

ELEKTRONIKA

dla wszystkich

nr 4/2024 (339) • kwiecień • www.elportal.pl

DIY PLUS
tylko dla prenumeratorów

Jednookładowy cyfrowy odbiornik radiowy FM/AM/SW

PROJEKTY dla elektroników

- ▶ Przetestuj do czterech serwo mechanizmów oddzielnych lub w systemie
- ▶ Programowany, hybrydowy zasilacz laboratoryjny ze sterowaniem Wi-Fi, część 2

DIY dla wszystkich

- ▶ Wyłącznik czasowy
- ▶ Alarm kuchenny z wykorzystaniem czujnika gazu MQ2
- ▶ Wydłużenie żywotności pompy wody wraz z oszczędnością energii

TUTORIALE

- ▶ Lutowanie SMD. Porady i wskazówki, sztuczki i triki
- ▶ Historia Uniwersalnego Złącza Szeregowego. Ekspansja USB
- ▶ Jak działa zasilanie USB-C
- ▶ Korzystanie z tanich modułów elektronicznych. Ładowarki USB
- ▶ Chirurgia obwodowa: Transformatory i LTSpice, część 2
- ▶ Audio OUT: Przedwzmacniacz mikrofonowy (do wokodera), część 1
- ▶ Ekscytacje Maxa: Migające diody LED i śliniacy się inżynierowie



**Super Tester
Serwomechanizmów**



EP.com.pl

Największy portal dla elektroników konstruktorów



**Król automatyki
jest w Tobie**

AutomatykaB2B.pl

FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

przełączniki
półprzewodniki
złącza
przekazniki
radiatory
obudowy
i wiele więcej...

www.piekarz.pl



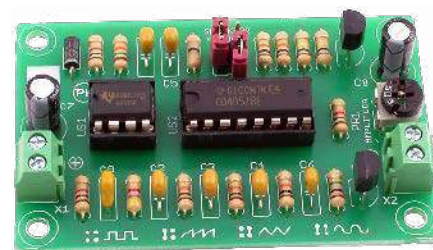
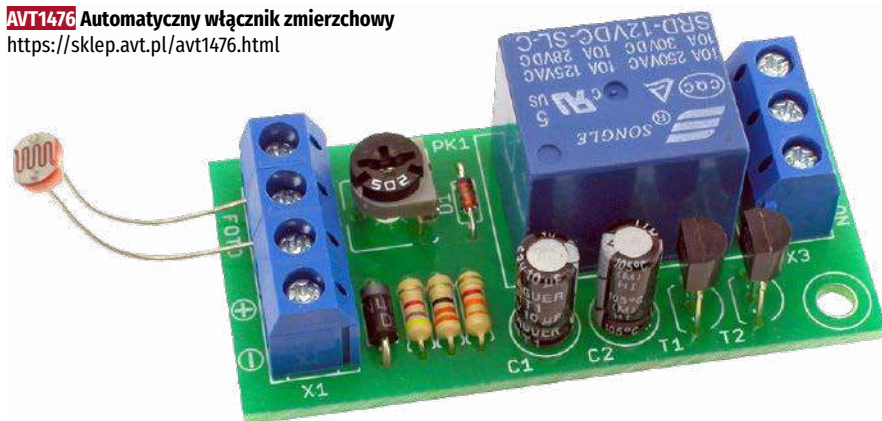


Najbardziej popularne kity AVT

Poznaj listę **TOP 100** na www.elportal.pl/kityavt

AVT1476 Automatyczny włącznik zmierzchowy

<https://sklep.avt.pl/avt1476.html>



AVT1327 Mini generator funkcyjny

<https://sklep.avt.pl/avt1327.html>



AVT735 Regulator mocy PWM 10 A

<https://sklep.avt.pl/avt735.html>



AVT5540 Radio FM z RDS

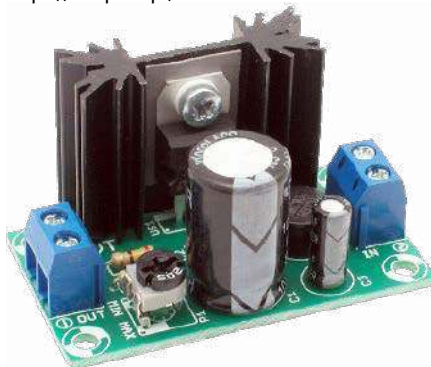
<https://sklep.avt.pl/avt5540.html>



AVT1597/3 Wzmacniacz audio z układem

TDA2050 35 W

<https://sklep.avt.pl/wzmacniacz-audio-z-ukladem-tda2050-zestaw-do-samodzielnego-montazu.html>



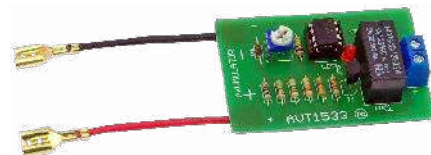
AVT1066 Miniaturowy zasilacz uniwersalny z LM317

<https://sklep.avt.pl/avt1066.html>



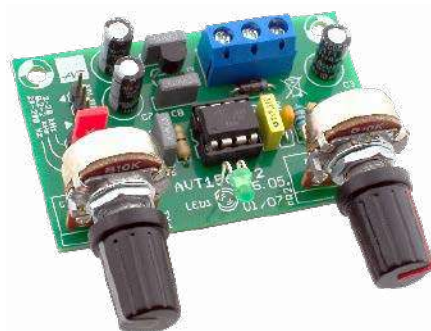
AVT1594 Wzmacniacz mocy 2x45 W z STK4182

<https://sklep.avt.pl/avt1594.html>



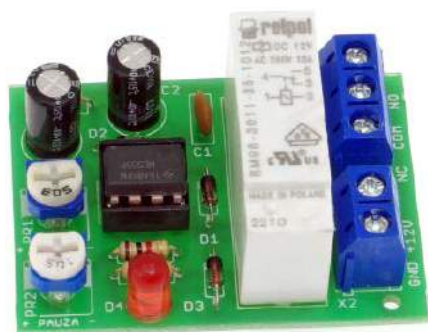
AVT1533 Zabezpieczenie akumulatora 12 V przed rozładowaniem

<https://sklep.avt.pl/avt1533.html>



AVT1569 Generator akustyczny 20 Hz...20 kHz

<https://sklep.avt.pl/avt1569.html>



AVT1459 Uniwersalny układ czasowy

<https://sklep.avt.pl/avt1459.html>



AVT1661 Elektroniczna kostka do gry

<https://sklep.avt.pl/avt1661.html>



Pełna oferta na: sklep.avt.pl

obejrzyj filmy na <https://www.youtube.com/@serwisAVT>

Wiosenne

-50%

na wybrane roczne
prenumeraty
drukowane



101,40 zł
12 wydań w roku



113,40 zł
12 wydań w roku



89,40 zł
12 wydań w roku

Zaprenumeruj wybrane czasopismo z rabatem aż 50%!

Promocja wiosenna dotyczy rocznych prenumerat drukowanych czasopism:

- Elektronika dla Wszystkich (12 wydań w roku) • Elektronika Praktyczna (12 wydań w roku) • Młody Technik (12 wydań w roku)

Zamów prenumeratę na www.UlubionyKiosk.pl/prenumerata lub poprzez dokonanie przelewu na konto AVT-Korporacja sp. z o.o., ul. Leszcynowa 11, 03-197 Warszawa, ING BANK ŚLĄSKI 18 1050 1012 1000 0024 3173 1013 (w tytule wpłaty podaj nazwę czasopisma).

Masz opłaconą bieżącą prenumeratę? Już teraz przedłuż ją z rabatem 50%.

Promocja trwa do 31.05.2024 i nie łączy się z innymi promocjami Wydawnictwa AVT.

Koszt wysyłki (list standardowy nierejestrowany) pod wskazany adres w Polsce ponosi wydawnictwo.

8



Projekty dla elektroników:

- Jednokładowy cyfrowy odbiornik radiowy FM/AM/SW wg Silicon Labs..... 8
- Przetestuj do czterech serwomechanizmów oddzielnych lub w systemie. Super Tester Serwomechanizmów..... 18
- Programowany, hybrydowy zasilacz laboratoryjny ze sterowaniem Wi-Fi, część 2..... 24

Tutoriale:

- Lutowanie SMD. Porady i wskazówki, sztuczki i triki..... 39

Wszystko o USB:

- Historia Uniwersalnego Złącza Szeregowego. Ekspansja USB..... 48
- Jak działa zasilanie USB-C 54
- Korzystanie z tanich modułów elektronicznych. Ładowarki USB 60

Edukacja w EdW dla szkół i uczelni:

- Wykład 17: Generatory sygnału na wzmacniaczach operacyjnych..... 64

Ekscytacje Maxa:

- Migające diody LED i śliniacy się inżynierowie, część 7 74
- Sprytne porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania .. 77

- Chirurgia obwodowa: Transformatory i LTspice, część 2..... 78

- Audio OUT: Przedwzmacniacz mikrofonowy (do wokodera), część 1..... 82

18



24

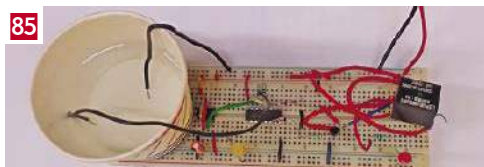


DIY dla wszystkich:

- Wydłużenie żywotności pompy wody wraz z oszczędnością energii..... 85
- Alarm kuchenny z wykorzystaniem czujnika gazu MQ2 87
- Wyłącznik czasowy 89

- PCBWay zmontuje Twoje płytki drukowane..... 36

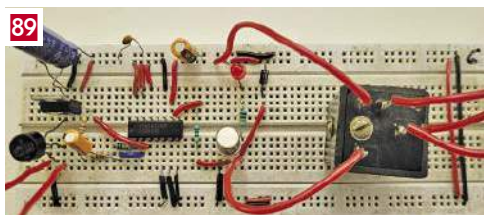
85



DIY PLUS

- 16-kanałowy sterownik serwomechanizmów RC z interfejsem I²C 91
- Sterowanie silnikiem DC za pomocą joysticka..... 91

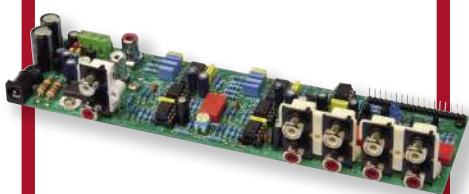
89



Rubryki stałe:

- Prenumerata 3
- Od wydawcy 5
- Poczt..... 6

A za miesiąc w majowym EdW



* Cyfrowy przedwzmacniacz z ekranem dotykowym i zdalnym sterowaniem, część 1

Przedwzmacniacz łączy w sobie najlepsze cechy techniki cyfrowej i analogowej, a więc cyfrowe sterowanie z intuicyjnym interfejsem wykorzystującym ekran dotykowy, możliwość zapisywania własnych ustawień, zdalne sterowanie, a także niski poziom szumów i zniekształceń obwodów analogowych. Zastosowano klasyczne obwody regulacji głośności i tonów ze wzmacniaczami operacyjnymi wysokiej jakości i potencjometrami cyfrowymi.

* Autotransformatorowy regulator napięcia sieciowego

To urządzenie (variac) zostało zbudowane w celu uzyskania stabilizowanego źródła zasilania 115 V AC, ale można je łatwo dostosować do zapewnienia stabilizowanego napięcia 220...240 V AC do zasilania dowolnego sprzętu sieciowego, na przykład starych urządzeń lampowych wymagających zasilania z sieci 220 V.

* Sterowany napięciem filtr syntezatora 24 dB/okt. z użyciem fotosyntezatorów

Prosty układ przestrzajanego filtra, niewymagającego użycia żadnych specjalnych układów scalonych, a jedynie dwóch zwykłych wzmacniaczy operacyjnych, które mogą być zasilane z pojedynczego źródła 5 V. Pomimo swej prostoty filtr oferuje bardzo niskie zniekształcenia (THD < 0,01%), przetwarzanie sygnałów o dużej amplitudzie i szeroki zakres dynamiki (> 90 dB).

* Tester kabli USB

Gdy masz szufladę pełną kabli USB, a część z nich nie działa prawidłowo, to używanie tych kabli staje się frustrujące. To urządzenie pozwoli Ci zaprowadzić porządek w tej sprawie.

* Plus kolejna porcja intrygujących projektów DIY.

* Plus wiele artykułów w Twoich ulubionych cyklach Tutoriali.

**W kioskach
od 30 kwietnia**

Komputery analogowe jak historia koła zatacza

Temat wykładu w tym wydaniu EdW, dotyczącego generatorów na bazie wzmacniaczy operacyjnych, pobudził mnie do wspomnień i refleksji historycznej. Termin „wzmacniacz operacyjny” pojawił się w latach czterdziestych minionego wieku w związku z pracami nad komputerami analogowymi na lampach elektronowych. Powstała wówczas koncepcja wzmacniacza, który poprzez podłączenie odpowiednich elementów zewnętrznych może wykonywać operacje algebraiczne oraz całkowanie i różniczkowanie.

Pojęcie komputer analogowy w szerokim znaczeniu można wywodzić od czasów starożytnych Greków, gdy nie było jeszcze elektroniki, ale zmyślnie konstrukcje mechaniczne potrafiły nawet symulować ruchy planet, księżyca i słońca. Swego rodzaju mechanicznym komputerem analogowym był też suwak logarytmiczny, znany od czasów Newtona, z którym jeszcze 50 lat temu nie rozstawał się żaden inżynier. Jednak dopiero elektronika, najpierw lampowa, a potem tranzystorowa, otworzyła drogę do rozwoju komputerów rozwiązujących równania różniczkowe z zastosowaniem wzmacniaczy operacyjnych całkujących i różniczkujących bez przetwarzania wielkości analogowych w zera i jedynki.

W latach pięćdziesiątych i sześćdziesiątych możliwości obliczeniowe komputerów cyfrowych były zbyt skromne dla ich efektywnego zastosowania do wielu obliczeń naukowych, opartych na rozwiązywaniu równań różniczkowych. Szczególnie w krajach rozwijających atomistykę (USA, Francja, Wielka Brytania, ZSRR) powstawały coraz doskonalsze konstrukcje komputerów analogowych, w których w czasie rzeczywistym, poprzez odpowiednią obróbkę przebiegu analogowego uzyskiwano rozwiązania zagadnień, na które ówczesne komputery cyfrowe potrzebowałyby kilkudziesięciu godzin.

Zaskakująco wysoki poziom osiągnięto w tej dziedzinie w Polsce, a ściślej w Wojskowej Akademii Technicznej. Komendant WAT, generał Sylwester Kaliski, realizował ambitne prace naukowe w zakresie fuzji termojądrowej metodą „fokus plazma”, których teoria opiera się na rozwiązaniach nieliniowych równań matematyki fizycznej, nie mających rozwiązań analitycznych. Do rozwiązywania tych równań doskonale nadają się komputery analogowe, dlatego świetny zespół konstruktorów, pod kierownictwem charyzmatycznego pułkownika Józefa Kapicy, tworzył w WAT nowoczesne wersje tych komputerów. Epizodycznie brałem udział w pracach tego zespołu.

W roku 1967, opromieniony lokalną sławą twórcy pierwszego w kraju pulsometru tranzystorowego EKG, zostałem zaproszony przez pułkownika Kapicę do jego zespołu. Moje zadanie dotyczyło opracowania układu serwomechanizmu do wstępnych nastaw potencjometrów wieloobrotowych. Otóż komputer analogowy miał ze sto, albo i więcej precyzyjnych potencjometrów wieloobrotowych, których nastawy określały parametry początkowe określonego zadania. Pulpit sterowniczy z setkami gałek nie robił najlepszego wrażenia. Lepszym rozwiązaniem było nastawianie tych potencjometrów za pomocą podłączanego do nich kolejno jednego serwomechanizmu – zamiast setek gałek była tylko jedna. Skonstruowałem prototyp takiego serwomechanizmu na tranzystorach germanowych, które silnie zmieniały parametry w funkcji temperatury i pierwsza demonstracja działania prototypu obnażyła ten feler. Absolutnie charyzmatyczny pułkownik Kapica ocenił krótko moją konstrukcję: „to całkiem dobrze działa, a teraz weź tę płytkę i wyrzuć ją w krówskie gówno”. To była niezła szkoła. Dopiero zastosowanie wzmacniacza operacyjnego (prototyp wykonany w Instytucie Tele-Radiotechnicznym w technologii hybrydowej, bo wzmacniacz $\mu A 741$ miał się pojawić w USA dopiero za rok) dało konstrukcję o zadowalającym działaniu.

Trzeba jednak zauważyć, że dokładność rozwiązań otrzymywanych przy użyciu komputerów analogowych pozostawia wiele do życzenia. Błędy sięgają kilku do kilkunastu procent, dlatego z czasem, gdy zdolności obliczeniowe komputerów cyfrowych wzrosły o kilka rzędów wielkości, zaniechano prac nad komputerami analogowymi. Okazuje się, że nie na zawsze. W ostatnich latach pojawiło się zainteresowanie procesorami analogowymi do wspomagania sztucznej inteligencji. Pracuje nad tym firma Mythic i IBM.

Okazuje się, że w rozwiązaniach zagadnień metodą iteracyjnego zbliżania się do zbieżnego wyniku, można znakomicie przyspieszać rozwiązywanie, jeśli proces iteracji rozpoczyna się od dobrze „odgadniętego” rozwiązania. Zatem przebieg rozwiązania jest dwuetapowy – najpierw komputer analogowy znajduje rozwiązanie z błędem kilku do kilkunastu procent, które staje się „odgadniętym” rozwiązaniem początkowym dla procesu iteracyjnego wykonywanego w drugim etapie przez komputer cyfrowy. W jakimś sensie można powiedzieć, że historia zataczyła koło.

Wiesław Marciniak

W rubryce „Począ” zamieszczamy fragmenty listów od Czytelników. Szczególnie chętnie publikujemy komentarze do artykułów w bieżących wydaniach EdW oraz propozycje tematów artykułów, zadań i quizów.

Pomoc ze źródła

Czytelnicy EdW niekiedy ulegają stereotypowym przekonaniom, że łatwość kontaktu z krajowym autorem projektu jest zasadniczą przewagą publikacji krajowych na artykułami zagranicznymi. Z naszej wieloletniej praktyki wiemy, że z krajowymi autorami bywa różnie, natomiast zawsze możemy liczyć na pomoc naszych redakcji źródłowych. Dla ilustracji jak to działa, publikujemy przykładowy list naszego Czytelnika do redakcji EdW i odpowiedź redakcji Silicon Chip na pytania Czytelnika przesłane przez nas do SC.

[...]Szczególnie zainteresował mnie projekt generatora AM/FM opublikowany w ramach licencji Silicon Chip, w numerach 7/2022 i 8/2022. Do własnych celów postanowiłem złożyć to urządzenie.

Posiadam zakupione w Silicon Chip Australia płytkę PCB oraz zaprogramowany mikroprocesor. Jestem w trakcie montażu elementów, ale niestety brakuje mi bardzo istotnej informacji, stąd ta wiadomość z prośbą o pomoc.

W opisie jest podane jak należy nawinąć autotransformator na rdzeniu typu „nosek świnki” (balun). Z powodu braku informacji o spodziewanej wartości indukcyjności muszę rozwiązać niejednoznaczności.

W artykule jest mowa o nawinięciu bifilarnych zwojów drutem o średnicy 0,315 mm, ale załączony rysunek z oryginalnym opisem po angielsku pokazuje to nawinięcie drutem 0,15 mm. Dla parametrów indukcyjności i impedancji to znacząca różnica.

Dodatkowo, co do wielkości tego rdzenia mowa jest o „długości” (long) wynoszącej 7 mm. Taki też rdzeń udało mi się zamówić z granicą, bo w Polsce nigdzie nie znalazłem tego typu rdzeni. Niestety wg nadruku i rozstawu wyprowadzeń, rdzeń o długości 7 mm jest dwukrotnie za mały. Być może to, co autor miał na myśli, to wysokość tego rdzenia, ale wówczas jego długość to ok. 15 mm (tak aby pasował na obrys na PCB oraz mniej więcej z tego co widać na zamieszczonym zdjęciu w artykule).

Takie różnice co do parametrów autotransformatora noskowego muszą być wyjaśnione.

Z poważaniem
A.W.

Ten list po kilku dniach (przebywałem na urlopie i nie czytałem pocztą do 10.03. – W.M.) wysłał mi redakcja Silicon Chip i następnego dnia otrzymaliśmy odpowiedź:

Hello Wiesław,

This is at least partially answered by the Notes & Errata available from our website (under 2019 as that is when we published it): <https://www.siliconchip.com.au/Articles/Errata>.

AM/FM/CW HF/VHF RF Signal Generator, June-July 2019: (1) The second article describes the core for transformer T1 as 7mm long in some places and 14mm long in others. It should be 7mm, while a 14mm core will work, it's harder to fit.

He told me that he used 0.15mm wire for winding this but he says that it will also work with 0.315mm wire.

Nicholas Vinen

Okazuje się, że w tym samym czasie Autor listu skontaktował się niezależnie z redakcją Silicon Chip i przysłał nam następującą korespondencję:

Dziękuję za informację, ale czekając nieco długo, skontaktowałem się osobiście z redakcją Silicon Chip Magazine i otrzymałem odpowiedź po jednym dniu, dodatkowo mam kontakt z autorem rozwiązania i jestem po wymianie z nim szeregu informacji.

On opublikował poprawki na swojej stronie dotyczące rezystora (akurat w EdW było to już na schemacie uwzględnione) z 1 kΩ na 10 kΩ oraz encoderów.

Ja mogę dodać od siebie, że stała czasowa nawet przy zmianie rezystora była na tyle mała, że nie dało się uzyskać efektu wyłączenia układu „miękiego załączania”. Pomogło dopiero zwiększenie pojemności kondensatora emiterowego z 1 μF na 3,3 μF (4,7 μF też jest OK.). Uwagi dotyczące starego typu wyświetlaczy są ok, ale teraz to pewnie każdy chętniej skorzysta z tanich i lepszych LCD 1602 (przynajmniej, że jeszcze nie montowałem wyświetlacza, ale mając na uwadze niepewność, sprawdź jeszcze czy i tu nie ma błędnych połączeń na druku). Biorąc pod uwagę obudowę i drobne przeróbki, elementy typu przełączniki suwakowe, enkoder, potencjometr i wyświetlacz będą montować na samym końcu żeby by odpowiednio wyzozonować do płyty czołowej obudowy (zamierzam też wbudować zasilacz wewnątrz urządzenia, ekranowany blachą stalową, aby ograniczyć wpływ transformatora – nie chcę stosować impulsowego zasilania). Jak widać, nie uruchomiłem jeszcze generatora, czekam na dostawę ferrytu do tego autotransformatora. W artykule błędnie podano zarówno jego wymiar jak i średnicę drutu (inaczej w treści i inaczej na oryginalnym rysunku). Biorąc pod uwagę obrys nadrukowany na płytce (PCB kupiłem w Silicon Chip), wymiar ferrytu to ok. 14 mm długości i 7 mm wysokości, a drut o średnicy 0,15 mm (przy czym Andrew Woodfield potwierdził, że 0,315 mm też powinien być ok).

A więc warto sprostować dodatkowo te informacje również.

Z poważaniem
A.W.

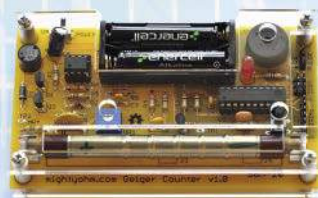
Patronat AVT

Poniżej prezentujemy listę szkół biorących udział w programie PATRONAT AVT, który jest całkowicie bezpłatny, a szkoły objęte tym patronatem korzystają z różnych benefitów, takich jak bezpłatne prenumeraty, darmowe pakiety próbne kitów AVT, itp. Szkoły, które dopiero teraz dowiadują się o naszej akcji PATRONAT AVT, prosimy o przeczytanie listu w EdW 09/2022 (wydanie dostępne na www.ulubionykiosk.pl) i zgłoszenie akcesu do PATRONATU AVT. Zgłoszenia prosimy wysłać na adres: prenumerata@avt.pl.

- Centrum Edukacji Zawodowej, 82-200 Malbork, De Gaulle'a 75a
- Centrum Edukacji Zawodowej i Biznesu, 66-400 Gorzów Wielkopolski, Pomorska 67
- Gminny Zespół Szkolno-Przedszkolny nr 4 w Więckach, 42-110 Popów, Więcki, Szkolna 1
- Górnośląskie Centrum Edukacyjne im. Marii Skłodowskiej-Curie w Gliwicach, 44-100 Gliwice, Okrzei 20
- Noworudzka Szkoła Techniczna w Nowej Rudzie, 57-401 Nowa Ruda, Stara Droga 4
- Regionalne Centrum Edukacji Zawodowej w Biłgoraju, 23-400 Biłgoraj, Kościuszki 98
- Regionalne Centrum Edukacji Zawodowej w Lubartowie, 21-100 Lubartów, 1 Maja 82
- Technikum nr 4 im. Marii Skłodowskiej-Curie, 41-902 Bytom, Katowicka 35
- Zespół Placówek Edukacyjno-Wychowawczych w Gołdapi, 19-500 Gołdap, Wojska Polskiego 18
- Zespół Placówek Oświatowych w Rudniku, 32-440 Sułkowice, Rudnik, Szkolna 55
- Zespół Szkolno-Przedszkolny nr 2 w Wiśle, 43-460 Wiśła, Malinka 53
- Zespół Szkolno-Przedszkolny nr 3 w Gliwicach, 44-122 Gliwice, Żwirki i Wigury 85
- Zespół Szkolno-Przedszkolny nr 4 w Rybniku, 44-207 Rybnik, Komisji Edukacji Narodowej 29
- Zespół Szkolno-Przedszkolny w Choceniu, 87-850 Chocień, Sikorskiego 12
- Zespół Szkolno-Przedszkolny w Ostrożnicy, 47-280 Pawłowiczki, Ostrożnica, Kościelna 42
- Zespół Szkół Budowlano-Elektrycznych im. Jana III Sobieskiego w Świdnicy, 58-100 Świdnica Śląska, Wałbrzyska 35-37
- Zespół Szkół Centrum Kształcenia Ustawicznego w Gronowie, 87-162 Lubicz Dolny, Gronowo 128
- Zespół Szkół Elektronicznych i Telekomunikacyjnych w Olsztynie, 10-144 Olsztyn, Bałtycka 37a
- Zespół Szkół Elektronicznych im. I. Domeyki w Bolesławcu, 59-700 Bolesławiec, Tyran-kiewiczów 2
- Zespół Szkół Elektronicznych w Rzeszowie, 35-078 Rzeszów, Hetmańska 120
- Zespół Szkół Elektronicznych, Elektrycznych i Mechanicznych, 43-300 Bielsko-Biała, Słowackiego 24
- Zespół Szkół Elektrycznych nr 2 w Krakowie, 31-977 Kraków, Os. Szkolne 26
- Zespół Szkół Elektrycznych w Kielcach, 25-317 Kielce, Kaczorowskiego 8
- Zespół Szkół im. Bolesława Prusa, 42-207 Częstochowa, Prusa 20
- Zespół Szkół im. ks. dra Jana Zwierza w Ropczycach, 39-100 Ropczyce, Mickiewicza 14
- Zespół Szkół im. Ks. Stanisława Staszica, 39-400 Tarnobrzeg, Kopernika 1
- Zespół Szkół nr 1 w Przysietnicy, 36-200 Brzozów, Przysietnica 198
- Zespół Szkół nr 10 im. Prof. Janusza Groszkowskiego w Zabrze, 41-807 Zabrze, Chopina 26
- Techniczne Zakłady Naukowe w Dąbrowie
- Górnicej, 41-300 Dąbrowa Górnicza, Zawidzkiej 10
- Zespół Szkół nr 2 im. Eugeniusza Kwiatkowskiego w Dębicy, 39-200 Dębica, Lisa 2
- Zespół Szkół nr 2 im. Gen. Józefa Bema, 05-822 Milanówek, Wójtowska 3
- Zespół Szkół nr 2 im. Ks. Prof. Józefa Tischnera w Żorach, 44-240 Żory, Boryńska 2
- Zespół Szkół nr 2 w Pabianicach im. prof. Janusza Groszkowskiego, 95-200 Pabianice, św. Jana 27
- Zespół Szkół nr 4 w Nowym Sączu, 33-300 Nowy Sącz, św. Ducha 6
- Zespół Szkół nr 40 im. Stefana Starzyńskiego, 03-771 Warszawa, Objazdowa 3
- Zespół Szkół Politechnicznych im. Bohaterów Monte Cassino we Wrześni, 62-300 Września, Wojska Polskiego 1
- Zespół Szkół Ponadgimnazjalnych nr 1 w Jarocinie, 63-200 Jarocin, Franciszkańska 1
- Zespół Szkół Ponadpodstawowych nr 2 im. E. Kwiatkowskiego w Jarocinie, 63-200 Jarocin, Franciszkańska 2
- Zespół Szkół Ponadpodstawowych nr 3 im. Armii Krajowej w Zamościu, 22-400 Zamość, Zamoyskiego 62
- Zespół Szkół Powiatowych im. Stanisława Staszica w Opocznie, 26-300 Opoczno, Kossaka 1a
- Zespół Szkół Publicznych w Szewnie, 27-400 Ostrowiec Świętokrzyski, Szewna, Langiewicza 3
- Zespół Szkół Spożywczych i Hotelarskich w Radomiu, 26-600 Radom, św. Brata Alberta 1
- Zespół Szkół Techniczno-Informatycznych w Elblągu, 82-300 Elbląg, Ryckarska 2
- Zespół Szkół Technicznych i Licealnych w Piechowicach, 58-573 Piechowice, Przemysłowa 21
- Zespół Szkół Technicznych i Ogólnokształcących nr 3 im. Ł. Abramowskiego, 40-659 Katowice, Harcerzy Września 1939 2
- Zespół Szkół Technicznych im. Armii Krajowej w Skarżysku-Kamiennie, 26-110 Skarżysko-Kamienna, Tysiąclecia 22
- Zespół Szkół Technicznych im. Ignacego Mościckiego w Tarnowie, 33-101 Tarnów, E. Kwiatkowskiego 17
- Zespół Szkół Technicznych w Kolbuszowej, 36-100 Kolbuszowa, Bytnara 2
- Zespół Szkół w Błażowej, 36-030 Błażowa, Kowala 3
- Zespół Szkół w Gościńcu, 78-120 Gościńcu, Kościuszki 5
- Zespół Szkół w Zarzeczu, 37-205 Zarzecze, św. Jana Pawła II 7
- Zespół Szkół Zawodowych nr 1 im. gen. F. Kleeberga w Dęblinie, 08-530 Dęblin, Tysiąclecia 3
- Zespół Szkół Samochodowych im. inż. Tadeusza Tańskiego, 33-300 Nowy Sącz, Rejtana 18a
- Szkoła Podstawowa im. Rodzimych Bohaterów II Wojny Światowej w Zatałkowie, 83-342 Kamienica Królewska, Zatałkowo 6

Elektor Bestsellers

SAVE UP TO
26% NOW!



www.elektor.com/sale/deals





Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/dm44c>
Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania

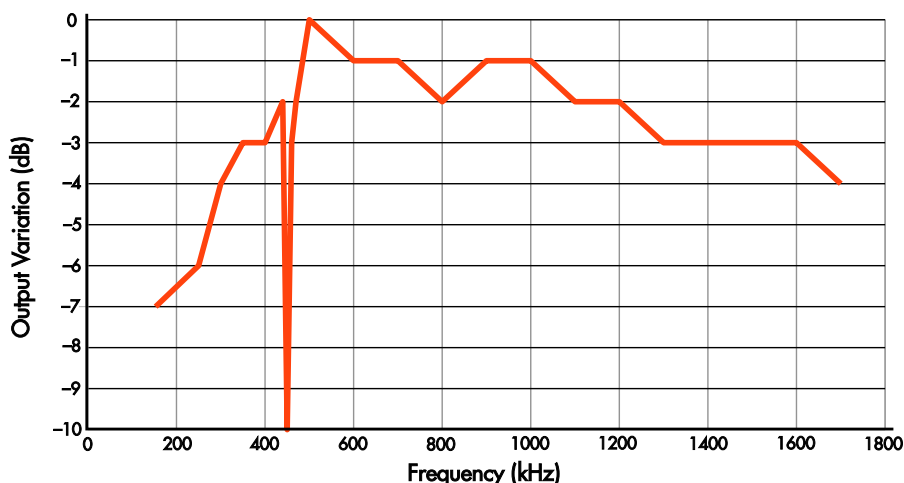


Jednoulkowy cyfrowy odbiornik radiowy FM/AM/SW wg Silicon Labs

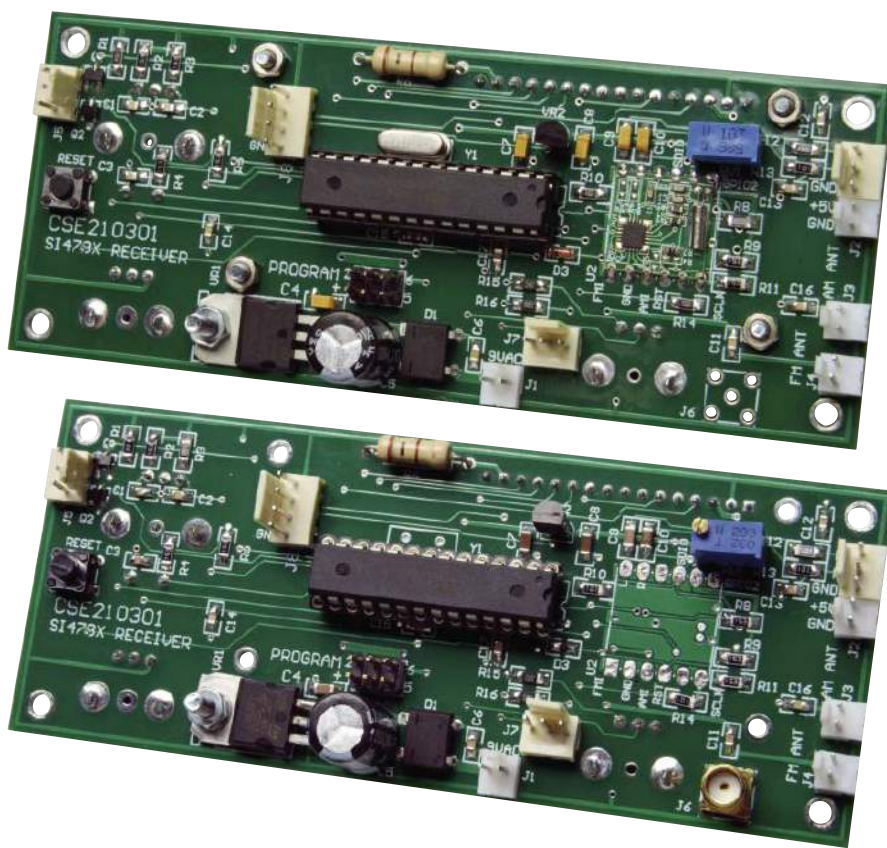
Najnowsza technologia umożliwia budowanie radioodbiorników FM/AM w oparciu o pojedyncze, dedykowane układy scalone. Wszystko, co musisz zrobić, to podłączyć kilka anten do wejść takiego układu, wystać mu kilka poleceń przez port szeregowy, a na drugim końcu pojawi się dźwięk stereo. W rezultacie układy Silicon Labs odpowiadające takiej koncepcji sprawiają, że zbudowanie sprawnego odbiornika radiowego jest dziecinnie proste. Taki radioodbiornik jest łatwy w konfiguracji i obsłudze, mieści się w kompaktowej obudowie i działa przy zastosowaniu zwykłego zasilacza sieciowego.

Autor był w miarę zadowolony ze swojego projektu odbiornika AM/FM/SW pokazanego w wydaniu Silicon Chip-a ze stycznia 2021 r. (siliconchip.com.au/Article/14704), przynajmniej pod względem łatwości budowy, łatwości użytkowania i pokrycia wielu pasm radiowych. W wersji polskiej tekst ukazał się w wydaniu EdW z września 2023 r. Ale projekt pozostawił po sobie niedosyt, ze względu na wrażenie, że jego ogólna atrakcyjność pozostawia trochę do życzenia. Autor nie był również zadowolony z faktu, że nie miał wystarczającej ilości informacji pozwalających na wdrożenie pełnego cyfrowego sterowania układem radiowym BK1198.

Chociaż konstrukcja tego radia była stosunkowo prosta, byłoby to o wiele prostsze, gdyby można było uzyskać działające sterowanie cyfrowe.



Rysunek 1. czułość radia w poszerzonym paśmie AM, od 153 kHz do 1,7 MHz. Z wyjątkiem spadku w okolicach 445...455 kHz (typowe częstotliwości pośrednie), wynik jest dość płaski. W standardowym paśmie nadawania AM 550...1720 kHz, różnica wynosi tylko około 4 dB



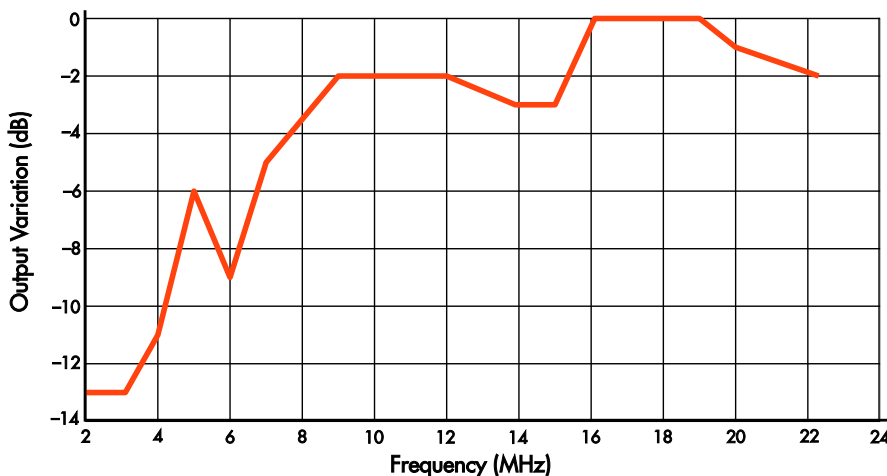
Te dwa zdjęcia pokazują, że górna część płytki drukowanej dla wersji z modułem Si4730 (góra) tego projektu niewiele różni się od wersji Si4732 (dół). Zignoruj dodatkowe śruby/nakrętki na górnej fotografii, ponieważ służą one tylko do montażu wyświetlacza

W ciągu ostatnich paru lat pojawiło się kilka nowych układów, które znacznie ułatwiają projektowanie odbiorników radiowych. Wiele z nich pochodzi z firmy Silicon Labs; istnieje około 34 odmian układów z rodziny Si473x, do których notę katalogową można pobrać ze strony siliconchip.com.au/link/ab7y.

Mają one podobną architekturę do układu BK1198, użytego w projekcie Silicon Chip-a ze stycznia 2021 roku. Jedną

z głównych zalet układów Silicon Labs jest dokumentacja; podczas gdy informacje na temat BK1198 są co najmniej skąpe, nota aplikacyjna dla układów SiLabs liczy 321 stron! (Patrz siliconchip.com.au/link/ab7z).

Płytką, zaprojektowaną w Silicon Chip-ie, jest odpowiednia dla wstępnie zmontowanego modułu z układem Si4730 lub samodzielnego układu Si4732. Oba są dostępne (patrz uwagi dalej!) na AliExpress w dość niskich



Rysunek 2. Podobny wykres „odpowiedzi częstotliwościowej” dla zakresu SW od 2 MHz do 22,3 MHz

cenach. Si4730 obsługuje tylko standardowe pasma AM i FM, podczas gdy Si4732 można zaprogramować tak, aby obejmował odbiór fal długich i krótkich. Oba mogą dekodować FM stereo. Specyfikacje podają następujące pasma:

- Obsługa pasm FM całego świata: 64...108 MHz
- Obsługa pasma AM w pełnym zakresie: 520...1710 kHz
- Obsługa pasma SW (Si4734/32/35): 2,3...26,1 MHz
- Obsługa pasma LW (Si4734/32/35): 153...279 kHz

Ale co z przerwaniami pomiędzy pasmami? Odrobina eksperymentów pozwoliła ustawić częstotliwości odbioru w tych lukach. I co za niespodzianka; dzięki układowi Si4732 można wybrać dowolną częstotliwość od 153 kHz do 30 MHz, wysyłając odpowiedni kod do układu. Żadnych luk! To, czy w tych lukach jest coś interesującego, to inna sprawa.

W rezultacie mamy zaprogramowane pasmo AM od 153 kHz do 1730 kHz, a pasmo SW od 2 MHz do 30 MHz.

Efektywność

Dla pasma FM, krótki kawałek drutu wewnątrz obudowy pozwalał na odbiór większości stacji w Melbourne z dobrym poziomem sygnału do zakłóceń (SNR – Signal Noise Ratio).

Z zewnętrzną długą anteną przewodową podłączoną bezpośrednio do wejścia anteny AM, Autor mógł uzyskać odbiór wielu stacji z SNR 25 dB lub lepszym bez żadnej anteny ferrytowej. W ten sposób w obwodzie odbiornika nie jest wymagana ani jedna cewka indukcyjna! Używając pręta ferrytowego, słabsze stacje również były odbierane, ale razem z dużą ilością zakłóceń powodowanych przez całą elektronikę w laboratorium Autora.

Pokazany na **rysunku 1** wykres przedstawia czułość w paśmie AM od 153 kHz do 1700 kHz. Zwraca uwagę ostry spadek czułości przy 450 kHz. Trudno zgadnąć, dlaczego tak się dzieje, ale występuje to w pobliżu częstotliwości pośredniej większości odbiorników superheterodynowych, więc nie ma to znaczenia.

Na falach krótkich czułość jest porównywalna z pasmem AM (**rysunek 2**). Nie jest to doskonały wynik, ale można go uznać za wystarczający. Podczas pomiarów było trochę „ćwierkania” na niektórych częstotliwościach, np. 8 MHz, 14 MHz i 16 MHz, które utrudniały pomiar SNR. Powyżej 22 MHz, wyświetlacz SNR raczej nie dawał sensownych odczytów, chociaż efektywność aż do 30 MHz wydawała się taka sama jak przy 20 MHz.

Nota katalogowa jest dosyć uboga w informacje na temat parametrów wyjścia audio. Eksperymentalnie udało się ustalić,

że minimalna rezystancja obciążenia na wyjściu słuchawkowym wynosi 1,6 kΩ. Jeśli będzie mniejsza, nastąpi przesterowanie.

Maksymalny poziom napięcia wyjściowego przy takim obciążeniu wynosi dla fali sinusoidalnej 250 mV międzyszczytowo lub około 88 mV RMS, co daje mniej niż 1 mW mocy wyjściowej. Układ działa ze słuchawkami o niskiej impedancji, chociaż przy maksymalnej głośności wystąpią pewne zniekształcenia. Słuchawki Sennheiser o impedancji 60 Ω zapewniały w ciłym otoczeniu akceptowalny poziom odsłuchu.

Słuchawki Panasonic z redukcją szumów o impedancji wejściowej 330 Ω (z włączoną redukcją szumów) dawały znacznie wyższy poziom dźwięku. Podanie sygnału do zewnętrznych głośników ze wzmacniaczem dało dźwięk dobrej jakości.

Z powodu tak słabego sygnału wyjściowego, Autor zdecydował się dodać bufor w postaci wzmacniacza operacyjnego, który umożliwia wysterowanie słuchawek o niskiej impedancji, zapewniając jednocześnie wystarczający poziom napięcia dla niskieffektywnych słuchawek o wysokiej impedancji. Jest to również przydatne w przypadku podawania dźwięku do przedwzmacniacza lub wzmacniacza, ponieważ sygnał jest bliższy typowemu „poziomowi liniowemu”.

Gdy pokrętko strojenia jest obracane, każdy impuls wału enkodera wysyła sześć bajtów przez szynę I²C, a następnie odbiera siedem bajtów stanu. Zajmuje to dużo czasu, więc zbyt szybkie obracanie pokrętką strojenia powoduje pominięcie impulsów enkodera i niewielką zmianę częstotliwości. Wystarczy spowolnić obroty.

Opis układu

Pełny schemat ideowy pokazano na rysunku 3. Moduł Si4730 zawiera rezonator kwarcowy 32 768 kHz i powiązane kondensatory. Antena FM jest podłączona do wejścia modułu FM za pomocą szeregowego kondensatora 1 nF, podczas gdy pasmo AM wymaga pręta ferrytowego, zwykle z cewką o indukcyjności 400 μH.

Opcjonalny kondensator 10 nF łączy dwa wejścia antenowe, umożliwiając użycie pojedynczego przewodu antenowego zapewniającego w obszarach metropolitalnych odbiór zarówno FM, jak i AM. Linia SEN jest powiązana wewnętrznie z modułem Si4730.

Wyjście audio jest podłączone do złącza CON4. Poziom wyjściowy z samego układu radiowego jest wystarczający do zasilania słuchawek 60 Ω. Jak wspomniano powyżej, w zależności od słuchawek, poziom głośności może być nieco niski, a zniekształcenia mogą być wyższe niż byśmy chcieli.



Wersja Si4732 jest nieco odmienna ze względu na instalację dwóch kondensatorów 22 pF, kwarcu (X2) i samego układu na spodzie płytki drukowanej

Podwójny wzmacniacz operacyjny (IC3) obecny w ostatecznej wersji, nie był montowany w pokazanych na fotografiach prototypach. Wzmacniacz operacyjny daje czterokrotne wzmocnienie napięciowe i wyjście o niskiej impedancji, wystarczające do wysterowania prawie wszystkich słuchawek nausznych lub dousznych do przyzwoitego poziomu głośności (nawet w przypadku niskieffektywnych typów), a być może nawet bardzo efektywnych głośników biernych.

Pragniemy jednak uspokoić Czytelników – płytka PCB zaprojektowana jest pod wzmacniacz operacyjny i tylko od Ciebie zależy, czy użyjesz go w swojej konstrukcji.

Alternatywnie można użyć zewnętrznego wzmacniacza audio, takiego jak głośniki komputerowe, z lub bez wzmacniacza operacyjnego na płytce radioodbiornika. Jeśli nie potrzebujesz wzmacniacza operacyjnego, możesz po prostu zmostkować pary styków 1/3 i 5/7, aby doprowadzić wyjście układu radiowego do złącza CON4.

CON4 ma również styki +5 V i GND. Może to być wykorzystane do zasilania małego modułu wzmacniacza zamontowanego w tej samej obudowie, z głośnikami o impedancji 8 Ω. Nie polecamy w tym miejscu wzmacniaczy klasy D, ponieważ mogą one generować impulsy, które będą zakłócać odbiór radiowy, podobnie jak zasilacz impulsowy.

Sterowanie odbywa się za pomocą standardowej magistrali szeregowej I²C i linii reset. Zastosowano układ mikrokontrolera ATmega328P z pamięcią 32 KB w obudowie DIL, chociaż w prototypie Autora użyto ATmega168 z pamięcią o rozmiarze 16 KB; program zajmuje tylko 68% z 16 KB pamięci FLASH, a w przyszłowiowych szufladach elektroników leży mnóstwo tych układów z poprzednich projektów. Poza rozmiarem pamięci FLASH, są one zasadniczo identyczne.

Wyświetlacz to standardowy alfanumeryczny moduł LCD 16×2. Istnieje możliwość podłączenia zewnętrznego rezonatora kwarcowego dla układu ATmega, ale zdaniem Autora, wewnętrzny oscylator RC 8 MHz jest zupełnie wystarczający.

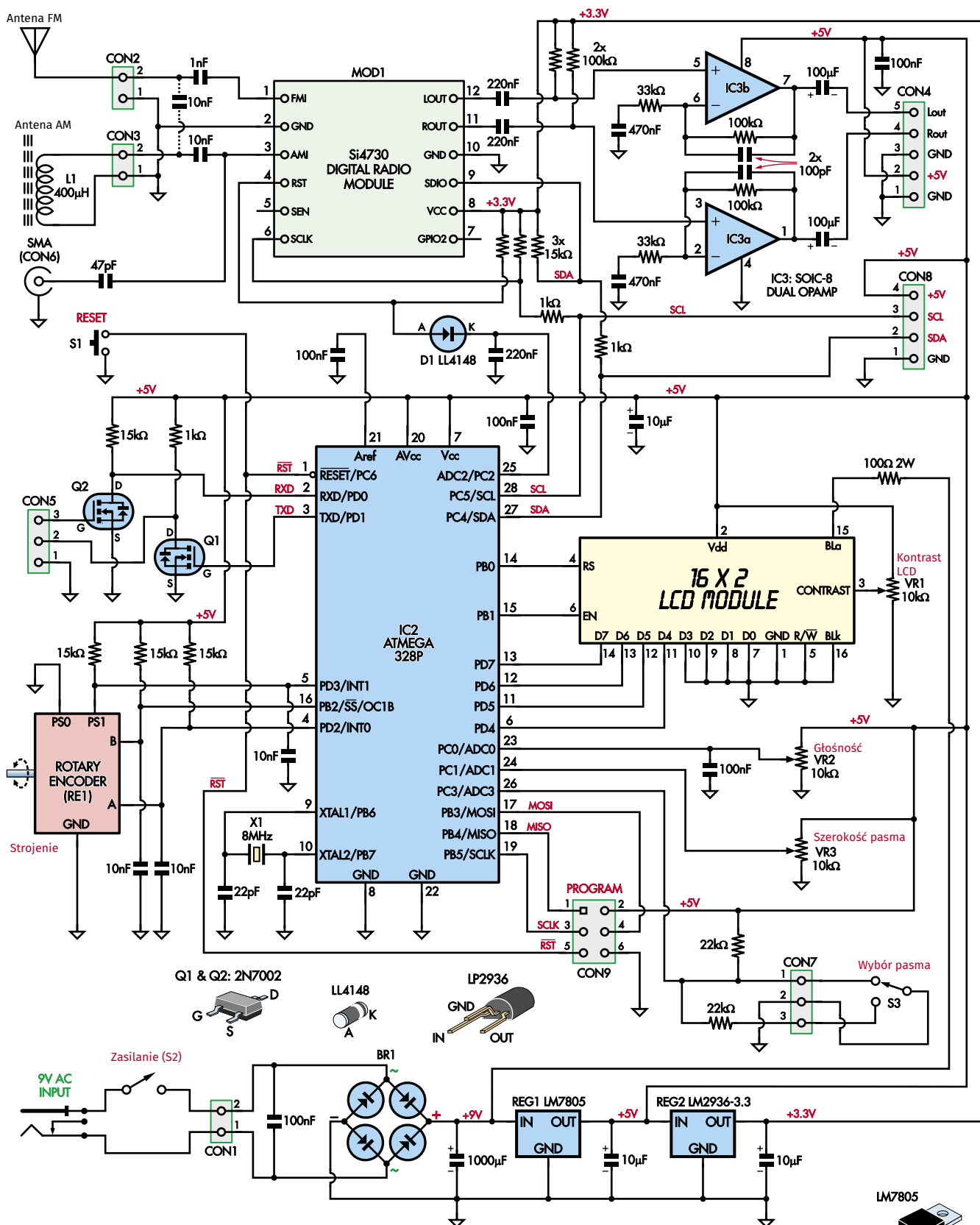
Procesor zasilany jest napięciem 5 V, podczas gdy układ SiLabs wymaga napięcia 3,3 V. Nie stanowi to problemu dla interfejsu I²C, ponieważ wyjście jest typu „otwarty kolektor” (open-drain), a rezystory polaryzujące o rezystancji 15 kΩ ustawiają na szynach napięcie 3,3 V. Istnieją również dwa szeregowo rezystory ograniczające prąd o rezystancji 1 kΩ między wyjściami I²C mikrokontrolera i wejściami modułu radiowego jako zabezpieczenie przed nieprawidłowym zaprogramowaniem linii I²C.

Typowa wartość rezystora ustawiającego napięcie na szynie I²C wynosi 4,7 kΩ, ale linie SCL i SDA w układzie SiLabs mają ograniczone możliwości wysterowania. Praca z rezystorami polaryzującymi o rezystancji 4,7 kΩ może uniemożliwić sterowanie szynami I²C przez mikrokontroler, zwłaszcza biorąc pod uwagę szeregowo rezystory zabezpieczające o rezystancji 1 kΩ.

Stąd zastosowanie rezystorów polaryzujących 15 kΩ; niższe wartości dałyby krainowo niskie napięcie na obu szynach przy zewnętrznej polaryzacji przez rezystory 1 kΩ. Nie wystąpiły żadne problemy z rezystorami polaryzującymi o wyższej wartości (np. wrażliwość na EMI).

Strojenie odbywa się za pomocą standardowego enkodera z przyciskiem monostabilnym (RE1). Przełącznik przełącza na pasmach różne wielkości kroków. Zewnętrzny przełącznik pasma, S3, przełącza między trybami AM i FM.

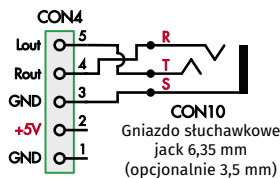
Zastosowany został przełącznik typu ON-OFF-ON, aby zapewnić trzy pasma. Daje to trzy różne napięcia, które można odczytać



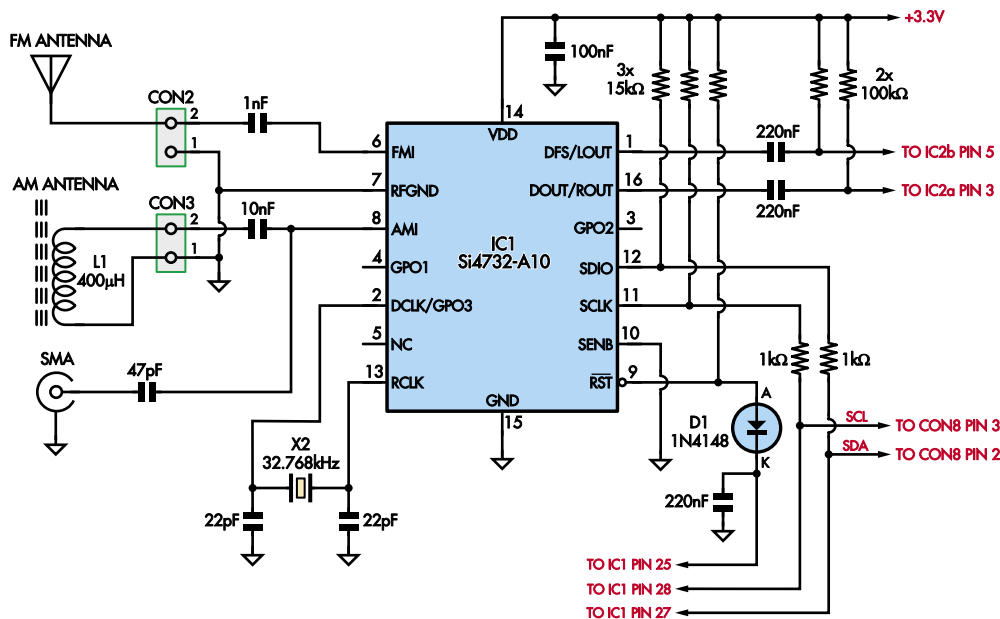
Cyfrowy odbiornik radiowy AM/FM

Rysunek 3. Obwód radiowy nie jest zbyt rozbudowany, dzięki modułowi radiowemu Si4730. Anteny po lewej stronie są po prostu sprzężone z modułem za pomocą kondensatorów, podczas gdy wyjścia audio po prawej stronie zasilają parę stopni buforowych/wzmacniających wzmacniacza operacyjnego, które są efektywniejsze wysterowaniu słuchawek niż sam moduł. Układ IC2 steruje radiem za pośrednictwem magistrali I²C, jednocześnie monitorując dane wejściowe użytkownika ustawiane za pomocą enkodera obrotowego RE1 i wyświetlając informacje o strojeniu i sile sygnału na dwuwierszowym wyświetlaczu LCD

► Rysunek 4(a). Jeśli chcesz odbierać pasmo SW, wszystko co musisz zrobić, to zrezygnować z modułu Si4730 (MOD1) i zamiast tego zamontować IC1, kondensator odsprężający 100 nF, kwarc X2 i dwa kondensatory 22 pF. Wszystkie pozostałe elementy pokazane tutaj znajdowały się w oryginalnym układzie (rysunek 3) i zostały zduplikowane tylko w celu wyjaśnienia, w jaki sposób IC1 jest podłączony do reszty układu



▲ Rysunek 4(b). Pokazane jest podłączenie do złącza CON4 montowanego na panelu gniazda słuchawkowego. Sprawdź dokładnie podłączenia końcówek odpowiadających poszczególnym stykom gniazda



przez wejście przetwornika analogowo-cyfrowego (ADC) w ATmega, PC3 (pin 26). Jeśli używany jest moduł Si4730, nie ma pasma SW, więc zamiast tego należy użyć przełącznika dwupozycyjnego.

Kolejne wejście ADC, PC0 (pin 23), monitoruje napięcie na wyjściu potencjometru VR2, który reguluje głośność. Odczyt jest skalowany i wysyłany przez linie I²C w celu sterowania głośnością układu SiLabs.

Trzecie wejście ADC na PC1 (pin 24) odczytuje pozycję potencjometru VR3; odczyt jest skalowany i wysyłany do układu SiLabs

w celu regulacji szerokości pasma w zakresie AM. Można użyć przełącznika wielopozycyjnego, ale jest to prostsze i tańsze rozwiązanie.

Dostępne szerokości pasma to 1,0; 1,8; 2,0; 2,5; 3,0; 4,0 i 6,0 kHz. Użyty potencjometr ma środkowy odczep, który daje szerokość pasma 2,5 kHz, ale jest to opcjonalne. Nie ma opcji szerokości pasma dla FM.

Korzystanie z układu Si4732

Dla tych, którzy chcą włączyć pasmo SW lub LW, istnieje możliwość użycia układu Si4732 zamiast modułu Si4730. Jest on dostarczany

w obudowie SOIC SMD, która nie jest specjalnie trudna do manualnego montażu. Istnieją tylko niewielkie zmiany w układzie, jak pokazano na rysunku 4.

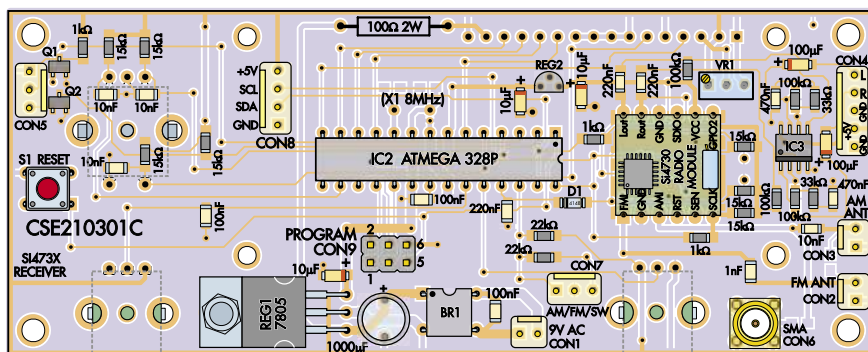
Pin SENB w Si4732 jest podłączony do masy, co daje mu inny adres I²C niż w przypadku modułu Si4730. Wymaga również dodatkowego rezonatora kwarcowego i trzech kondensatorów. Adresy I²C modułu Si4730 to C6hex dla zapisu i C7hex dla odczytu. W przypadku układu Si4732 odpowiednie adresy to 22hex i 23hex.

Nie należy stosować równocześnie modułu Si4730, jak i układu Si4732. Chociaż mają one różne adresy I²C, obciążenie wejść RF jest takie, że poważnie pogarsza czułość.

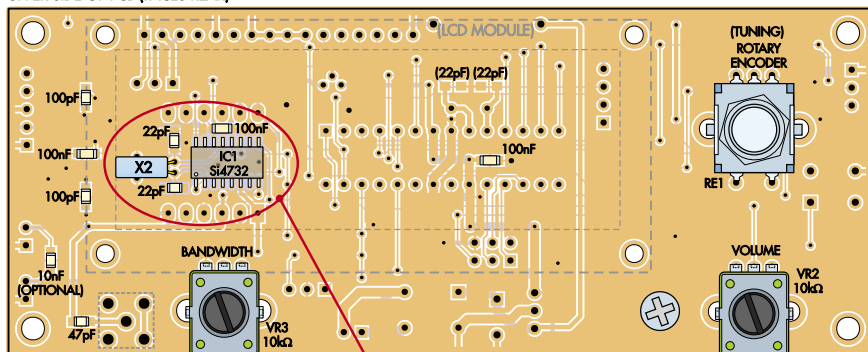
Warto zauważyć, że magistrala I²C jest dostępna z zewnątrz przez CON8, wraz z zasilaniem +5 V. Może to być przydatne do rozbudowy układu w przyszłości lub jako pomoc w debugowaniu.

Rysunek 5 i 6. Większość komponentów montuje się na górnej stronie płytki drukowanej; oprócz kilku części SMD, jedynymi częściami na dole są dwa potencjometry, enkoder obrotowy i kwarc X2 (jeśli zamontowany jest IC1). Najlepiej jest zamontować wszystkie elementy SMD na spodzie, następnie elementy SMD na górze, kolejno elementy przewlekane na górze, a następnie na spodzie. Upewnij się, że podzespoły o wymagającym kierunku montażu, takie jak moduł radiowy, wszystkie układy scalone, aluminiowe i tantalowe kondensatory elektrolityczne, mostek prostowniczy BR1, dioda D1 i potencjometr nastawny VR1 są zorientowane zgodnie z rysunkiem

Errata: jeśli używasz określonej w spisie części, stabilizator napięcia 3,3 V REG2 powinien być zamontowany do góry nogami w stosunku do schematu montażowego. Alternatywnie można go zamontować na spodniej stronie płytki drukowanej, upewniając się, że nie dotyka panelu przedniego. Wynika to z zamiany wyprowadzeń wejściowych i wyjściowych na płycie drukowanej. W każdym razie, po wlutowaniu REG2, upewnij się, że działa on prawidłowo i dostarcza napięcie 3,3 V



UPPER SIDE OF PCB (FACES REAR)



LOWER SIDE OF PCB (FACES FRONT)

Elementy użyte zamiast modułu Si3730

Można zastosować zasilanie 9 V AC lub 9...12 V DC, poprzez CON1. Jeśli używane jest zasilanie DC, nie może to być zasilacz impulsowy, ponieważ wytwarzają one dużo zakłóceń, które mogą zakłócać a nawet stłumić pasmo AM.

Regulator/stabilizator 7805 zasilą układ ATmega i moduł LCD, podczas gdy mały regulator liniowy w obudowie TO-92 zapewnia napięcie 3,3 V dla układu SiLabs.

Interfejs debugowania

MOSFET-y Q1 i Q2 zapewniają szeregowy interfejs debugowania. Było to nieocenione w fazie tworzenia i rozwijania układu, ale nie jest wymagane, jeśli chcesz po prostu korzystać z radia. Interfejs jest ustawiony na 38 400 bps, osiem bitów danych, jeden bit stopu i brak parzystości.

Mikrokontroler IC2 jest programowany za pomocą standardowego 6-stykowego złącza CON9. Do resetowania IC2 służy przycisk monostabilny.

Wybór komponentów

Chociaż Redakcja Silicon Chip-a starała się zastosować komponenty powszechnie dostępne, kilka głównych trzeba będzie zakupić u zagranicznych dostawców. Podczas opracowywania artykułu istniało kilku dostawców modułu Si4730-V2.0 na platformie AliExpress, którzy sprzedawali go po około 5 USD. Upewnij się, że jest to wersja z sześcioma

podłączeniami po każdej stronie. Istnieją wersje z tylko pięcioma złączami po każdej stronie, które nie będą pasować do PCB. Podobnie jak w przypadku większości zamówień z Chin, należy przygotować się na dość długi, około miesięczny, czas dostawy. **Od Red. EdW: wg naszego rozeznania, ta wersja modułu jest aktualnie niedostępna na znanych platformach handlowych. Sprawę znalezienia modułu pogarsza fatalnej jakości wyszukiwarka na portalu AliExpress. Nigdy nie ma też pewności, że otrzymamy to, co zamówiliśmy.**

Układ Si4732 jest produkowany w obudowie SOIC-16. Jest dostępny na AliExpress, po około 3 USD za sztukę. Jest również dostępny w Digi-Key i Mouser w nieco wyższej cenie, ale dobrą wiadomością jest to, że można go zamówić wraz z innymi częściami (o wartości około 60 USD) z bezpłatną ekspresową dostawą.

Oprócz kondensatora elektrolitycznego 1000 µF i rezystora 2 Ω, wszystkie inne rezystory i kondensatory są w wersji do montażu SMD, w rozmiarze 1206 lub 0805, i nie ma żadnych podzespołów o mniejszym rastrze wyprowadzeń, o których montaż należałoby się martwić.

Dostępne są różne kolory podświetlenia modułu LCD. Zdecydowanie polecamy wersję białą-niebieską, w przeciwieństwie do staromodnej wersji żółto-zielonej (wersja białą-niebieską jest również bardziej czytelna). Ten typ jest dostępny u kilku australijskich dostawców w serwisie eBay. Ale jeśli nie

masz nic przeciwko czekaniu, moduł LCD może kosztować u chińskich dostawców zaledwie około 2,50 USD.

Ponieważ używamy interfejsu równoległego, nie będzie potrzebna płytki interfejsu szeregowego I²C dostarczana z niektórymi modułami LCD.

Wyświetlacz LCD jest montowany do płytki za pomocą wsporników i podłączany za pomocą dostarczonej standardowej listwy kołkowej wciskanej do niskoprofilowego jednorzędowego gniazda prostego. Montaż wyświetlacza LCD nad płytką oraz sama wysokość tego wyświetlacza powodują, że dwa potencjometry i enkoder obrotowy wymagają osi o długości 25...30 mm. Lista części zawiera sugerowane komponenty.

Budowa

Słowo przestrogi. Rezonator kwarcowy na małym module „4730” nie jest zbyt mocno utwierdzony i można go łatwo wygiąć w jedną stronę i uszkodzić. Autor ostrzega przed tym na podstawie własnego doświadczenia! Zaleca więc użycie kleju kontaktowego super-glue, aby mocno unieruchomić rezonator kwarcowy na płytce. W każdym razie lepiej zamówić dwa takie moduły, na wszelki wypadek.

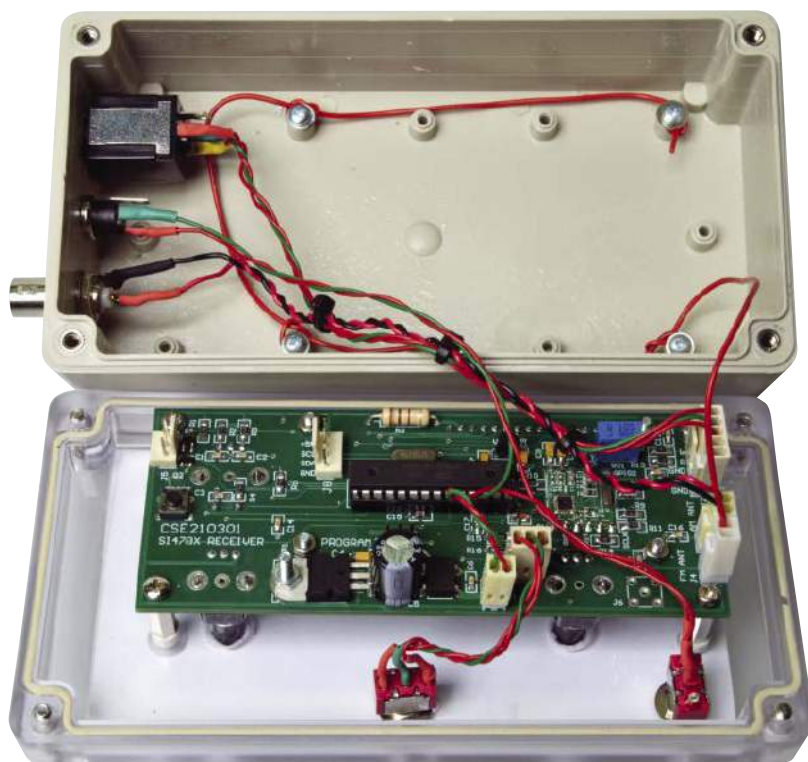
Płytką drukowaną (kod CSE210301C) jest dwustronna, z komponentami po obu stronach. Jej wymiary to 123×49,5 mm. Obie wersje wykorzystują tę samą płytkę drukowaną: albo montuje się moduł Si4730 po jednej stronie, albo układ Si4732 po drugiej. Zapoznaj się ze schematami montażowymi na rysunkach 5 i 6 i upewnij się, że albo zamontujesz moduł, jak pokazano na rysunku 5, albo komponenty pokazane w czerwonym owale na rysunku 6; nie montuj obu razem.

Zacznij od zamontowania 16-pinowego układu scalonego. Jest to typ SOIC-16 z wyprowadzeniami rozmieszczonymi w rastrze 1,27 mm, a więc na tyle szeroko, że można je lutować pojedynczo za pomocą lutownicy z cienką końcówką.

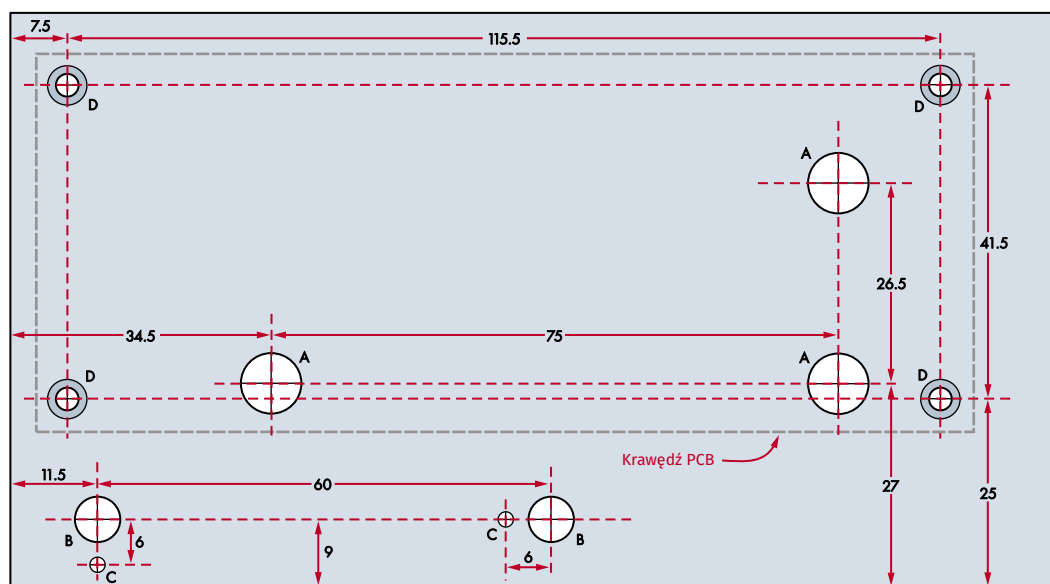
Najpierw nałóż topnik na pola lutownicze, aby zmniejszyć ryzyko tworzenia się mostków między wyprowadzeniami. Jeśli podczas lutowania utworzą się plecionki, użyj więcej topnika i miedzianej plecionki lutowniczej, aby je usunąć.

Następnie zamontuj kondensatory SMD na spodzie płytki. Zauważ, że dwa kondensatory 22 pF (wartości w nawiasach) są potrzebne tylko wtedy, gdy chcesz użyć oscylatora kwarcowego dla układu ATmega168/ATmega328. Nie jest to konieczne, więc sugerujemy aby ich nie montować.

Druga strona płytki zawiera resztę komponentów. Następnie zamontuj pozostałe



Sposób okablowania prototypowego radia z modułem Si4730



- Otwory A: $\varnothing$ 8,0 mm
- Otwory B: $\varnothing$ 6,0 mm (pod przełączniki)
- Otwory C: $\varnothing$ 2,0 mm (pod znaczniki przełączników)
- Otwory D: $\varnothing$ 3,0 mm (montaż PCB, szfawowane)

Wszystkie wymiary podano w milimetrach

Rysunek 7. Jeśli używasz obudowy z przezroczystą pokrywą, musisz tylko wywiercić okrągłe otwory, jak pokazano tutaj. Możesz przykleić taśmę maskującą do panelu, zmierzyć i zaznaczyć wymiary otworów lub po prostu skopiować/wydrukować ten schemat, wyciąć go i użyć jako szablonu. Aby uzyskać najlepszy efekt, należy pogłębić otwory oznaczone literą D po zewnętrznej stronie panelu

komponenty do montażu powierzchniowego. Jeśli używasz modułu Si4730, upewnij się, że jest on dokładnie umieszczony. Potrzebuje on sporej ilości lutu, aby mógł wpłynąć do „pół-otworów” po obu stronach (patrz zdjęcie na kolejnej stronie).

Upewnij się, że kondensatory tantalowe 10 μ F i 100 μ F są umieszczone z prawidłową polaryzacją. Koniec z paskiem jest dodatni, więc skieruj końce oznaczone paskami w stronę symboli „+” na płytce drukowanej. Następnie należy zamontować elementy przewlekane, w tym opcjonalny rezonator kwarcowy 8 MHz.

Przewidziano również miejsce na gniazdo SMA CON6, które nie zostało użyte. Jest to alternatywne wejście antenowe dla pasm AM, LW i SW.

Autor sugeruje, aby moduł LCD był wyjmowany; dlatego został podłączony do gniazda. Pasujące ze względu na wysokość złącza nie są łatwe do znalezienia, ale lista części wspomina o dostawcach. LCD jest następnie mocowany za pomocą 9 mm kołków dystansowych (Jaycar HP0862 lub Altronics H21362) oraz śrub M2,5 $\times$ 15 i nakrętek.

Ostatnimi elementami do przymocowania są dwa potencjometry (VR2 i VR3)

oraz enkoder obrotowy RE1 po stronie LCD. Na koniec należy dobrze umyć płytkę z obu stron środkiem do czyszczenia płytek drukowanych.

Przygotowanie obudowy

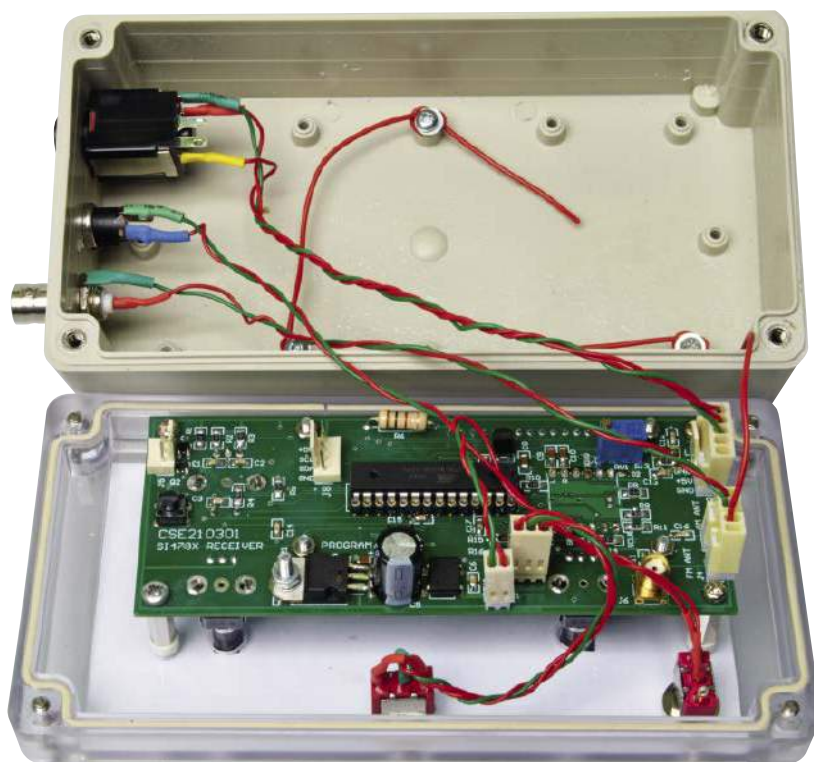
Prototypy radia zostały umieszczone w obudowie Hammond RP1175C, która ma przezroczystą pokrywę. Pozwala to uniknąć konieczności wykonywania prostokątnych wycięcia na wyświetlacz LCD, dzięki czemu można wywiercić wszystkie otwory. Jedynymi sklepami, w których można kupić tę obudowę, były Mouser i Digi-Key. Można użyć większej obudowy, która jest dostępna lokalnie, ale sprawiłoby to, że radio byłoby nieco mniej wygodne w użyciu.

Złącze zasilania, gniazdo słuchawkowe i złącze antenowe BNC można umieścić na dowolnej wygodnej powierzchni. W prototypach jest to prawa strona obudowy.

Gniazdo słuchawkowe stanowi pewien problem. Grubość obudowy jest zbyt duża dla łatwo dostępnych gniazd stereo 3,5 mm. Najprostszym rozwiązaniem jest użycie gniazda 6,35 mm, a w razie potrzeby adaptera na 3,5 mm, takiego jak Jaycar PA3590. Jakimś rozwiązaniem jest też wklejenie do obudowy gniazda 3,5 mm odciętego, wraz z przewodem o odpowiedniej długości, z przedłużacza 3,5 mm – 3,5 mm.

Szczegóły wiercenia pokazano na rysunku 7; użyj go jako początkowego szablonu do zlokalizowania otworów montażowych płytki drukowanej (D) i otworów przełącznika (B). Ponieważ wymagana jest duża dokładność, goła płytka drukowana może być wstępnie użyta jako szablon do wiercenia otworów montażowych.

Użyj narzędzia do pogłębiania, aby były śruby były wyrównane z panelem przednim.



Podobnie, przykład okablowania tego projektu dla wersji z układem Si4732

Wykaz elementów:

- 1 dwustronna płytką drukowaną o kodzie CSE210301C i wymiarach 123×49,5 mm
- 1 zasilacz wtyczkowy 9 V AC typu transformatorowego z wtyczką cylindryczną 2,1/2,5 mm ID
- 1 plastikowa obudowa z przezroczystą pokrywą [np. Altronics H0326, Hammond RP1175C: Digi-Key; Mouser]
- 1 etykieta panelu, pasująca do budowanej wersji
- 1 alfanumeryczny moduł LCD 16×2 z niebieskim podświetleniem (LCD1)
- 1 wąska podstawa DIL28 IC
- 3 2-stykowe spolarzowane kompletne złącza 402-2/403-2 z pasującymi stykami (CON1-3) [Jaycar HM3412/02, Altronics P5492/72 + 2× P5470A]
- 1 5-stykowe spolarzowane kompletne złącze 402-5/403-5 z pasującymi stykami (CON4) [Jaycar HM3415/05, Altronics P5495/75 + 5× P5470A]
- 2 3-stykowe spolarzowane kompletne złącza 402-3/402-3 z pasującymi stykami (CON5, CON7) [Jaycar HM3413/03, Altronics P5493/73 + 3× P5470A]
- 1 4-stykowe spolarzowane kompletne złącze 402-4/403-4 z pasującymi stykami (CON8; opcjonalnie) [Jaycar HM3414/04, Altronics P5494/74]
- 1 gniazdo BNC do montażu na panelu obudowy [Jaycar PS0658, Altronics P0516A]
- 1 montowane na płycie drukowanej gniazdo zasilania DC, 2,1/2,5 mm ID, pasujące do zestawu wtyczek [np. Jaycar PS0522/4, Altronics P0620/1A]
- 1 gniazdo stereo jack 6,35 mm do montażu na panelu obudowy lub patrz opis w tekście [np. Jaycar PS0182, Altronics P0065]
- 1 16-stykowa jednorzędowa niskoprofilowa listwa kołkowa z pasującym jednorzędowym gniazdem 16-stykowym (dla LCD) *
- 1 wieloobrotowy potencjometr nastawny 10 kΩ (VR1)
- 2 pionowe potencjometry 10 kΩ 9 mm z osiami w rozmiarze D (VR2, VR3) [np. Bourns PTV09A-4030F-B103-ND; lub użyj Altronics R1946 z karbowaną osią]
- 1 pionowy enkoder obrotowy z osią w rozmiarze D i zintegrowanym przełącznikiem przyciskowym (RE1) [np. Bourns PEC11R-4225F-S0024]
- 3 małe lub średnie pokręta pasujące do VR2, VR3 i RE1
- 1 mały dotykowy przełącznik przyciskowy (S1) do montażu na płycie drukowanej [np. Jaycar SP0601 lub Altronics S1120]
- 1 miniaturowy przełącznik SPDT z przylutowanym znacznikiem (S2) [np. Jaycar ST0335]
- 1 antena z prętem ferrytowym i cewką 400 μH (L1) [np. Jaycar LF1020]
- 4 kołki dystansowe 9 mm (do montażu LCD) [Jaycar HP0862, Altronics H1362]
- 4 śruby z łbem walcowym M3×9-10 i nakrętki (do REG1)
- 4 śruby z łbem stożkowym M3×12
- 4 śruby z łbem walcowym M3×6
- 4 gwintowane kołki dystansowe M3×12
- 4 kołki dystansowe bez gwintu o długości 6 mm i średnicy wewnętrznej 3,25 mm
- 4 śruby z łbem walcowym M2,5×15 i nakrętki (do montażu wyświetlacza LCD)
- różnej długości przewody połączeniowe o średniej wytrzymałości
- różne krótkie odcinki rurki termokurczliwej w zależności od rozmiaru przewodu
- * niektóre opcje obejmują Semtronics SBU400Z (listwa kołkowa) + MH1519-140 (gniazdo), Mouser 200-BBL116GF (listwa kołkowa) + Mouser 200-SL116T10 (gniazdo), element14 1667454 (listwa kołkowa) + Jaycar PI6470 (gniazdo) lub Altronics P5400 (gniazdo)

Półprzewodniki:

- 1 ATmega168 lub ATmega328, 8-bitowy mikroprocesor zaprogramowany wsadem CSE210301.hex (IC2)
- 1 wzmacniacz operacyjny 5 V „rail-to-rail”, SOIC-8 (IC3) [np. LME49721, dostępny w Digi-Key, Mouser, eBay, AliExpress]
- 1 regulator/stabilizator liniowy 7805 5 V 1 A, TO-220 (REG1)
- 1 regulator/stabilizator liniowy LM2936-3.3 3,3 V o niskim spadku napięcia, TO-92 (REG2)
- 2 MOSFET-y N-kanalowe małej mocy 2N7002, SMD SOT-23 (Q1, Q2)
- 1 mostek prostowniczy DB104 (BR1) [Jaycar ZR1308]
- 1 dioda małej mocy LL4148, SMD DO-80 MELF (D1) [Jaycar ZR1103]

Kondensatory: (wszystkie w rozmiarze SMD 0805, o ile nie podano inaczej)

- 1 kondensator elektrolityczny 1000 μF/16 V do radialnego montażu przewlekane
- 5 kondensatorów ceramicznych 100 nF/50 V X7R
- 2 kondensatory tantalowe 100 μF/6 V SMD, rozmiar SMA
- 5 kondensatorów ceramicznych 10 nF/50 V X7R
- 3 kondensatory tantalowe 10 μF/6 V SMD, rozmiar SMA
- 1 kondensator ceramiczny 1 nF/50 V X7R
- 2 kondensatory ceramiczne 470 nF/50 V X7R
- 2 kondensatory 100 pF/50 V ceramika COG/NPO
- 3 kondensatory ceramiczne 220 nF/50 V X7R
- 1 kondensator 47 pF/50 V ceramika COG/NPO

Rezystory: (wszystkie 1% SMD 1206, o ile nie podano inaczej)

- 4 szt. 100 kΩ 2 szt. 33 kΩ 2 szt. 22 kΩ 7 szt. 5 kΩ 3 szt. 1 kΩ
- 1 szt. 100 Ω 5%/2 W osiowy

Dodatkowe części dla wersji zawierającej układ Si4732

- 1 układ scalony Si4732, SOIC-16 (IC1) [AliExpress, eBay]
- 1 miniaturowy przełącznik on-off-on (środek wyłączony) z przylutowanym znacznikiem (S3) [np. Jaycar ST0336]
- 1 kwarc zegarkowy 32 768 Hz (X2)
- 1 kondensator ceramiczny 100 nF/50 V X7R, SMD 0805
- 2 kondensatory ceramiczne 22 pF/50 V ceramika COG/NPO, SMD 0805

Dodatkowe części dla wersji zawierającej moduł Si4730

- 1 moduł Si4730 do montażu powierzchniowego, z sześcioma stykami po obu stronach (MOD1) [AliExpress, eBay]
- 1 miniaturowy przełącznik SPDT z przylutowanym znacznikiem (S3) [np. Jaycar ST0335].

Części opcjonalne

- 1 pionowe gniazdo SMA (CON6) (wejście zewnętrznej anteny AM)
- 1 odcinek 2×3 dwurzędowej listwy kołkowej (CON9) (do programowania IC2 w obwodzie)
- 1 rezonator kwarcowy 8 MHz (X1) (patrz tekst)
- 2 kondensatory ceramiczne 22 pF/50 V ceramika COG/NPO, SMD 0805

Należy zauważyć, że w płycie PCB na miejscu środków osi enkodera i dwóch potencjometrów znajdują się małe otwory.

Po wywierceniu czterech otworów montażowych (D), przymocuj płytkę do panelu za pomocą śrub 3 mm i wywierć otwory 1 mm przez środek otworów w miejscu montażu dwóch potencjometrów i enkodera, aby dokładnie zaznaczyć środki otworów 8 mm (oznaczenie A).

Autor wydrukował etykietę panelu przedniego o wymiarach 139×76 mm na grubym papierze fotograficznym, który idealnie pasuje do szczeliny na przezroczystym panelu.

Rysunek 8 przedstawia etykietę panelu dla wersji zawierającej moduł Si4730, natomiast rysunek 9 przedstawia etykietę dla wersji z modulem Si4732. Jedyna różnica polega na oznaczeniu przełącznika zmiany pasma, dodając opcję SW dla układu Si4732. Etykiety te można również pobrać ze strony internetowej Silicon Chip-a i wydrukować.

Użyj ostrego noża modelarskiego, aby wyciąć w etykiecie otwór na wyświetlacz LCD i pięć otworów na osie potencjometrów, enkodera i przełączniki, a następnie wytnij panel i wsuń go do wnętrza przezroczystej pokrywy. Powinien być dobrze dopasowany.

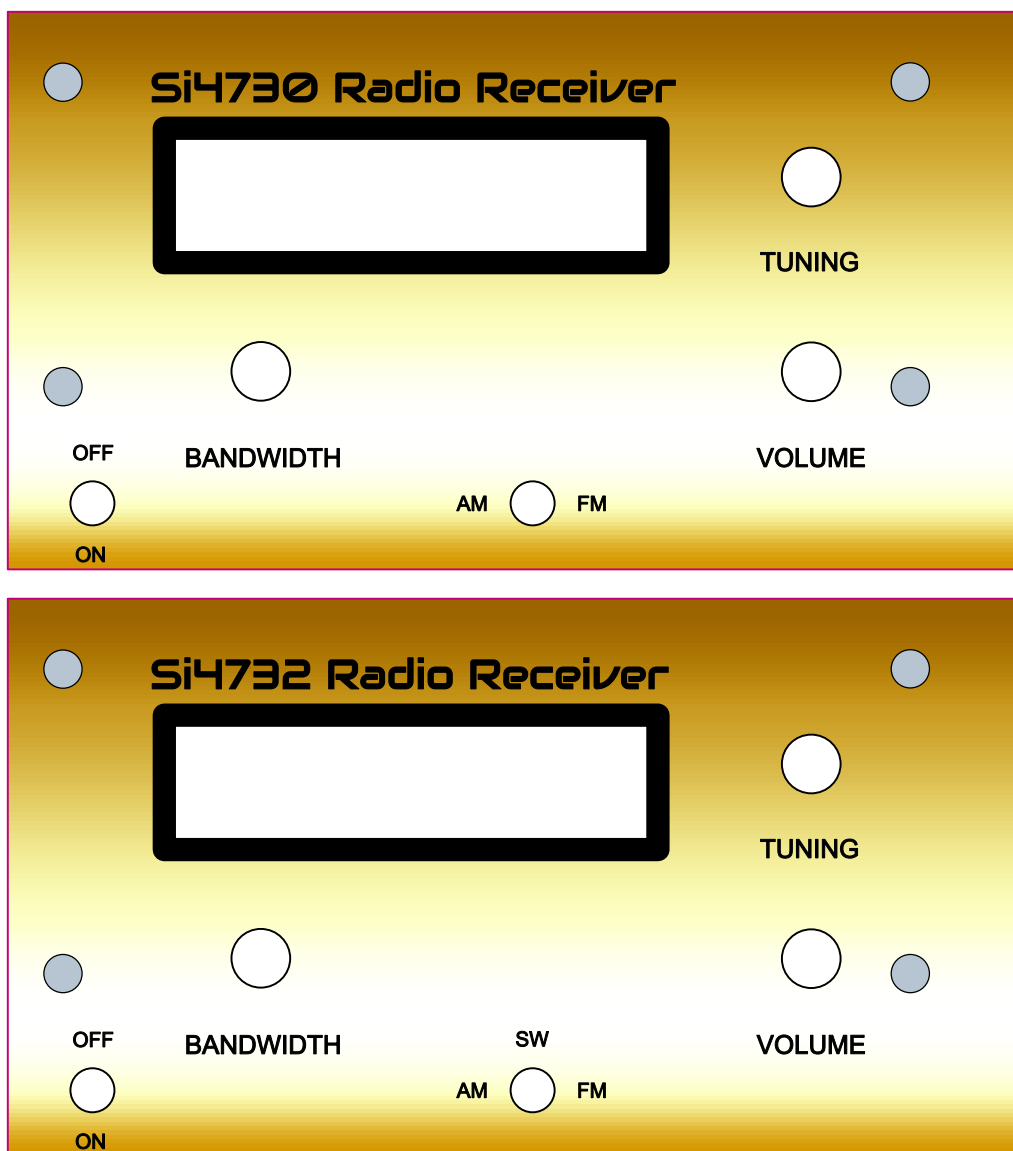
Używając gwintowanych kołków dystansowych (gwintowanych tulei mosiężnych), przymocuj płytkę drukowaną od spodu panelu przedniego za pomocą śrub M3 z łbem stożkowym o długości 12 mm z przodu i śrub M3×6 z tyłu. Potrzebne są kołki dystansowe o długości 18 mm, które można wykonać z gwintowanego kołka dystansowego 12 mm oraz niewykorzystanego kołka dystansowego 6 mm. Mogą istnieć inne kombinacje elementów dystansowych, aby uzyskać wymagane 18 mm. Szczegóły w spisie materiałów.

Wałki potencjometru i enkodera mają średnicę po 6 mm. Zachowaj ostrożność, jeśli używasz pokręteł metrycznych, ponieważ niektóre z nich mogą nie pasować do wałków. Wybierz typy z wkrętem dociskowym, ponieważ będą one pasować do szerokiej gamy typów wałków.

Pozostaje jeszcze wewnętrzne okablowanie do przełączników i złączy na obudowie. Jest to stosunkowo proste i pokazane na zdjęciach (patrz rysunki 3...6).

Programowanie mikroprocesora

Autor napisał oprogramowanie sterujące przy użyciu BASCOM, kompilatora BASIC dla mikroprocesorów AVR. Posiadanie dokumentacji aplikacji i notatek programistycznych dostarczonych przez SiLabs



Rysunek 8 i 9. Te etykiety na panel przedni są również dostępne do pobrania ze strony internetowej Silicon Chip-a, dzięki czemu można je wydrukować, wyciąć i przymocować do wewnętrznej (lub zewnętrznej) strony pokrywy obudowy

sprawiło, że kod był dość prosty. Zarówno kod źródłowy .bas, jak i plik wykonywalny .hex są dostępne do pobrania ze strony internetowej Silicon Chip. Należy pamiętać, że do skompilowania pliku .bas jest potrzebna płatna wersja programu BASCOM.

Złącze programowania na płytce jest przeznaczone dla programatora AVRISP Mk2. Można go używać w połączeniu z darmowym programem Atmel (obecnie Microchip) Studio dostępnym do pobrania ze strony www.microchip.com.

Sterowanie układem SiLabs odbywa się poprzez magistralę I²C za pomocą poleceń szeregowych, a naprawdę, jest ich mnóstwo. Istnieje wiele funkcji, takich jak skanowanie, które można włączyć do projektu, ale Autor zdecydował się zachować maksymalną prostotę („keep it simple, stupid” – KISS).

Oczywiście możesz Czytelniku rozwinąć oprogramowania Autora.

Jak wspomniano powyżej, przycisk monostabilny zintegrowany z enkoderem strojenia przełącza między krokami strojenia, umożliwiając precyzyjny wybór lub szybkie strojenie w całym paśmie. W paśmie AM krok wynosi 1 kHz, 9 kHz lub 100 kHz. Pasmo FM wynosi od 87 MHz do 108 MHz i ma krok 100 kHz lub 1 MHz.

W paśmie SW (jeśli jest używane) krok wynosi 1 kHz, 10 kHz, 100 kHz lub 1 MHz. Około pół sekundy po wybraniu częstotliwości, jest ona zapisywana wraz z rozmiarem kroku, w pamięci EEPROM. Oznacza to, że przy następnym włączeniu zasilania wartości EEPROM są odczytywane i wybierana jest ta częstotliwość.

Górny wiersz wyświetlacza LCD 16×2 pokazuje częstotliwość, a w pasmach AM i SW

także szerokość pasma. Druga linia pokazuje wielkość kroku i stosunek sygnału do szumu (SNR). Układ Si jest próbkowany raz na sekundę, aby zaktualizować wartość SNR.

Moduł Si4730 nie podaje jednak odczytów SNR w paśmie FM. Słabsze sygnały dają wyjście mono, a nie stereo, jak można by oczekiwać.

Początkowa konfiguracja

Projekt nie ma oddzielnego programu sterującego dla układów Si4730 i Si4732, więc typ układu jest automatycznie identyfikowany po włączeniu zasilania. Nie trzeba nic robić.

W drugim prototypie Redakcji Silicon Chip-a okazało się, że strojenie było odwrotne. Obrót w prawo zmniejszał częstotliwość! Wygląda na to, że enkodery różnią się działaniem osi między sobą. Autor sugeruje w tym miejscu metodę wyboru prawidłowego kierunku strojenia za pomocą istniejącego interfejsu radiowego. Jeśli okaże się, że działanie enkodera jest odwrócone, należy wykonać następujące czynności:

1. Przekręć pokrętkę Bandwidth (Szerokość pasma) do końca w prawo.
2. Dostrój pasmo AM do 500 kHz. Wyświetlacz pokaże „Toggle Direction”

(Kierunek strojenia) w górnym wierszu i „Direction 1” lub „Direction 2” w dolnym wierszu. Nie trzeba naciskać przycisku, ponieważ automatycznie wybierany jest alternatywny kierunek.

3. Dostrój do innej częstotliwości i upewnij się, że kierunek strojenia jest prawidłowy.

Tę konfigurację należy wykonać tylko raz, ponieważ parametry są zapisywane w pamięci EEPROM i przywracane po włączeniu zasilania. ■

Charles Kosina

Adaptacja do wydania polskiego – Andrzej Nowicki

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. www.siliconchip.com.au

Subscribe to Elektor's newsletter and get the chance to

WIN

a Raspberry Pi Pico W board



www.elektor.com/eda



Subscribe to Elektor's newsletter, get a €5 coupon code and get the chance to WIN a Raspberry Pi Pico W board



Be one of the 10 fortunate winners!



elektor
design > share > earn

Prezentowany tutaj tester serwomechanizmów opracowałem, ponieważ dron, nad którym pracuję, czasami podczas lotów testowych, wydziela nieprzyjemny zapach spalonej elektroniki i spalonego lakieru miedzianego. Potrzebowałem narzędzia, które pozwoli mi zdiagnozować problem i pomoże zrozumieć przyczynę problemów.

Przetestuj do czterech serwomechanizmów oddzielnych lub w systemie

Super Tester Serwomechanizmów

W Internecie można znaleźć wiele projektów testerów serwomechanizmów służących do testowania pojedynczego serwa. Pozwalają sprawdzić, czy właśnie zakupione serwo działa prawidłowo. Prezentowany tutaj tester idzie w tej koncepcji o kilka kroków dalej. Można go używać nie tylko do jednoczesnego testowania maksymalnie czterech serwomechanizmów, ale umożliwia także testowanie modułu elektronicznej kontroli prędkości (ESC) lub sterownika serwomechanizmów wraz z serwami, w wielu konfiguracjach.

Dron do którego potrzebowałem testera pokazany jest na **rysunku 1**. Jego konstrukcję oparto na projekcie opisanym w [1].

Sygnały serwa 101

Dla przypomnienia, serwomechanizm modelarski jest sterowany sygnałem w postaci impulsów o zmiennej szerokości, o częstotliwości powtarzania wynoszącej 50 Hz. Szerokość impulsu (czas, przez który sygnał jest wysoki) powinna mieścić się w zakresie od jednej do dwóch milisekund (niektóre serwa akceptują czas impulsu od 900 μ s do 2100 μ s), przez resztę czasu sygnał ma stan niski. Serwo znajduje się w pozycji środkowej, gdy szerokość

impulsu wynosi 1,5 ms. Gdy impuls jest krótszy, serwo obróci się w jednym kierunku (np. w lewo); gdy jest dłuższy, skręci w drugą stronę (np. w prawo).

W typowym systemie lotu należy sprawdzić pięć parametrów: czas trwania czterech impulsów sterujących serwomechanizmami (ciąg silnika, przechylenie lewo/prawo, nachylenie przód/tył i obrót wokół osi pionowej) oraz napięcie zasilania. W przypadku skomplikowanego projektu, np. opracowania kontrolera lotu (drona lub innego typu modelu RC) lub dowolnego miksera impulsowego (aparatura sterownicza), tester ten pozwala zobaczyć wynik działania miksera.

Należy zauważyć, że w pozostałej części tego artykułu termin „serwo” można zastąpić terminem ESC lub dowolnym innym urządzeniem wytwarzającym i/lub rozpoznającym sygnały serwo.

Funkcje i możliwości

Super Serwo Tester (SST) mierzy czas trwania impulsów sterujących maksymalnie czterema serwami i dostarcza informacji o jakości zasilania. Można go umieścić pomiędzy odbiornikiem zdalnego sterowania a serwami, pomiędzy odbiornikiem a kontrolerem

lotu drona lub pomiędzy kontrolerem lotu a serwami. **Rysunki 2...6** przedstawiają możliwe konfiguracje pracy testera.

Tryby pracy

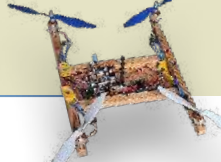
SST posiada dwa tryby pracy, wybierane za pomocą przełącznika suwakowego:

Ręczny: W tym trybie SST generuje impulsy dla czterech serwomechanizmów lub kontrolera lotu. Szerokość impulsów jest ustawiana za pomocą czterech potencjometrów. W tym trybie to SST dostarcza zasilanie do serw (lub ESC) i zasilanie modelu nie może być podłączone do żadnego z nich. Napięcie zasilania SST musi wynosić od 7,5 V do 12 VDC.

Wejścia: W tym trybie mierzone są długości impulsów pochodzących z odbiornika lub sterownika. Sygnały są również przekazywane do wyjść SST w celu sterowania serwami (lub kontrolerem lotu). W tym trybie tester i serwa (lub ESC) są zasilane z zasilacza modelu. Napięcie zasilania nie może przekraczać 7,49 V, a tester nie powinien być podłączany do własnego źródła zasilania. Ponadto wszystkie cztery kanały wejściowe muszą być podłączone, w przeciwnym razie dioda LED i brzęczyk zasygnalizują usterkę.



Rysunek 1. Quadcopter, dla uruchamiania którego zaprojektowano opisywany tester



Tryby wyświetlania

Wyświetlacz pokazuje graficznie czas trwania impulsów na czterech wykresach słupkowych wraz z ich wartością liczbową w mikrosekundach (od 1000 μ s do 2000 μ s). Wykresy słupkowe są ograniczone do zakresu od 1000 μ s do 2000 μ s, ale wartość liczbową ograniczona nie jest. Gdy wartość wykryta poza zakres, wokół jej wartości liczbowej rysowana jest ramka. W tym przypadku dioda LED również się zaświeci, a brzęczyk, jeśli jest włączony, wyda sygnał dźwiękowy.

Wświetlacz pokazuje także wartość napięcia zasilania serwo mechanizmów w dowolnym trybie pracy (**Ręcznym** lub **Wejścia**). Rzeczywiście, jakość zasilania modelu jest ważna zarówno dla kontrolera lotu, jak i dla bezpieczeństwa pilota oraz osób podziwiających jego umiejętności latania. Zmierzone napięcie zasilania zostanie zaznaczone ramką, gdy jego wartość będzie niższa niż 4,5 V. Dodatkowo zaświeci się też dioda LED, a brzęczyk wyda sygnał dźwiękowy, dzięki czemu nie trzeba cały czas obserwować testera.

SST oferuje dwa tryby wyświetlania, **stos** i **kwadrat** (patrz rysunek 7). Drugi typ lepiej nadaje się do quadkoptera. Położenie potencjometru P4 przy włączeniu zasilania określa tryb wyświetlania. Gdy przed włączeniem zasilania P4 zostanie obrócony całkowicie w lewo, SST będzie wyświetlał kanały jeden nad drugim (stos). Obrócenie potencjometru P4 w prawo przed załączeniem zasilania powoduje wybranie ułożenia kanałów na wyświetlaczu w kwadrat, po uruchomieniu się urządzenia.

Układ elektryczny

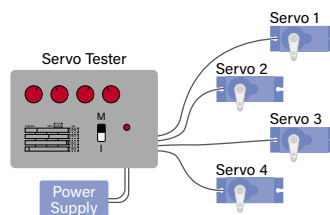
Skoro już wiemy jak korzystać z testera, przyjrzyjmy się układowi. Pokazano go na rysunku 8. Nie jest on zbyt skomplikowany, gdyż składa się głównie ze złączy. Sercem układu jest mikrokontroler ATmega328P, taki sam jak na płytce Arduino UNO. Jego częstotliwość zegara jest określana przez X1, rezonator kwarcowy o częstotliwości 16 MHz, wspomagany przez C5 i C6.

Cztery z jego końcówek GPIO (PB0 do PB3) są podłączone do wejść sygnałów serwo na złączu K1. Wyjścia sygnałów do serw (PD5, PD6, PD7 i PB4) są podłączone do złączy K2. Oba złącza są okablowane w taki sposób, że można bezpośrednio podłączyć standardowe kable serwo mechanizmów. Innymi słowy, trójki pinów 1-3-5, 2-4-6, 7-9-11 i 8-10-12 odpowiadają kolejnym serwo mechanizmom. Połączenia te często kodowane są za pomocą kolorów pomarańczowo-czerwono-brązowego. Pomarańczowy (sygnał impulsowy) łączy się z pinem 1 (lub 2, 7 lub 8), czerwony (VCC)

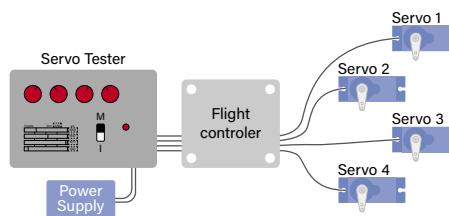
z pinem 3 (lub 4, 9 lub 10) i brązowy (GND) z pinem 5 (lub 6, 11 lub 12). Oczywiście będą wyjątki od tej reguły, dlatego sprawdź to najpierw przed podłączeniem czegośkolwiek (od Red. EdW: Serwo mechanizmy HiTec mają przewód masy w kolorze czarnym, plus (środkowy) w kolorze czerwonym i trzeci, sygnałowy w kolorze różnym od dwóch pozostałych).

Wejścia analogowe

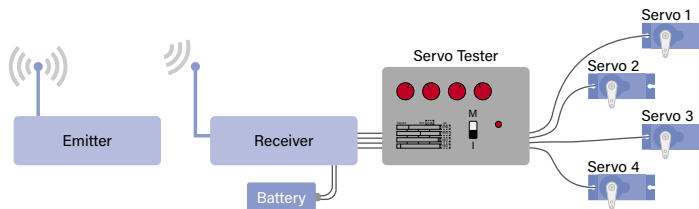
Cztery potencjometry są podłączone do wejść analogowych MCU na portach PC0 do PC3. Napięcie zasilania podawane jest poprzez dzielnik napięcia (R4, R5 i R6) na wejście analogowe PC4. Stosunek pomiędzy (R4+R5) i R6 musi wynosić 2:1, ale ich wartości bezwzględne nie są krytyczne. Użycie trzech rezystorów o tej samej wartości upraszcza ich sortowanie w celu zapewnienia precyzji (dobrze jest zastosować rezystory o tolerancji 1%).



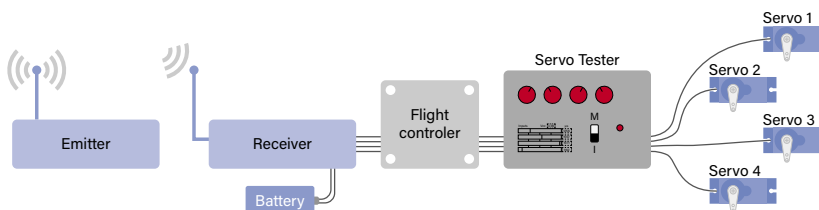
Rysunek 2. Konfiguracja 1, prosty test czterech serwo mechanizmów (tryb ręczny), pozwala na testowanie czterech serw w tym samym czasie. Serwo mechanizmy są zasilane z testera i sterowane potencjometrami



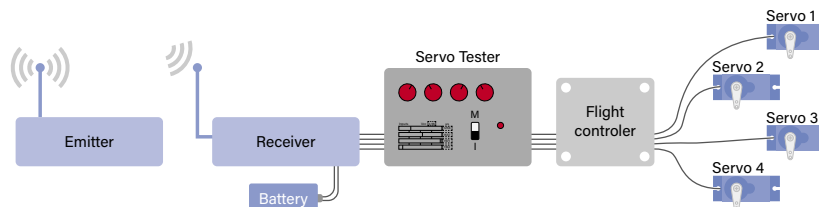
Rysunek 3. Konfiguracja 2, test kontrolera lotu po stronie odbiorczej (tryb ręczny), pozwala na przetestowanie kontrolera serw drona, bez nadajnika lub odbiornika. Serwo mechanizmy są zasilane z testera i sterowane potencjometrami



Rysunek 4. Konfiguracja 3, test nadajnika i odbiornika (tryb wejścia) w celu sprawdzenia nadajnika i odbiornika. Tester, serwa i ich sterownik są zasilane z baterii odbiornika



Rysunek 5. Konfiguracja 4, test nadajnika, odbiornika i sterownika serw (tryb wejścia), służy do sprawdzania poprawności działania nadajnika i odbiornika ze sterownikiem serw (w przypadku drona). Tester, serwa i ich sterownik są zasilane przez baterię odbiornika



Rysunek 6. Konfiguracja 5, test nadajnika i odbiornika (tryb wejścia), pozwala na testowanie sterownika serw z nadajnikiem i odbiornikiem. Tester, serwa i ich sterownik są zasilane przez baterię odbiornika



jest dokładniejszy niż zwykła 2-zaciskowa dioda Zenera, co poprawia tolerancję pomiaru napięcia zasilania serwa: 1% do 2% zamiast około 5% dla zwykłej diody stabilizacyjnej.



1



Należy pamiętać, że D2 jest dostarczany w wysokiej obudowie. Aby utrzymać konstrukcję jako niską, można go ułożyć na płytce drukowanej. Z drugiej strony możliwe jest również wykorzystanie go jako kółka podpierającego wyświetlacz.

Pomiar napięcia zasilania nie może być wyższy niż 7,49 V, ponieważ maksymalne napięcie wejścia analogowego MCU wynosi 2,5 V. Dlatego też zasilanie odbiornika i serwomechanizmów nie może nigdy przekraczać tej wartości.

Kwestie związane z wyświetlaniem

Wyświetlacz OLED I²C oparty na kontrolerze SSD1306 jest podłączony do złącza K4. Port I²C wyświetlacza nie jest podłączony do portu I²C MCU, ale do PD0 i PD1. Oprogramowanie emuluje magistralę I²C. Dzieje się tak dlatego, że w ATmega328 magistrala I²C jest współdzielona z wejściem analogowym PC4, które jest już używane do pomiaru napięcia zasilania. Dlatego nie można tutaj zastosować wbudowanego sprzętowego urządzenia peryferyjnego I²C.

R9 i R10 to rezystory podciągające dla magistrali I²C. Oficjalnie powinny mieć wartość około 4,7 kΩ, ale 10 kΩ też się sprawdza i oszczędza pozycję na liście materiałów, ponieważ większość pozostałych rezystorów również ma wartość 10 kΩ.

Należy pamiętać, że K4 jest narysowane jako złącze 8-pinowe z dwoma rzędami, ale na płytce PCB należy zamontować jednorzędowe 4-pinowe gniazdo w pozycji „A” lub „B”, a nie w obu. Pozycja zależy od wyświetlacza. Użyte wyświetlacze mają wyprowadzenia GND i VCC na stykach 3 i 4, ale nie zawsze w tej samej kolejności. K4A i K4B umożliwiają użycie dowolnego typu wyświetlacza. Użycie podwójnej podstawy zamiast (lutowanych) zworek zajmuje mniej miejsca na płycie i pozwala zaoszczędzić dwie (lutowane) zworki. Niedogodnością jest oczywiście to, że położenie otworu wyświetlacza w obudowie zależy od wyświetlacza, gdyż K4A i K4B nie znajdują się w tym samym miejscu. Ponadto w zależności od producenta wymiary tego wyświetlacza mogą się różnić. Dlatego przed obróbką obudowy należy wybrać swój wyświetlacz.

Wybór trybu pracy

Przełącznik suwakowy S1, dwubiegunowy, dwupołożeniowy (DPDT), wybiera tryb pracy super testera. W trybie **ręcznym** podłącza napięcie zasilania 5 V do złącza K1 i K2 serw. W trybie **wejść** napięcie zasilania testera wynosi VBATT, więc jego własne

zasilanie 5 V musi zostać odłączone, aby uniknąć konfliktów. Właściwie w tym trybie powinien być odłączony zasilanie testera, ale ponieważ lepiej być pewnym, niż żałować, S1 odłączy je za Ciebie. Dioda Zenra D3 wraz z R8 zapewniają, że zasilanie reszty układu pochodzące z VBATT nie przekroczy wartości maksymalnych obsługiwanych przez pozostałe elementy. Pozycja S1 jest odczytywana przez port PD4, więc MCU może wybrać odpowiedni tryb pracy.

Ścisłe mówiąc, S1 nie musi być typu DPDT, ale musi być w stanie przewodzić prąd pobierany przez maksymalnie cztery serwa, a nawet więcej, jeśli do wyjść podłączony jest również ESC lub kontroler lotu. Odpowiednie przełączniki suwakowe DPDT są często tańsze niż typy jednobiegunowe, co wyjaśnia, dlaczego są tutaj stosowane.

Inne szczegóły

Port GPIO PD2 MCU steruje diodą alarmową i brzęczykiem. Ponieważ są one połączone równolegle, MOSFET T1 zapewnia dodatkową moc do sterowania nimi obydwojema bez przeciążania MCU. Drugi przełącznik suwakowy, S2 służy do wyciszenia brzęczyka, gdyż czasami może on być nieco irytujący. Ten przełącznik jest typem o małej mocy.

Zasilanie 5 VDC dla SST w trybie **ręcznym** jest uzyskiwane z klasycznego liniowego regulatora napięcia 7805 (IC1). Upewnij się, że używasz komponentu w obudowie typu TO220, ponieważ musi on zasilac maksymalnie cztery

serwa. Aby utrzymać małą wysokość konstrukcji, C1 i C2 można zamontować w pozycji leżącej.

Wreszcie dostępne jest złącze K3, typu IDC-6 w obudowie z kodowaniem, umożliwiające programowanie mikrokontrolera bez wyjmowania go z układu. Jest okablowane w taki sam sposób, jak złącze Arduino ICSP (programowanie szeregowo w układzie).

Kilka słów o oprogramowaniu

Do tego projektu użyłem Arduino IDE, co udowadnia, że można go wykorzystywać do skomplikowanych projektów. Najpierw sprawdziłem koncepcję z Arduino UNO i płytką prototypową. Gdy wszystko zadziałało zgodnie z założeniami, załadowałem program do urządzenia końcowego poprzez złącze ICSP. Należy pamiętać, że podczas programowania mikrokontrolera poprzez złącze ICSP K3 nie można nic podłączać do złącza K1 i K2, ponieważ współdzielą one sygnał z K3.

Wszystkie biblioteki potrzebne do projektu znajdują się w pliku do pobrania [3] z wyjątkiem biblioteki Servo, która jest częścią Arduino IDE. Projekt jest podzielony na trzy pliki: *Display.ino*, *Inputs.ino* i główny plik szkicu. Nazwy plików dość dobrze wyjaśniają, co zawierają. Główny plik szkicu zawiera funkcje `setup()` i `loop()`, ale także procedury przerwań używane do pomiaru czasu impulsów wejściowych. Zmiany poziomu sygnałów wejściowych wywołują przerwania związane

Wykaz elementów:

Rezystory: (THT, 0,25 W, 5%)

R1, R3: 1 kΩ
R2, R9, R10: 10 kΩ
R4, R5, R6: 10 kΩ 1%
R8: 22 Ω
P1, P2, P3, P4: 10 kΩ, liniowe, pionowe

Kondensatory: (THT)

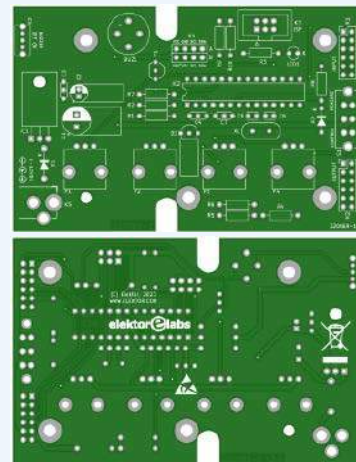
C1: 100 μF, 35 V, średnica 8 mm, raster 3,5 mm
C2: 10 μF, 16 V, średnica 5 mm, raster 2 mm
C3, C4, C7: 100 nF, raster 2,5 mm lub 5 mm
C5, C6: 22 pF, raster 2,5 mm

Półprzewodniki: (THT)

D1: 1N5817 dioda Schottky'ego
D2: LM385Z-2.5, TO-92
D3: BZX79 – C5V1
IC1: stabilizator 7805, TO220
IC2: mikrokontroler ATmega328-P
LED1: LED, 3 mm, czerwona
T1: 2N7000, TO-92

Różne: (THT)

BUZ1: Brzęczyk piezoo z oscylatorem, średnica 12 mm, raster 6,5 mm lub 7,6 mm
K1, K2: listwa kołkowa 2×6, kątowna, raster 2,54 mm
K3: Wtyk IDC, do druku, prosty (np. AVT Z-WS6G)
K4*: gniazdo jednorzędowe 1×4, raster 2,54 mm (wlutowane w pozycji A lub B w zależności od wyświetlacza)
K5: Gniazdo zasilania, beczka
S1: przeł. suwakowy DPDT, kątowny (np. MFS201N-16-Z, Mouser)
S2: przełącznik suwakowy SPDT, kątowny (np. OS102011MA1, Mouser)
X1: rezonator kwarcowy, 16 MHz, HC-49S (niskoprofilowy)
Wyświetlacz OLED, 0,96", 128×64 piksele, 4-pin I²C (np. BOTLAND MOD-08866)
Płytką PCB Elektora 220069-1
Sugerowana obudowa: Hammond 1593N (dostępne w SOSELECRONIC)
* = patrz tekst



ze zmianą stanu pinów. W ten sposób można użyć funkcji `micros()` do pomiaru długości impulsów. Zmierzone szerokości impulsów są kopiowane na wyjście i przekazywane do funkcji wyświetlacza w celu zobrazowania. **Rysunek 9** przedstawia ogólny zarys przepływu sygnałów wewnątrz mikrokontrolera.

Korzystanie z biblioteki OLED_I2C

Jak wspomniano wcześniej, wyświetlacz OLED jest sterowany za pośrednictwem magistrali I²C. W ATmega328P magistrala ta jest współdzielona z przetwornikiem ADC. PC4 może być albo pinem SDA, albo wejściem analogowym A4. Aby obejść ten problem,

magistrala I²C jest emulowana programowo. Robi to biblioteka `OLED_I2C` firmy Rinky-Dink Electronics [2], która pozwala nam przypisać dowolny port GPIO do magistrali I²C.

Musiałem nieco zmodyfikować plik `OLED_I2C.cpp` tej biblioteki, ponieważ niepoprawna dyrektywa `#include` uniemożliwiała jej kompilację. Zastąp linię 25:

```
#include "hardware/avr/HW_AVR.h"
przez
#include "HW_AVR.h"
```

Ponadto w pliku `HW_AVR.h` w funkcji `OLED::update()` skomentowałem linię 24, jak poniżej:

```
// noInterrupts();
```

Jest to konieczne, ponieważ biblioteka Servo wykorzystuje przerwania do generowania sygnałów PWM 50 Hz na wyjściach SST. Muszą one także pozostać włączone do pomiaru przychodzących impulsów.

Przed włączeniem zasilania

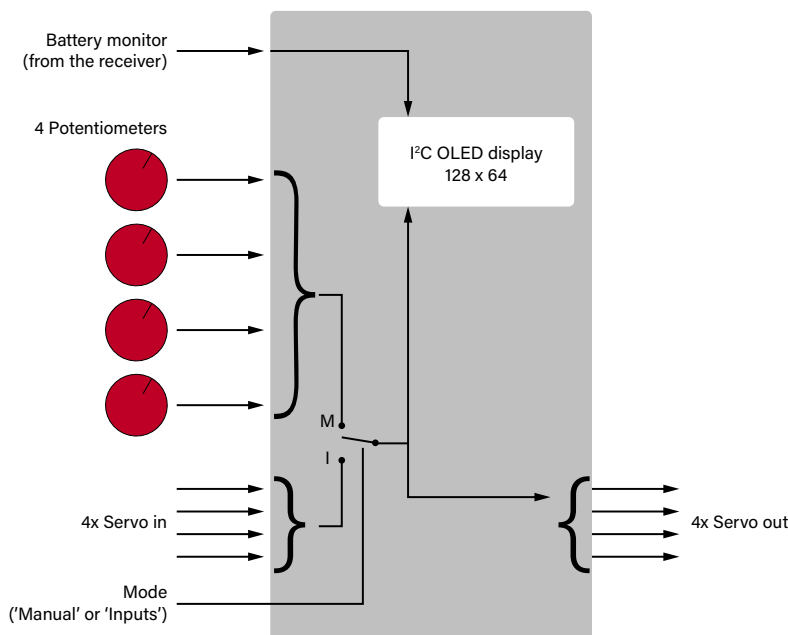
Jeżeli po włączeniu zasilania potencjometr P1 znajduje się w pozycji minimalnej (obrócony w lewo), na ekranie wyświetlana jest biała ramka. Ułatwia to ustawienie wycięcia w obudowie pod ekran i umożliwia poprawienie otworu w razie potrzeby. Nie każdy ma w domu maszynę CNC. Obróć P1 do środka, aby powrócić do normalnej pracy.

Podobnie, jeśli P4 znajduje się w pozycji minimalnej po włączeniu zasilania, wyświetlacz będzie wyświetlał w trybie **stos**. Jeśli po włączeniu zasilania P4 znajduje się w maksymalnej pozycji, aktywowany jest tryb wyświetlania **kwadrat**. W tym drugim przypadku oznacza to, że silnik podłączony do wyjścia 4 będzie miał pełną prędkość podczas rozruchu (z powodu położenia P4).

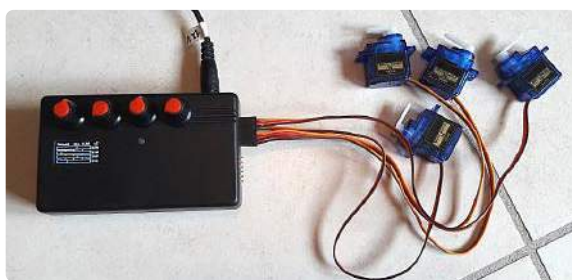
Czas wystartować

Projektowanie oraz rozwijanie testera było świetną zabawą i jest on przydatnym narzędziem, jeśli interesujesz się dronami czy quadkoopterami lub ogólnie modelami sterowanymi radiowo. Projekt ten nie tylko pomógł mi zwiększyć moje umiejętności programowania, ale teraz mogę wreszcie zabrać ze sobą mojego drona, który niecierpliwie czeka na start! ■

Olivier Croiset (Francja)



Rysunek 9. Schemat blokowy testera pokazuje przepływ sygnałów wewnątrz mikrokontrolera



Rysunek 10. Prototyp wykonany przez autora sterujący czterema serwami w trybie ręcznym



Rysunek 11. Prototyp wykonany w Laboratoriach Elektora zamontowany w swojej obudowie

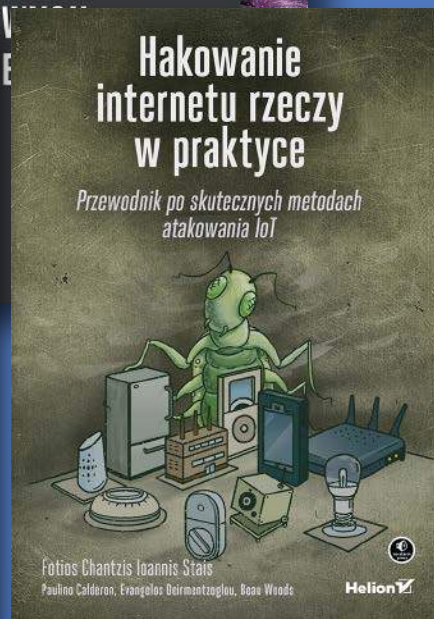
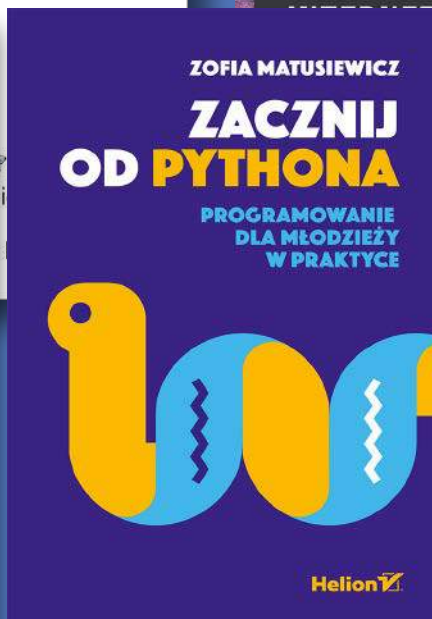
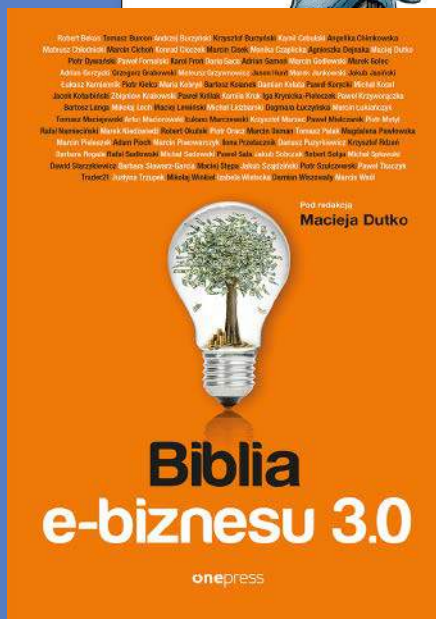
Linki internetowe:

- [1] Strona internetowa Joopa Brokkinga o dronach: http://www.brokking.net/ymfc-a_main.html
- [2] Biblioteka `OLED_I2C`: <http://www.rinkydinkelectronics.com/index.php>
- [3] Pobieranie projektów w Elektor Labs: <https://www.elektormagazine.com/labs/super-servo-tester>

Pytania lub komentarze?

Jeśli masz pytania techniczne, wyślij e-mail do autora na adres albot.dumitru@hsrw.org, do zespołu redakcyjnego Elektora na adres editor@elektor.com lub z redakcją EdW redakcja@elportal.pl

KSIĄŻKI W ULUBIONYM KIOSKU Z RABATEM DO 30%



Zobacz pełną ofertę! PONAD 500 TYTUŁÓW

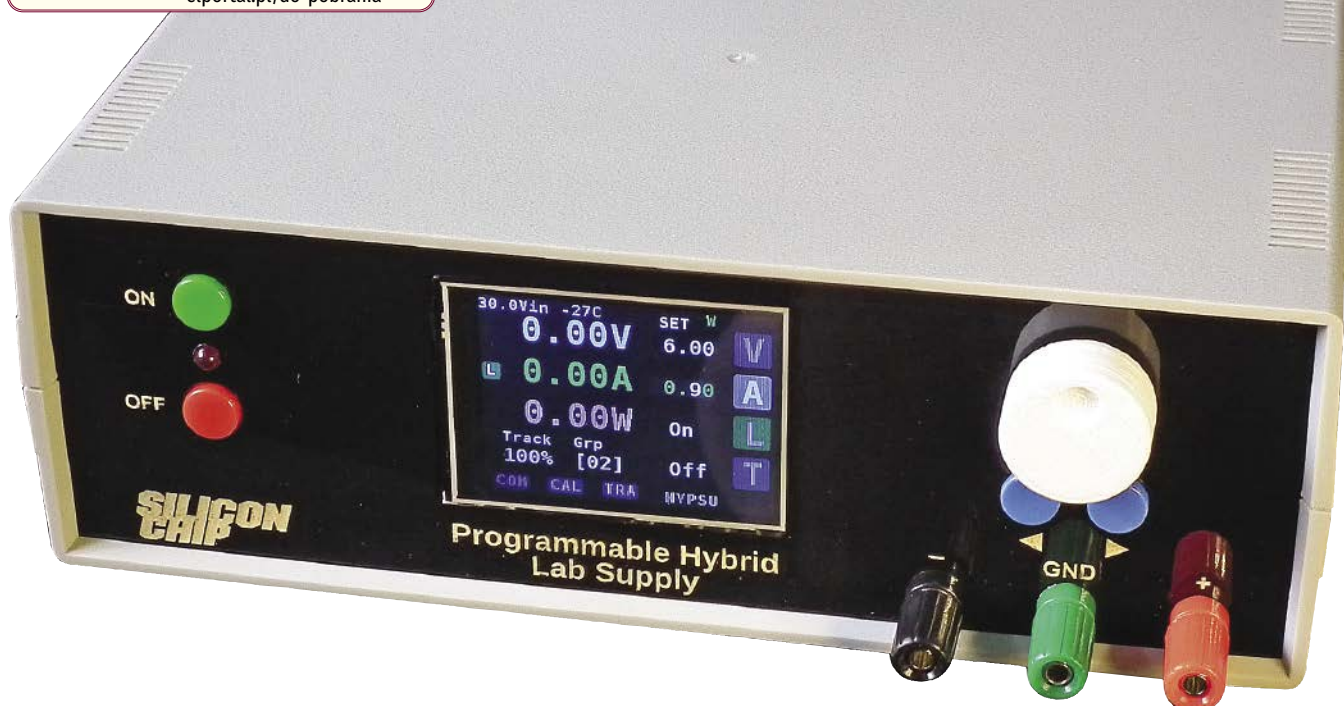
Zamów wygodnie na www.UlubionyKiosk.pl

epresa.pl 40635921f8



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/dhstf>

Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania



Programowany, hybrydowy zasilacz laboratoryjny ze sterowaniem Wi-Fi, część 2

Nasz nowy zasilacz laboratoryjny dostarcza napięcie w zakresie 0...27 V, o natężeniu do 5 A przy napięciu do 16 V i niewiele mniejszym (natężeniu) przy napięciu powyżej 16 V. Może być zdalnie sterowany przez Wi-Fi. Można nawet zbudować kilka sztuk takiego zasilacza, a zbudowane zasilacze połączyć szeregowo lub równolegle i zdalnie nimi zarządzać. Po opisanu w zeszłym miesiącu konfiguracji, zasady działania i schematu obwodów, ten kolejny artykuł pokazuje, jak zmontować dwie płytki drukowane i podłączyć starannie wszystko w skromnej, niewielkiej plastikowej obudowie przyrządu.

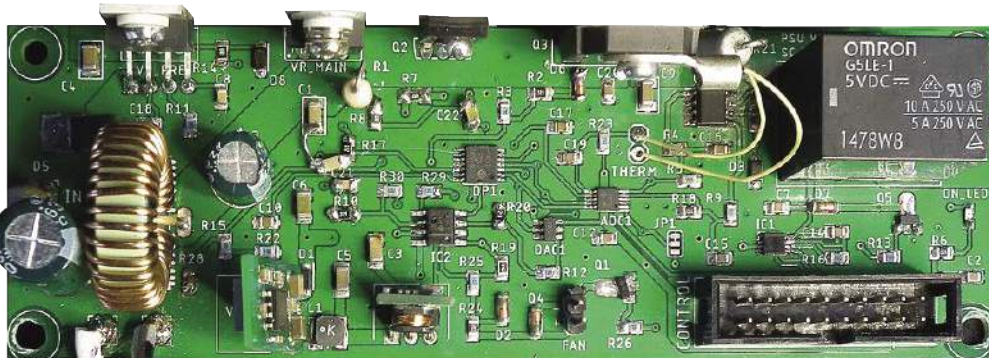
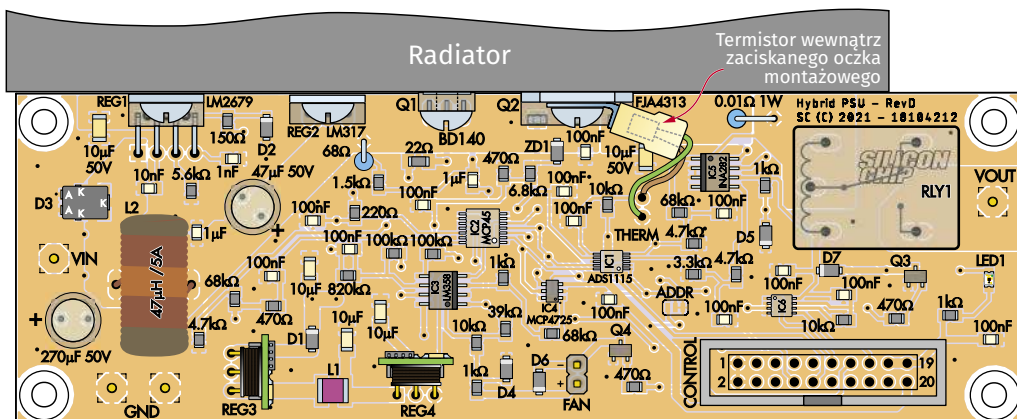
Jak już wcześniej wyjaśniliśmy, zasilacz ten wykorzystuje trzystopniowy układ hybrydowy, z dwoma zasilaczami impulsowymi i końcowym stopniem liniowym. Taka konstrukcja zapewnia doskonałą wydajność i sprawia, że całość jest kompaktowa i lekka, a jednocześnie oferuje bardzo dobrą sprawność.

Nasz zasilacz ma wiele przydatnych funkcji, takich jak łagodny rozruch i szybki czas ustalania napięcia z minimalnym przeregulowaniem.

Dzięki tym cechom, a także możliwości programowania, może on generować impulsy o zadanej mocy lub sekwencje napięć w celu

Cechy i specyfikacja płytki sterującej:

- Dwurdzeniowy, 32-bitowy mikrokontroler ESP-32 240 MHz
- Wbudowany kolorowy ekran dotykowy LCD o przekątnej 2,8 cala lub opcjonalnie 3,5 cala
- 520 kB pamięci RAM, 4 MB pamięci FLASH EEPROM
- Gniazda kart: pełnowymiarowe SD i mikro SD
- Interfejs dotykowy oraz odłączane przełączniki, dioda LED i enkoder obrotowy
- 20-pinowe złącze rozszerzeń z 2×I²C, SPI, 2×DAC, 2×ADC, komunikacją szeregową i GPIO
- Możliwość użycia maksymalnie 17 styków GPIO/PWM
- Wi-Fi (802.11 b/g/n) o przepustowości 150 Mb/s
- Obsługa Bluetooth i BLE (opcjonalnie)
- Port szeregowy USB
- Funkcje serwera WWW i klienta WWW
- Aktualizacja oprogramowania over-the-air (OTA) lub przez USB



Rysunek 6. Wszystkie komponenty montuje się na górnej stronie płytki regulatora w podanych miejscach. Generalnie najłatwiej jest zamontować wszystkie elementy SMD przed lutowaniem części do montażu przewlekanych. Podspęły, które mocuje się do radiatora, najlepiej pozostawić wzdłuż górnej krawędzi płytki, do czasu przestawienia podstawowych funkcji. Należy zauważyć, że dioda SMD D3 ma dwa zaciski anodowe i trzy zaciski katodowe, z których dwa znajdują się po bokach.

Errata: REG4 został nieprawidłowo wymieniony na liście części w zeszłym miesiącu jako VXO7803, podczas gdy powinna to być wersja 5V oznaczona jako VXO7805. Jeśli został zakupiony, część z przystrojką „7803” nadal będzie działać. Ponadto IC4 powinien być MCP4725A0T-E/CH. (Podano poprawne nazwy w spisie części) Końcówki bazy i emitera tranzystorów Q3 i Q4 są na PCB regulatora wzajemnie zamienione, dlatego tranzystory MUSZĄ być przylutowane „plecami” w stronę PCB, odwrotnie w stosunku do schematu montażowego. Mało starannie przylutowane tranzystory Q3 i Q4 w tej „nadmierzającej” pozycji widać na fotografii gotowej płytki, przy okazji z niewłaściwymi opisami Q1 i Q5.

testowania, jak zasilane urządzenia radzą sobie ze stanami przejściowymi.

Zasilacz impulsowy AC-DC jest on generować impulsy o zadanej mocy lub sekwencję napięć, ale pozostałe dwa moduły w urządzeniu muszą zostać zmontowane, zanim całość będzie można umieścić w obudowie i połączyć ze sobą. Przejdźmy więc do budowy tych dwóch płytek.

Budowa

Pierwszym krokiem jest montaż płytek. Rysunek 6 przedstawia schemat montażowy PCB płytki regulatora, natomiast rysunek 7 przedstawia schemat montażowy płytki sterowania.

Wszystkie części na płytce regulatora są montowane z jednej strony, a większość z nich jest montowana powierzchniowo. Płytką sterującą ma komponenty po obu stronach, ale tylko kilka

elementów SMD do zamontowania po tej samej stronie. Najlepiej jest najpierw przylutować wszystkie elementy SMD, a następnie przejść do montażu elementów przewlekanych.

Jeśli masz piec do lutowania rozpliwowego (lub jesteś w stanie zrobić własny! Zobacz, jak to zrobić w wydaniu kwiecień/maj 2020 Silicon Chip-a – siliconchip.com.au/series/343, lub wersja po polsku w wydaniu EdW maj/czerwiec 2023 r.), możesz przylutować wszystkie elementy SMD jednocześnie. Ręcznie nakładasz pastę lutowniczą na wszystkie pady lutownicze SMD, następnie ostrożnie rozmieszczasz podspęły na płytce PCB, zgodnie z rysunkiem montażowym, a na koniec poddajesz obie płytki cyklowi lutowania rozpliwowego.

Po ostygnięciu musisz dokładnie sprawdzić wszystkie układy scalone, aby upewnić się,

że nie ma mostków między wyprowadzeniami lub niepolutowanych styków.

Niepolutowane styki możesz naprawić, dodając odrobinę pasty topnikowej, a następnie odrobinę lutu przy pomocy lutownicy z cienkim grotem. Zmostkowane styki możesz rozdzielić, dodając odrobinę pasty topnikowej, a następnie nakładając i przyciskając rozgrzanym grotem lutownicy miedzianą plecionkę lutowniczą, następnie odciągając ją, gdy tylko nadmiar lutu zostanie wchłonięty do plecionki.

Oczywiście użyte części mogą być lutowane ręcznie, a jedynymi, które smogą okazać się nieco trudnymi w montażu, są IC1, IC2 i IC6 na płytce regulatora.

Reszta montażu powinna być prosta, jednak należy uważać na części o zdefiniowanej

Opcje budowy i montażu:

Oba panele boczne płytki sterującej można oddzielić, zapewniając elastyczność układu.

Wyświetlacz LCD o przekątnej 2,8 cala można zamienić na typ 3,5 cala, o ile jest na to miejsce (potrzebna będzie większa obudowa niż ta podana w specyfikacji).

Jeśli to zrobisz, upewnij się, że zaprogramowałeś układ za pomocą alternatywnego pliku binarnego, ponieważ 3,5-calowy wyświetlacz LCD ma inny sterownik niż typ 2,8-calowy.

Można użyć bardziej kompaktowego zasilacza impulsowego AC-DC 75 W (np. MeanWell LRS-75-24), co zmniejszyłoby ogólne wytwarzanie ciepła, chociaż ograniczyłoby również maksymalny prąd wyjściowy.

Podczas gdy można zbudować dwa oddzielne zasilacze i połączyć je razem jako zasilacze współpracujące, możliwe byłoby

również podłączenie dwóch płytek regulatora napięcia do pojedynczej płytki sterującej, aby stworzyć uniwersalny zasilacz regulowany, który można również skonfigurować tak, aby zapewniał dwa razy większy prąd (z wyjściami połączonymi równolegle) lub dwa razy większe napięcie (wyjścia połączone szeregowo).

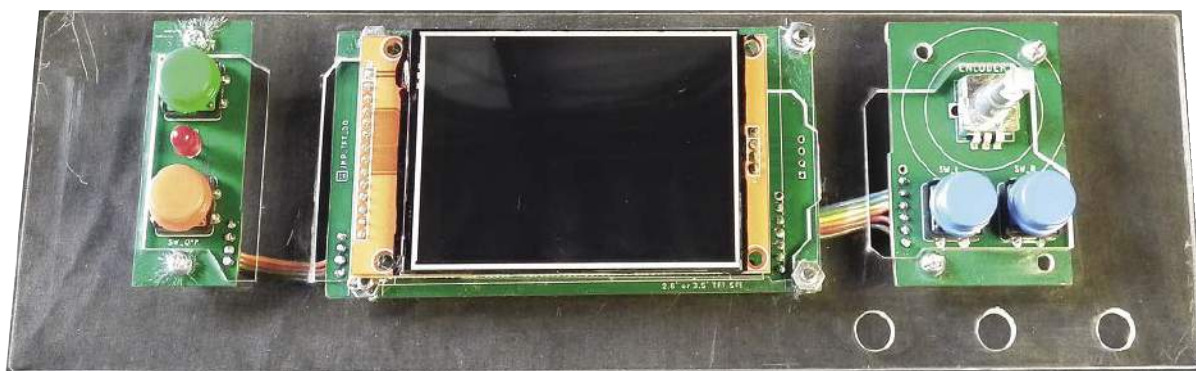
Wymagałoby to dodatkowego izolatora, aby dwie płytki regulatora były niezależne napięciowo względem siebie (rozdzielone masy!), a także dwóch oddzielnych zasilaczy AC-DC. Ten dwukanałowy projekt będzie wymagał większej obudowy, takiej jak Jaycar Cat HB5556. Będzie również wymagać zmodyfikowanego oprogramowania.

Mamy nadzieję przedstawić w przyszłości zmiany wymagane dla tej opcji.



Po przyłutowaniu wszystkich elementów SMD przejdź do elementów z montażem przewlekany. Najlepiej zacząć od dwóch wtyków IDC Z-WS20; upewnij się, że są zorientowane tak, jak pokazano (mają wycięcie na klucz gniazda montowanego na końcu kabla połączeniowego).

Po zamontowaniu enkodera obrotowego po drugiej stronie płytki, pozostaje już



Prototyp wykorzystywał panel wsporczy z pleksiglasu do montażu na panelu przednim obudowy, aby uniknąć dodatkowych otworów montażowych. Podczas budowania go w sposób opisany w tym artykule, konieczne będzie użycie elementów dystansowych, aby zamontować płytkę sterowania bezpośrednio na panelu przednim

tylko montaż wyświetlacza LCD. W przypadku pozostawienia płytki w całości (bez separacji obszarów z enkoderem obrotowym i przyciskami) zaznaczonych szarym kolorem połączeń przewodowych nie trzeba realizować.

Wyrównanie wysokości wyświetlacza z przełącznikami ma zasadnicze znaczenie dla schludnego wyglądu panelu przedniego. Zapoznaj się z dolną częścią rysunku 7, aby zobaczyć, jak powinien wyglądać ostateczny montaż. Aby uzyskać dobry wynik, ustaw górną część wyświetlacza 2...3 mm niżej niż górne części przycisków przełączników.

Powinno to oznaczać, że ekran dotykowy będzie znajdował się 0...1 mm od powierzchni panelu, a przyciski powinny wystawać o około 1,5 mm. Długość styków lutowanych fabrycznie w wyświetlaczach jest różna, więc może być konieczne usunięcie istniejących wyprowadzeń i wlutowanie dłuższych, jeśli fabryczne są zbyt krótkie.

Przerwane linie pokazane na rysunku 7 wskazują, gdzie byłyby podłączone przewody, gdybyś rozciął płytkę wzdłuż szczelin, ale nie zalecamy tego robić, chyba, że masz konkretne plany zamontowania panelu sterowania w innej obudowie niż ta wybrana przez nas. Jak to wygląda w takiej sytuacji, możesz zobaczyć dalej na zdjęciach naszego prototypu.

Wykończenie płytki regulatora

Na płycie regulatora zamontuj wtyk zasilania wentylatora, a następnie pionowe rezystory osiowe i kondensatory elektrolityczne, przestrzegając oznaczeń biegunowości tych ostatnich. Postępuj podobnie z REG3 i REG4, upewniając się, że nie pomylisz dwóch różnych typów, ponieważ mają one odmienne wyprowadzenia. Następnie możesz zamontować przekładnik i cewkę toroidalną.

Kółki lutownicze pokazane dla VIN, GND i VOUT są opcjonalne. Istnieją zalety lutowania przewodów do kołków montażowych (łatwiej jest wykonać dobre połączenie i istnieje mniejsze prawdopodobieństwo awarii związanych z naprężeniami), ale można przylutować przewody bezpośrednio do płytki.

Pozostają tylko komponenty, które montujesz na radiatorze: REG1, REG2, Q1, Q2 i termistor NTC. Nie zapomnij odizolować radiatorów tych podzespołów i śrub montażowych od radiatora chłodzącego za pomocą silikonowych podkładek i tulejek.

Uruchomienie płytki sterującej

W pierwszym etapie wystarczy sam moduł ESP32 i kabel USB. Montaż modułu na płycie sterującej nastąpi później.

Zakładamy, że jesteś już zaznajomiony ze środowiskiem programistycznym Arduino. Jeśli nie masz jeszcze zainstalowanego Arduino IDE (zintegrowanego środowiska programistycznego), możesz je pobrać ze strony www.arduino.cc/en/software.

Będziesz musiał dodać obsługę płytki ESP32 do IDE, jeśli jeszcze tego nie zrobiłeś. Aby to zrobić, przejdź do File → Preferences i dodaj https://dl.espressif.com/dl/package_esp32_index.json do adresów URL Additional Boards Manager-a. Następnie otwórz Boards Manager-a (Tools → Board → Board Manager), wyszukaj ESP32 i kliknij „Install”.

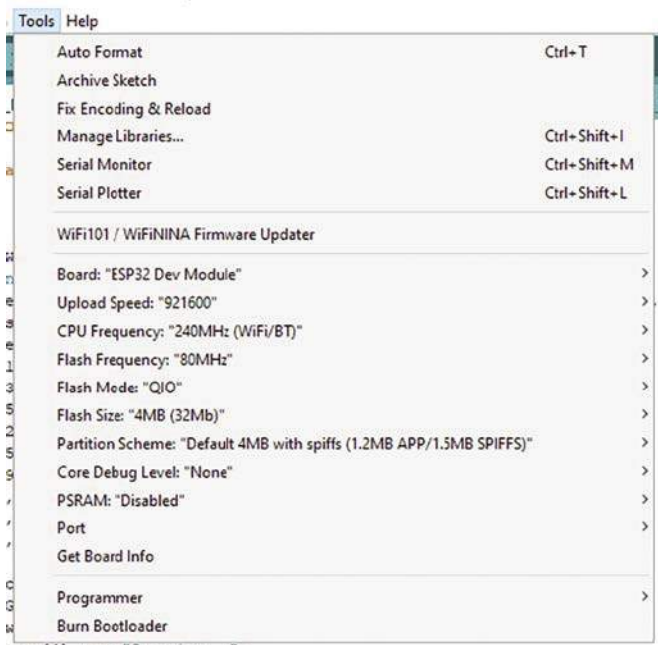
Spowoduje to skonfigurowanie środowiska programistycznego i dodanie do niego obszernej listy przykładowych programów. Ustaw płytkę na „ESP32 Dev Module” za pomocą menu (ekran 1). Pozostałe ustawienia można pozostawić jako domyślne.

Podłącz moduł ESP32 i wybierz nowy port komunikacyjny, który pojawi się w menu.

Aby sprawdzić, czy działa poprawnie, otwórz przykład Communication → ASCII Table i załaduj go (CTRL+U w Windows). Otwórz Serial Monitor, ustaw szybkość transmisji na 9600 bpd, a ekran powinien wypełnić się danymi wyjściowymi ASCII ze szkicu testowego.



Aby zapewnić lepszy układ panelu przedniego, płytkę sterowania w prototypie Silicon Chip-a została podzielona na trzy części i połączona tęczowymi kablami. Przedstawione tutaj rozwiązania montażowe wykorzystują kawałek przezroczystego pleksiglasu (na górze), który nie jest wymagany do ukończenia projektu wg opisu w tekście artykułów



Ekran 1. Po wybraniu właściwej płytki („Board”) w menu „Narzędzia Arduino IDE” („Arduino IDE Tools”), ustawienia powinny wyglądać następująco

Bezprzewodowe ładowanie oprogramowania

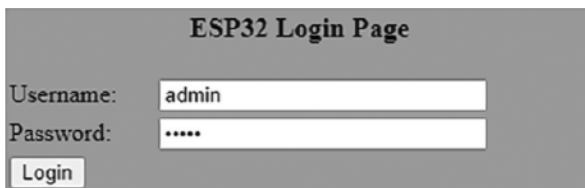
Aby zademonstrować inne możliwe zastosowania płytki sterującej, stworzyliśmy wersję aplikacji pogodowej Wi-Fi używanej jako program demonstracyjny dla modułu D1 Mini Backpack-a (październik 2020; siliconchip.com.au/Article/14599, opis po polsku: EdW z kwietnia 2023 r.). Jest to również dobry sposób na niezależne przetestowanie płytki sterującej.

```
#include <WiFi.h>
#include <WiFiClient.h>
#include <WebServer.h>
#include <ESPmDNS.h>
#include <Update.h>

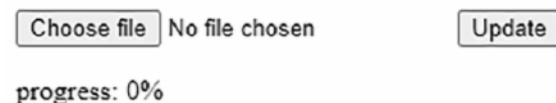
const char* host = "esp32";
const char* ssid = "PalHome";
const char* password = "AF18181F9C4516EE36AF9D79D3";

WebServer server(80);
```

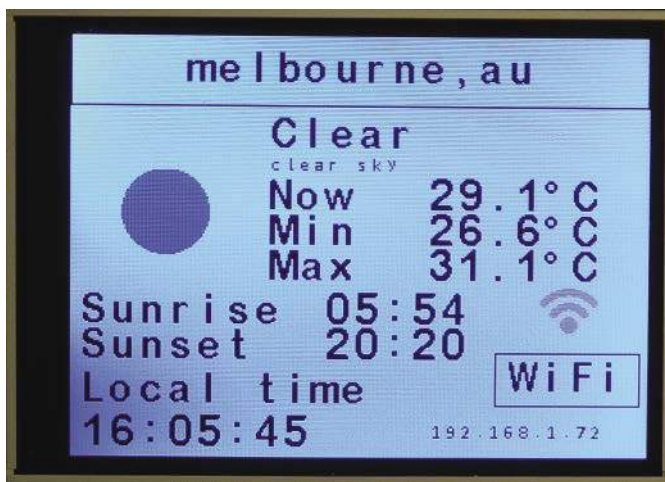
Rysunek 8. Aby przesać kod do ESP-32 przez Wi-Fi (aktualizacja OTA), musisz dodać swoje uwierzytelnienie sieciowe w górnej części programu, jak pokazano tutaj. Nazwę hosta można pozostawić bez zmian lub zmienić zgodnie z własnymi wymaganiami



Rysunek 9. Po wyświetleniu strony logowania modułu ESP-32 należy użyć domyślnych nazw użytkownika („admin”) i hasła („admin”). Nie ma potrzeby ich zmieniać, ponieważ są one używane tylko raz



Rysunek 10. Po zalogowaniu się na stronie OTA można wybrać plik, a następnie przesać go zdalnie do pamięci FLASH ESP-32 za pomocą przycisków „Wybierz plik” („Choose file”) i „Aktualizuj” („Update”)



Ekran 2. Jeśli Twój moduł ESP-32 został poprawnie zmontowany i zaprogramowany, po podłączeniu do sieci Wi-Fi powinien podawać lokalne dane pogodowe, jak pokazano tutaj. Przypisany adres IP znajduje się w prawym dolnym rogu

Udostępniliśmy do pobrania ze strony Silicon Chip-a plik ZIP, przystosowany do dwóch opcji ekranu dotykowego, wyświetlacza 2,8-calowego lub 3,5-calowego (można również pobrać plik ze strony siliconchip.com.au/link/ab72). Wersja 2,8 cala kończy się na -28.bin, podczas gdy druga wersja kończy się na -35.bin. Załaduj plik za pomocą procesu aktualizacji OTA (over-the-air) opisanego poniżej. Aplikacja Weather ma wbudowaną funkcję OTA, która upraszcza ładowanie kodu sterownika zasilacza.

Programowanie ESP32 w trybie over-the-air jest procesem dwuetapowym. Najpierw ładujemy prosty szkic z aktualizacją over-the-air (OTA) przez USB. Załaduj przykład ArduinoOTA (File → Examples → ArduinoOTA → OTAWebUpdater). Wprowadź dane uwierzytelniające Wi-Fi (SSID i hasło) w górnej części programu (rysunek 8).

Otwórz Serial Monitor i zmień prędkość transmisji na 115 200 bpdów. Zapisz szkic Arduino, ponieważ będziemy go używać ponownie. Skompiluj i załaduj szkic i zanotuj adres IP wyświetlany przez Serial Monitor.

Teraz możesz odłączyć moduł ESP-32 i podłączyć go do płytki sterującej, upewniając się, że jest zorientowany jak na poniższym zdjęciu. Podłączenie go w niewłaściwy sposób może być katastrofalne w skutkach! Nie podłączaj jeszcze płytki sterującej do płytki regulatora, ale upewnij się, że ekran dotykowy TFT jest zamontowany na płytce sterującej.

Zasil to złożenie, używając kabla USB lub (jeśli zamontowałeś CON1 i REG1) zasilania DC 9...12 V. Kabel USB nie musi być podłączony do komputera, choć może być.

Otwórz przeglądarkę internetową na komputerze i wpisz adres IP ESP32. Powinien zostać wyświetlony ekran logowania (rysunek 9). Nazwa użytkownika i hasło to „admin”. Nie ma sensu zmieniać ich na coś bezpieczniejszego, ponieważ będziemy używać tego szkicu tylko raz.

Po zalogowaniu wybierz pobrany plik oprogramowania za pomocą przycisku „Choose file” (rysunek 10), a następnie „Update”. Strona internetowa będzie śledzić postęp przesyłania; następnie, po krótkiej chwili, moduł ESP32 uruchomi się ponownie, otwierając aplikację pogodową (ekran 2).

Po sprawdzeniu, że płytka sterująca działa poprawnie, można załadować program zasilacza. Jest on częścią tego samego pakietu ZIP, który zawiera aplikację pogodową. Istnieje tylko jeden plik binarny dla programu zasilacza, ponieważ ten projekt został opracowany dla wyświetlacza 2,8 cala (ergo, użyj pliku o nazwie „Control_C_ADS_Touch_2_8_V2.ino.bin”).



Ekran 3. Główny ekran, który pojawia się po włączeniu. Aktualne napięcie, prąd i moc są pokazane po lewej stronie, a napięcie wejściowe i temperatura radiatora powyżej. Ustawienia napięcia i prądu znajdują się po prawej stronie, a przyciski do włączania/wyłączania ograniczania prądu i śledzenia poniżej. Status urządzenia wyświetlany jest w prawym górnym rogu ekranu

Po załadowaniu tego programu przy użyciu tej samej procedury aktualizacji OTA (lub przesłaniu go bezpośrednio przez USB), odłącz zasilanie DC (jeśli jest obecne) i podłącz kabel USB do komputera (zarówno w celu zasilania, jak i komunikacji).

Otwórz monitor szeregowy Arduino ustawiony na prędkość 115 200 bodów. Powinienesz zobaczyć kilka poleceń startowych, kończących się monitem „SCPI Command?”.

Jeśli wpiszesz „*IDN?” w polu poleceń i klikniesz „Send”, oprogramowanie powinno odpowiedzieć czymś w rodzaju „SiliconChip,PSU01,PS01-01,1.0,NONE”.

Nieco później omówimy konfigurację Wi-Fi i inne opcje konfiguracji zasilacza.

Obracanie i kalibracja ekranu

Niektóre ekrany TFT są dostarczane z obróconą o 180° pozycją startową ekranu dotykowego w stosunku do wyświetlacza. Jeśli ekran

dotykowy wydaje się nie działać, może to być przyczyną.

Spróbuj dotknąć ekranu w pobliżu legendy SET w prawym górnym rogu. Jeśli spowoduje to przejście do ekranu kalibracji, wystarczy dotknąć przycisku ROT na środku ekranu (ekran 4). Liczba pod nim powinna zmienić się z 3 na 1. Poczekaj, aż żółty wskaźnik [E] zgaśnie (po około 60 sekundach), a nowa wartość zostanie trwale zapisana w pamięci EEPROM ESP32.

Aby dokładnie zgrać ekran dotykowy z wyświetlaczem, dotknij przycisku TCH w lewym dolnym rogu ekranu kalibracji. Postępuj zgodnie z instrukcjami, dotykając każdego z dwóch symboli + sześć razy. Podobnie jak powyżej, wartości zostaną trwale zapisane po 60 sekundach.

Pobrane oprogramowanie zasilacza zawiera również instrukcje PDF dla dwóch płytek, z informacjami wykraczającymi poza te zawarte w naszych dwóch artykułach.



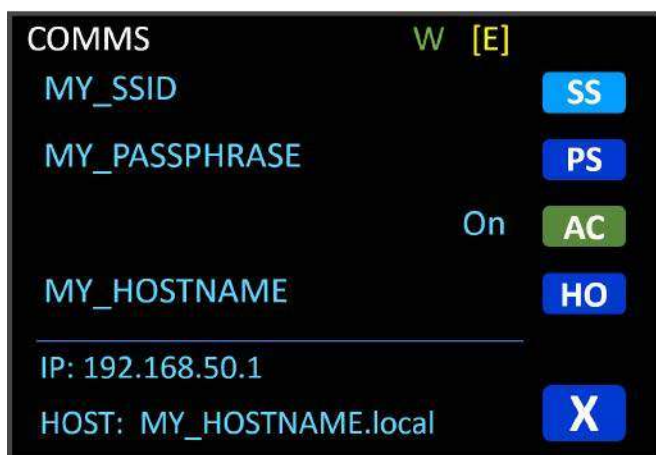
Ekran 4. Ekran kalibracji pokazuje odczyty napięcia i prądu urządzenia w lewym górnym rogu, z regulowanymi przesunięciami kalibracji po ich prawej stronie. Przyciski zapisywania i anulowania znajdują się w prawym dolnym rogu, przycisk obracania ekranu pośrodku, a przycisk menu kalibracji dotykowej w lewym dolnym rogu. (Dostęp do wszystkich pozycji menu uzyskuje się poprzez naciśnięcie przycisków, które w razie potrzeby pojawiają się na dole)

Konfiguracja sieci Wi-Fi

Po zaprogramowaniu płytki sterowania i włączeniu zasilania powinno pojawić się menu sterowania (ekran 3) z nałożonym zielonym polem. Program spróbuje połączyć się z lokalną siecią Wi-Fi LAN i po 10 sekundach wyłączy się, ponieważ nie przekazaliśmy mu jeszcze danych uwierzytelniających.

Następnie powinno nastąpić kolejne 10-sekundowe opóźnienie, podczas gdy program będzie szukał istniejącej sieci ESPINST. Wreszcie powinien niemal natychmiast stać się punktem dostępu dla sieci ESPINST. W tym momencie zielone pole powinno zniknąć, pozostawiając wyświetlone menu główne. Małe zielone „W” w prawym górnym rogu wskazuje, że Wi-Fi działa.

Przycisk „On” powinien zapalić diodę LED, a przycisk „Off” powinien ją ponownie wyłączyć. Dotknięcie dowolnego przycisku menu po prawej stronie wyświetlacza powinno podświetlić wartość ustawienia obok niego lub



Ekran 5. Ekran ustawień Wi-Fi umożliwia ustawienie nazwy hosta urządzenia, identyfikatora SSID sieci i hasła, a także pokazuje bieżący adres IP i nazwę hosta urządzenia. Ustawienie „AC” oznacza automatyczne połączenie



Ekran 6. Ekran śledzenia pozwala przypisać zasilacz do grupy śledzenia (GRP), a następnie ustawić, czy śledzi on napięcie, prąd lub obie wartości innych urządzeń w grupie

zmienić tryb funkcji. Jak opisano w zeszłym miesiącu, po wybraniu ustawienia na ekranie dotykowym, obroty enkodera powinny zmieniać wartość wybranej cyfry, a użyte przyciski „SW_L” i „SW_R” powinny przesunąć podświetloną cyfrę na ekranie w lewo lub w prawo.

Dalsze testowanie

Teraz nadszedł czas, aby wyłączyć zasilanie płytki sterującej i podłączyć ją do płytki regulatora za pomocą 20-żyłowego kabla taśmowego o długości około 10 cm, z gniazdami IDC na obu końcach.

Jeśli jeszcze nie przygotowałeś tego kabla, zrób to teraz, upewniając się, że wskaźnik styku 1 na każdej wtyczce IDC (zwykle trójkąt uformowany w plastiku na jednym końcu) wskazuje ten sam przewód w kablu. Szary kabel taśmowy ma zwykle jeden czerwony przewód wskazujący pin 1. Jeśli używasz kabla wielokolorowego, użyj czarnego przewodu (czarny = 0 w schemacie kolorów rezystora).

Upewnij się, że złącza IDC są zaciśnięte wystarczająco mocno, aby wszystkie ostrza całkowicie przebiły izolację kabla taśmowego. Zazwyczaj można to stwierdzić, ponieważ dwie (lub trzy) części każdego gniazda IDC będą całkowicie zlicowane i równoległe.

Częściowo zaciśnięte gniazda IDC zwykle mają przerwę na jednym lub obu końcach, widoczną po dokładnym sprawdzeniu. Jest to najczęstsza przyczyna nie działania kabli taśmowych.

Połącz ze sobą obie płytki. Nie powinno być możliwe ich nieprawidłowe połączenie ze względu na gniazda z kluczem. Mimo to, dla pewności, warto sprawdzić za pomocą omomierza lub testera ciągłości, czy szyny GND, 5 V i 3,3 V są prawidłowo podłączone na obu końcach. Sprawdź również, czy żadna z tych szyn nie jest zwarta z inną.

Teraz podłącz kabel USB z powrotem do modułu ESP32. Ponieważ obecnie nie ma innego źródła zasilania, jest to całkiem bezpieczne. Port USB może również zapewnić wystarczający prąd do przetestowania podstawowych funkcji zasilacza, innych niż wentylator.

Sprawdź, czy napięcia szyn 5 V i 3,3 V na płycie regulatora są prawidłowe. Katoda (paskowany koniec) diody D7 jest dogodnym punktem do pomiaru zasilania 5 V, podczas gdy styk złącza termistora najbliższy tranzystorów mocy powinien pokazywać 3,3 V.

Na ekranie monitora szeregowego Arduino, autotest po włączeniu zasilania (POST) powinien zgłosić, że wykryto trzy urządzenia I²C. Jeśli żadne się nie pojawi, najbardziej prawdopodobnym winowajcą jest mostek lutowniczy na stykach jednego z układów scalonych.

Dioda LED na płycie zasilacza, gdy są obsługiwane przełączniki włączania/wyłączania wyjścia, powinna podążać za diodą na płycie sterowania. Awaria w tym przypadku jest najprawdopodobniej spowodowana odwrotnym włączeniem diody LED lub mostkiem lutowniczym na wyprowadzeniach układu scalonego IC6.

Kilka sekund po włączeniu prąd wyjściowy na ekranie panelu sterowania powinien wskazywać 0 A, po zakończeniu funkcji automatycznego zerowania. Napięcia wejściowe i wyjściowe powinny wynosić mniej niż jeden wolt, ponieważ istnieją do tych szyn pewne ścieżki prądowe z zasilaczy 3,3 V i 5 V.

Odczyt temperatury będzie poza zakresem do czasu wlutowania termistora. Ustaw napięcie wyjściowe na 2,0 V, ponieważ będzie ono potrzebne do kalibracji.

Następnie odłącz panel sterowania od płyty regulatora i podłącz zasilanie 7...12 V DC między zaciskiem VIN i GND na płycie drukowanej regulatora.

Sprawdź napięcia wyjściowe na regulatorach +5 V i -5 V. Nominalna szyna +5 V powinna wskazywać około 4,5 V, ponieważ dioda zabezpieczająca przed odwrotnym podłączeniem zasilania (D1) jest połączona szeregowo z tym zasilaniem.

Niskoprądowe testowanie samego regulatora można bezpiecznie kontynuować bez regulatora. Używając tego samego zasilania 7...12 V DC podłączonego jak poprzednio do VIN i przy odłączonej płycie sterowania, sprawdź napięcie wyjściowe REG1, które pojawia się na ZD1. Powinno ono wynosić między 3,6 V a VIN, z wartością o około 3,6 V wyższą niż napięcie na wyjściu REG2 (środkowy pin).

Wyłącz zasilanie, podłącz ponownie płytkę sterującą i włącz ponownie zasilanie. Przekaznik powinien teraz włączać się i wyłączać, wraz z diodą LED, po naciśnięciu przełączników panelu sterowania. Ustaw napięcie wyjściowe na 2 V, włącz wyjście i sprawdź napięcia na Vout (2 V) i Vpre (około 5,6 V). Wyreguluj napięcie wyjściowe i sprawdź, czy napięcie Vpre wynosi około (Vout + 3,6 V).

Następnie podłącz na wyjściu rezystor 47 Ω/1 W (lub o nieco większej wartości). Ustaw napięcie na 5 V i upewnij się, że prąd wyjściowy wynosi około 100 mA.

Przed wyłączeniem panelu sterowania ustaw napięcie wyjściowe ponownie na 2 V i poczekaj, aż wartość ta zostanie zapisana po około 60 sekundach w pamięci FLASH, gdy wskaźnik [E] w prawym górnym rogu ekranu zgaśnie. Następnie odłącz kabel USB. Pozostają nam jeszcze tylko testy po zmontowaniu całego zasilacza.

Przygotowanie panelu

Wywierć i wytnij otwory w przednim plastikowym panelu obudowy przyrządu, jak

Zdalne sterowanie przez SCPI

Protokół Standard Commands for Programmable Instruments (SCPI) wykorzystywany w tym projekcie został opracowany na początku lat 90., w celu zapewnienia standardowej składni i struktury poleceń dla programowalnych przyrządów, od zasilaczy po oscyloskopy i nie tylko.

Został zaprojektowany jako protokół master-slave (urządzenie nadrzędne-podrzędne), z komputerem sterującym zawsze jako master.

Początkowo został zaimplementowany na magistrali GPIB (IEEE 488), ale obecnie powszechnie wykorzystywane są inne kanały komunikacyjne, takie jak transmisja szeregową (w tym szeregowy interfejs USB) i TCP (protokół internetowy).

Polecenia SCPI składają się ze słów kluczowych z uwzględnieniem wielkości liter, oddzielonych dwukropkami. Polecenia zakończone znakiem zapytania są zapytaniami, a urządzenie zwraca wartość lub zestaw wartości dla każdego zapytania. Każde słowo kluczowe może mieć powiązane z nim parametry, np:

„SET:VOLTage 350 mV”

„MEASure:VOLTage?”

Parametry mogą być liczbami całkowitymi,

zmiennoprzecinkowymi lub ciągami znaków, w zależności od polecenia.

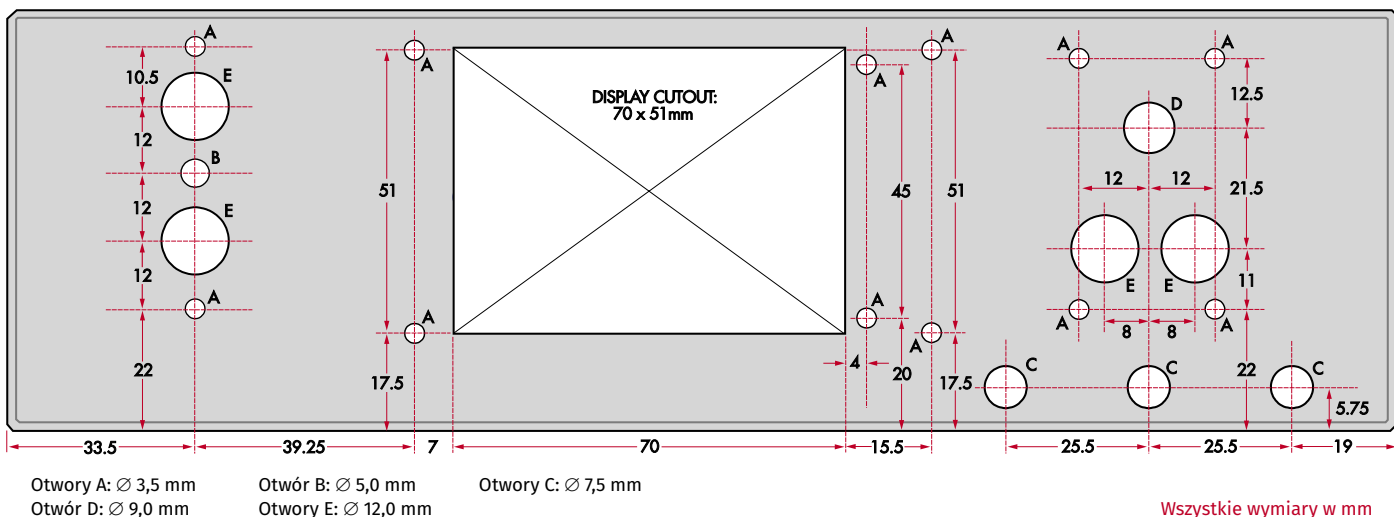
Po poleceniach numerycznych mogą występować jednostki, takie jak V, mV, A lub mA. Pełne SCPI rozumie wszystkie mnożniki od yotta (10²⁴) do yocto (10⁻²⁴).

Ten przyrząd akceptuje tylko „gołe” jednostki lub jednostki miliamperowe, unikając problemów związanych z ustawianiem megaamperów, gdy zamierzano użyć miliamperów!

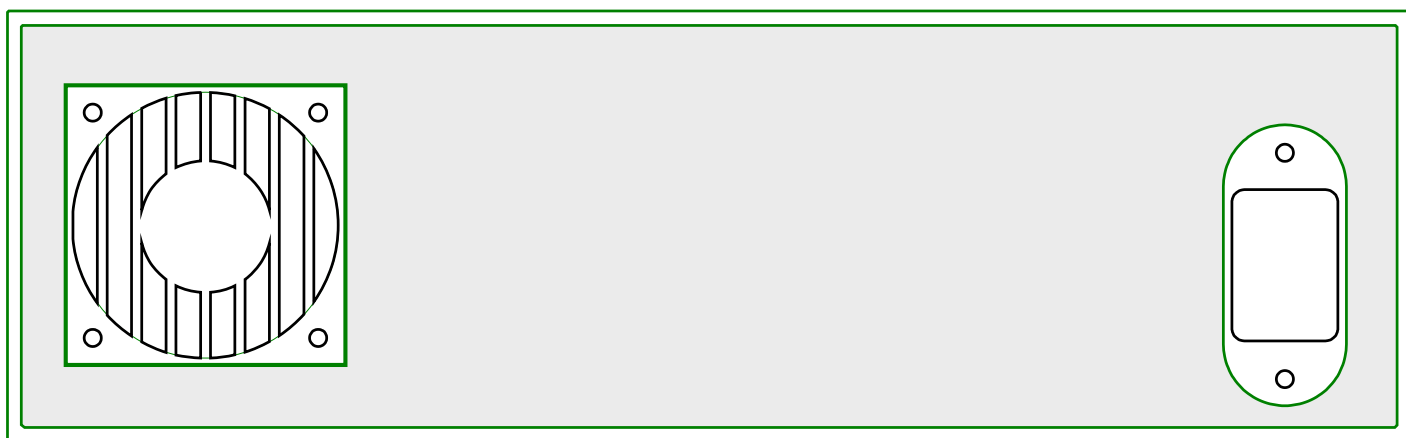
Każde polecenie, takie jak „MEASure”, może być wydane w pełnej formie lub w formie skrótu, który jest zawsze pisany wielkimi literami i prawie zawsze składa się z czterech znaków. Tak więc „:MEAS:VOLT?” jest równoważne „:MEASure:VOLTage?”.

Fundacja IVI, która jest następcą konsorcjum SCPI non-profit, posiada stronę internetową pod adresem www.ivifoundation.org/specifications/default.aspx, z wyczerpującą dokumentacją na temat SCPI i niedawno opracowanych i bardziej elastycznych protokołów komunikacyjnych przyrządów, takich jak VISA i VXI.

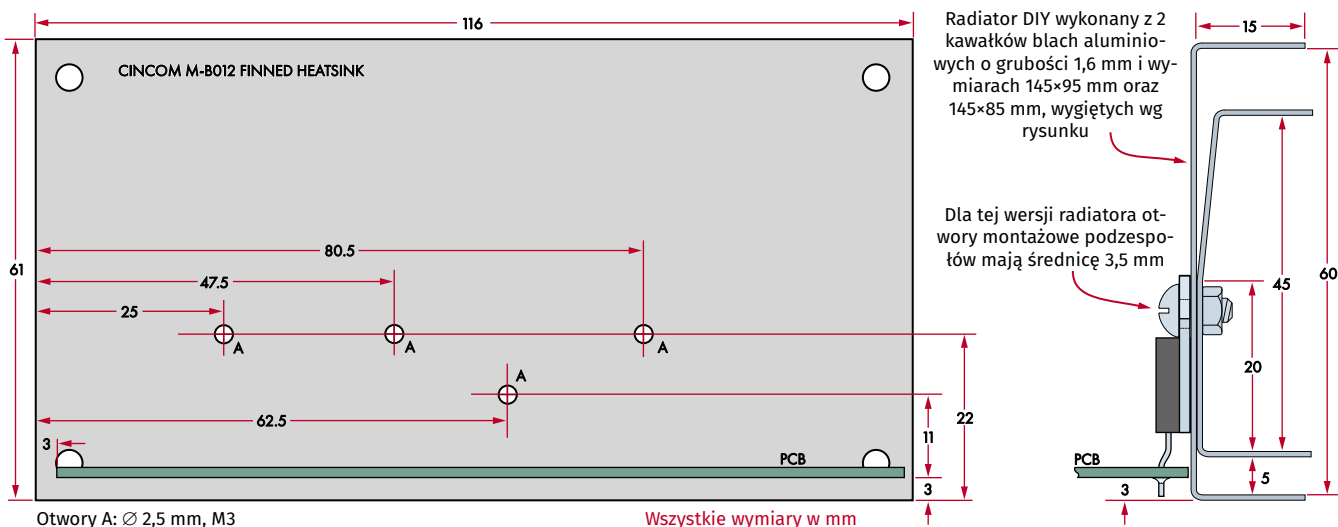
Polecenia SCPI używane dla tego programowanego zasilacza są szczegółowo opisane w podręczniku dołączonym do plików do pobrania dla tego projektu.



Rysunek 11. Sztampon wiercenia i frezowania panelu przedniego przeskalowany w rozmiarze 75% wielkości rzeczywistej. Prostokątny otwór na ekran dotykowy można wykonać, kopiując ten schemat, dołączając go tymczasowo do panelu, a następnie wierząc serię małych (2...3 mm) otworów tuż wewnątrz konturu. Użyj narzędzia tnącego, takiego jak frez obrotowy lub, w razie potrzeby, pary nożyc bocznych, aby połączyć wszystkie otwory razem, aż panel wypadnie, a następnie spiłuj krawędzie na gładko, aż ekran dotykowy będzie pasował



Rysunek 12. Wiercenie i frezowanie tylnego panelu jest stosunkowo proste, ponieważ potrzebne są tylko otwory do zamontowania złącza wejściowego zasilania sieciowego IEC i wentylatora chłodzącego. Choć pokazałismy szczeliny na wylot wentylatora, znacznie łatwiej byłoby po prostu wywiercić serię otworów o średnicy 5 mm w pokazanym obszarze. Nie należy robić większych otworów, aby nie można było włożyć małych palców. Dobrym pomysłem jest użycie osłony wentylatora



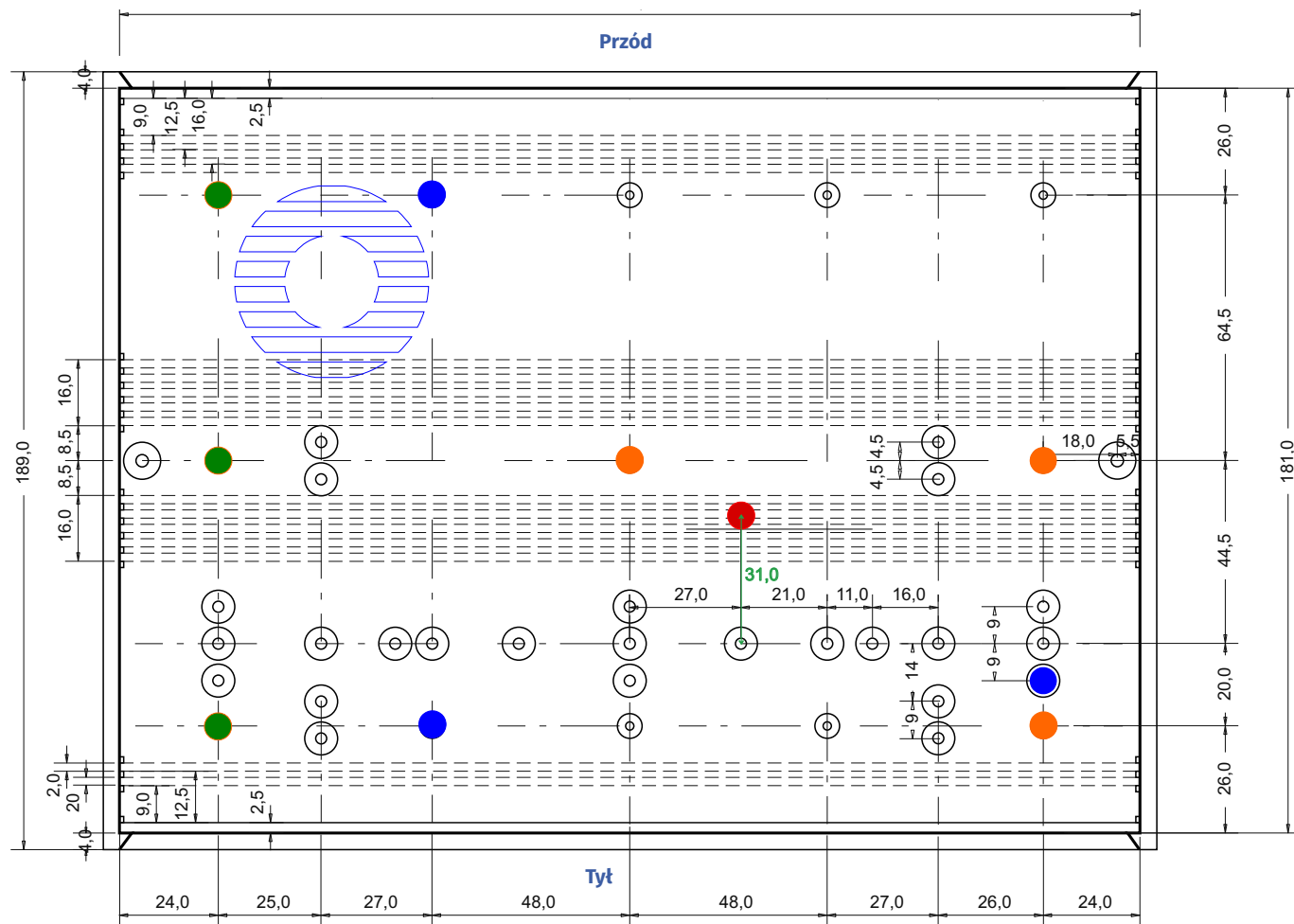
Rysunek 13. Szczegóły wiercenia radiatora oraz plan wersji radiatora DIY wykonanej samodzielnie z blachy aluminiowej (po prawej stronie). W radiatorze komercyjnym wszystkie otwory mają średnicę 2,5 mm w celu nagwintowania do M3. W wersji DIY wywierć otwory 3,5 mm w złożonych blachach. Nie należy wiercić otworów montażowych w podstawie radiatora, dopóki komponenty nie zostaną do niego przymocowane. Otwory można następnie rozplanować, wierząc je w dolnej części obudowy. Radiator DIY wykorzystuje dwa kawałki blachy aluminiowej o grubości 1,6 mm i wymiarach 145×95 mm oraz 145×85 mm. Dolna krawędź radiatora znajduje się 3 mm poniżej dolnej krawędzi płytki drukowanej

Niebieskie. Użyj do montażu śrubami samogwintującymi 4G×6

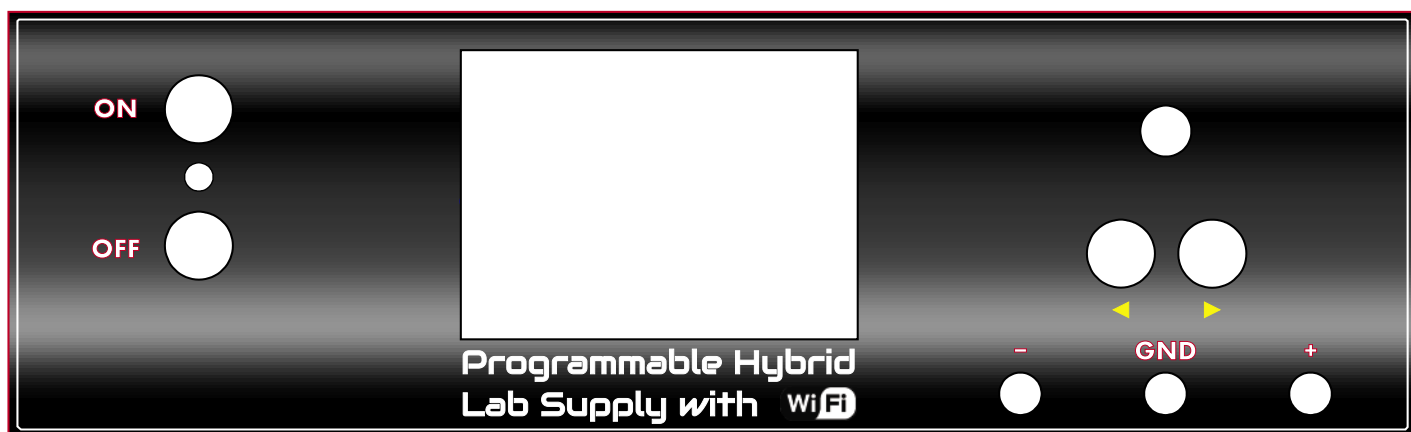
Zielone. Dopasuj radiator i wierć przez otwory 3,5 mm dla śrub mocujących M3

Pomarańczowe. Dopasuj

Czerwony. Nowy otwór montażowy pod kotek dystansowy M3×12



Rysunek 14. Pokazuje, jak przygotować spód obudowy. Otwory montażowe są zaznaczone na zielono i należy je przewiercić przez słupki na wylot (średnica 3,5 mm) i pogłębić od spodu. Górną część tych słupków należy spiłować do wysokości dolnych słupków. Należy wywiercić tylko dwa z zielonych otworów, w zależności od tego, czy używany jest radiator komercyjny czy DIY. Pomarańczowe występy również należy spiłować, ale nie wiercić. Płytke PCB montuje się bezpośrednio na dwóch niebieskich wystęпах za pomocą wkrętów samogwintujących 4G×9 z okrągłym łbem, podobnie jak konwerter AC-DC MeanWell. Drugi otwór montażowy jest wymagany dla zasilacza MeanWell. Znajduje się on 51 mm przed niebieskim otworem montażowym i 73,5 mm do wewnątrz



Rysunek 15. Grafika panelu przedniego urządzenia w 75% naturalnej wielkości – innymi słowy, przed skopiowaniem należy ją powiększyć do 133% wydrukowanego rozmiaru. Alternatywnie, wzór panelu można pobrać w pełnym rozmiarze ze strony internetowej Silicon Chip-a, aby wydrukować wysokiej jakości wersję do naklejenia na przyrząd



Ukończony prototypowy moduł regulatora przymocowany do radiatora, który został wykorzystany ponownie z innego projektu (stąd dodatkowe otwory). Zwróć uwagę na podkładki izolacyjne pod obudowami podzespołów i plastikowe tuleje pod śrubami montażowymi. Jeśli używasz podkładek mikowych, dodaj pastę termoprzewodzącą po obu stronach. Przed włączeniem zasilania należy sprawdzić, czy istnieje izolacja między każdym radiatorem a odstąpionym metalnym na radiatorze

pokazano na rysunku 11. Otwory powinny być pokrywane się z częściami na płytce sterowania (szczegóły montażu znajdują się również na dole rysunku 7).

Otwór „B” po lewej stronie jest przeznaczony dla diody LED, podczas gdy 12 otworów oznaczonych „A” odpowiada śrubom montażowym wyświetlacza i płyty sterowania z tyłu panelu przedniego. Trzy otwory „C” w prawym dolnym rogu są przeznaczone dla montowanych na panelu gniazd wyjściowych i uziemienia (via przewód ochronny zasilania).

Rysunek 12 pokazuje frezowanie i wiercenie wymagane dla tylnego panelu, co jest stosunkowo proste. Po zakończeniu należy zamontować wentylator i od wewnątrz gniazdo sieciowe IEC (z gniazdem sieciowym włożonym od zewnątrz).

Etykiety w tym samym rozmiarze, co panel przedni, można pobrać ze strony internetowej Silicon Chip-a, wydrukować, zaalaminować i przykleić do przedniej części panelu (należy pamiętać, że rysunek 15 jest niewymiarowy – w przypadku kopiowania należy powiększyć go do 133% wielkości).

Wytnij otwory ostrym nożem modelarskim, a następnie możesz przymocować płytkę sterującą i zamontować pokrętła.

Wykonanie/przymocowanie radiatora

Wszystkie podstawowe funkcje zostały sprawdzone, więc można zamontować radiator. Rysunek 13 pokazuje, gdzie należy wywiercić otwory w dostępnym na rynku, a więc „komercyjnym” radiatorze, a także szczegóły dotyczące stworzenia własnego („DIY”).

Dolna krawędź obu typów radiatorów wystaje 3 mm poniżej dolnej krawędzi płytki drukowanej, jak pokazano na rysunku. Koniec radiatora najbliżej tylnego panelu również wystaje 3 mm poza koniec płytki drukowanej, aby umożliwić ustawienie otworu montażowego od końca.

Gdy to zrobisz, zamontuj wentylator i złącze zasilacza na tylnym panelu obudowy i podłącz przewody po stronie AC konwertera AC-DC. Kable AC muszą mieć tylko około 7 cm długości, ponieważ konwerter będzie zamontowany dość blisko złącza zasilania.

Przewód ochronny (żółto-zielony) musi sięgać do zacisku GND na panelu przednim. Zaizoluj końce kabli sieciowych przy gnieździe zasilania rurką termokurczliwą.

Przytnij kilka słupków montażowych na spodzie obudowy urządzenia i wywierć trzy otwory montażowe – patrz rysunek 14.

W zależności od wybranej opcji radiatora, wiercone są uchwyty środkowe (CINCON) lub przednie (DIY). Wynika to z faktu, że radiator CINCON jest nieco za krótki, aby osiągnąć przedniego występu montażowego.

Wywierć otwór 3,5 mm w czerwonym punkcie, aby zamocować konwerter AC-DC. Znajduje się on bezpośrednio w linii z jednym z istniejących (niezmodyfikowanych) słupków, ale 31 mm bliżej linii środkowej obudowy.

Zamontuj konwerter AC-DC w obudowie za pomocą wkrętu samogwintującego 4G×9 z okrągłym łbem przez otwór obok zacisków i gwintowanego wkrętu M3 z łbem stożkowym przyciętego na długość, z podkładką dystansową, przez wywiercony otwór.

Założ plastikową osłonę zacisków AC i gniazda zasilania i zamocuj ją pod krawędzią konwertera. Autor swoją zrobił z czerwonej polipropylenowej maty do cięcia.

Teraz podłącz wyjście konwertera AC-DC do pól/kołków VIN i GND na płytce regulatora oraz przyłutuj przewód ochronny (żółto-zielony z gniazda zasilania sieciowego) do zacisku GND na panelu przednim. Przewód od drugiego pola/kołka GND idzie do ujemnego gniazda zasilania na panelu przednim. Nawiń 4...5 zwojów przewodu

Możliwości rozbudowy

20-szypkowy wtyk CON2, wraz z dwoma opcjonalnymi złączami związanymi z enkoderem obrotowym i przelaznikami przyciskowymi, oferuje szeroki zakres wejść i wyjść do rozbudowy, a także, gdy płytka sterująca jest używana do innych celów.

Do tych złącz podłączonych jest łącznie 17 styków modułu ESP32, oprócz magistrali SPI, która jest współdzielona z kartą SD i ekranem dotykowym (tabela 1).

W tym projekcie zasilacza jest używanych kilka portów I/O ogólnego przeznaczenia (GPIO) i magistrala I²C; jednak magistrala SPI, port szeregowy, port USB, kanały DAC i ADC są nieużywane i dlatego są dostępne do ewentualnej rozbudowy.

Magistrala I²C obsługuje wszystkie tryby do 5 MHz z 7-bitowym lub 10-bitowym adresowaniem. Najlepiej jednak trzymać się trybu 400 kHz/7-bit, ponieważ wiele starszych układów I²C nie obsługuje bardziej zaawansowanych trybów.

Na płytce znajdują się rezystory polaryzujące magistrali I²C. Druga magistrala I²C jest dostępna jako jedna z opcji konfiguracji dla styków 13 i 14 CON2, jako alternatywa dla GPIO0 i drugiego kanału DAC.

Magistrala SPI została włączona do 20-stykowego złącza rozszerzeń; jeden styk GPIO będzie musiał zostać przydzielony jako linia wyboru układu (CS) dla każdego dodatkowego używanego urządzenia SPI. Ponieważ sygnały SPI przechodzą przez kabel taśmowy, najlepiej jest trzymać się częstotliwości magistrali 10 MHz lub niższych.

Obsługiwane jest przechowywanie plików na kartach SD. Podobnie jak w przypadku modułu LCD Mini D1 opartego na ESP8266, oprócz pełnowymiarowego gniazda na module LCD zapewniono wbudowane gniazdo karty mikro SD. Oba mogą być używane, ale nie razem, ponieważ pojedyncza linia wyboru układu jest współdzielona między nimi.

Opcjonalnie, przelaznik wykrywania karty (CD) w gnieździe może być podłączony do GPIO3. Jest on uziemiony po włożeniu karty i będzie wymagał skonfigurowania prądu polaryzującego w oprogramowaniu tego styku.

Dwukanałowy przetwornik ADC ma rozdzielczość 12 bitów, a maksymalna częstotliwość próbkowania wynosi około 27 kHz, pod nadzorem oprogramowania. Pomiędzy tymi stykami i masą GND znajdują się pola kontaktowe w celu zmniejszenia szumów wejściowych, gdy styki GPIO 34 lub 35 są używane jako wejścia ADC. Zalecane kondensatory po 100 nF zapewniają znaczne filtrowanie nawet przy umiarkowanych częstotliwościach, ponieważ wejście pobiera zaledwie prąd o natężeniu 50 nA.

Dostępne są dwa 8-bitowe kanały DAC o praktycznej przepustowości około 200 tys. próbek na sekundę. Dostępny jest interfejs szeregowy na poziomie logicznym, zdolny do nadawania i odbierania z prędkością do 5 Mb/s. Obsługiwany jest również interfejs szeregowy USB. Jak zauważono w tekście, niez izolowane zasilanie lub komunikacja USB nie są zalecane dla tego zasilacza.

Poza sygnałami I²C i SPI, pozostałe porty są wielofunkcyjne. Każdy styk GPIO można skonfigurować jako wejście przerwania lub wyjście PWM.

Większość wyspecjalizowanych portów (ADC, DAC i szeregowy) może być również używana do obsługi cyfrowych wejść/wyjść, zwiększając całkowitą liczbę portów obsługujących GPIO do 8 lub 17, jeśli enkoder obrotowy i przelazniki przyciskowe nie są wymagane.

Jeśli liczba styków GPIO jest niewystarczająca dla danego projektu, można dodać ekspander I²C I/O, taki jak MCP23008.

Komunikacja

ESP32 oferuje szeroki zakres opcji Wi-Fi; może łączyć się z istniejącą siecią LAN Wi-Fi 2,4 GHz lub tworzyć sieć lokalną w trybie „soft-AP”. Oba tryby są włączone do naszego projektu zasilacza.

Najpierw sterownik próbuje połączyć się z istniejącą siecią LAN, z uwierzytelnieniem wprowadzonym w podmenu COMMS. Jeśli to się nie powiedzie, próbuje dotrzeć do istniejącej sieci

Tabela 1. Mapowanie styków CON2

Nr styku złącza CON2	Funkcja ESP-32	Styk ESP-32	Funkcja zasilacza
1	GND		GND
2	GND		GND
3	SPI:MISO	GPIO19	–
4	SPI:SCK	GPIO18	–
5	I ² C1:SDA/GPIO	GPIO21	Sterowanie I ² C dla IC1, IC2, IC4
6	SPI:MOSI	GPIO23	–
7	I ² C1:SCL/GPIO	GPIO0	Sterowanie I ² C dla IC1, IC2, IC4
8	I ² C2:SDA/GPIO	GPIO22	–
9	COM2:TX/GPIO	GPIO17	Detekcja sygnału DRDY z IC1
10	COM2:RX	GPIO16	–
11	GPIO	GPIO2	–
12	GPIO	GPIO4	Detekcja naciśnięcia SW_ON
13	DAC1/GPIO	GPIO25	–
14	DAC2/I ² C2:SCL/GPIO	GPIO26	Sterowanie wentylator wł./wył.
15	ADC1-7/GPIO	GPIO35	–
16	GPIO	GPIO12	Detekcja naciśnięcia SW_OFF
17	ADC1-6/GPIO	GPIO35	–
18	+5 V		+5 V
19	+3,3 V		+3,3 V
20	+5 V		+5 V

z SSID ESPINST. Jeśli to się nie powiedzie, tworzy sieć ESPINST dla innych przyrządów.

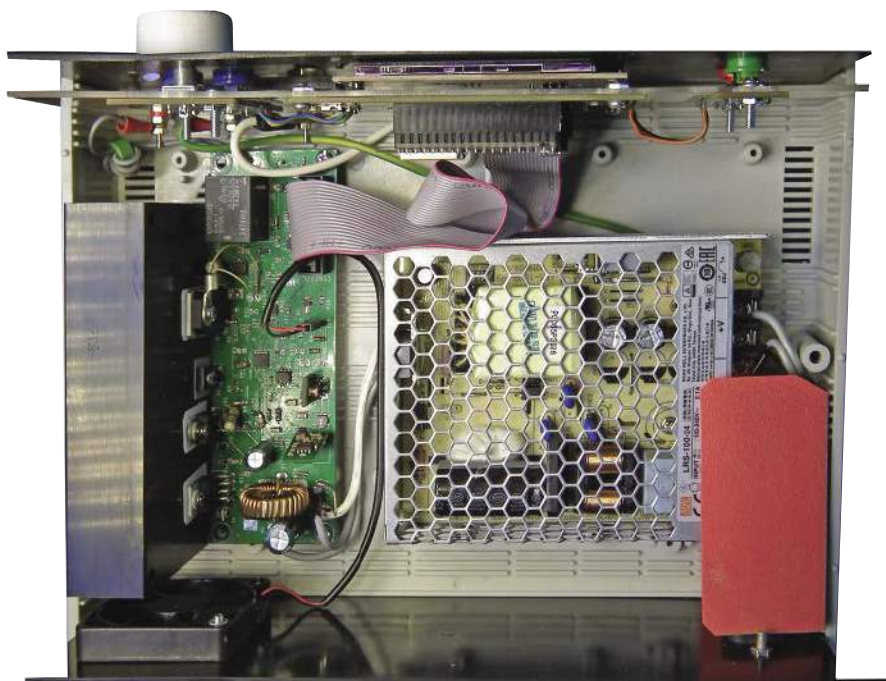
Chociaż karta sterowania obsługuje zarówno tradycyjne tryby Bluetooth, jak i BLE, nie zostały one wykorzystane w projekcie zasilacza. Drugi port szeregowy jest dostępny na 20-stykowym złączu rozszerzeń. Jest on również nieużywany w projekcie zasilacza.

Komunikacja szeregową USB jest dostępna za pośrednictwem gniazda mikro USB w module ESP-32, zapewniając gotowy sposób programowania urządzenia i debugowania kodu za pomocą jednej z dostępnych zintegrowanych platform programistycznych, takich jak Arduino. Port USB zapewnia również jeden z interfejsów sterowania SCPI dla tego projektu zasilacza.

Zdecydowanie zaleca się stosowanie izolatora USB w projekcie zasilacza, aby uniknąć pętli uziemienia, która mogłaby zniszczyć ESP-32 lub port USB komputera.

Izolatory te są dostępne w serwisie eBay lub AliExpress za około 15 USD. Działają w trybie pełnej prędkości (11 Mb/s) lub wysokiej prędkości (480 Mb/s). Autor z powodzeniem korzystał z odmiany przedstawionej na poniższym zdjęciu.





Gotowy hybrydowy zasilacz laboratoryjny powinien wyglądać podobnie do prototypu na tym zdjęciu, z drobnymi różnicami wynikającymi z innego rozmieszczenia elementów panelu sterującego

połączeniowego wokół toroidalnego rdzenia dla cewki filtra wyjściowego, a następnie przylutuj jeden koniec do pola/kołka VOUT na płycie drukowanej regulatora i zamontuj płytkę drukowaną/radiator w obudowie. Podłącz drugi przewód filtra toroidalnego do dodatniego zacisku zasilania na panelu przednim.

Autor użył zaciskanych końcówek oczkowych, aby umożliwić łatwe zakładanie/zdejmowanie przewodów z gwintowanych końcówek gniazd wyjściowych; można jednak również przylutować przewody bezpośrednio do końcówek tych gniazd. Wyjściowy dławik toroidalny (zielony) schowany jest w rogu obudowy, pomiędzy radiatorem a panelem przednim.

Wykończenie

W tym momencie montaż laboratoryjnego zasilacza hybrydowego jest zasadniczo zakończony i przechodzimy do dalszych testów.

Jeśli nie pamiętałeś przed wyłączeniem o ustawieniu napięcia wyjściowego na panelu sterowania na 2 V, odłącz zasilanie, podłącz ponownie kabel USB i postępuj zgodnie z instrukcjami powyżej. Odłącz kabel USB.

Ustaw napięcie wyjściowe konwertera AC-DC na najniższe ustawienie za pomocą potencjometru nastawnego (całkowicie w lewo),

a następnie włącz zasilanie. Wartość VIN powinna mieścić się w zakresie kilku woltów wokół 20 V (minimalne ustawienie konwertera). Zwróć uwagę na napięcia na VOUT i VPRE (około 3,6 V wyższe). Jeśli wszystko jest w porządku, można bezpiecznie obrócić potencjometr nastawny do najwyższego ustawienia, co podniesie VIN do około 30 V.

Podstawowe testy są teraz zakończone i można zacząć używać przyrządu do zasilania różnych budowanych urządzeń.

Kalibracja

Kalibracja zasilacza jest jedną z opcji zaimplementowanych w GUI, ponieważ pomiary prądu i napięcia będą podawane z błędem kilku procent, głównie w zależności od tolerancji rezystorów.

Aby skalibrować pomiary napięcia, ustaw napięcie wyjściowe na 25 V (lub inne ustawienie o kilka woltów poniżej wartości maksymalnej) i wybierz menu CAL na ekranie (w lewym dolnym rogu ekranu 3). Bez podłączonego obciążenia włącz wyjście zasilacza i zmierz napięcie za pomocą multimetru.

Za pomocą przycisków numerycznych wprowadź różnicę między odczytem multimetru a wartością wyświetlaną na ekranie po lewej

stronie. Jeśli odczyt multimetru jest wyższy, wprowadź wartość dodatnią. W przykładzie pokazanym na ekranie 4 zasilacz odczytuje 25,00 V, ale multimetr referencyjny odczytuje 25,12 V, więc 0,12 V jest ustawione jako przesunięcie w prawym górnym rogu.

Dotknij przycisku ZAPISZ, aby zapisać wynik. Spowoduje to również wyjście z menu kalibracji. Poczekaj przed wyłączeniem urządzenia, aż wskaźnik [E] zgaśnie, aby nowa wartość kalibracji została trwale zapisana w pamięci FLASH (EEPROM).

Powtórz ten sam proces kalibracji dla prądu, z włączonym wyjściem, używając rezystora obciążenia, który pobiera 1 A lub więcej przy dowolnym napięciu wyjściowym. Rezystor mocy 10 Ω /10 W lub większy będzie działał prawidłowo. Nie ma potrzeby kalibracji prądu zerowego, ponieważ wartość ta kalibruje się automatycznie po krótkim czasie, gdy wyjście jest wyłączone.

Jeśli chcesz korzystać z połączenia bezprzewodowego Wi-Fi z przyrządem, przejdź do podmenu „COMMS” (ekran 5). Wprowadź swoje dane uwierzytelniające sieci Wi-Fi i naciśnij przycisk automatycznego połączenia (AC), jeśli nie jest jeszcze zielony. Spowoduje to zainicjowanie protokołu połączenia Wi-Fi. Pojawi się zielony prostokąt wskazujący postęp połączenia.

Po zakończeniu wskaźnik „W” w górnej części ekranu powinien być zielony, a adres IP wyświetlany w dolnej części ekranu „COMMS”.

Podsumowanie

Laboratoryjny zasilacz hybrydowy, w przedstawionej tutaj formie, jest rzeczywiście bardzo przydatnym urządzeniem. Mimo to można rozszerzyć jego możliwości o jeszcze więcej funkcji, dzięki opcjom modułu Wi-Fi i mikrokontrolera ESP-32.

Należy również pamiętać, że nasza płyta sterująca w stylu Backpack-a jest portowna i wszechstronna sama w sobie i może być używana do współpracy z różnymi innymi projektami. ■

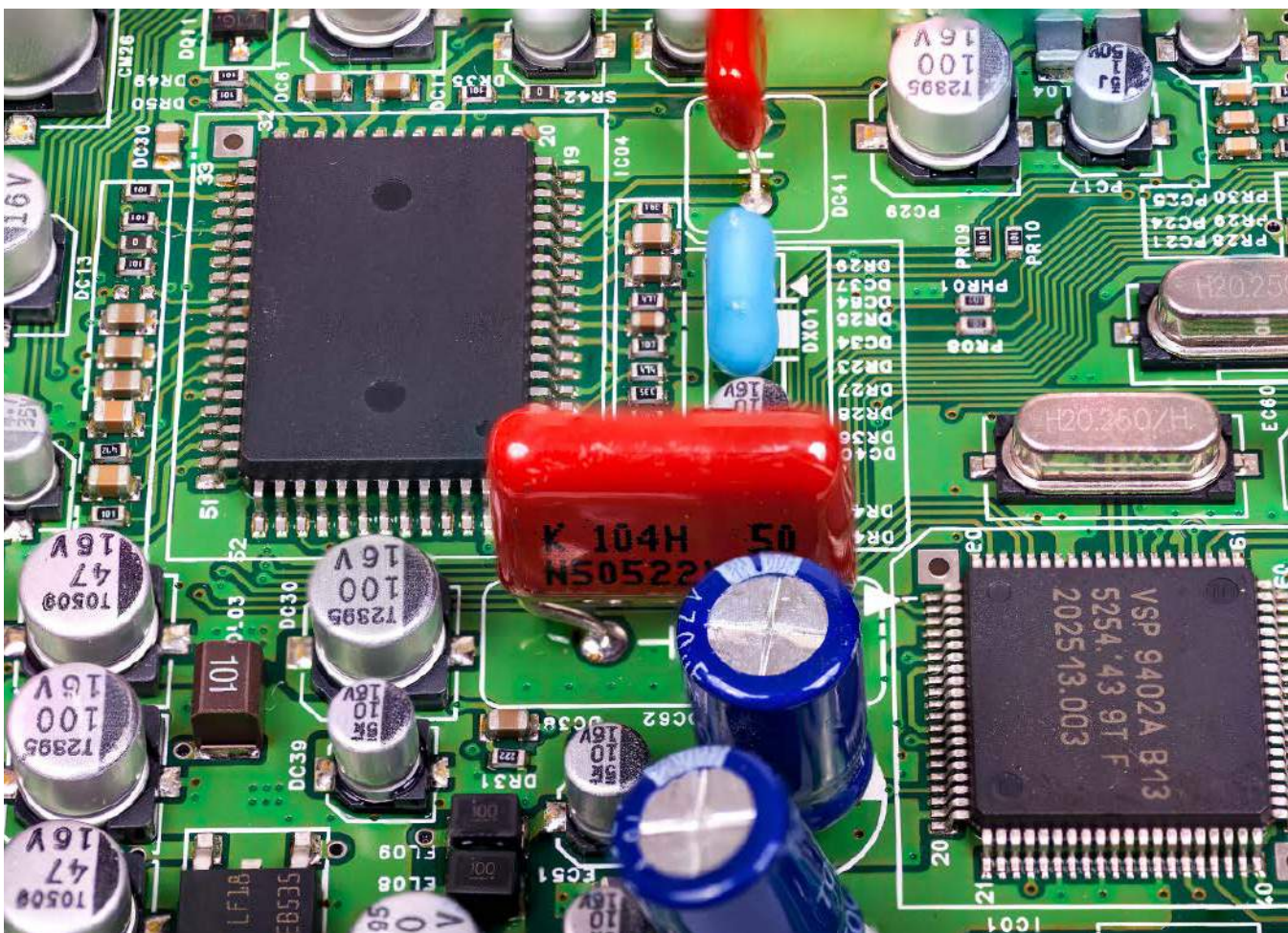
Richard Palmer

*Adaptacja do wydania
polskiego – Andrzej Nowicki*

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022.
www.siliconchip.com.au

REKLAMA

numery archiwalne • prenumerata • książki
www.UlubionyKiosk.pl



PCBWay zmontuje Twoje płytki drukowane

Jakość montażu PCB ma kluczowe znaczenie dla niezawodności docelowego urządzenia. Nawet najmniejsze błędy, często trudne do zauważenia lub wręcz niewidoczne gołym okiem, mogą z łatwością przekreślić szanse na prawidłowe uruchomienie układu. Mało tego – trudność montażu współczesnych urządzeń elektronicznych rośnie wraz z miniaturyzacją układów scalonych i elementów biernych. A gdyby tak zlecić całość prac zewnętrznemu producentowi? Nic prostszego – wykwalifikowany zespół PCBWay jest do Twojej dyspozycji!

Montaż montażowi nierówny

Pojęcie „montaż PCB” obejmuje zarówno klasyczną technologię THT, jak i powierzchnię (SMT), choć w większości przypadków konieczne okazuje się połączenie obu tych metod. PCBWay oferuje usługi dostosowane do rodzaju i objętości produkcji – w przypadku małych zamówień preferowany jest montaż ręczny, podczas gdy większe serie produkcyjne realizowane są z użyciem nowoczesnych automatów: drukarek pasty lutowniczej,

maszyn pick&place oraz systemów automatycznej kontroli optycznej (AOI).

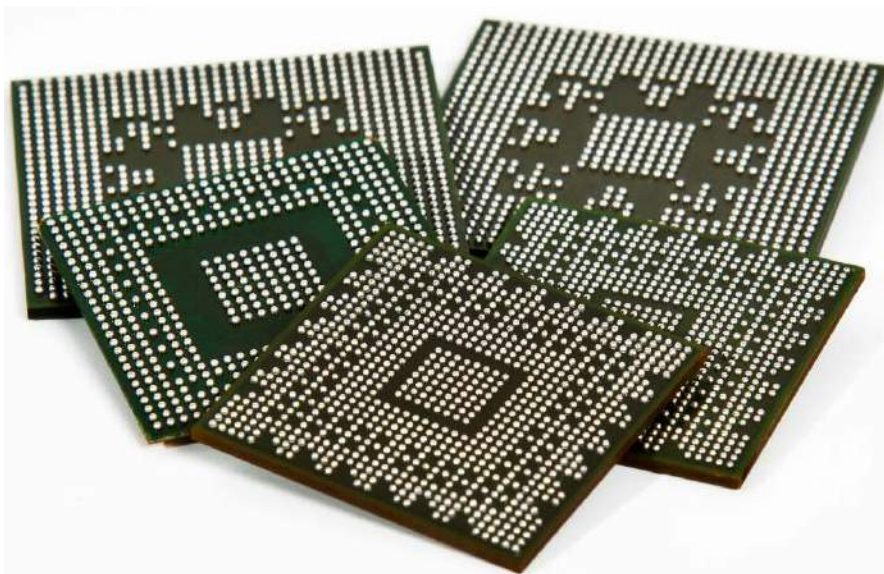
Montaż SMT w ofercie PCBWay

Rozbudowany park maszynowy PCBWay umożliwia montaż układów w technologii SMT w niezwykle szerokim zakresie rozmiarów i typów obudowy komponentów. Obsługiwane są układy w obudowach:

- Ball Grid Array (BGA – rysunek 1), w tym:
 - Plastic BGA (PBGA),

- Ceramic BGA (CBGA),
- Micro Fine Line BGA (MBGA),
- Stack BGA (Package-On-Package, PoP),
- Ultra-Fine Ball Grid Array (uBGA),
- Quad Flat Pack No-Lead (QFN),
- Quad Flat Package (QFP),
- Small Outline Integrated Circuit (SOIC),
- Plastic Leaded Chip Carrier (PLCC).

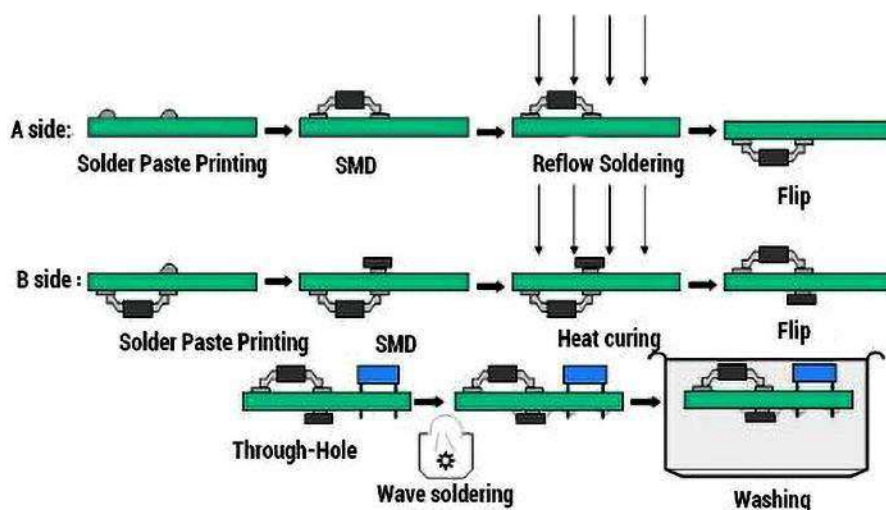
Warto dodać, że firma jest w stanie zapewnić niezawodny montaż elementów w obudowach o niezwykle małym rozstawie wyprowadzeń



Rysunek 1. Przykładowe układy BGA



Rysunek 2. Schemat procesu montażu SMT w PCBWay



2 layers board Mixed Assembly

Rysunek 3. Przebieg procesu montażu mieszanego (SMT+THT) w PCBWay



Fotografia 1. Automat pick&place w czasie montażu płytki drukowanej

– 0,25 mm dla układów BGA (!) oraz 0,2 mm dla innych rodzajów podzespołów. Silnie zminiaturyzowane urządzenia, w których pracują tak niewielkie układy scalone, wymagają rzecz jasna również kompaktowych komponentów pasywnych – PCBWay ma możliwość pracy nawet z elementami w obudowach 01005 oraz 0201, co pokrywa praktycznie wszystkie potrzeby współczesnej elektroniki, także w zakresie urządzeń mobilnych czy elektroniki ubieralnej.

Kluczowym elementem procesu seryjnego montażu elektroniki pozostaje skrupulatna kontrola jakości. W PCBWay jest ona realizowana na wszystkich etapach produkcji, przy czym w większości punktów kontrolnych pracują nie ludzie, a wyspecjalizowane automaty (**rysunek 2**). Wspomniana już automatyczna inspekcja optyczna pasty lutowniczej (SPI) opiera się na trójwymiarowych skanerach wysokiej rozdzielczości, zdolnych do wykrywania nawet najmniejszych ubytków lub nadmiarów pasty lutowniczej na padach PCB. Płytki opuszczające piec do lutowania rozplwowego są następnie kontrolowane przez kolejny automat AOI, tym razem sprawdzający poprawność pozycji i oznaczeń komponentów, co umożliwia szybkie zidentyfikowanie ewentualnych błędów (np. brakujących elementów, błędnej polaryzacji etc.). Należy podkreślić, że producent wdrożył nie tylko uniwersalny system kontroli jakości ISO9001, ale także uzyskał certyfikaty ISO 13485 (System Zarządzania Jakością dla Wyrobów Medycznych) oraz IATF 16949 (analogiczny system dla branży motoryzacyjnej). Wyroby dostarczane przez PCBWay spełniają ponadto wymogi amerykańskich standardów UL oraz normy IPC-610, określającej tzw. dopuszczalność pakietów elektronicznych.

Montaż THT i mieszany

Obwody drukowane przeznaczone do pracy z samymi tylko komponentami przewlekany (lub przewlekany i powierzchniowymi jednocześnie), wymagają odpowiedniej modyfikacji procesu montażu. O ile bowiem układy bazujące wyłącznie na komponentach THT mogą być w całości lutowane na fali, to płytki wykonane w technologii mieszanej przechodzą przez kilka dodatkowych etapów. Przebieg procesu stosowanego w fabryce PCBWay zobrazowano schematycznie na **rysunku 3**. Na początku montowane są komponenty tylko na warstwie górnej (TOP), przy użyciu standardowego procesu układania elementów (pick&place – **fotografia 1**) oraz lutowania rozplwowego. Następnie płytka jest odwracana w osi poziomej i zamontowane

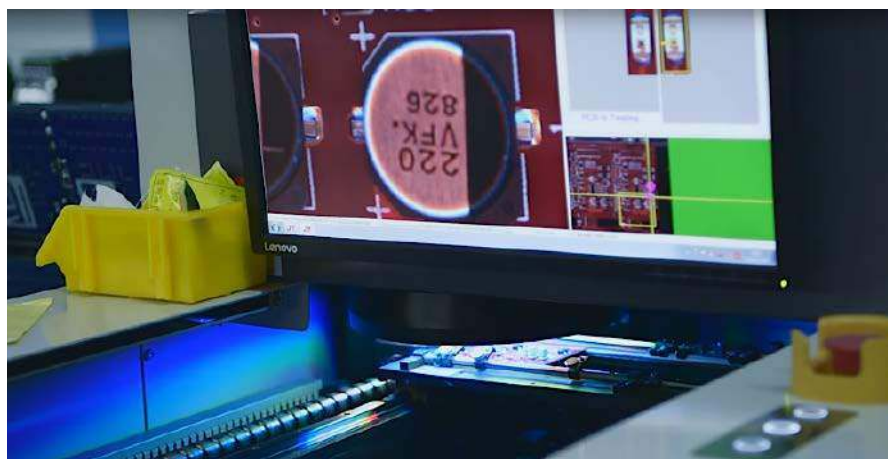
zostają podzespoły umiejscowione na warstwie dolnej (BOTTOM). Po zakończeniu drugiego etapu produkcji płytką powraca do położenia początkowego, po czym zostaje „uzbrojona” elementami przewlekanyymi. Tak przygotowany pakiet jest kierowany do maszyny odpowiedzialnej za lutowanie na fali, co pozwala zamontować komponenty THT bez konieczności angażowania montażyistów wykwalifikowanych w lutowaniu selektywnym. Rzecz jasna, jeżeli zajdzie taka potrzeba, elementy przewlekane mogą być także lutowane ręcznie, jednak rozwiązanie takie jest efektywne tylko pod warunkami niewielkiej objętości partii produkcyjnej, a także małego udziału komponentów THT w BOM urządzenia.

Nie tylko montaż – oferta usług dodatkowych

W portfolio PCBWay można znaleźć szereg usług dodatkowych, które znakomicie przyspieszają produkcję kontraktową urządzeń elektronicznych i upraszczają proces logistyczny po stronie klienta. Firma świadczy outsourcing m.in. w zakresie inspekcji rentgenowskiej układów BGA oraz testów funkcjonalnych prowadzonych na docelowym układzie (ICT – In-Circuit Testing) przy pomocy wyspecjalizowanych, zautomatyzowanych stacji testowych (**fotografia 2**). Tego rodzaju badania pozwalają wykryć błędy w działaniu zmontowanych układów elektronicznych, manifestujące się np. poprzez niewłaściwe wartości rezystancji, pojemności, czy nawet indukcyjności. Mało tego – odpowiednie scenariusze testowe mogą obejmować także wartości napięć w poszczególnych miejscach obwodu,



Fotografia 2. Stacja testowa ICT



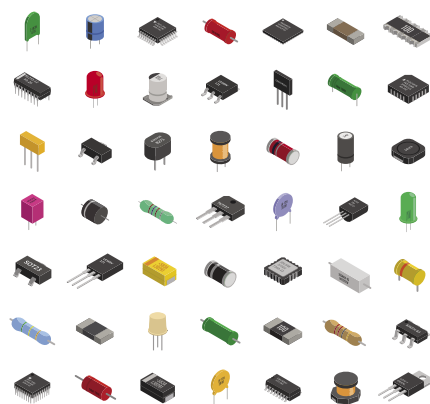
Fotografia 3. Stanowisko automatycznej inspekcji wizyjnej (AOI)

czy też poprawność odczytów kontrolnych z układów cyfrowych (np. poprzez zastosowanie skanu JTAG). Jest to zatem dalece bardziej zaawansowany sposób kontroli jakości, niż wszelkiego rodzaju testy optyczne – o ile bowiem działanie AOI (**fotografia 3**) sprawdza się do wykrywania uszkodzeń, braków lub błędnie zamontowanych komponentów, o tyle testy funkcjonalne są w stanie wykryć awarie niewidoczne przy zastosowaniu metod optycznych lub nawet rentgenowskich, a wynikające np. z głębokich uszkodzeń elementów, przekroczenia dopuszczalnych tolerancji produkcyjnych PCB, błędów pamięci i wielu innych.

Obsługa łańcucha dostaw komponentów

Ważną zaletą PCBWay jako producenta kontraktowego jest duża elastyczność w zakresie obsługi łańcucha dostaw komponentów przeznaczonych do montażu (**rysunek 4**). Firma oferuje zasadniczo trzy możliwości, określone jako:

- **Consigned/kitted** – w tym scenariuszu klient sam kompletuje BOM projektu i przesyła cały komplet elementów wraz ze szczegółową dokumentacją do PCBWay;
- **Turn-key** – wykwalifikowany zespół PCBWay ds. zaopatrzenia przejmuje od klienta całość wysiłków związanych z kompletacją zestawu elementów do produkcji, konsultując z nim ceny i dostępność podzespołów (co ważne – producent nie zarabia dodatkowo na komponentach, stosuje się do oryginalnych cen dostawców).
- **Partial Turn-key/Combo** – najczęściej wybierane rozwiązanie polega na dostarczeniu kluczowych komponentów przez klienta i pozostawieniu zespołowi PCBWay jedynie uzupełnienia brakujących pozycji (np. drobnych elementów pasywnych czy też prostych elementów półprzewodnikowych).



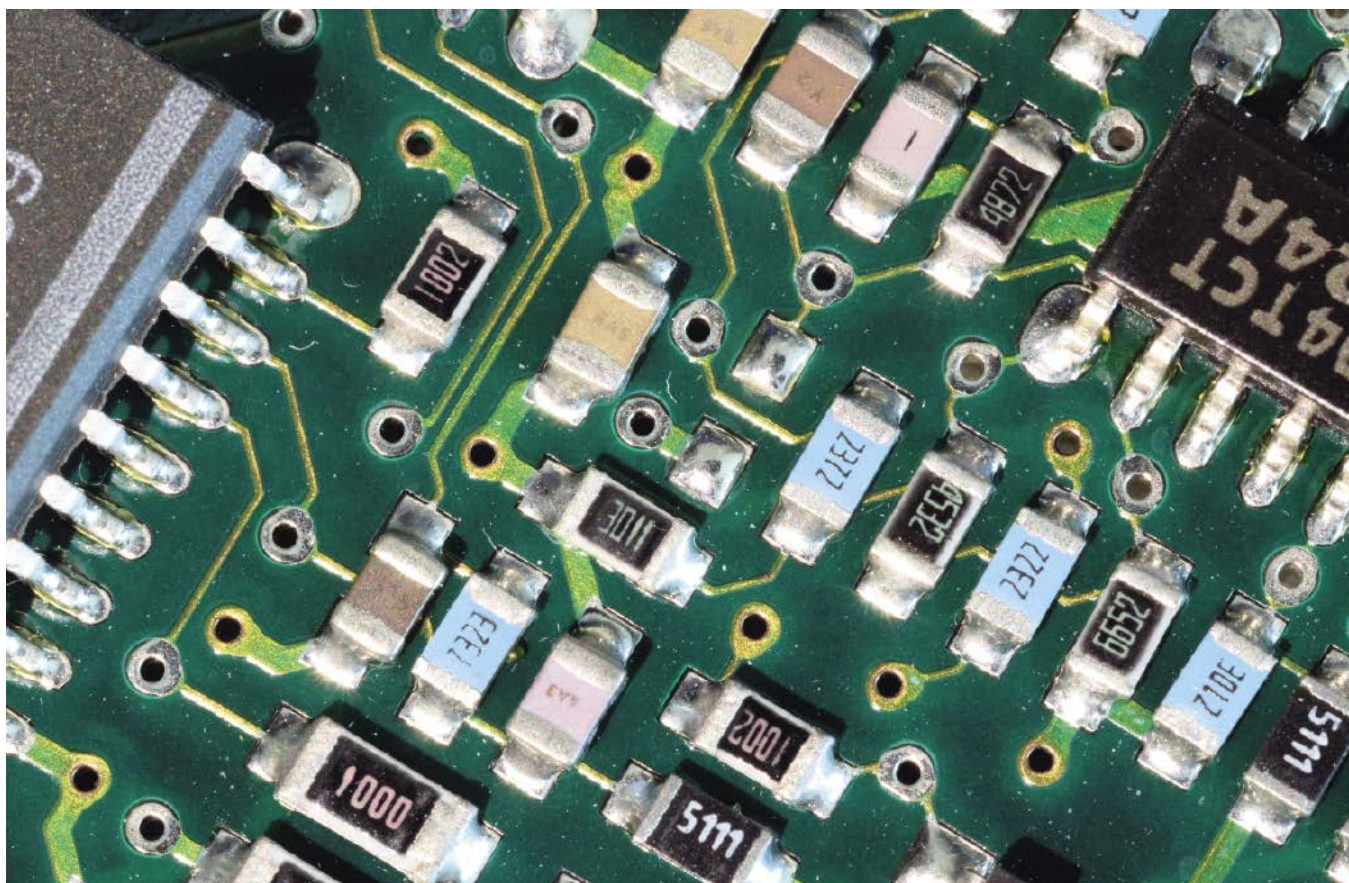
Rysunek 4. PCBWay oferuje szerokie możliwości w zakresie kompletacji BOM produkowanych urządzeń

Dodatkowo PCBWay może realizować zakupy elementów zgodnie z wytycznymi klienta, tj. wybierać potrzebne elementy tylko z oferty wskazanych przez niego dostawców. Takie rozwiązanie pozwala odbiorcom znacznie uprościć proces logistyczny, niweluje bowiem konieczność ponoszenia podwójnych opłat transportowych, skraca czas realizacji produkcji kontraktowej, a nade wszystko – zapewnia pełną identyfikowalność łańcucha dostaw.

Podsumowanie

Produkcja urządzeń elektronicznych jest złożonym, wieloetapowym procesem, w którym fundamentalne znaczenie zyskuje nie tylko doskonałe wyposażenie w postaci nowoczesnego parku maszynowego, ale także wykwalifikowany zespół montażyistów i operatorów, wspierany przez sprawną infrastrukturę logistyczną. Szeroki wachlarz usług EMS w portfolio PCBWay pozwala efektywnie zrealizować nawet największe wyzwania w zakresie produkcji zaawansowanych urządzeń na miarę XXI wieku. ■

www.pcbway.com



Zdjęcie: www.pxfuel.com/en/free-photo-qhfan

Lutowanie SMD

Porady i wskazówki, sztuczki i triki

Podczas gdy jedyną różnicą między komponentami SMD i przewlekаныmi jest sposób ich lutowania do płytki drukowanej, istnieje wiele nieświadomości, a nawet nieporozumień, otaczających elementy SMD i nowe techniki wymagane do pracy z nimi, zwłaszcza z częściami o mniejszych rozmiarach. Ten artykuł towarzyszy naszemu projektowi płytki treningowej SMD (w poprzednim numerze EdW) i zawiera wiele szczegółów, które pomogą stać Ci się mistrzem lutowania SMD.

Niewątpliwie niektórzy ludzie wolą nauczyć się lutowania SMD, po prostu zabierając się za montaż uprzednio nabytej płytki treningowej (patrz tekst w poprzednim numerze EdW) wraz z niezbędnymi elementami. Jednak lutowanie SMD jest o wiele łatwiejsze, jeśli zna się niektóre wskazówki i sztuczki.

Informacje zawarte w tym artykule mogą okazać się pomocne, nawet jeśli nie planujesz lutowania modułu płytki treningowej SMD. Artykuł zawiera wiele ogólnych porad i wskazówek, więc choćby tylko z tego względu warto go przeczytać. Należy jednak pamiętać, że ten tekst ma towarzyszyć nauce lutowania na płycie szkoleniowej, zatem nie opisuje mniej popularnych komponentów i obudów SMD, które nie pojawiają się na tej płycie drukowanej.

Jeśli masz już pewne doświadczenie w pracy z SMD, jest szansa, że nadal możesz się czegoś nauczyć, czytając ten artykuł i pomijając sekcje dotyczące tematów, które już rozumiesz.

Rozmiary i obudowy komponentów SMD

Wiele komponentów użytych w naszym projekcie płytki treningowej (w tym rezystory, kondensatory i diody) ma dwa wyprowadzenia i jest w tak zwanych obudowach typu „chip”. Są one małe, płaskie i z grubsza prostokątne. Komponenty w takich obudowach występują najliczniej w każdym projekcie opartym na elementach do montażu powierzchniowego.

Niektóre elementy pasywne występują w różnych typach obudów SMD. Na przykład,

często spotyka się małe kondensatory elektrolityczne umieszczone na niewielkiej plastikowej podstawie z wystającymi wyprowadzeniami w stylu SMD. Chociaż są one mniejsze niż większość kondensatorów elektrolitycznych, to nadal są większe niż większość elementów pasywnych do montażu powierzchniowego, więc nie są trudne do lutowania.

Części w obudowach „chip” są często opisywane za pomocą cztero- lub sześciocyfrowego kodu, przy czym istnieją zarówno calowe, jak i metryczne wersje tego kodu. Na przykład, część o rozmiarze metrycznym 3216 będzie zamiennie znana jako 1206 w systemie calowym. Co myłące, istnieją części o tych samych kodach w obu systemach (w tym 1206), ale mają one bardzo różne rozmiary! W tym

tekście tradycyjnie będziemy preferować rozmiary calowe.

Jednym ze sposobów rozróżnienia obu typów oznaczeń jest użycie prefiksu „M” dla rozmiarów metrycznych; to jest to, co znajdziemy na listach części Silicon Chip-a i zwykle cytujemy oba kody, aby rozwiązać niejednoznaczność. Na przykład, na listach części często można zobaczyć (M3216/**1206**). Jest to największy rozmiar rezystora i kondensatora, jaki zastosowaliśmy na płytce treningowej SMD. Dostępne są jednak większe części; następnym krokiem jest zwykle M3226/**1210**, a potem M4532/**1812**.

Pierwsze dwie cyfry określają długość elementu, podczas gdy pozostałe cyfry określają szerokość. Większość części ma zwykle większą długość niż szerokość, więc pierwsze dwie cyfry będą większe, ale nie zawsze tak jest. Zwykle wyprowadzenia znajdują się wzdłuż krótszych boków, ale w przypadkach, gdy wyprowadzenia obejmują dłuższe boki, liczby mogą być odwrócone (np. M1632/**0612**).

Wymiary metryczne podane są w dziesiątych częściach milimetra, więc część M3216 mierzy 3,2 mm długości i 1,6 mm szerokości. Należy również pamiętać, że dwa wyprowadzenia będą znajdować się na przeciwnych końcach, na krótszych bokach.

W systemie calowym każda para cyfr odpowiada 1/100 cala, więc część **1206** ma wymiary **0,12** cala na **0,06** cala, zbliżone do odpowiednika metrycznego.

Tabela 1 podsumowuje niektóre z bardziej popularnych rozmiarów elementów dwukońcówkowych. Zwróć uwagę na ostatni wiersz pokazujący pięciocyfrowy kod calowy (z wymiarem poniżej **1/100** cala lub 0,25 mm!). Można również zauważyć, że niektóre kody (takie jak **0603** i **0402**) są obecne w obu wierszach.

Na płytce szkoleniowej SC wszystkie części wokół IC1 mają rozmiar M3216/**1206**. Jest to jeden z największych rozmiarów, w którym

Tabela 1. Typowe rozmiary pasywnych elementów SMD

Metryczne	M3216	M2012	M1608	M1005	M0603	M0402
Długość	3,2 mm	2,0 mm	1,6 mm	1,0 mm	0,6 mm	0,4 mm
Szerokość	1,6 mm	1,2 mm	0,8 mm	0,5 mm	0,3 mm	0,2 mm
Calowe	1206	0805	0603	0402	0201	01005
Długość	0,12 cala	0,08 cala	0,06 cala	0,04 cala	0,02 cala	0,01 cala
Szerokość	0,06 cala	0,05 cala	0,03 cala	0,02 cala	0,01 cala	0,005 cala

dostępna jest szeroka gama części, więc jest to dobry wybór do stosowania części SMD, gdy nie ma potrzeby stosowania mniejszych.

Diody LED wokół IC2 mają różne rozmiary, poczynając od M3216/**1206** przez M2012/**0805**, M1608/**0603** i M1005/**0402** aż do najmniejszych M0603/**0201**. Każda z tych diod ma odpowiedni rezystor szeregowy o tym samym rozmiarze.

Inną dwuwyprowadzeniową obudowę, którą możemy tutaj znaleźć, jest często używana dioda i jest znana jako SOD-123 („small outline diode” – „mała dioda”). Są one podobne w wyglądzie do obudów tranzystorów, które opiszemy poniżej, ale mają tylko dwa wyprowadzenia.

Komponenty z trzema lub więcej wyprowadzeniami

Na płytce treningowej zastosowano układy IC1 oraz Q1 również w bardzo popularnych dla elementów SMD obudowach.

W przypadku części z więcej niż dwoma wyprowadzeniami, często występują warianty z różną ich liczbą, ale identycznym rozstawem pomiędzy nimi i odstępami między rzędami. Części nazywane SOIC lub SOP („small outline IC” lub „small outline package”) czyli „mały układ scalony” lub „mała obudowa”) mają zazwyczaj styki o rozstawie 1,27 mm czyli 0,05 cala.

Jest to dokładnie połowa rastra styków większości elementów przewlekanych DIL („dual in-line”). Układ IC1 znajduje się w obudowie

SOIC-8 o szerokości 3,9 mm. Podobnie jak w przypadku części DIL, szerokość ma tendencję do zwiększania się wraz ze wzrostem liczby wyprowadzeń, aby zapewnić miejsce na wewnętrzne połączenia oraz większe struktury krzemowe.

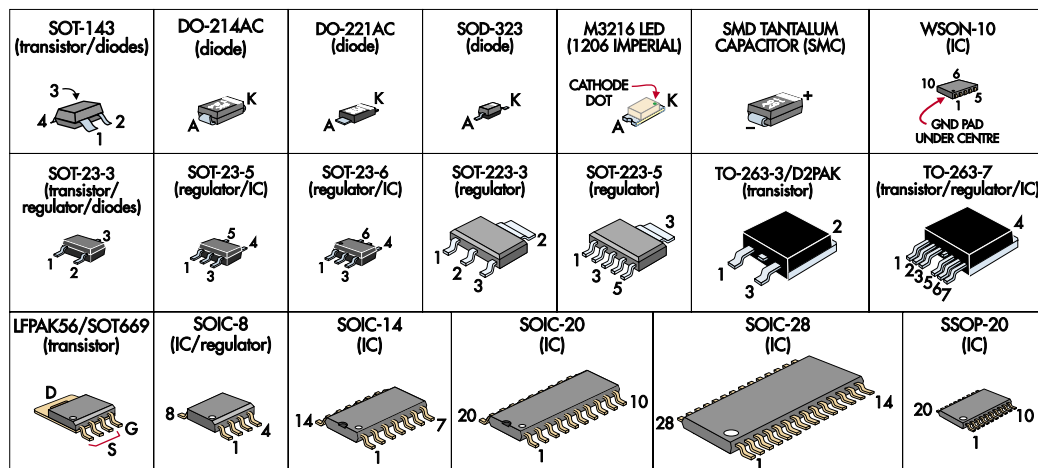
Obudowa, którą wybraliśmy dla tranzystora Q1, nazywa się SOT-23 („small-outline transistor” – „mały tranzystor”). Istnieją również warianty z dodatkowymi wyprowadzeniami naprzeciwko każdego z nich, zwane SOT-23-6, a także SOT-23-5, który jest taki sam jak SOT-23-6, ale nie ma środkowego styku po jednej stronie (**rysunek 1**).

Typowe elementy w obudowach SOT-23 (MOSFET-y, tranzystory małej mocy, podwójne diody itp.) są dość łatwe w montażu, ponieważ pasują do swoich pól lutowniczych tylko w jeden sposób, a wyprowadzenia są dość luźno rozmieszczone i dostępne. Jednak ich rozmiar sprawia, że przy nieostrożnym obchodzeniu się z nimi istnieje większe prawdopodobieństwo, że zostaną zgubione.

Czysta przestrzeń robocza o jednolitym kolorze jest najlepszą strategią zapobiegającą zgubieniu małych elementów.

Obudowy układów scalonych

Rozmiar obudowy IC2 na naszej płytce szkoleniowej jest kolejnym krokiem w miniaturyzacji. Obudowa nosi nazwę SSOP („small shrink outline package” – „mała zminiaturyzowana obudowa”). Można je również zobaczyć z innymi opisami, takimi jak TSSOP („thin



Rysunek 1. Po lewej stronie pokazano niektóre z bardziej popularnych obudów i wyprowadzeń („footprintów” – „odcisków stopy”) komponentów do montażu powierzchniowego (np. SOT-23, SOIC-8, SSOP-16, M3216/1206), wraz z numeracją styków

shrink small-outline package” – „cienka mała zminiaturyzowana obudowa”). Tak czy inaczej, będą one miały raster wyprowadzeń 0,65 mm, czyli o połowę mniejszy niż w przypadku SOIC. Oprócz tego, że są cieńsze, obudowy TSSOP są również węższe niż SSOP, więc uważaj – niektóre pola lutownicze będą pasować do obu typów, ale nie wszystkie.

Inną popularną obudową układów scalonych, która nadaje się do lutowania ręcznego, jest QFP („quad flat pack” – „płaska obudowa kwadratowa”) i wiele jej wariantów, takich jak TQFP („thin quad flat pack” – wersja o zmniejszonej grubości). Są one dostępne z różnymi rastrami styków, od 0,4 mm do 0,8 mm.

Są one często używane tam, gdzie potrzeba więcej wyprowadzeń na małej przestrzeni, np. w mikroprocesorach. Chociaż obudowy te nie są dużo mniejsze, mogą być trudniejsze do prawidłowego wycentrowania wyprowadzeń względem padów lutowniczych na płytce PCB, ze względu na wyprowadzenia rozmieszczone wokół czterech boków.

W celach demonstracyjnych umieściliśmy pole lutownicze dedykowane obudowie QFP-44 (10×10) z tyłu PCB; ma 44 wyprowadzenia (11 po każdej stronie), podczas gdy 10×10 odnosi się do wymiarów plastikowej obudowy w milimetrach. Raster wyprowadzeń wynosi 0,8 mm. Możesz przetestować swoje umiejętności, jeśli masz odpowiednią część, chociaż do działania układu elektronicznego nic to nie wniesie – styki nie są z niczym połączone. Rysunek pół lutowniczych może być również przydatny jako punkt odniesienia do sprawdzania wymiarów i rastra wyprowadzeń.

Zazwyczaj to niewielki rozmiar części SMD utrudnia lutowanie ręczne, ale istnieją również

inne powody. W przypadku chęci zastosowania elementów mniejszych niż SSOP, projektant może wybrać QFN („quad flat no-lead” – „płaska obudowa kwadratowa bez wyprowadzeń”), BGA („ball grid array” – „układ siatki kulek”), VTLA („very thin leadless array” – „bardzo cienka matryca bez wyprowadzeń”) lub WLCSP („wafer level chip scale packaging” – „obudowa na poziomie struktury krzemowej”).

Części te nie są przeznaczone do lutowania ręcznego, a ich prawidłowe lutowanie wykonywane jest w procesie rozpliwowym lub podobnym. Nie oznacza to, że w ogóle nie można ich lutować ręcznie, ale jest to bardzo trudne.

Niektóre części mogą mieć również na spodzie obudowy duże „pady termiczne”, które należy przylutować. O ile płytka drukowana nie jest zaprojektowana z otworem przelewym przez PCB, aby umożliwić doprowadzenie lutu z drugiej strony, nie jest praktyczne lutowanie ich ręcznie (choć z dużym powodzeniem może być używana ręczna stacja lutownicza na gorące powietrze).

Pakiety i części opisane do tej pory są do pewnego stopnia standardowe. Istnieje również wiele elementów SMD, które są dostarczane w unikalnych obudowach. Płytkę treningową SMD Silicon Chip-a ma dwie takie części: gniazdo mini USB i uchwyt na baterię pastylkową.

Oznaczenia elementów SMD

Oznaczenia na częściach SMD mogą nie być oczywiste, nawet jeśli są obecne, ale te dla rezystorów (powyżej pewnego rozmiaru) są na szczęście dość proste.

Zamiast kodu koloru, są one po prostu drukowane (lub wykonywane laserowo) z numerycznym odpowiednikiem kodu koloru. Rezystor 10 kΩ do montażu przelotowego miałby paski

w kolorze brązowym, czarnym, pomarańczowym lub brązowym, czarnym, czarnym, czerwonym, wskazujące 10, po którym następują trzy zera lub 100, po którym następują dwa zera.

Rezystor SMD 10 kΩ będzie po prostu oznaczony jako „103” lub „1002”. Należy pamiętać, że nie ma kodu tolerancji.

Niestety, zwykle kondensatory SMD z dielektrykami ceramicznymi typowo nie są w ogóle oznaczone. W takim przypadku wszystko, co można zrobić, to upewnić się, że części są dobrze oznakowane w opakowaniu i pracować tylko z jedną wartością na raz.

Układy scalone również mogą być trudne do identyfikacji, ponieważ zwykle mają tajemnicze kody wykonane na niewielkiej przestrzeni dostępnej na ich górnej powierzchni. Części SOIC mogą być wystarczająco duże, aby mieć sensowny kod, ale części SOT-23 są na to zbyt małe. Niektórzy producenci mogą nawet używać tego samego kodu, którego inny producent użył dla innej, niekompatybilnej części. Dane katalogowe części zwykle wskazują, jakie kody zostały użyte.

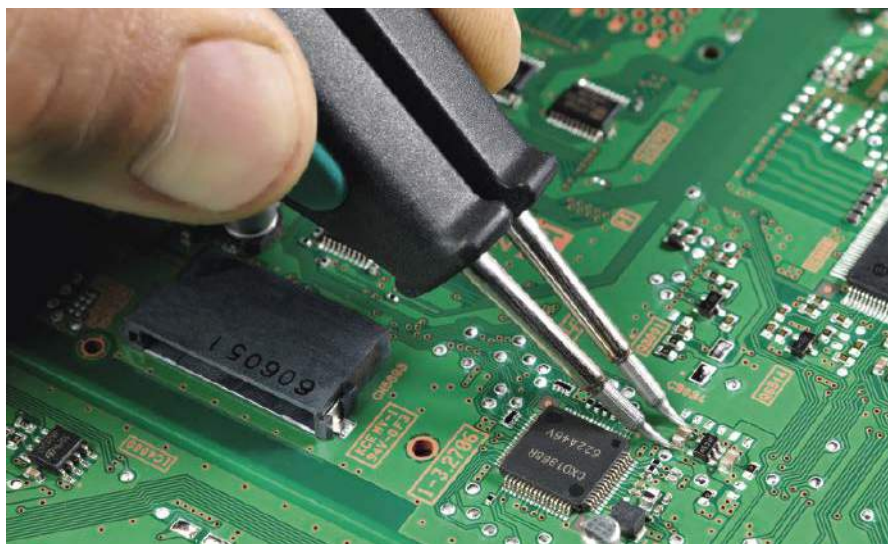
Układy scalone mają również oznaczenia wskazujące ich orientację. Zazwyczaj oznaczenie ma na celu wyróżnienie wyprowadzenia o numerze 1. Może to być wgłębienie w plastikowej obudowie lub skos wzdłuż jednej krawędzi. Może to być również symbol wykonany na górnej powierzchni elementu.

Wgląd do noty katalogowej jest najlepszym sposobem, aby dowiedzieć się, jakiego rodzaju będzie to oznaczenie. Na warstwie opisowej (legend, silk screen) PCB zazwyczaj oznaczamy lokalizację wyprowadzenia 1 małą kropką lub «1».

Niektóre elementy SOIC będą miały również wycięcie i skos zaznaczone na warstwie opisowej PCB, odpowiadające tym cechom, które mogą występować na układzie scalonym. Należy jednak pamiętać, że różni producenci zamiennych elementów mogą stosować różne metody oznaczania styku 1.

Ponieważ najmniejsze komponenty SMD nie są przeznaczone do ręcznego montażu, zazwyczaj nie mają wyraźnych oznaczeń. Zamiast tego skomputeryzowana maszyna do pobierania i umieszczania ich na PCB jest zaprogramowana tak, aby znać ich orientację na taśmie ze szpuli, na której są dostarczane; nota katalogowa powinna to definiować.

Ponieważ diody LED są spolaryzowane, one również zwykle mają oznaczenie elektrod. Może się ono różnić, ale zwykle jest to zielona kropka lub kształt litery T oznaczający katodę albo mały trójkąt, który odpowiada kierunkowi trójkąta w symbolu diody, a tym samym również wskazuje na katodę.



Podczas lutowania przydatna do przytrzymywania elementów jest pęseta. Można również zakupić pęsety z rdzeniami grzewczymi, które mogą być używane do rozlutowywania, jak pokazano na tym zdjęciu. Źródło zdjęcia: https://commons.wikimedia.org/wiki/File:Soldering_a_0805.jpg



Do nakładania topnika można użyć wielu różnych przyborów, takich jak powyższy długopis. Generalnie zalecamy używanie strzykawki z topnikiem w postaci żelu zamiast długopisu lub pojemnika z pastą, ponieważ strzykawka jest łatwa w użyciu do nakładania a sam topnik w jej wnętrzu nie wysycha

Narzędzia i materiały eksploatacyjne

Ten artykuł jest przeznaczony dla względnie początkujących, więc zakładamy, że posiadasz podstawowe narzędzia przeznaczone do lutowania części przewlekanych. Oznacza to lutownicę (najlepiej z regulacją temperatury) i trochę cyny lutowniczej w formie drutu. Z odrobiną ostrożności możesz użyć tych narzędzi do zlutowania pierwszej sekcji płytki treningowej SMD, chociaż pomocnych będzie kilka dodatkowych przyborów.

Pęseta

Będziesz potrzebował czegoś do przytrzymania części podczas lutowania. Niewielki rozmiar oznacza, że nie można używać palców; nawet gdyby były wystarczająco małe, bardzo szybko by się poparzyły! Idealna będzie pęseta z cienkimi końcówkami.

Zestawy takie jak TH1752 firmy Jaycar lub T2374 firmy Altronics są całkowicie wystarczające, chociaż w przypadku mniejszych części mogą być pomocne precyzyjne końcówki. Prawie wszystko, co można opisać jako pęsetę, będzie lepsze niż nic.

Topnik

Praktycznie każdy lut w postaci drutu zawiera topnik lub kalafonię, zwykle wystarczające do lutowania połączeń przelotowych. Ale prawdopodobnie nie zdasz sobie sprawy z korzyści, jakie może przynieść oddzielny topnik, dopóki nie zaczniesz go używać.

Chociaż możesz być przyzwyczajony do tego, że drut lutowniczy „po prostu działa”, w rzeczywistości odpowiedzialny jest za to w dużej mierze rdzeń żywiczny (żywica, np. kalafonia z niektórych drzew iglastych, stanowi doskonały topnik). Istnieją inne, bardziej nowoczesne, a nawet syntetyczne topniki, ale żywice (zwane po oczyszczeniu kalafoniami) są nadal używane, ponieważ są wystarczająco skuteczne.

Jeśli kiedykolwiek próbowałeś ponownie użyć lutowni z odzysku, wiesz, że nie lutuje ono tak dobrze, jak nowe. Nie jest to spowodowane jego wiekiem, ale zużyciem topnika.

Wynika to przede wszystkim z powstawania na lutowanych elementach tlenków metali, które gromadzą się z czasem, gdy metale reagują z tlenem w powietrzu. Jedną z cech topnika jest to, że jest on środkiem redukującym; prostym wyjaśnieniem tej właściwości jest to, że może on odwrócić proces utleniania.

Topnik reaguje z tlenkami, pozostawiając czysty metal, który lepiej wiąże. Wiele topników tworzy również warstwę chroniącą przed tlenem i zapobiegającą dalszemu utlenianiu, co dotyczy również samego lutowni, ścieżek PCB i wyprowadzeń komponentów.

Inną cechą topnika jest to, że powinien być aktywowany ciepłem i działać tylko w pobliżu temperatury lutowania. Zapobiega to jego przedwczesnemu zużyciu.

Topnik może również poprawić transfer ciepła. Ponieważ wszystkie powierzchnie muszą być podgrzane powyżej punktu topnienia lutowni, aby umożliwić jego dobre związanie, topnik może pomóc w dostarczeniu ciepła tam, gdzie jest ono wymagane. Topnik może być nakładany bezpośrednio na części i płytkę drukowaną podczas montażu powierzchniowego, ułatwiając przenoszenie ciepła z lutownicy zarówno do wyprowadzeń komponentów, jak i padów lutowniczych na płytce PCB.

Topnik reaguje również z różnymi tlenkami i zanieczyszczeniami, neutralizując ich negatywny wpływ na proces lutowania. Produkty reakcji nazywane są (w dosłownym tłumaczeniu) żużłem. Rezultatem jest często ciemna, lepka substancja, osad, który zbiera się na grotcie lutownicy.

Topnik może być również silnie żrącą substancją chemiczną i może uszkodzić płytkę, jeśli zostanie na niej dłużej. W nocie katalogowej topnika ten aspekt powinien zostać wyczerpująco opisany. Topniki sprzedawane jako „nie wymagające czyszczenia” rzadziej pozostawiają żrące pozostałości.

Dostępne są topniki w płynie, długopisy z topnikiem i pasty z topnikiem; preferujemy

pastę lub żel, ponieważ łatwiej je nakładać i regulować ich ilość oraz dłużej się utrzymują (topniki w płynie szybciej odparowują). Przy tej ilości lutowania, jaką wykonujemy, nawet dość mała strzykawka wystarcza na lata (lub przynajmniej do momentu gdy minie jej okres przydatności do użycia), więc nie ma potrzeby kupowania dużej ilości pasty topnikowej.

Aby ułatwić posługiwanie się takim topnikiem, zalecamy zakup małej strzykawki, takiej jak H1650A Flux Gel Syringe firmy Altronics. Strzykawka pozwala na precyzyjną aplikację niewielkich ilości topnika.

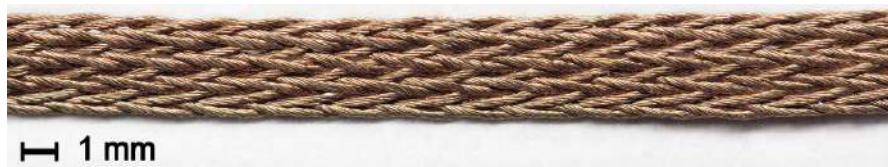
Czyszczenie

Ważne jest, aby wyczyścić płytkę PCB po lutowaniu, zwłaszcza jeśli używasz dużej ilości topnika (co nie jest złym pomysłem, ponieważ skutkuje bardziej niezawodnym połączeniem). Prawdopodobnie konieczne będzie czyszczenie grot lutownicy.

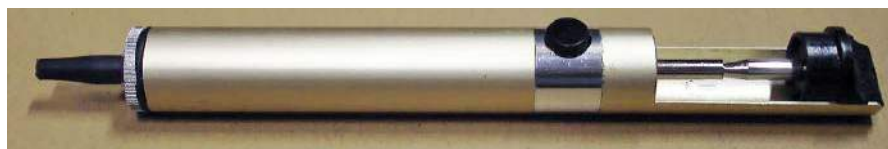
Najczęstszym wyborem jest tutaj gąbka czyszcząca; lekko ją zwilż, tylko na tyle, aby grot nie przypalił gąbki. Widzieliśmy mosiężne czyściki, które działają całkiem dobrze, ale nie wydają się mieć zdolności do wychwytywania wszystkich pozostałości. W razie potrzeby dobrze sprawdza się lekko zwilżony ręcznik papierowy. **Od Red. EdW: Mosiężne lub stalowe czyściki (nawet te kuchenne, do zmywania naczyń użyte w miejsce tych oryginalnych) sprawdzają się całkiem dobrze i są zazwyczaj wygodniejsze w użyciu i stabilniejsze niż nawilżana gąbka.**

Rozpuszczalniki

Dla większości topników zalecany jest również po lutowaniu środek czyszczący (nawet dla tak zwanych topników nie wymagających czyszczenia). Alkohol izopropylowy (izopropanol) to rozsądny, uniwersalny wybór. Niektóre topniki i ich osady są lepkie i mogą wymagać



To jest zbliżenie miedzianej plecionki lutowniczej. Jest ona zwykle sprzedawana na rolce i służy do usuwania nadmiaru lutowni. Źródło zdjęcia: https://commons.wikimedia.org/wiki/File:Solder_wick_close_up.jpg



Lutowniczy odsysacz cyny jest lepszy do usuwania większej ilości lutowni, podczas gdy plecionka jest lepsza do mniejszych zadań, takich jak usuwanie zwarć w postaci nadmiaru lutowni pomiędzy wyprowadzeniami komponentów SMD. Źródło zdjęcia: https://commons.wikimedia.org/wiki/File:Solder_sucker.jpg



Chociaż nie musi to być zestaw typu „wszystko w jednym”, szkło powiększające, uchwyt do PCB i dobre oświetlenie pomogą ułatwić lutowanie małych elementów. Jest to wspomniane poniżej stanowisko Jaycar TH1987

bardziej energicznego działania w celu prawidłowego oczyszczenia płytek PCB.

Z uwagi na powyższe jeszcze lepszą opcją jest specjalistyczny środek do usuwania resztek topników, taki jak Kleanium Deflux-It G2 Flux Remover firmy Chemtools (siliconchip.com.au/link/abad).

Zachowaj ostrożność z tymi rozpuszczalnikami. Wiele z nich, w tym alkohol izopropylowy, jest łatwopalnych, a niektóre są trujące lub mogą uszkodzić skórę. Karta charakterystyki rozpuszczalnika to najlepsze miejsce, w którym można znaleźć porady i informacje na ten temat.

Obecność topnika nie powinna utrudniać testowania większości obwodów niskiego napięcia, ale należy przed podłączeniem zasilania usunąć go z obwodów napięcia sieci 230 V. Zanieczyszczenia wychwycone przez topnik mogą utworzyć ścieżkę przewodzącą, która byłaby niebezpieczna przy takich napięciach.

Należy również oczyścić płytkę PCB z topnika, aby móc ją prawidłowo sprawdzić. Topnik i osady mogą zasłaniać mostki lutownicze i „zimne luty”. Najlepiej jest czyścić na bieżąco, zamiast zostawiać wszystko na koniec, ponieważ topnik jest łatwiejszy do usunięcia, gdy jest ciepły.

Czyść przy użyciu odpowiednich środków chemicznych. Najlepiej używać nylonowych szczotek i/lub niestrzępiących się ściereczek, ponieważ nie chcesz pozostawić włókien na płytce. Nie rozpylaj ani nie wylewaj

roztworu czyszczącego na płytkę; musisz go usunąć, gdy będzie miał szansę rozpuścić topnik. Czasami pozostawienie roztworu do splukania spowoduje usunięcie dużej części topnika, ale nadal będziesz musiał wytrzeć go do sucha.

Może się okazać, że proces czyszczenia jest niedoskonały lub, co gorsza, ujawnia awarię lutowania. Nie pozostaje nic innego, jak wrócić i naprawić problem, a następnie wyczyścić miejsce lutowania i sprawdzić ponownie jakość wykonanego połączenia.

Plecionka lutownicza

Jest to miedziana taśma z drobno tkanego drutu miedzianego, bez izolacji, który zwykle został zaimpregnowany jakiegoś

rodzaju topnikiem. Służy do odprowadzania (lub pochłaniania) nadmiaru lutowni.

Typowym zastosowaniem jest usuwanie nadmiaru lutu, który utworzył mostek między dwoma wyprowadzeniami lub usuwanie lutu z pola lutowniczego po usunięciu wadliwej części, przed zamontowaniem nowej.

Plecionka jest dość tania; można kupić małą rolkę o długości ponad 1 m za kilka dolarów od Jaycar (Cat NS3020) lub Altronics (Cat T1206A) (i oczywiście w sklepie AVT). Typowe użycie może pochłonąć kilka milimetrów opłotu, więc on również wystarczy na dość długo.

Uchwyt PCB

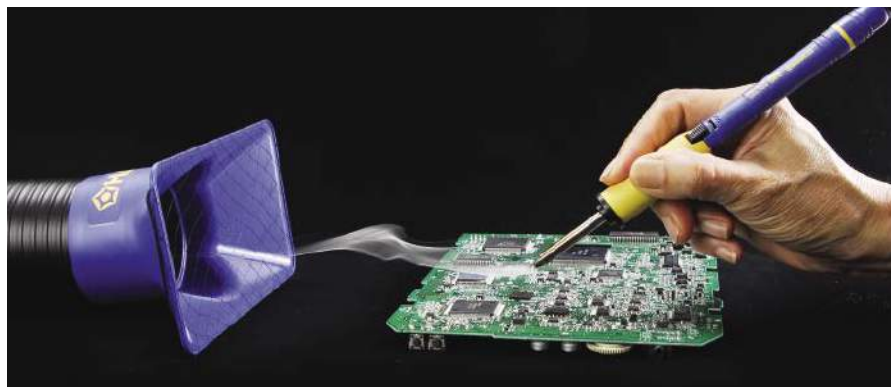
Dla niewielkich płytek PCB zawierających małe komponenty SMD bardzo przydatny może okazać się uchwyt unieruchamiający płytkę na czas montażu. Przydatna jest również możliwość manipulowania trzymaną płytką w dowolnym kierunku w celu uzyskania dostępu do określonego komponentu pod określonym kątem.

Narzędzia takie jak TH1982 Third Hand PCB Holder firmy Jaycar lub T2356 Spring Loaded PCB Holder firmy Altronics, popularnie zwane „trzecią ręką”, są idealne. Płytkę PCB jest utrzymywana w miejscu, ale można ją dowolnie ustawiać lub obracać cały uchwyt, aby umożliwić dostęp pod różnymi kątami.

Chociaż narzędzia te nie są drogie, nawet coś takiego jak Blu-Tack lub podobna masa unieruchamiająca wielokrotnego użytku może być przydatnym prowizorycznym zamiennikiem. Chociaż ciepło z lutownicy prawdopodobnie zmiękczy i zmatowi Blu-Tack-a, nigdy nie mieliśmy żadnych problemów z użyciem go do przytrzymania płytki drukowanej na miejscu.

Lupy

Możliwość wyraźnego zobaczenia drobnych części i elementów związanych z projektami SMD jest najważniejsza. Istnieją dwa ważne



Jeśli pracujesz w zamkniętym pomieszczeniu, ważna jest jakaś forma odsysania oparów. Chociaż ten odcąg Hakko FA-430 (Silicon Chip z sierpnia 2011; siliconchip.com.au/Article/1121) może być poza budżetem niektórych hobbystów, można zamiast tego użyć matego wentylatora, aby zdmuchnąć opary

sposoby na poprawę jakości widzenia: powiększenie i oświetlenie.

Jeśli masz bystre oczy i pracujesz z niektórymi większymi częściami w obudowach SOIC i M3216/1206, możesz sobie poradzić bez powiększenia. Jednak nadal ważne jest, aby przyjrzeć się bliżej swojej pracy i sprawdzić, czy wszystko jest tak, jak powinno.

Na szczęście istnieje szeroka gama przybórów, których można użyć do powiększenia, a niektóre z nich możesz już mieć, na przykład proste ręczne szkło powiększające.

Niektóre uchwyty do PCB zawierają lupę, w tym uchwyt do PCB z lupą LED TH1987 firmy Jaycar. Zawiera on również podstawkę pod lutownicę.

Drugą skrajnością jest mikroskop. Choć z pewnością nie jest tak tani, nie nadszarpnie Twojego budżetu, nie wymaga bowiem dużego powiększenia. Wiele mikroskopów zapewnia również doskonałe oświetlenie. Obecnie dostępnych jest wiele mikroskopów USB i cyfrowych.

Aparat fotograficzny w smartfonie to też odpowiedni sprzęt, który większość Czytelników ma już w kieszeni. Podobną opcją jest aparat cyfrowy z wizjerem LCD i opcją „Makro”.

Może być konieczne użycie funkcji zoomu (nawet zoom cyfrowy będzie bardzo pomocny), aby

zobaczyć rozsądną ilość szczegółów. Jeśli urządzenie posiada tryb „Makro”, będzie ono również lepiej dostosowane do obserwacji z bliska.

Zazwyczaj okazuje się jednak, że przydatna jest stacjonarna lupa, którą można zamocować nad płytką drukowaną, a także mała lupa ręczna, którą można wziąć do ręki i użyć w razie potrzeby.

Oświetlenie

Dobre oświetlenie jest najważniejsze dla udanej pracy z elementami SMD. Najlepsze jest rozproszone źródło światła, ponieważ punktowe źródła mogą powodować cienie, które zasłaniają części płytki drukowanej, zwłaszcza między wyprowadzeniami komponentów, gdzie mogą tworzyć się mostki.

Jeśli masz tylko punktowe źródła światła, skieruj je z przeciwnych stron, aby zniwelować cienie. Możesz rozproszyć światło, odbijając je od czegoś białego, takiego jak ściana, sufit lub arkusz papieru.

Jeśli dobrze widzisz to, co chcesz zobaczyć, prawdopodobnie masz wystarczającą ilość światła.

Wyciągi oparów

Należy pamiętać, że topnik generuje również dym, którego wdychanie jest niezdrowe.

Zalecany sposób radzenia sobie z tym problemem jest wyciąg oparów, ale może on być kosztowny. Niewielki wentylator (taki jak wentylator komputerowy) może również działać, ustawiony tak, aby kierował dym z dala od użytkownika.

Jeśli nie możesz poradzić sobie z aktywną kontrolą oparów, inną opcją jest praca na zewnątrz (lub w pobliżu dużego otwartego okna).

Najlepszy sprzęt

Jeśli jeszcze ich nie masz, wszystkie elementy, o których wspomnieliśmy do tej pory, są dostępne w dość niskich cenach. Pokróćce omówimy również kilka elementów, które mogą jeszcze bardziej poprawić komfort pracy z SMD.

Jak zauważyliśmy wcześniej, podstawowa lutownica jest prawdopodobnie wystarczająca do pracy z większymi częściami SMD. Kiedy zaczynasz zajmować się mniejszymi częściami, niektóre opcjonalne funkcje stają się niezbędne.

Pomocne będą dwa aspekty. Cienki grot pozwoli na dokładniejsze lutowanie, ponieważ zazwyczaj chcesz operować tylko z jednym wyprowadzeniem na raz (ale zobacz sekcję poniżej dotyczącą lutowania z przeciąganiem; większe grotu mogą być lepsze w przypadku



Końcówka punktowa lub stożkowa (średnia)



Średnie grotu stożkowe są w miarę uniwersalne. Dobrze sprawdzają się przy lutowaniu elementów przewlekanych i większych elementów SMD. Ich zaletą jest możliwość użycia pod praktycznie dowolnym kątem.



Końcówka punktowa lub stożkowa (mała)



Drobniejsze grotu stożkowe pozwalają lutować mniejsze wyprowadzenia, więc są bardziej odpowiednie do lutowania dużych i średnich elementów SMD, a także mniejszych elementów przewlekanych.



Końcówka płaska ścięta typu wkrętak/dłuto



Szeroki obszar styku grotów w kształcie dłuta sprawia, że są one przydatne do rozgrzewania i prowadzenia plecionki lutowniczej w celu usunięcia lutu, a także podgrzewania wyprowadzeń SMD lub ponownego lutowania wyprowadzeń po jednej stronie podzespołu.



Końcówka typu ostrze noża, do SOIC, PLCC etc.



Podobnie jak dłuto, grot w kształcie noża może stykać się z dużym obszarem płytki. Kąt nachylenia końcówki sprawia, że jest ona wygodniejsza do pracy po bokach układów scalonych.



Końcówka ścięta duża



Grotu skośne mogą stykać się z jeszcze większym obszarem, ale większe końcówki, takie jak ta, są zazwyczaj zbyt duże, aby dostać się w pobliże mniejszych komponentów.



Końcówka ścięta mała



Mniejsze grotu kątowe są nie tylko łatwiejsze w operowaniu, ale można je również ustawić pod kątem, aby stykały się tylko z jedną krawędzią lub w razie potrzeby z całą powierzchnią.



Końcówka z wgłębieniem na lut



Grot SMD typu flow jest podobny do grotu kątowego, ale ma wgłębienie, które może utrzymać porcję stopionego lutu. Dzięki temu idealnie nadaje się do lutowania jednocześnie wielu styków blisko siebie.

Metryczne		Calowe
0402	-	01005
0603	-	0201
1005	-	0402
1608	-	0603
2012	-	0805
2520	-	1008
3216	-	1206
3225	-	1210
4516	-	1806
4532	-	1812
5025	-	2010
6332	-	2512

■	1×1 mm
■	0,1×0,1 cala
■	1×1 cm

Ten diagram pokazuje w rzeczywistych rozmiarach typowe komponenty SMD. Metryczny komponent M0402 jest tak mały, że ledwo go widać!

tych technik). Krawędź grotu-dłuta może być wystarczająco wąska, aby pracować z relatywnie małymi rozmiarami.

Stacja lutownicza z regulowaną temperaturą jest nieoceniona podczas pracy z większymi częściami. Wiele z tych stacji jest wyposażonych w stojaki i gąbki, które również są pomocne.

Wreszcie, dysza do lutowania gorącym powietrzem ze stacji lutowniczej może być bardzo przydatna do odlutowywania części SMD lub lutowania rozplwowego niektórych trudniejszych części. Stacje są dostępne w zaskakująco niskich cenach, szczególnie używane, i warto mieć taką, jeśli planujesz dużo pracy z komponentami do montażu powierzchniowego.

Korzystanie z narzędzi

Podsumowując powyższe porady, upewnij się, że masz pastę topnikową, gąbkę do czyszczenia grotu lutownicy i odpowiedni



Przykład lutowania na fali pokazujący płytkę PCB opuszczającą część grzewczą maszyny i przenoszoną do wanny z falą lutowniczą. Źródło zdjęcia: <https://youtu.be/VVH58QrprVc>

rozpuszczalnik do wybranego topnika. Używaj topnika obficie i utrzymuj grot lutownicy w czystości.

Techniki lutowania

Jeśli czytałeś już któryś z naszych artykułów na temat budowy układów SMD, to poniższe informacje będą Ci znane. Omówimy nawet dość szczegółowo sposób korzystania z narzędzi, o których właśnie wspomnieliśmy. Możesz również śledzić postępy dzięki dołączonym zdjęciom.

Nałóż topnik na pola stykowe danych komponentów. Dobrym pomysłem jest praca w małych grupach podobnych komponentów. Na przykład, możesz zaplanować pracę ze wszystkimi rezystorami 10 kΩ, jeśli jest ich wiele.

Jeśli istnieje niewielka liczba różnych wartości, można pracować z nimi równolegle. Jednym z wyjątków są kondensatory, które, jak zauważyliśmy wcześniej, zwykle nie mają żadnych wyróżniających je oznaczeń. W takim przypadku zalecamy trzymanie się jednej wartości na raz.

Z grubsza umieść komponenty na ich miejscach. Żel topnikowy lub pasta mają na ogół wystarczającą lepkość, aby utrzymać je na miejscu. Może się okazać, że pęseta zbierze niewielkie ilości topnika, a następnie przyklei się do komponentów. To kolejny powód, aby utrzymywać wszystko w czystości.

Dopasuj element za pomocą pęsety, aby wyprowadzenia układu były wycentrowane względem padów lutowniczych na płycie PCB. Ilość widocznego dedykowanego pola lutowniczego PCB będzie decydować o łatwości operowania lutownicą, więc symetryczne umieszczenie podzespołu jest nie tylko schludne, ale kluczowe dla sprawnego lutowania.

W przypadku małych wyprowadzeń pomocne może być również nałożenie topnika na górną część wyprowadzenia. Oczyszcz grot lutownicy i nałóż na niego niewielką ilość świeżego lutu.

Delikatnie przytrzymaj element płasko na PCB za pomocą pęsety i dotknij grotem zarówno pola lutowniczego, jak i wyprowadzenia. Przytrzymaj przez chwilę, aby



Niektóre popularne grotu lutownicze; większość z nich nadaje się do pracy z elementami SMD



Podczas lutowania metodą przeciągania lutu zazwyczaj używa się grotu typu flow lub grotu skośnego. Najłatwiejszym sposobem na nauczenie się lutowania metodą przeciągania byłoby obejrzenie filmu, takiego, jakich wiele można znaleźć na YouTube Źródło: <https://youtu.be/nyele3CIs-U>

Tabela 2. Popularne rodzaje lutów

Typ lutu	Skład/Nazwa	Temperatura topnienia	Komentarz
Na bazie ołowiu	SnPb 60/40%	188°C	Wyższe stężenie cyny (Sn) prowadzi do większej wytrzymałości
	SnPb 63/37%	183°C	Eutektyka topi się/zestala w jednej temperaturze
Bezołowiowy	Sn100C	227°C	Bezsrebrowy; może zawierać miedź, nikiel, german
	SAC305	217...220°C	Zawiera cynę, srebro i miedź; stosowany w lutowaniu na fali
	SnCu	217...232°C	Zawiera cynę i miedź; lutowie bezołowiowe na bazie cyny jest często używane do lutowania rozplwowego i na fali
	SAC387	217...219°C	Zawiera cynę, srebro i miedź
	Sn42Bi52	138°C	Niskotopliwy, do wrażliwych części
Topnik	Kalafonia	NA	Ułatwia lutowanie
	Nie zawiera kalafonii	NA	Często zawiera halogenki metali, takie jak chlorek cynku, kwas solny, kwas cytrynowy itp.; może być żrący
Lut twardy	Srebro, miedź, mosiądz, brąz itp.	>450°C	Często używane w biżuterii i zaprojektowane tak, aby miały temperaturę topnienia tuż poniżej temperatury topnienia odpowiedniego metalu

wyprowadzenia elementu oraz pad lutowniczy mogły się nagrzać i połączyć z lutowiem. Powinieneś zobaczyć, jak lut przepływa z grotu na wyprowadzenie i pad.

Odsuń grot i kontynuuj przytrzymywanie części na miejscu, aż lut się nieco ostudzi i przyjmie formę stałą. Jedna sekunda będzie wystarczająca dla małych części z cienkimi wyprowadzeniami, więcej czasu może być potrzebne dla większych komponentów.

Jeśli część przesunęła się lub nie leży płasko na płytce drukowanej, chwyć ją pęsetą i podgrzej, aby stopić lutowie. Dostosuj pozycję elementu, aż będziesz zadowolony.

Jeśli wygląda na to, że część jest ponownie dobrze wyrównana i płaska względem płytki drukowanej, nałóż trochę świeżego lutowia na grot i przeprowadź przez pozostałe pady.

W przypadku bardzo wąskich lub drobnych padów, w pierwszej kolejności przykładaj grot lutownicy do padów lutowniczych. Maski lutownicze na płytce drukowanej pomoże zapobiec przepływowi lutu tam, gdzie nie powinien. Staramy się wydłużać pady lutownicze w wielu projektach SMD opracowanych w Redakcji Silicon Chip-a, aby ułatwić lutowanie, chociaż nie we wszystkich projektach jest to widoczne.

W zależności od lutownicy, rozmiaru ścieżki i topnika, lut może zostać przyciągnięty do pola lutowniczego i wyprowadzenia wyłącznie przez napięcie powierzchniowe lutu. Zjawisko

trudniej zaobserwować stosując lutownicę kolbową, gdy każde z wyprowadzeń jest lutowane niemal oddzielnie. Przy korzystaniu ze stacji lutowniczej na gorące powietrze, przyciągany do padów lutowniczych komponent (samocentrumujący się komponent) jest zjawiskiem każdorazowym. Zachowanie się lutu i jego napięcie powierzchniowe przy pracy z SMD w małych skalach jest krytyczne, więc powinno Ci to pomóc w poprawnej pracy.

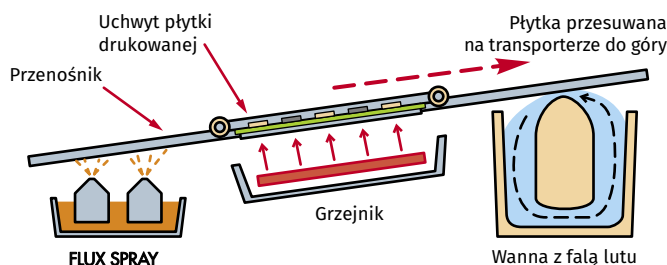
Być może widziałeś części lutowane pastą lutowniczą w piecu rozplwowym; gdy lut upłynnia się, element przyciągany jest do pól lutowniczych i wręcz zatraskuje się na właściwym miejscu. Wynika to z napięcia powierzchniowego, lokalizacji pól lutowniczych i oddziaływania maski lutowniczej.

Napięcie powierzchniowe przyciąga również lut dokładnie tam, gdzie jest potrzebny.

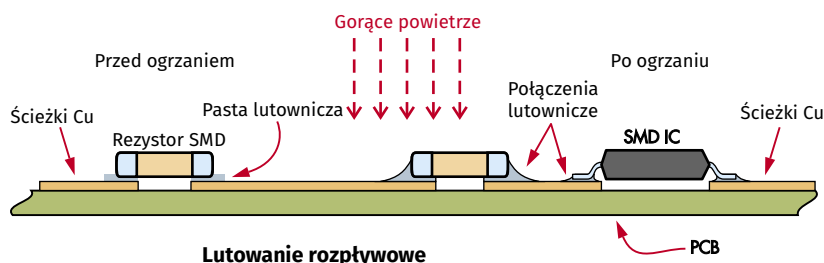
Jeśli części płasko przylegają do PCB, wymagana jest tylko niewielka ilość lutu. Jeśli widzisz czyste, zakrzywione luty, które dobrze rozplnęły się wokół padu lutowniczego i wyprowadzenia, jest to dobry znak, że połączenie jest dobrze uformowane.

Możesz wykorzystać napięcie powierzchniowe, aby nałożyć dużą ilość lutu i zapewnić mocne połączenie. Wybrzuszone, ale czyste i błyszczące połączenie z pewnością będzie bardziej funkcjonalne i solidne niż niewidoczna cienka warstewka, pod warunkiem, że nie będzie się ono stykało z żadną inną częścią!

Ruchy lutownicą powinieneś wyćwiczyć metodą prób i błędów. Czas lutowania będzie również zależał od takich rzeczy, jak temperatura grotu i wybór lutowia (ołowiowego lub bezołowiowego).



Podstawowe zasady lutowania na fali. Płytkę PCB jest przenoszona nad kąpielą lutowniczą przez przenośnik. W pewnym momencie lut jest wypychany w górę „falą”, tak, że spód płytki przechodzi przez nią. Komponenty i miedziane pady są lutowane, a następnie płytka wychodzi z kąpieli



◀ Lutowanie rozplwowe w ogóle nie wykorzystuje lutownicy – do roztopienia pasty lutowniczej nałożonej na pady lutownicze i wyprowadzenia komponentu, które mają zostać przylutowane, wykorzystuje się gorące powietrze o kontrolowanej temperaturze lub podczerwień. Płytkę przechodzi przez podczerwień. Płytkę przechodzi „po krzyku” – lutowane połączenie gotowe

Jeśli wystąpi mostek lutowniczy, dopóki część jest prawidłowo wyrównana, kontynuuj lutowanie pozostałych wyprowadzeń. Następnie pozbadź się nadmiarowych połączeń.

Użyj techniki opisanej wcześniej, aby usunąć lut z mostków na wyprowadzeniach. W razie potrzeby nałóż więcej lutu (zwłaszcza jeśli nie masz łatwego dostępu do mostka). Nałóż topnik, plecionkę (patrz poniżej), a następnie użyj lutownicy. Pozwól plecionce wchłonąć nadmiar lutu, a następnie ostrożnie odsuń lutownicę i plecionkę.

Dokładnie obejrzyj część za pomocą lupy. Jeśli połączenie wydaje się luźne lub nieczyste, nałóż świeży topnik i delikatnie dotknij czystą końcówką lutownicy po kolei każdego wyprowadzenia. Przekonasz się, że nawet ten krok poprawiania każdego wyprowadzenia pomoże rozprzecznić lut tam, gdzie jest potrzebny.

Lutowanie z przeciąganiem

Gdy komponenty SMD mają wyprowadzenia, które są bardzo blisko siebie, niepraktyczne staje się lutowanie ich pojedynczo. Jedynym komponentem na płycie treningowej PCB SMD, który uważamy za posiadający tak ciasne odstępy między stykami, jest IC2, w obudowie SSOP, z rastrem wyprowadzeń 0,65 mm. Niektóre układy mają jeszcze ciaśniejszy raster styków, zmniejszający się do około 0,4 mm (np. TQFP-144).

W takich przypadkach łatwiej jest lutować takie układy scalone metodą przeciągania lutownicy. **Po przyklejeniu układu (!!!)** i nałożeniu topnika na styki, niewielką ilość lutowni nakładasz na grot lutownicy, a następnie delikatnie przeciągasz wzdłuż rzędu styków. Napięcie powierzchniowe przeciąga niewielką ilość lutowni z grota na pady lutownicze i leżące na nich wyprowadzenia lutowanego układu. Wykonane prawidłowo

przeciąganie tworzy idealne połączenia za pierwszym razem.

Generalnie lepiej jest nałożyć za dużo lutu niż za mało, ponieważ mostki są łatwiejsze do zauważenia niż zimne luty, czyli połączenia z niewystarczającą ilością lutu. Mostki możesz łatwo usunąć za pomocą plecionki (patrz poniżej) i większej ilości topnika.

Można stosować specjalne groty lutownicy do lutowania metodą przeciągania, z „dołkami” (włgłębieniami) do przytrzymania lutowni, ale można również użyć, gdy masz wprawę, standardowego grotu. Trzeba tylko częściej dodawać do niego nieco więcej lutu (np. co 5...10 lutowanych styków, zamiast co 30...40 wyprowadzeń).

Nawet układy scalone o większym rastrze, takie jak typy SOIC, mogą być lutowane przy użyciu tej techniki; może to być szybsze (i schludniejsze) niż lutowanie ich wyprowadzeń pojedynczo.

Używanie miedzianej plecionki lutowniczej

Miedziana plecionka lutownicza jest najlepszym sposobem do usuwania małych ilości lutu, podczas gdy odsysacz cyny jest lepszy do usuwania dużych ilości. Jeśli więc masz dużo lutu do usunięcia, np. przy lutowanych elementach przewlekanych, zacznij od odsysacza, aby usunąć większość spoiwa, i zakończ plecionką, aby precyzyjnie usunąć jego resztki.

Jednak przy niewielkich rozmiarach połączeń lutowniczych związanych z częściami SMD, odsysacz cyny staje się nieporęczny i mogą po prostu wciągać części SMD, razem z lutem, który próbujesz usunąć. Ilości lutu, które trzeba usunąć, również będą niewielkie.

Przed użyciem plecionki warto dodać topnik. Słowo „topnik” pochodzi od łacińskiego słowa „fluere”, oznaczającego płynąć; chcemy zachęcić lut do spłynięcia na plecionkę.

Przyciśnij czystą część plecionki grotom do lutu i pozwól, aby wszystko rozgrzało się na tyle, aby stopić lut; powinien zacząć wsiąkać w plecionkę. Wykonana z miedzi plecionka dobrze przewodzi ciepło, więc należy ostrożnie uchwycić ją palcami lub użyć pęsety.

Po tym, jak plecionka wchłonie lut, ostrożnie odsuń zarówno grot, jak i plecionkę, przesuwając je w poprzek płytki drukowanej. Nie powinienesz najpierw usuwać lutownicy, bo zostaniesz z plecionką przylutowaną do płytki drukowanej!

Czasami może pomóc dodanie większej ilości lutu w miejscu, w którym chcesz go usunąć, zwłaszcza jeśli jest to mostek lutowniczy schowany głęboko między dwoma wyprowadzeniami. Dodatkowa objętość lutu może dać plecionce więcej powierzchni do kontaktu. **Od Red. EdW: metoda z dodawaniem lutu przydaje się szczególnie przy wylutowywaniu podzespołów z PCB, zwłaszcza lutowanych w technice bezołowiowej. Dodanie w takiej sytuacji większej ilości lutu ołowiowego lub niskotopliwego Sn42Bi58 radykalnie obniża temperaturę topnienia spoiwa i zdecydowanie ułatwia wylutowywanie.**

Jeśli po użyciu plecionki na płycie drukowanej pozostaje ciemny osad, jest to prawdopodobnie produkt uboczny działania topnika. W przypadku takich obszarów można użyć wacika nasączonego rozpuszczalnikiem do czyszczenia topnika, aby oczyścić małe obszary przed kontynuowaniem pracy. ■

Tim Blythman

Adaptacja do wydania polskiego – Andrzej Nowicki

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. www.siliconchip.com.au

REKLAMA



MECHANIC
SHOW THE STYLE OF MASTER

TEMPERATURA TOPNIENIA	SKŁAD PASTY
183°C	SN63/PB37
143°C	SN63/PB37
138°C	SN42/BI58

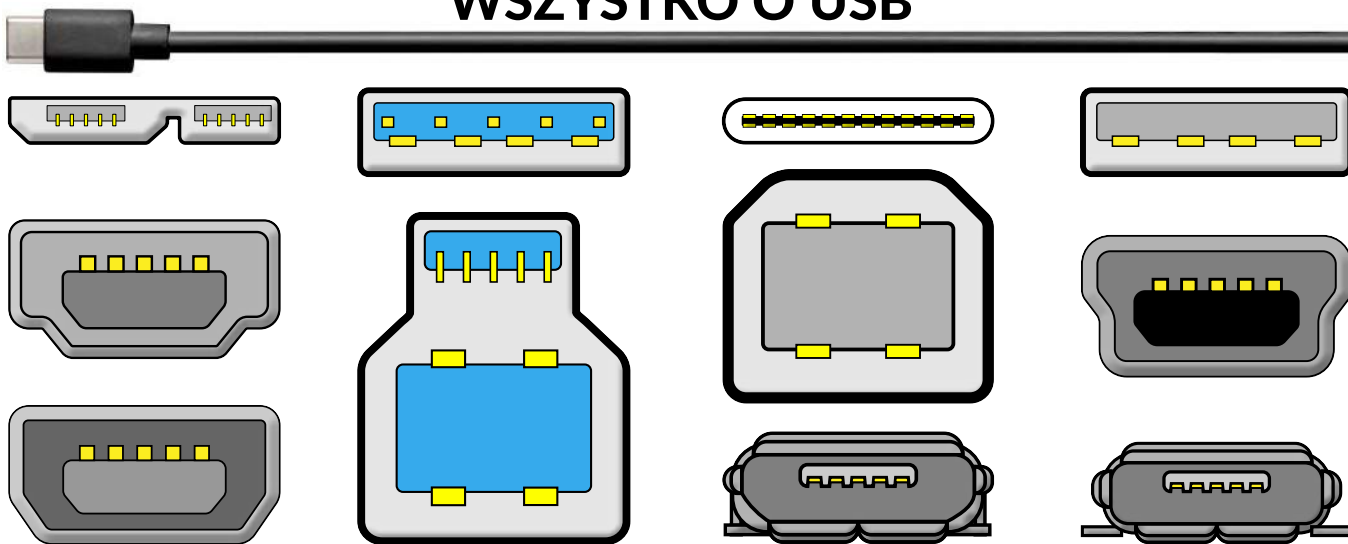


🌐 sklep.avt.pl
☎ +48 456 455 240
📍 Leszczynowa 11, 03-197 Warszawa

eprasa.pl 40635921f8

Pasty lutownicze





Historia Uniwersalnego Złącza Szeregowego. Ekspansja USB

Około 30 lat temu grupa firm, aby ułatwić podłączanie urządzeń zewnętrznych do komputerów PC, opracowała uniwersalną magistralę szeregową USB, zastępującą mnóstwo złączy i interfejsów, które były używane wcześniej. Znacznie zwiększyło to również szybkość komunikacji, w porównaniu z istniejącymi protokołami szeregowymi (i równoległymi). Od tego czasu wydajność interfejsu USB i zakres jego zastosowania znacznie wzrosły.

Kiedy w latach 70. pojawiła się pierwsza generacja komputerów PC lub komputerów osobistych – maszyny takie jak MITS Altair, Commodore PET, Tandy TRS-80 i Apple II – były nieco ograniczone pod względem możliwości łączenia się z urządzeniami peryferyjnymi, takimi jak drukarki, modemy i zewnętrzne napędy taśmowe lub dyskowe.

Ale kiedy IBM wypuścił w 1981 roku swój pierwszy komputer PC (model 5150), wszystko zaczęło się zmieniać. Komputer IBM 5150 był dostępny z maksymalnie dwoma wbudowanymi napędami dyskietek, 16 KB pamięci RAM i kolorową kartą graficzną CGA (dla której dostępny był kolorowy monitor). Co ważne, posiadał również sloty z tyłu płyty głównej na karty interfejsów, które udostępniały port równoległy drukarki Centronics, Game Port oraz jeden lub dwa porty szerebowe RS-232C.

Wkrótce można było również podłączyć do płyty głównej komputera dysk twardy o pojemności 10 MB.

Następnie zaczęło pojawiać się wiele nowych komputerów PC, z których większość oferowała podobne funkcje. Około 1990 roku każdy dostępny komputer miał zwykle 640 KB pamięci RAM, wbudowany dysk twardy o pojemności 20...40 MB, kolorową kartę graficzną oraz port drukarki Centronics i kilka portów szeregowych RS-232C. Wiele z nich posiadało również kartę Ethernet, dzięki czemu można je było podłączyć do sieci LAN (sieć lokalna).

Zaczęły pojawiać się również różne bardziej wyspecjalizowane interfejsy; na przykład do podłączenia magistrali GPIB w celu sterowania z komputera zewnętrznymi urządzeniami czy przyrządami naukowymi i badawczymi. Pojawiło się również

złącze Fire-Wire (IEEE1394), szeregową magistralę o wysokiej przepustowości zaprojektowaną do wydajnego podłączania urządzeń peryferyjnych, takich jak szybkie dyski twarde czy pamięci taśmowe. Wkrótce z tyłu komputerów znalazło się wiele różnych złączy interfejsów, umożliwiających podłączenie wielu urządzeń peryferyjnych, często dedykowanych do danego interfejsu.

Narodziny USB

Rozwój USB rozpoczął się w 1994 roku, gdy grupa firm mocno zaangażowanych w branżę PC (Compaq, DEC, IBM, Intel, Microsoft, NEC i Nortel) porozumiała się i postanowiła ułatwić podłączanie urządzeń zewnętrznych do komputerów PC.

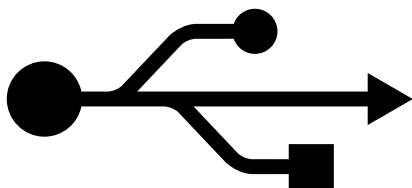
Polegało to na zastąpieniu wszystkich różnych złączy interfejsów grupą prostszych,



Logo „Certified USB” pojawia się na urządzeniach USB, które zostały sprawdzone i uznane za działające w sposób akceptowalny przez USB Implementers Forum



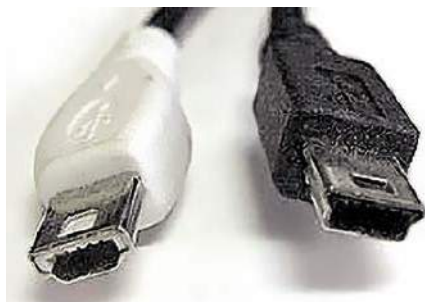
Oryginalny kabel USB do podłączania urządzeń peryferyjnych, takich jak drukarki, z pełnowymiarową wtyczką typu A na końcu komputera (po prawej) i wtyczką typu B na końcu do podłączenia urządzenia (po lewej)



Ta ikona USB lub jej odmiana zwykle pojawia się na wszystkich lub prawie wszystkich wtyczkach i gniazdach kompatybilnych z USB



Odmiana logo certyfikacji USB-IF, które pojawia się na urządzeniach zgodnych z USB 2.0+ pracujących przy transmisji 480 Mb/s



Kwadratowa wtyczka typu B jest zbyt duża dla smukłych urządzeń, takich jak smartfony i tablety, więc zaprojektowano mniejsze wtyczki mini typu A (po lewej) i typu B (po prawej). Są one nadal używane przez niektóre tańsze urządzenia, zwłaszcza te, które używają USB 5 V do ładowania, ale w większości zostały zastąpione przez gniazda mikro typu B lub typu C



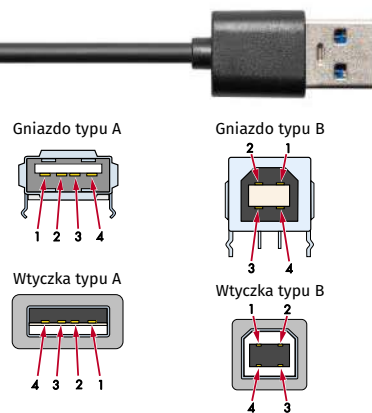
Mikrowtyczka typu B (i odpowiednie, nie pokazane tutaj gniazdo USB) rozwiązały kilka problemów z gniazdem mini typu B; oprócz tego, że są znacznie smuklejsze, zostały również zaprojektowane tak, aby wtyczka zużywała się znacznie wcześniej niż gniazdo, zapobiegając przedwczesnemu zużyciu gniazda i ostatecznie konieczności jego wymiany

identycznych złączy wielofunkcyjnych, z których każde mogłoby być konfigurowane przez oprogramowanie do wykonywania różnych zadań związanych z transmisją danych. Tak narodziła się uniwersalna magistrała szeregową, niemal natychmiast identyfikowaną przez akronim USB.

Oficjalna specyfikacja USB 1.0 została opublikowana w styczniu 1996 roku i definiowała dwie szybkości transmisji danych: 1,5 Mb/s (187,5 KB/s), zwaną Low Speed lub Low Bandwidth (przeznaczoną dla urządzeń peryferyjnych, takich jak klawiatury, myszy i joysticki) oraz 12 Mb/s (1,5 MB/s), zwaną Full Speed (do obsługi urządzeń peryferyjnych o wyższej prędkości, takich jak drukarki, napędy dyskowe itp.). Dla wyjaśnienia, skrót „B” oznacza „bajty”, skrót „b” oznacza „bity”.

Intel wyprodukował w 1995 roku pierwsze układy scalone interfejsu zaprojektowane do obsługi USB, ale do sierpnia 1998 roku, kiedy to wydano specyfikację USB 1.1, niewiele urządzeń USB lub komputerów wyposażonych w porty USB pojawiło się na rynku. Wkrótce specyfikacja USB została powszechnie przyjęta, prowadząc do tego, co Microsoft nazwał „komputerem bez obciążenia z przeszłości”.

Złącza USB używane w tych początkowych implementacjach to płaski prostokątny wtyk (lub gniazdo) typu A, dla portów „downstream” (od serwera do klienta) w komputerze PC oraz kwadratowe gniazdo typu B dla portu „upstream” (od klienta do serwera) w urządzeniu peryferyjnym, takim jak drukarka (rysunek 1).



Styk	Nazwa	Kolor izolacji	Opis
1	Vbus	Czerwony lub pomarańczowy	+5 V
2	D-	Biały lub złoty	Dane-
3	D+	Zielony	Dane+
4	GND	Czarny lub niebieski	Masa

Rysunek 1. Podłączenia typu A i typu B
Złącza USB, od których wszystko się zaczęło. Prostokątne złącze typu A pojawia się w urządzeniach głównych, takich jak komputery, zasilacze, huby itd. Natomiast kwadratowe złącze typu B jest używane w urządzeniach peryferyjnych, takich jak drukarki. Utrudnia to przypadkowe podłączenie dwóch portów hosta lub urządzenia, co w najlepszym przypadku nic by nie dało, a w najgorszym spowodowałoby uszkodzenie

Oba złącza miały tylko cztery styki, dwa do przesyłania danych i dwa do dostarczania zasilania 5 V DC do urządzenia peryferyjnego. Większość urządzeń zewnętrznych była podłączana do komputera za pomocą kabla USB wyposażonego we wtyczkę typu A na końcu podłączanym do komputera i wtyczkę typu B na końcu podłączanym do urządzenia. Obudowa obu typów wtyczek była oznaczona charakterystycznym logo USB znanym jako „trójkąb” (patrz wyżej).

Taka sytuacja miała miejsce do kwietnia 2000 roku, kiedy to wydano poprawioną specyfikację USB 2.0. Dodano trzecią szybkość transmisji 480 Mb/s (60 MB/s), nazwaną High Speed lub High Bandwidth, oprócz Low Speed i Full Speed. Umożliwiło to również stosowanie złączy i kabli mini A i mini B, a wkrótce także złączy i kabli mikro USB.

Zarówno złącza mini USB, jak i mikro USB były wyposażone w pięć kontaktów, a dodatkowy styk służył do identyfikacji typu urządzenia, peryferyjnego lub hosta, do którego były podłączone („ID urządzenia”).

Specyfikacja USB 2.0 pozwoliła również dwóm urządzeniom peryferyjnym na bezpośrednią komunikację, z pominięciem hosta PC – funkcja zwana USB On-The-Go lub USB-OTG.

Specyfikacja została również rozszerzona o wsparcie dla dedykowanych ładowarek baterii, a także zezwolenie na zwiększony przepływ prądu w kablu USB między komputerem a urządzeniem peryferyjnym, w porównaniu

z pierwotnym limitem 500 mA (lub 100 mA dla nieskonfigurowanych urządzeń).

W czasach, gdy urządzenia miały mieszankę portów USB 1.1 i USB 2.0, porty USB 1.1 były zwykle oznaczone wewnątrz za pomocą białego plastiku, podczas gdy porty USB 2.0 używały wewnątrz czarnego plastiku.

Pojawienie się USB 3.0

Specyfikacja USB 3.0 została wydana w listopadzie 2008 roku, kiedy to ogólne zarządzanie standardem USB przeniesiono z USB 3.0 Promoter Group do USB Implementers Forum (USB-IF). USB 3.0 dodał kolejny tryb transferu: SuperSpeed, zapewniający nominalną szybkość przesyłania danych 5,0 Gb/s (625 MB/s) oprócz trzech uprzednio istniejących szybkości transferu.

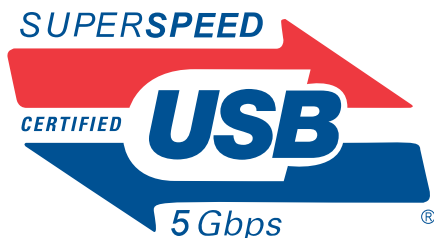
Komunikacja w trybie SuperSpeed odbywa się w pełnym duplexie, podczas gdy w trzech wcześniejszych trybach był to półdupleks.

USB 3.0 wprowadziło również protokół UASP, który zapewnia ogólnie większą prędkość transferu niż protokół Bulk-Only-Transfer (BOT) zapewniany przez USB 1.X i USB 2.0.

Specyfikacja USB 3.0 dodała szereg kompatybilnych wstecznie wtyczek, gniazd i kabli. Wtyczki i gniazda SuperSpeed mają łącznie dziewięć kontaktów (4+5) i są oznaczone wyraźnym logo oraz wewnętrzną warstwą izolacyjną w kolorze niebieskim, w przeciwieństwie do czarnego lub białego koloru używanego w złączach USB 1.1 i USB 2.0.



Wyraźnie widać pięć nowych styków (z przodu) i niebieski kolor izolacji tego gniazda USB 3.0



Logo certyfikatu USB-IF dla urządzeń zgodnych ze specyfikacją USB 3.0 lub nowszą i prędkości transmisji 5 Gb/s



Kabel Thunderbolt 1/2 wykorzystuje to samo złącze co Mini DisplayPort i zastępuje złącze FireWire w komputerach Apple. Łączy on sygnały PCI Express i DisplayPort oraz zapewnia zasilanie prądem stałym (fot. <https://w.wiki/o26>)



Thunderbolt 3 (na fotografii pokazane gniazdo) wykorzystuje kable z wtyczką USB typu C, ale zapewnia więcej funkcji niż tylko przenoszenie danych; oferuje również większą moc zasilania niż USB (np. do ładowania laptopów) i może transmitować wideo, a nawet sygnały złącza PCI Express. USB 4.0 zasadniczo implementuje funkcje Thunderbolt do standardu USB (fot. <https://w.wiki/o27>)

Urządzenia o niskiej i wysokiej prędkości nadal działają w USB 3.0, ale urządzenia SuperSpeed mogą korzystać ze wzrostu dostępnego prądu od 150 mA do 900 mA. Różne złącza i kable USB opiszemy szczegółowo nieco później.

W lipcu 2013 roku wydana została specyfikacja USB 3.1, zapewniająca dwie odmiany trybu USB 3.0 SuperSpeed: USB 3.1 Gen1, podobny do oryginalnej specyfikacji USB 3.0, oraz USB 3.1 Gen2, który wprowadza nowy tryb transferu SuperSpeed+.

SuperSpeed+ podwaja maksymalną szybkość transmisji danych do 10 Gb/s (1,25 GB/s), jednocześnie zmniejszając obciążenie linii kodowaniem do zaledwie 3% poprzez zmianę protokołu kodowania na 128 b/132 b.

Następnie we wrześniu 2017 r. wydano specyfikację USB 3.2. Wprowadziła ona dwa dodatkowe tryby transferu SuperSpeed+, zaprojektowane w celu wykorzystania 24 kontaktów (2x12) w nowo opracowanych złączach USB typu C.

Chociaż dwa rzędy 12 styków zostały początkowo zaimplementowane symetrycznie, aby umożliwić podłączenie złączy typu USB-C w dowolny sposób, specyfikacja USB 3.2 wykorzystuje je do zapewnienia pracy wielopasmowej (przy użyciu dodatkowych przewodów w kablu), aby umożliwić przesyłanie danych z prędkością 10 Gb/s lub 20 Gb/s (2,5 GB/s).

Gdy komputery mają zarówno porty USB 2.0/3.0, jak i USB 3.1/3.2, zazwyczaj porty USB 3.1/3.2 będą oznaczone kolorem turkusowym lub żółtym, porty USB 3.0 pozostaną niebieskie, a starsze porty będą miały czarny (USB 2.0) lub biały (USB 1.1) plastik.

Thunderbolt 1, 2 i 3

Pod koniec 2008 roku Apple wprowadziło zminiaturyzowaną wersję cyfrowego interfejsu audio-wideo DisplayPort, nazwaną Mini DisplayPort lub MiniDP. Zastąpił on port DVI w większości modeli Apple, takich jak MacBook, MacBook Air, MacBook Pro, iMac, Mac Mini i Mac Pro. Port MiniDP zaczął również pojawiać się w notebookach Asus, Microsoft, MSI, Lenovo, Toshiba, HP, Dell i innych producentów.

Następnie, na początku 2011 roku, Intel i Apple ogłosiły uzgodnienie interfejsu sprzetowego Thunderbolt, który łączył funkcje PCI Express i MiniDP i zastąpił FireWire (IEEE1394). Kable Thunderbolt łączą w sobie transmisję elektryczną i światłowodową, przy czym miedziane przewody są zwykle używane do przesyłania energii, podczas gdy światłowody przesyłają z dużą prędkością dane. Wykorzystują one 20-stykowe złącze MiniDP.

W czerwcu 2013 roku Intel ogłosił specyfikację Thunderbolt 2, która wykorzystywała te same złącza co Thunderbolt 1, ale podwoiła szybkość transmisji danych do 20 Gb/s (2,5 GB/s) poprzez połączenie dwóch kanałów 10 Gb/s.

Pierwszymi produktami konsumenckimi wykorzystującymi Thunderbolt 2 były płyta główna Asus Z87-Deluxe/Quad i MacBook Pro Retina firmy Apple, oba wyroby wypuszczone na rynek w drugiej połowie 2013 roku.

Następnie w 2015 roku Intel ogłosił specyfikację Thunderbolt 3, która ponownie podwoiła maksymalną szybkość transmisji danych do 40 Gb/s (5 GB/s), jednocześnie zmniejszając o połowę zużycie energii. Wykorzystując miedziane kable i 24-stykowe złącza USB-C, które zostały wprowadzone w 2014 roku, Thunderbolt 3 może zawierać USB Power Delivery i przesyłać do 100 W mocy wraz z szybką transmisją danych.

Urządzenia z portami Thunderbolt 3 stały się dostępne w listopadzie 2015 roku, w tym notebooki firm Acer, Asus, Clevo, HP, Dell, Dell Alienware, Lenovo, MSI, Razer i Sony z systemem Microsoft Windows, a także płyty główne firmy Lintex Technology. Następnie, w październiku 2016 roku, Apple wypuściło zmodyfikowany MacBook Pro, wyposażony w dwa lub cztery porty Thunderbolt 3, w zależności od modelu.

USB-C

Specyfikacja złącza USB typu C lub USB-C została sfinalizowana przez USB-IF w sierpniu 2014 roku i jest przede wszystkim związana z miniaturowymi 24-stykowymi (2x12-stykowymi) złączami USB-C.

Początkowo złącza te były używane z interfejsami USB 3.1, dzięki czemu można je było wkładać do gniazd w dowolny sposób (są one również znacznie bardziej wytrzymałe niż wtyczki mini lub mikro typu B). Kiedy jednak w 2015 roku pojawił się Thunderbolt 3, zaczęto ich używać również w tym interfejsie. A kiedy pod koniec 2017 roku pojawiło się USB 3.2, zyskały one również możliwość transmisji SuperSpeed+.

Należy jednak pamiętać, że urządzenie wyposażone w złącze USB-C niekoniecznie implementuje USB, USB Power Delivery lub którykolwiek ze zdefiniowanych trybów alternatywnych. Tryb alternatywny dedykuje niektóre fizyczne przewody w kablu USB-C 3.1 do bezpośredniej transmisji między urządzeniem a hostem innych protokołów danych, takich jak DisplayPort.

Cztery szybkie linie, dwa styki pasma bocznego (tylko w przypadku zadokowanych, odłączanych urządzeń i stałych kabli), dwa styki danych USB 2.0 i jeden styk konfiguracyjny

mogą być używane do transmisji w trybie alternatywnym. Tryby są konfigurowane za pomocą komunikatów VDM (Vendor Defined Messages – Wiadomości Definiowane/Konfigurowane przez Nadawcę) poprzez kanał konfiguracyjny.

Ostatnio, na forum europejskim, złącze USB-C zostało wybrane jako uniwersalne złącze do ładowania telefonów komórkowych, smartfonów etc.

USB4

Kolejna implementacja złącza USB-C została zdefiniowana w sierpniu 2019 r., kiedy USB-IF opublikowało specyfikację USB4. USB4 opiera się na protokole Thunderbolt 3. Obsługuje przepustowość danych 40 Gb/s (5 Gb/s), jest kompatybilny z Thunderbolt 3 i wstecznie kompatybilny z USB 3.2 i USB 2.0.

Architektura definiuje metodę dynamicznego współdzielenia pojedynczego szybkiego łącza z wieloma urządzeniami końcowymi, zaprojektowaną w celu optymalizacji transferu danych według typu i aplikacji.

Thunderbolt 4

Thunderbolt 4 został zapowiedziany w styczniu 2020 roku na targach CES (Consumer Electronics Show), a ostateczna specyfikacja została wydana w lipcu 2020 roku.

Główne ulepszenia to obsługa protokołów USB 4 i szybkości transmisji danych, minimalna wymagana przepustowość 32 Gb/s dla łącza PCIe, obsługa dwóch wyświetlaczy 4K (lub jednego 8K) oraz ochrona bezpośredniego dostępu do pamięci (DMA) oparta na Intel VT-d w celu zapobiegania fizycznym atakom DMA.

Maksymalna przepustowość pozostaje na poziomie 40 Gb/s, takim samym jak w przypadku Thunderbolt 3 i czterokrotnie większym niż USB 3.2 Gen2x1. Mimo to, minimalna przepustowość, którą producenci są zobowiązani wdrożyć, została podwojona z 16 Gb/s, wcześniej dozwolonych w specyfikacji Thunderbolt 3.

USB Power Delivery (USB PD)

W lipcu 2012 roku, USB Promoter Group sfinalizowała specyfikację USB Power Delivery (USB PD rev.1), aby umożliwić uniwersalne zasilanie lub ładowanie laptopów, tabletów, dysków zasilanych przez USB i podobnej elektroniki użytkowej o wyższej mocy. Jest to logiczne rozszerzenie istniejących europejskich i chińskich standardów ładowania telefonów komórkowych.

Rozszerzenie USB PD rev.1 określa użycie certyfikowanych kabli USB „PD aware” ze standardowymi złączami USB typu A i typu B, aby dostarczyć zwiększoną moc (ponad 7,5 W) do urządzeń o większym zapotrzebowaniu na energię.

Urządzenia mogą żądać wyższych prądów i napięć od zgodnych hostów – do 2 A przy 5 V (10 W) i opcjonalnie do 3 A albo 5 A przy 12 V (36 W lub 60 W) lub 20 V (60 W lub 100 W).

We wszystkich przypadkach obsługiwane są zarówno konfiguracje host-to-device, jak i device-to-host (gospodarz-gość w obu kierunkach). Protokół konfiguracji zasilania wykorzystuje binarny kanał transmisji FSK (kluczowanie z przesunięciem częstotliwości) 24 MHz na linii Vbus.

Wersja 2.0 specyfikacji USB PD (USB PD Rev.2.0) została wydana w sierpniu 2014 roku jako część specyfikacji USB 3.1.

Obejmuje ona użycie kabli i złącza USB-C z czterema parami przewodów zasilania/uziemia i oddzielnym kanałem konfiguracyjnym, z wykorzystaniem kodowania danych BMC (Biphase Mark Code lub Differential Manchester) o niskiej częstotliwości ze sprzężeniem stałoprądowym w celu zmniejszenia możliwości wystąpienia zakłóceń RF.

Od tego czasu pojawiły się kolejne wersje USB PD Rev.2.0. W marcu 2016 r. wydano wersję 1.2, tworząc nowe zasady zasilania USB PD, które definiują cztery nominalne poziomy napięcia (5 V, 9 V, 12 V i 20 V) i poziomy mocy wyjściowej w zakresie od 0,5 W do 100 W.

Następnie, w styczniu 2017 roku, USB-IF wydało wersję 3.0 USB PD, która definiuje protokół programowalnego zasilania (PPS), który pozwala regulować napięcie Vbus w krokach co 20 mV, aby ułatwić ładowanie baterii zarówno stałym prądem (CC), jak i stałym napięciem (CV).

W styczniu 2018 r. wprowadzono logo „Certified USB Fast Charger” dla ładowarek wykorzystujących programowalny protokół zasilania USB PD 3.0.

Złącza i kable USB

Obecnie w użyciu jest tak wiele różnych złączy USB, że nie jest możliwe szczegółowe omówienie ich wszystkich. Przygotowaliśmy jednak kilka informacji, które pomogą Ci rozpoznać najpopularniejsze typy złączy i kabli.

Jak wspomniano wcześniej, rysunek 1 pokazuje te, które są prawdopodobnie najbardziej znane: gniazdo i wtyczka typu A oraz gniazdo i wtyczka typu B. Są to oryginalne czterostykowe złącza USB, ze złączami typu A przeznaczonymi do użytku na końcu hosta/komputera, oraz złączami typu B na wejściu urządzenia peryferyjnego/końcu zewnętrznego kabla USB łączącego oba urządzenia.

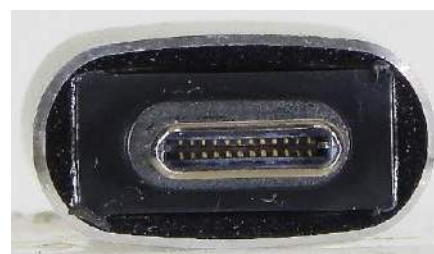
Poniższa tabela przedstawia nazwy zwykle nadawane czterem stykom, wyznaczony kolor izolacji dla każdego przewodu oraz opis funkcji.



Logo używane na najszybszych obecnie dostępnych urządzeniach USB 40 Gb/s



Wtyczki USB typu C mają więcej styków rozmieszczonych symetrycznie, dzięki czemu wtyczkę można włożyć dowolną stroną i połączenie nadal będzie działać



Zbliżenie na wtyczkę USB typu C wyraźnie pokazuje wszystkie 24 styki



Ten symbol pojawia się na wtyczkach i w pobliżu gniazd, które obsługują najszybszą prędkość transmisji 40 Gb/s w standardzie USB 4



Logo certyfikatu USB-IF dla urządzenia obsługującego USB-PD o mocy do 45 W. Obejmuje to dostarczanie na żądanie urządzenia zarówno wyższych napięć, jak i prądów, niż normalne 5 V/500 mA



Mini USB-B gniazdo



Mini USB-B wtyczka



Mikro USB-B gniazdo



Mikro USB-B wtyczka



powyższe złącza pokazano w rzeczywistej wielkości

Styk	Nazwa	Kolor izolacji	Opis
1	Vbus	Czerwony	+5 V
2	D-	Biały	Dane-
3	D+	Zielony	Dane+
4	ID	Bez przewodu	NC we wtyku USB-B
5	GND	Czarny	Masa

Rysunek 2. Podłączenia mini USB-B i mikro USB-B
Te zminiaturyzowane wersje oryginalnego złącza typu B są znacznie bardziej odpowiednie dla mniejszych urządzeń, takich jak telefony komórkowe i tablety. Dodają piąty styk Device ID i co ważne, wtyczka mikro typu B jest zaprojektowana tak, aby zużywała się ona, a nie gniazdo, więc wystarczy wymienić kabel, jeśli wtyczka się zużyje, a nie gniazdo lub całe urządzenie

Należy zauważyć, że w kablu USB przewody D- i D+ są „skręconą parą”, aby zmniejszyć ryzyko zakłóceń elektromagnetycznych (EMI) – zarówno pod względem zmniejszenia emisji, jak i uniknięcia problemów z odbiorem EMI.

Rysunek 2 przedstawia złącza mini USB i mikro USB, nadal używane do podłączania wielu kompaktowych urządzeń, takich jak tablety i PDA (osobiści asystenci cyfrowi czy osobisty menedżer informacji), smartfony i aparaty cyfrowe.

Chociaż istniały wtyczki i gniazda mini typu A, kiedy USB 2.0 i mini USB zostały wprowadzone w 2000 roku, zostały one oficjalnie „wycofane” w 2007 roku wraz z gniazdem mini typu AB. Dlatego też obecnie nie spotkasz ich zbyt wiele, a w rezultacie nie pokazaliśmy ich. To samo dotyczy wtyczek i gniazd mikro typu A.

Warto zauważyć, że chociaż funkcjonalna część wtyczek mikro USB ma podobną szerokość do wtyczek mini USB, to są one o około połowę cieńsze. Mimo to są one przystosowane do co najmniej 10 000 cykli podłączania i odłączania, czyli znacznie więcej niż wtyczki mini USB.

Rysunek 3 przedstawia szczegóły złączy USB 3.0 SuperSpeed, które zostały wprowadzone w 2008 roku. Są one zasadniczo zmodyfikowaną wersją oryginalnych złączy typu A i typu B, z dodanymi pięcioma stykami, aby sprostać wymaganiom protokołu SuperSpeed, przy jednoczesnym zachowaniu wstecznej kompatybilności z USB 1.x i USB 2.x.

Być może najbardziej widoczną na pierwszy rzut oka różnicą (zwłaszcza w przypadku złączy typu A) jest niebieski kolor plastikowej izolacji wewnątrz złączy, w porównaniu z białą lub czarną izolacją wewnątrz wcześniejszych złączy.

Wewnątrz złączy typu A dodatkowe pięć styków znajduje się w niewielkiej odległości od pierwszych czterech, równoległe do nich i nieco dalej od siebie.

Z drugiej strony, w złączach typu B górna część funkcjonalnej części złącza jest wydłużona o około 3 mm, a wszystkie dodatkowe pięć styków jest zamontowanych blisko siebie w węższej górnej części.

Dzięki temu gniazdo USB 3.x typu B może akceptować starsze wtyczki typu B, ale oczywiście wtyczka SuperSpeed typu B nie jest kompatybilna ze starszym gniazdem typu B.

Tak jak poprzednio, tabela pod schematami złączy pokazuje nazwy i przeznaczenie każdego z dziewięciu styków. Należy pamiętać, że styki 5, 6, 8 i 9 mają inną nazwę dla złącza A i złącza B.

Rysunek 4 przedstawia szczegóły charakterystycznych złączy USB 3.0 SuperSpeed mikro B, w których dodatkowe pięć styków znajduje się obok oryginalnych pięciu styków w ich „ściętym prostokącie” w linii z nimi, ale w oddzielnej grupie. Poniższa tabela przedstawia nazwy i przeznaczenie wszystkich dziesięciu styków.

Wreszcie dochodzimy do 24-stykowych złączy USB-C. **Rysunek 5** pokazuje szczegóły gniazda i wtyczki USB-C, dla przejrzystości w dwukrotnie większym rozmiarze. Styki znajdują się w dwóch rzędach, po 12, i są oznaczone jako A1-A12 i B1-B12.

Pierwotnie były one po prostu duplikatami, aby umożliwić podłączenie wtyczki do gniazda w dowolny sposób. Jednak obecnie, aby poradzić sobie z wieloma rozszerzonymi zastosowaniami USB-C, większość styków ma różne funkcje, pokazane w tabeli poniżej wtyczki i gniazda.

Styki A1, B12, B1 i A12 służą teraz „masie zasilania”, podczas gdy A4, B9, B4 i A9 służą zasilaniu Vbus, aby zapewnić dodatkowe

możliwości zasilania dla USB PD. B5 jest również dedykowany dla Vconn, aby zapewnić zasilanie dla kabli zasilających.

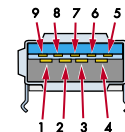
Inną rzeczą, na którą należy zwrócić uwagę w przypadku złączy USB-C jest to, że oprócz oryginalnej różnicowej pary danych USB (przypisanej do styków A6 i A7), zapewniają one również cztery pary ekranowanych par różnicowych dla transmisji danych SuperSpeed, SuperSpeed+, USB 4, Thunderbolt 3 i Thunderbolt 4 o wysokiej przepustowości.

Są one przypisane do styków A2 i A3, B11 i B10, B2 i B3, A11 i A10. Istnieje również linia konfiguracyjna przypisana do A5 i B5 i wreszcie dwie linie „pasma bocznego” przypisane do styków A8 i B8.

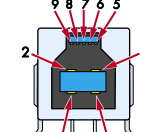
Jak to się rozwinęło...

Oczywiste jest, że USB rozwinęło się dramatycznie w ciągu ostatnich 30 lat, zarówno pod względem wydajności, jak i funkcji. Zmieniło się z systemu mającego na celu uproszczenie podłączania podstawowych urządzeń, takich jak klawiatury i myszy,

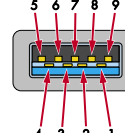
Gniazdo typu A



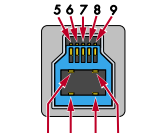
Wtyczka typu B



Wtyczka typu A

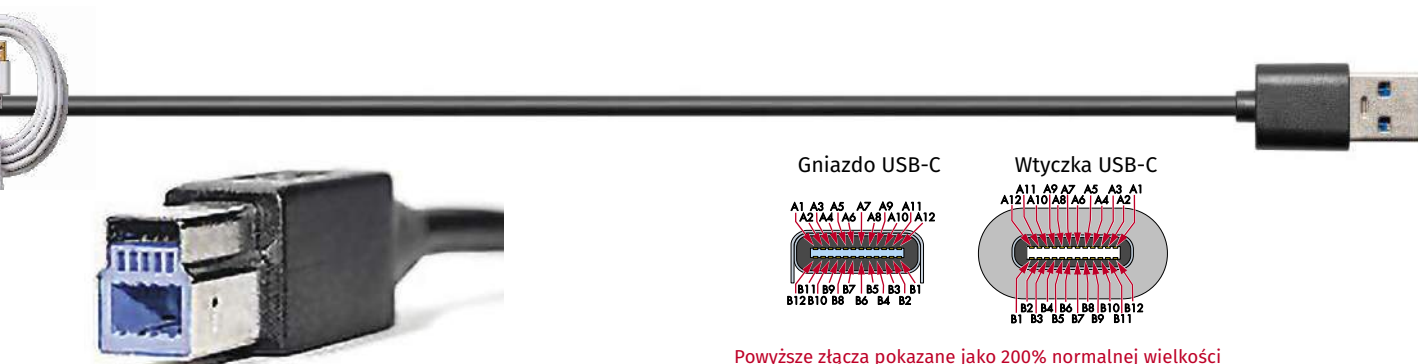


Wtyczka typu B



STYK	Kolor przewodu	NAZWA złącze A	NAZWA złącze B	OPIS
1	czerwony	VBUS	VBUS	zasilanie +5 V
2	biały	D-	D-	para różnicowa USB 2.0
3	zielony	D+	D+	
4	czarny	GND	GND	masa zasilania
5	niebieski	StdA_SSRX-	StdB_SSTX-	Para różnicowa odbiornika SuperSpeed
6	żółty	StdA_SSRX+	StdB_SSTX+	
7	N/A	Dren_GND	Dren_GND	masa sygnałowa
8	purpurowy	StdA_SSTX-	StdB_SSRX-	Para różnicowa nadajnika SuperSpeed
9	pomarańczowy	StdA_SSTX+	StdB_SSRX+	

Rysunek 3. Złącza USB 3.0 SuperSpeed typu A/B
Nowe wtyczki i gniazda USB 3.0/4.0 dodają pięć nowych styków do przenoszenia sygnałów o wyższej przepustowości. Zostały one zaprojektowane tak, aby urządzenia USB 1.0-2.0 mogły być nadal podłączane i działać normalnie na tych samych czterech stykach, z których zawsze korzystały. Urządzenia USB 3.0/4.0 mogą być podłączane do wcześniejszego gniazda typu A, ale dodatkowe styki nie będą się kontaktować, więc komunikacja będzie odbywać się z mniejszą prędkością



Powyższe złącza pokazane jako 200% normalnej wielkości

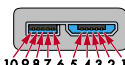
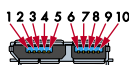
Zamiast dodawać wewnętrznie dodatkowe styki, jak to zrobiono dla USB 3.0 w przypadku wtyczek i gniazd typu A, pełnowymiarowe złącze USB 3.0 typu B rozbudowuje osłonę wtyczki. Ponieważ jednak dolna kwadratowa sekcja jest identyczna z wcześniejszą wtyczką USB 1.0/1.1 i USB 2.0 typu B, starsze kable mogą być nadal używane w urządzeniach z gniazdami akceptującymi tę nowszą wtyczkę



Podobnie, wtyczka USB 3.0 micro typu B dodaje po jednej stronie zupełnie nową sekcję z pięcioma dodatkowymi stykami. Po raz kolejny gniazda są kompatybilne ze starszymi (4-stykowymi) wtyczkami, ale nie na odwrót

Gniazdo mikro typu B

Wtyczka mikro typu B



STYK	NAZWA	OPIS
1	VBUS	zasilanie +5 V
2	D-	para różnicowa USB 2.0
3	D+	
4	ID	Identyfikacja urządzenia
5	GND	masa zasilania
6	SSTX-	Para różnicowa nadajnika SuperSpeed
7	SSTX+	
8	GND	masa sygnałowa
9	SSRX-	Para różnicowa odbiornika SuperSpeed
10	SSRX+	

Rysunek 4. Złącza SuperSpeed USB 3.0 mikro typu B

Złącza SuperSpeed micro są szersze, dodając oddzielną sekcję z pięcioma dodatkowymi stykami obok oryginalnego gniazda. Dlatego też gniazda te są również wstecznie kompatybilne ze starszymi kablami i hostami

w system z co najmniej dziewięcioma różnymi typami złączy – niektóre z aż 24 stykami – i dziewięcioma różnymi prędkościami przesyłania danych, od oryginalnego 1,5 Mb/s aż do 40 Gb/s SuperSpeed+ i Thunderbolt 3.

Styk	Nazwa	Kolor przewodu	Opis
A1, B12 B1, A12	GND	Cynowana plecionka	Masa zasilania
A4, B9 B4, A9	VBUS	czerwony	Zasilanie Vbus
A5	CC1	niebieski	Kanał 1 konfiguracji
B5	CC2	żółty	Kanał 2 konfiguracji
A6	Dp1	zielony	UTP_Dp (USB 2.0 D+)
A7	Dn1	biały	UTP_Dn (USB 2.0 D-)
A8	SBU1	czerwony	Sygnał pasma bocznego, strona A
B8	SBU2	czarny	Sygnał pasma bocznego, strona B
A2	SSTXp1	(żółty)	Ekranowana para różnicowa #1, TX, dodatni
A3	SSTXn1	(brązowy)	Ekranowana para różnicowa #1, TX, ujemny
B11	SSRXp1	(zielony)	Ekranowana para różnicowa #2, RX, dodatni
B10	SSRXn1	(pomarańczowy)	Ekranowana para różnicowa #2, RX, ujemny
B2	SSTXp2	(biały)	Ekranowana para różnicowa #3, TX, dodatni
B3	SSTXn2	(czarny)	Ekranowana para różnicowa #3, TX, ujemny
A11	SSRXp2	(czerwony)	Ekranowana para różnicowa #4, RX, dodatni
A10	SSRXn2	(niebieski)	Ekranowana para różnicowa #4, RX, ujemny

Rysunek 5. Złącza USB-C

Wtyczka i gniazdo USB typu C mają podobny rozmiar do wtyczki i gniazda mikro B, ale są w stanie zapewnić znacznie wyższą prędkość transmisji danych i większą moc zasilania. Są również symetryczne i znacznie bardziej wytrzymałe niż złącza mikro typu B

Zdolność do dostarczania zasilania do urządzeń za pośrednictwem kabla USB również znacznie wzrosła. Ze skromnych 100 mA lub 500 mA przy 5 V dostępnych przez USB 1.x i USB 2.x, USB PD pozwala teraz urządzeniom żądać jednego z czterech różnych poziomów napięcia (5 V, 9 V, 12 V lub 20 V), przy natężeniu prądu do 5 A.

Daje to możliwość uruchamiania urządzeń peryferyjnych o znacznie większej mocy, a także pozwala na szybkie ładowanie baterii większej liczby urządzeń zasilanych akumulatorowo, takich jak laptopy, tablety i telefony komórkowe, za pomocą kabla USB.

W dzisiejszych czasach można podłączyć pojedynczy kabel USB typu C między komputerem przenośnym a monitorem i nie tylko naładować komputer (z wewnętrznego zasilacza monitora), ale także przesyłać sygnały wideo o wysokiej rozdzielczości, a nawet podłączyć klawiaturę, mysz, drukarkę, szybkie urządzenia pamięci masowej i inne.

Zastanawiamy się, czy twórcy USB w ogóle pomyśleliby o takich możliwościach w 1994 roku, kiedy to wszystko się zaczęło! ■

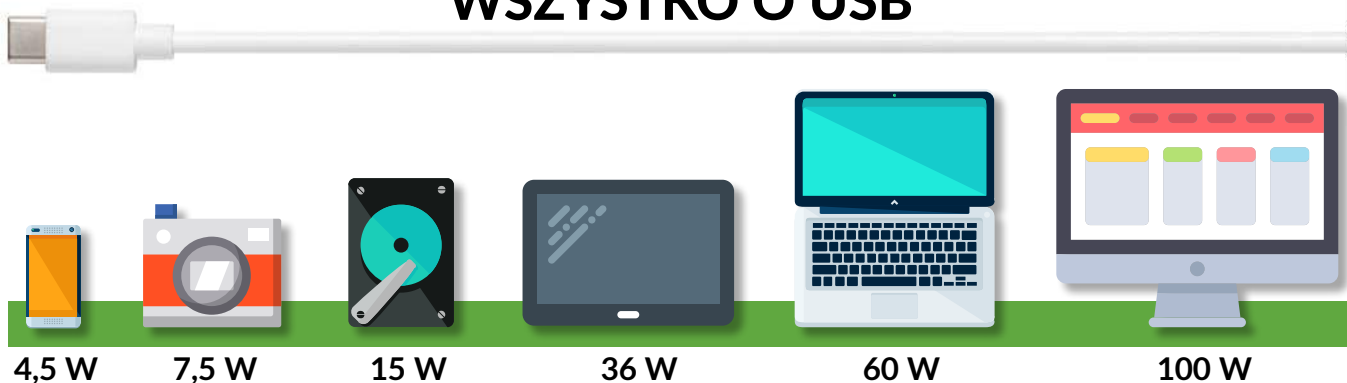
Jim Rowe

Adaptacja do wydania polskiego – Andrzej Nowicki

Więcej informacji

- Standard USB: <https://w.wiki/Usb>
- USB 3.0: <https://w.wiki/ntm>
- USB 4: <https://w.wiki/ntn>
- USB Type-C: <https://w.wiki/nto>
- Thunderbolt: <https://w.wiki/ntp>
- USB PD: <https://w.wiki/ntq>
- USB Implementers Forum: www.usb.org/
- <https://www.digikay.pl/pl/articles/decoding-the-usb-standards-from-1-to-4>

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. www.siliconchip.com.au



Jak działa zasilanie USB-C

Od czasu jego wprowadzenia w połowie lat 90. interfejs USB przeszedł długą drogę. Powszechne przyjęcie USB 3.2 wprowadziło złącze typu C oraz nową funkcję zarządzania przesyłaną energią (Power Delivery – PD), która umożliwia dostarczenie do 100 W (20 V @ 5 A). Jest to nieco bardziej skomplikowane niż poprzednie schematy zasilania USB, ale warto się temu przyjrzeć.

Poprzedni artykuł, pt. „Historia Uniwersalnego Złącza Szeregowego. Ekspansja USB”, opisuje dość szczegółowo złącze USB Typu-C (USB-C) i porusza kwestię nowego mechanizmu USB Power Delivery (siliconchip.com.au/Article/14883). W obu tematach jest jednak znacznie więcej do powiedzenia, więc ten artykuł uzupełni lukę.

Mamy również równoległy artykuł, który opisuje niektóre tanie źródła zasilania zgodne z USB Power Delivery (PD). Ale najpierw przeczytaj poniższe informacje, aby dowiedzieć się, dlaczego warto korzystać z tych urządzeń.

Standard USB 3.2 jest pierwszym, który oficjalnie pozwala źródłom zasilania i odbiornikom (oraz kablom!) negocjować dostarczane napięcie i prąd.

Zanim jednak przejdziemy do szczegółów tego, jak to wszystko działa, musimy wyjaśnić kilka pojęć. Sprawy stały się bardziej skomplikowane od czasów, gdy istniały tylko dwa możliwe urządzenia, do których można było podłączyć kabel USB: „gospodarz” (host), który dostarczał zasilanie i nadzorował

komunikację, oraz zasilane urządzenie odbiorcze, które zużywało energię i komunikowało się z „gospodarzem” – hostem.

Obecnie porty USB mogą mieć trzy możliwości przesyłania danych i trzy możliwości dostarczania zasilania. W przypadku danych są to:

- Porty skierowane do odbiornika (Downstream-facing ports – DFP) – zazwyczaj host lub koncentrator.
- Porty skierowane do nadajnika (Upstream-facing ports – UFP) – zazwyczaj urządzenie odbiorcze.
- Porty o podwójnej roli (Dual-role ports – DRP), które mogą przełączać się między rolą odbiornika i hosta.

Porty USB-C pełnią również rolę zasilania jako:

- Źródło, zasilanie w energię.
- Odbiornik, zużywający energię.
- Port zasilania o podwójnej roli.

Porty zasilania o podwójnej roli mogą przełączać się między rolami źródła i odbiornika. Dobrym tego przykładem jest port USB-C w laptopie, który może być obciążeniem

po podłączeniu do zasilacza (lub do monitora, który ma funkcję zasilacza) w celu zasilania lub ładowania urządzenia, lub źródłem po podłączeniu do portu urządzenia peryferyjnego, takiego jak dysk twardy z interfejsem USB; lub pendrive.

Po podłączeniu, DFP są portami źródłowymi, a UFP są portami odbiorczymi; można to jednak zmienić później po uzgodnieniu w ramach protokołu USB-PD.

Złącze USB-C

Standard USB 3.2 wprowadził nowe złącze typu C, które jest używane na obu końcach kabla USB-C. Złącze to można włożyć w dowolny sposób (jest symetryczne), a wtyczki złącza mogą opcjonalnie zawierać elektronikę, która pozwala systemowi zidentyfikować określone możliwości kabla.

Połączenie USB-C zawiera dedykowany kanał konfiguracyjny (CC), używany do wykrywania obecności kabla, orientacji wtyczki, możliwości kabla i negocjowania „umów o zasilanie” między źródłem a odbiornikiem.

Ponadto USB-C ma dwa superszybkie kanały różnicowe przesyłania danych typu full-duplex do szybkiej wymiany informacji oraz dwie linie boczne. W tej dyskusji skupimy się na CC.

Dwa rzędy styków w gnieździe są pokazane na środku **rysunku 1**, a dwie możliwe orientacje wtyczki (nieodwrócona i odwrócona) są pokazane odpowiednio powyżej i poniżej.

Styki GND i VBUS są symetryczne, więc styki te łączą się prawidłowo niezależnie od orientacji wtyczki. Klasyczne linie USB D+ i D– również działają poprawnie, ponieważ oba styki D+ i oba styki D– w gnieździe są wewnętrznie połączone.

Wtyk nieodwrócony

GND	TX1+	TX1–	VBUS	CC	D+	D–	SBU2	VBUS	RX2–	RX2+	GND
GND	RX1+	RX1–	VBUS	SBU1			VCON	VBUS	TX2–	TX2+	GND

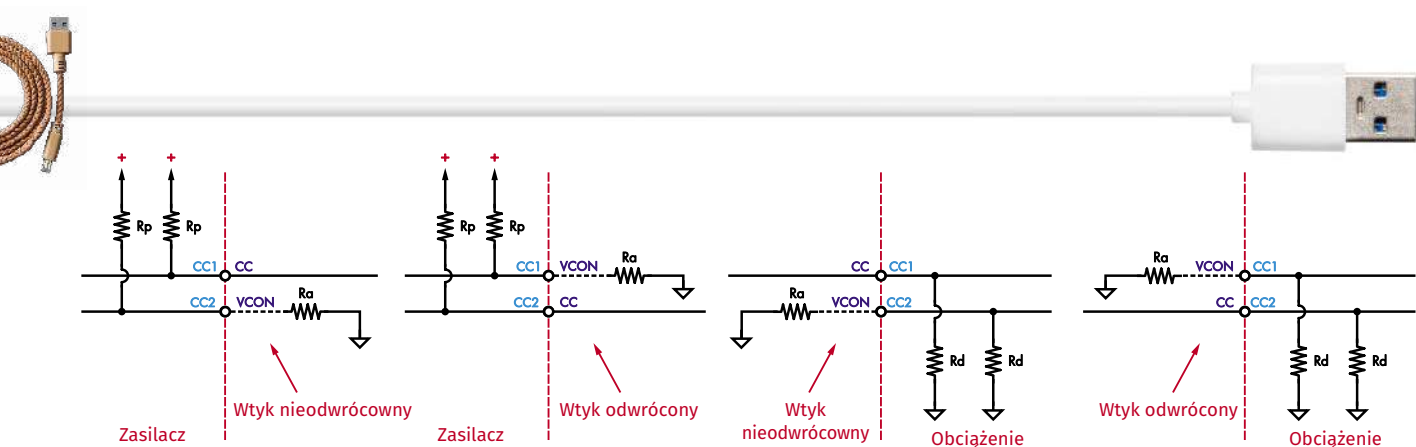
Gniazdo

GND	TX1+	TX1–	VBUS	CC1	D+	D–	SBU2	VBUS	RX2–	RX2+	GND
GND	RX1+	RX1–	VBUS	SBU1	D–	D+	CC2	VBUS	TX2–	TX2+	GND

Wtyk odwrócony

GND	TX2+	TX2–	VBUS	VCON			SBU1	VBUS	RX1–	RX1+	GND
GND	RX2+	RX2–	VBUS	SBU2	D–	D+	CC	VBUS	TX1–	TX1+	GND

Rysunek 1. Przypisanie styków gniazda USB typu C jest pokazane pośrodku. Powyżej i poniżej znajdują się dwie możliwe orientacje wtyczki. Styki CC1 i CC2 w gnieździe łączą się ze stykami CC i VCON w wtyczce; można stwierdzić, w którą stronę wtyczka jest włożona, monitorując, które ze styków: CC1 lub CC2 łączą się z CC



Rysunek 2. Po podłączeniu źródło ustawia za pośrednictwem R_p wysoki poziom logiczny linii CC1 i CC2, a odbiornik za pośrednictwem R_d ustawia ten poziom na niski. W kablu jest tylko jedna linia CC, więc zarówno źródło, jak i odbiornik mogą wykryć, czy wtyczka jest odwrócona. Aktywny kabel ustawia poziom logiczny za pośrednictwem R_a na styku VCON, aby zasygnalizować swoją obecność. Kable pasywne pozostawiają VCON otwarty

Dwie skręcone pary przewodów superszybkiej transmisji danych zostaną zamienione w zależności od orientacji wtyczki, więc urządzenie musi użyć szybkiego multipleksu („muxa”), aby je przekierować, jeśli wtyczka zostanie odwrócona. Podobnie, styki użycia pasma bocznego (SBU) są zamienione i muszą być uporządkowane w zależności od orientacji wtyczki.

Styki, którymi jesteśmy zainteresowani w przypadku zasilania energią, to dwa styki kanału konfiguracji (Configuration Channel – CC) w gnieździe oraz styki CC i VCON we wtyczce. Umożliwiają one portom na każdym końcu kabla określenie orientacji wtyczki, wykrycie podłączenia i określenie możliwości drugiego końca.

Wykrywanie podłączenia i orientacji kabla

Zapoznaj się z **rysunkiem 2**. Gdy połączony są dwa porty USB typu C, DFP domyślnie stanie się źródłem i ustawi, za pomocą rezystorów R_p , dwie linie CC na wysoki poziom logiczny, a koniec odbiorczy kabla ustawi linie CC na niski poziom logiczny za pomocą rezystorów R_d . Źródło wykrywa dołączenie obciążenia, gdy jeden z dwóch styków CC jest ustawiony przez R_d na niski poziom logiczny, a odbiornik wykrywa dołączenie, gdy jeden z jego styków CC jest ustawiany przez R_p na wysoki poziom logiczny.

Po podłączeniu, napięcie źródła jest domyślnie ustawiane na 5 V w celu zapewnienia

kompatybilności i bezpieczeństwa. Obciążenie (odbiornik) może ustalić, jaki prąd może dostarczyć źródło przy napięciu 5 V, mierząc wartość R_p (lub, dokładniej, pobierany prąd). **Tabela 1** pokazuje, jak to działa (dalej).

Źródło i odbiornik mogą określić orientację wtyczki, odnotowując, czy styk CC1 lub CC2 jest ustawiany przez R_p na wysoki poziom logiczny lub przez R_d na niski poziom logiczny.

Elektronicznie sygnowany zespół kabla (Electronically Marked Cable Assembly – EMCA)

Złącza kabla USB typu C mogą opcjonalnie zawierać mikrokontroler, umożliwiając kablowi zgłaszanie swoich możliwości urządzeniom źródłowym i odbiorczym. Przykładowo, standardowy kabel USB typu C może obsługiwać dostarczanie zasilania o natężeniu do 3 A.

Aby skorzystać z maksymalnego natężenia 5 A, należy użyć aktywnego kabla, który może poinformować źródło, że jest żądany prąd o takim natężeniu. Zasilacz nie pozwoli odbiornikowi pobierać więcej niż 3 A, jeśli kabel nie zgłosi, że jest w stanie obsłużyć to natężenie prądu.

Po podłączeniu, kabel wskazuje, że zawiera elektronikę, ustawiając styk VCON na niskim poziomie logicznym przez rezystor R_a , który zwykle ma rezystancję od 800 Ω do 1,2 k Ω . Kable bez logiki pozostawiają VCON na nieokreślonym poziomie (otwarty). Jeśli wykryty zostanie kabel z logiką,

urządzenie źródłowe jest odpowiedzialne za dostarczenie napięcia na odpowiednim styku CC do VCON w celu zasilania mikrokontrolera w kablu, jak pokazano na **rysunku 3**.

Inne tryby

Złącza typu C obsługują dodatkowe tryby, które wykorzystują kanały pasma bocznego (SideBand Use – SBU). Są to tryb akcesoriów audio (na przykład dla słuchawek) i tryb akcesoriów debugowania, w którym kanały SBU mogą być używane do przesyłania sygnałów debugowania.

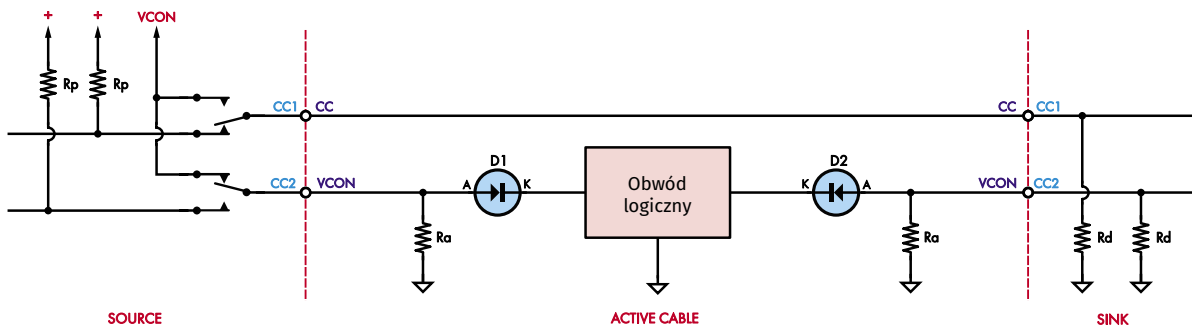
Tabela 2 podsumowuje różne kombinacje ustawiania styków CC na niskim lub wysokim poziomie logicznym, oraz ich znaczenie.

Zasilanie

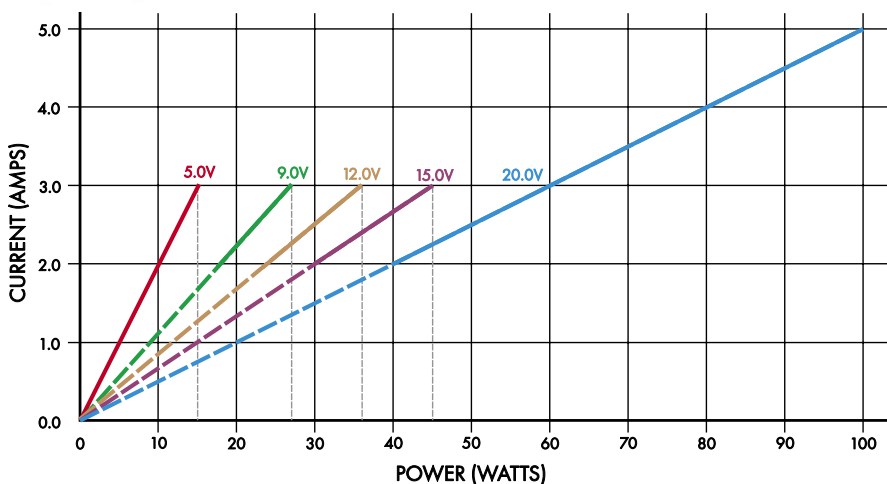
Po podłączeniu, źródło zasilania USB Power Delivery jest skonfigurowane do zasilania napięciem 5 V w celu zapewnienia kompatybilności ze starszymi urządzeniami. Będzie ono w stanie dostarczyć do 1,5 A dla maksymalnej mocy 7 W lub 3 A/15 W, w zależności od wartości R_p , jak opisano powyżej. Jeśli to wszystko, czego wymaga odbiornik, nie musi on robić nic więcej. Jest to znane jako „porozumienie domyślne”.

Jeśli odbiornik wymaga większej mocy lub wyższego napięcia, może zażądać więcej, negocjując zdefiniowaną umowę za pośrednictwem kanału konfiguracji.

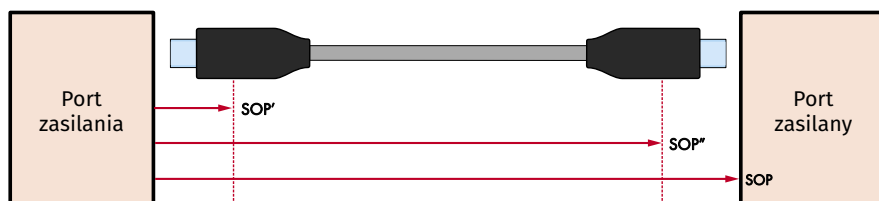
Zanim przejdziemy do szczegółów tego, jak to się robi, powinniśmy przyrzeć się napięciom



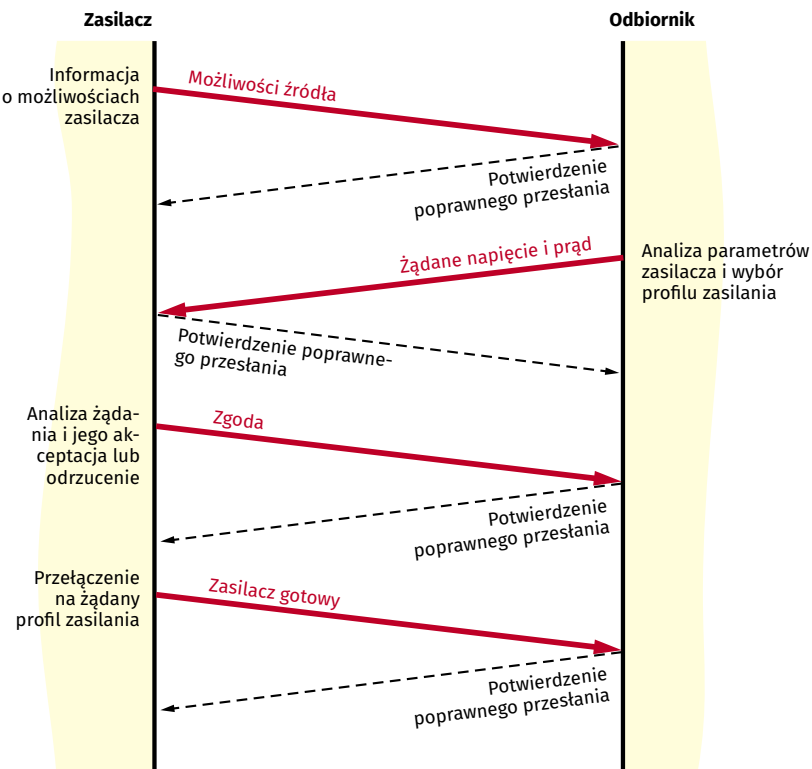
Rysunek 3. Jeśli źródło wykryje aktywny kabel, przetrąca zasilanie na styk VCON, aby zasilić mikrokontroler w złączu. Kabel musi być zdolny do zasilania z obu końców, ponieważ kabel USB typu C może być podłączony w obie strony



Rysunek 4. Moc znamionowa źródła USB Power Delivery zazwyczaj wskazuje zakres napięć, które może zapewnić zasilacz. Nie zawsze tak jest (to mały przytyk w kierunku Apple), więc oplota się to sprawdzić przed zakupem



Rysunek 5. Nagłówek „początku pakietu” (SOP) wskazuje, czy wiadomość jest przeznaczona dla UFP, czy dla konkretnego złącza aktywnego kabla. Są one mylączo nazywane odpowiednio pakietami SOP, SOP' i SOP''



Rysunek 6. Udane negocjacje kontraktu energetycznego rozpoczynają się od zgłoszenia przez źródło swoich możliwości. Następnie odbiornik żąda jednej z tych opcji. Jeśli żądanie zostanie zaakceptowane, źródło wysyła komunikat zgody („Accept”), zmienia swoje wyjście, a następnie wysyła komunikat gotowości („PS_RDY”). Wszystkie komunikaty są potwierdzane komunikatem poprawności przesłania informacji („GoodCRC”), jeśli zostały poprawnie odebrane

i prądom, jakie mogą dostarczać typowe źródła USB typu C.

Rysunek 4 pokazuje, że źródła USB Power Delivery są generalnie oceniane na podstawie mocy, jaką mogą dostarczyć. Do 15 W, zazwyczaj dostarczają tylko 5 V. Te o mocy powyżej 15 W zazwyczaj dostarczają 5 V lub 9 V; te o mocy powyżej 27 W powinny oferować 5 V, 9 V i 12 V, podczas gdy te o mocy powyżej 36 W powinny oferować 5 V, 9 V, 12 V i 15 V. Te o mocy powyżej 45 W dodają do listy napięcie 20 V.

Należy pamiętać, że choć na ogół tak jest, niektórzy dostawcy poszli własną drogą i oferują dziwne kombinacje.

Negocjowanie umowy dotyczącej zasilania

Dane w kanale konfiguracyjnym są kodowane w 32-bitowych pakietach przy użyciu dwufazowego kodowania (BMC) z prędkością 300 kilobitów, z korektą błędów CRC. Komunikacja jest inicjowana przez DFP (zwykle źródło).

Wiadomości zaczynają się od pakietu początkowego („start-of-packet” – SOP), który opisuje, dokąd wiadomość ma dotrzeć. Mylące jest to, że wiadomości są oznaczane jako pakiety SOP, SOP' i SOP''.

Jak pokazano na rysunku 5, wiadomości kierowane w pakietach SOP są przeznaczone dla UFP; te z pakietami SOP' są przeznaczone dla złącza na źródłowym końcu elektronicznie zaimplementowanego kabla (tj. odbierającego VCON); a te z pakietami SOP'' są przeznaczone dla złącza na końcu odbiorczym kabla.

Ważne jest, aby pamiętać, że chociaż transmisja SOP oznacza, że kanał konfiguracji wykorzystuje protokół wielopunktowy, ogólnie USB pozostaje połączeniem punkt-punkt. Transmisja na liniach D+, D- pozostaje niezmienną w stosunku do poprzednich generacji USB, a kanały „superspeed” również działają w trybie punkt-punkt. Nadal można podłączyć tylko jeden port DFP do jednego portu UFP i mieć jedno źródło i jeden odbiornik w danym połączeniu.

Podstawowy proces negocjowania sprzyzowanego kontraktu zasilania pokazano na rysunku 6. Do momentu wynegocjowania takiego kontraktu źródło będzie okresowo wysyłać komunikat „Source Cap”, aby ogłosić swoją możliwość dostarczania napięcia i prądu (nazwa procedury „Source Cap” pochodzi od „source capabilities”, czyli „możliwości źródła”). Jeśli zostanie on pomyślnie odebrany przez odbiornik, odpowie on komunikatem „odebrane-zrozumiane” („GoodCRC”).

Komunikat „Source_Cap” zawiera listę możliwych napięć i maksymalnych prądów, które

może dostarczyć zasilanie. Odbiornik może następnie poprosić źródło o dostarczenie określonego napięcia zidentyfikowanego w jednej z możliwości źródła i określić maksymalny wymagany prąd (do zgłaszanej mocy), wysyłając komunikat „Żądanie” („Request”) do źródła. Jest to również potwierdzane przez źródło komunikatem „GoodCRC”.

Źródło analizuje żądanie i określa, czy może je spełnić. Jeśli może, wysyła komunikat zgody („Accept”) do odbiornika. Zakładając, że zasilacz otrzymał komunikat „Good-CRC” potwierdzający odbiór przez zasilane urządzenie, źródło zmieni swoją moc wyjściową na uzgodniony poziom i wyśle komunikat gotowości („Power Source ready” – „PS_RDY”) do odbiornika. Jeśli wiadomość ta zostanie potwierdzona, negocjacje zostaną uznane za zakończone, a umowa na dostarczanie energii zawarta.

Rysunek 7 przedstawia zrzut ekranu tego procesu przy użyciu analizatora Total Phase USB Power Delivery. W górnej części ekranu znajdują się komunikaty przekazywane między źródłem a odbiornikiem w trakcie negocjowania umowy. Dolna połowa ekranu pokazuje rozwinięcie podświetlonego komunikatu „Source_Cap”.

Tabela 1. Wartości Rp i prąd źródła w zależności od możliwości prądowych

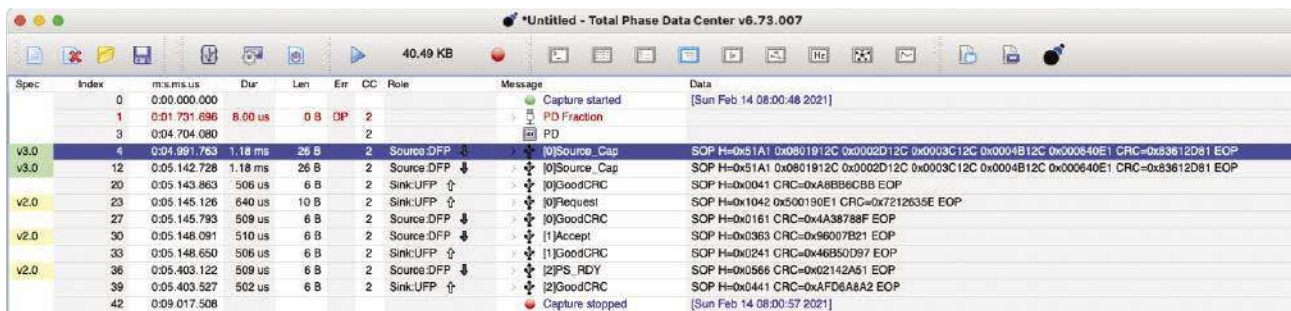
Maksymalny prąd źródła	Rp (ustawiony na 5 V)	Rp (ustawiony na 3,3 V)	Prąd źródła
Domyślny (0,5 A lub 0,9 A)	56 kΩ	36 kΩ	80 μA
1,5A	22 kΩ	12 kΩ	180 μA
3,0A	10 kΩ	4,7 kΩ	330 μA

Tabela 2. Dekodowanie stanów CC1 i CC2 (z perspektywy źródła)

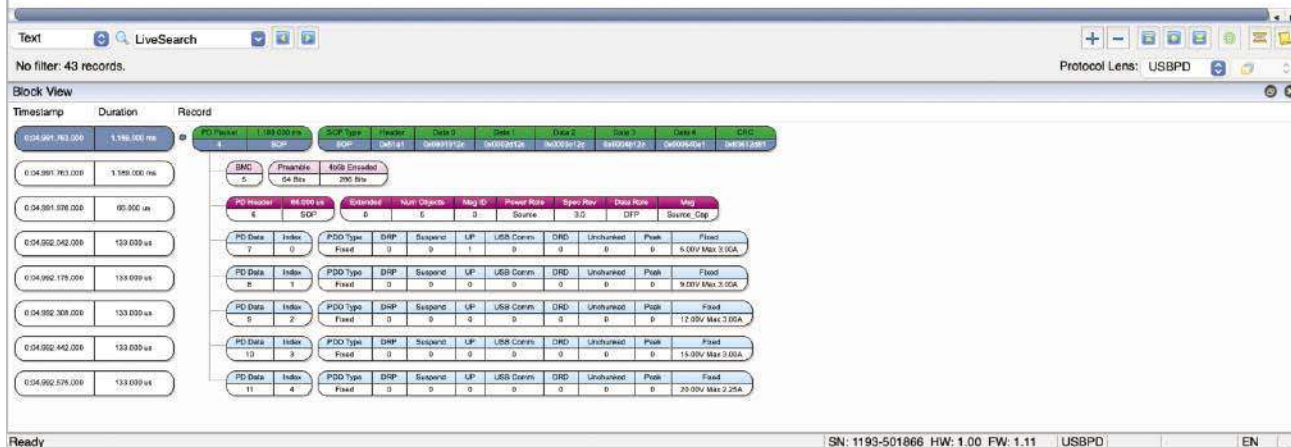
CC1	CC2	Podłączone?	Aktywny kabel?	Odwrócony?
otwarty	otwarty	nie	–	–
Rd	otwarty	tak	nie	nie
otwarty	Rd	tak	nie	tak
otwarty	Ra	nie	tak	nie
Ra	otwarty	nie	tak	tak
Rd	Ra	tak	tak	nie
Ra	Rd	tak	tak	tak
Rd	Rd	tryb debugowania akcesoriów		
Ra	Ra	tryb adaptera audio		

Widać, że to konkretne źródło (zasilacz Targus o mocy 45 W oznaczone jako APA95AU) może dostarczać napięcia 5 V, 9 V, 12 V, 15 V i 20 V, jak opisano w pięciu „sposobach zasilania” zawartych w komunikacie „Source_Cap”.

Rysunek 8 pokazuje dokładnie tę samą negocjację, ale podkreśla komunikat żądania odbiornika w odpowiedzi na komunikat „Source_Cap”. W tym przypadku odbiornik żąda 20 V, prosząc o „pozycję 5”, odpowiadającą piątemu poziomowi mocy w komunikacie



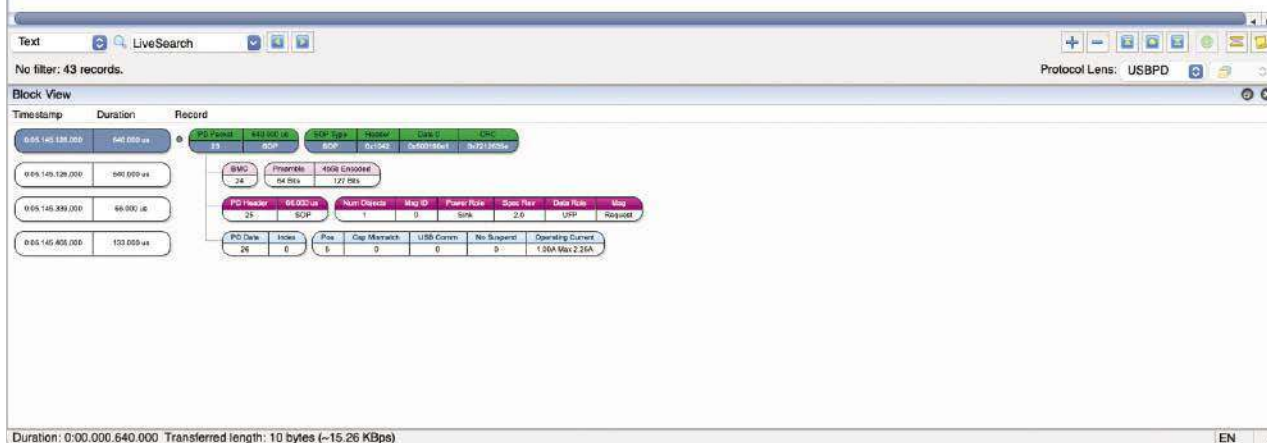
Rysunek 7. Górna połowa tego zrzutu ekranowego pokazuje komunikaty wymieniane między zasilaczem a obciążeniem podczas udanych negocjacji kontraktu 20 V. Dolna połowa pokazuje szczegóły podświetlonego komunikatu charakterystyki zasilacza („Source_Cap”). W tym przypadku istnieje pięć schematów opisu mocy, odpowiadających pięciu poziomom napięcia obsługiwanych przez to źródło



WSZYSTKO O USB

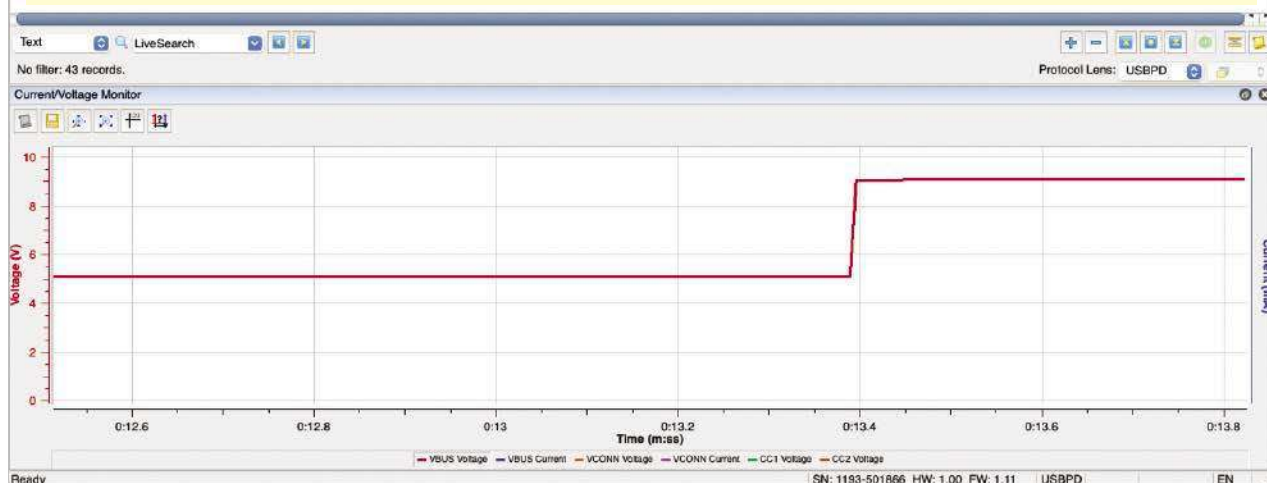
Spec	Index	m.s.ms.us	Dur	Len	Err	CC	Role	Message	Data
	0	0:00.000.000						Capture started	[Sun Feb 14 08:00:48 2021]
	1	0:01.731.696	8.00 us	0 B	DP	2		PD Fraction	
	3	0:04.704.080						PD	
v3.0	4	0:04.991.763	1.18 ms	26 B		2	Source:DFP	[0]Source_Cap	SOP H=0x51A1 0x0801912C 0x0002D12C 0x0003C12C 0x0004B12C 0x000640E1 CRC=0x83612D81 EOP
v3.0	12	0:05.142.728	1.18 ms	26 B		2	Source:DFP	[0]Source_Cap	SOP H=0x51A1 0x0801912C 0x0002D12C 0x0003C12C 0x0004B12C 0x000640E1 CRC=0x83612D81 EOP
	20	0:05.143.863	506 us	6 B		2	Sink:UFP	[0]GoodCRC	SOP H=0x0041 CRC=0xA8B86CBB EOP
v2.0	23	0:05.145.126	640 us	10 B		2	Sink:UFP	[0]Request	SOP H=0x1042 0x500190E1 CRC=0x7212635E EOP
	27	0:05.145.793	509 us	6 B		2	Source:DFP	[0]GoodCRC	SOP H=0x0161 CRC=0x4A38788F EOP
v2.0	30	0:05.148.091	510 us	6 B		2	Source:DFP	[1]Accept	SOP H=0x0363 CRC=0x96007B21 EOP
	33	0:05.148.650	506 us	6 B		2	Sink:UFP	[1]GoodCRC	SOP H=0x0241 CRC=0x46B50D97 EOP
v2.0	36	0:05.403.122	509 us	6 B		2	Source:DFP	[2]PS_RDY	SOP H=0x0566 CRC=0x02142A51 EOP
	39	0:05.403.527	502 us	6 B		2	Sink:UFP	[2]GoodCRC	SOP H=0x0441 CRC=0xAFD6A8A2 EOP
	42	0:09.017.508						Capture stopped	[Sun Feb 14 08:00:57 2021]

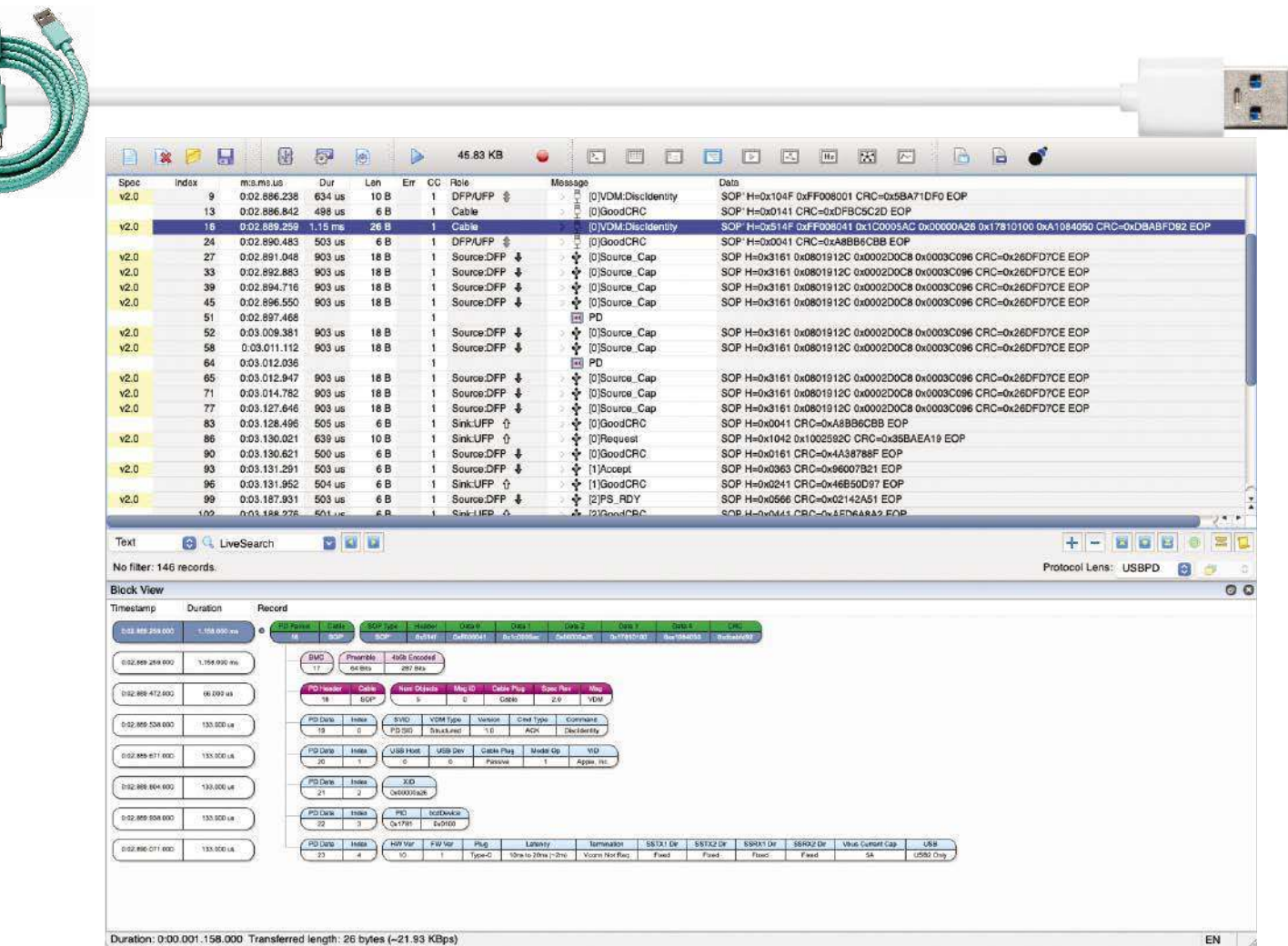
Rysunek 8. Ta sama transakcja, co na Rysunku 7, ale w tym przypadku dolny panel pokazuje szczegóły komunikatu „Żądanie” wystanego przez odbiornik. Żąda on „pozycji 5”, odpowiadającej informacji o możliwości zasilania napięciem 20 V w komunikacie charakterystyki zasilacza („Source_Cap”)



Spec	Index	m.s.ms.us	Dur	Len	Err	CC	Role	Message	Data
	0	0:00.000.000						Capture started	[Sun Feb 14 08:13:48 2021]
v2.0	1	0:13.326.631	501 us	6 B		1	Sink:UFP	[0]Soft_Reset	SOP H=0x004D CRC=0x040E23B7 EOP
	4	0:13.326.991	500 us	6 B		1	Source:DFP	[0]GoodCRC	SOP H=0x0161 CRC=0x4A38788F EOP
v2.0	7	0:13.327.562	500 us	6 B		1	Source:DFP	[0]Accept	SOP H=0x0163 CRC=0x780E1A0D EOP
	10	0:13.328.106	504 us	6 B		1	Sink:UFP	[0]GoodCRC	SOP H=0x0041 CRC=0xA8B86CBB EOP
v2.0	13	0:13.328.784	902 us	18 B		1	Source:DFP	[1]Source_Cap	SOP H=0x03361 0x0801912C 0x0002D0C8 0x0003C096 CRC=0x46B36285 EOP
	19	0:13.329.735	505 us	6 B		1	Sink:UFP	[1]GoodCRC	SOP H=0x0241 CRC=0x46B50D97 EOP
v2.0	22	0:13.330.799	635 us	10 B		1	Sink:UFP	[1]Request	SOP H=0x1242 0x200320C8 CRC=0x72853D24 EOP
	26	0:13.331.285	499 us	6 B		1	Source:DFP	[1]GoodCRC	SOP H=0x0361 CRC=0x4A3619A3 EOP
v2.0	29	0:13.331.959	500 us	6 B		1	Source:DFP	[2]Accept	SOP H=0x0563 CRC=0x7F63DE14 EOP
	32	0:13.332.594	501 us	6 B		1	Sink:UFP	[2]GoodCRC	SOP H=0x0441 CRC=0xAFD6A8A2 EOP
v2.0	35	0:13.388.578	503 us	6 B		1	Source:DFP	[3]PS_RDY	SOP H=0x0766 CRC=0xEC1A4B7D EOP
	38	0:13.389.205						PD	
	39	0:13.389.238	497 us	6 B		1	Sink:UFP	[3]GoodCRC	SOP H=0x0641 CRC=0x41D8C98E EOP
	42	0:17.538.533						Capture stopped	[Sun Feb 14 08:14:05 2021]

Rysunek 9. Ten zrzut ekranowy pokazuje przejście z kontraktu 5 V na kontrakt 9 V. Odbiornik wysyła komunikat miękkiego resetu („Soft Reset”), aby zmusić źródło do ponownego zadeklarowania swoich możliwości, w celu wynegocjowania nowej transakcji. Źródło nadal honoruje istniejący kontrakt, dopóki nie zostanie uzgodniony nowy





Rysunek 10. Pokazane negocjacje transakcji po dołączeniu, jeśli wykryty zostanie „inteligentny” kabel z układami elektronicznymi. Źródło najpierw wysyła zapytanie o kabel za pomocą komunikatu identyfikacji logiki kabla („DiscIdentity”). Dolny panel pokazuje odpowiedź kabla. Widzimy między innymi, że jest to kabel Apple i może obsługiwać prąd o natężeniu do 5 A

„Source_Cap”. Żąda również prądu roboczego 1 A i prądu szczytowego 2,25 A, maksymalnego dostępnego przy tym napięciu.

Renegocjacja

Jeśli odbiornik chce wynegocjować nowy kontrakt ze źródłem, może wysłać komunikat „miękkiego resetu” („Soft Reset”), który spowoduje ponowne wysłanie przez odbiornik komunikatu „Source_Cap”. Odbiornik może następnie zażądać innego kontraktu. Bieżący kontrakt pozostaje w mocy do momentu pomyślnego wynegocjowania nowego.

Rysunek 9 pokazuje jak to działa w praktyce. W tym przypadku odbiornik żąda zmiany napięcia z 5 V na 9 V. Dolna połowa zrzutu ekranu pokazuje zmianę VBUS odpowiadającą komunikatowi „PS_RDY”, tuż przed znakiem 13,4 sekundy.

Wszystkie pokazane do tej pory transakcje dotyczyły kabla pasywnego, więc nie wysłano żadnych pakietów SOP’ ani SOP”.

Rysunek 10 pokazuje komunikację przy połączeniu, gdy używany jest aktywny kabel. Początkowo źródło wysyła pakiet identyfikacji („DiscIdentity”) do kabla, który odpowiada podobnie oznaczoną wiadomością zawierającą

szczegóły dotyczące kabla. Ta odpowiedź jest rozwinięta w dolnej połowie ekranu. W dolnym rzędzie widać, że kabel zgłasza swoje opóźnienie i wskazuje, że ma zdolność przesyłania prądu 5 A.

Uwagi praktyczne

Jeśli Czytelnik chce zbudować projekt wykorzystujący zasilanie USB typu C, ma ku temu kilka opcji. Zdecydowanie najprostszą z nich jest praca z napięciem 5 V; tj. nic nie robić i zaakceptować domyślną umowę 5 V. Jest to identyczne z zasilaniem urządzenia ze złącza USB typu B, takiego jak Mini-B lub Micro-B.

Jeśli Czytelnik chce wynegocjować określoną, zdefiniowaną umowę, może użyć jednego z gotowych chipów, które zapewniają różne stopnie integracji. Zakładając, że Czytelnik buduje urządzenie jedynie pobierające zasilanie z portu USB, dobrą opcją, z której korzystał Autor, jest układ ST Microelectronics STUSB4500.

Może być on używany jako samodzielny sterownik (po zaprogramowaniu, jeśli domyślne ustawienia nie odpowiadają Czytelnikowi) lub może być używany z mikrokontrolerem, aby zapewnić pełen nadzór nad procesem negocjacji.

Autor użył tego układu do zbudowania prostego zasilacza. Wykorzystuje on liniowy regulator zapewniający niski poziom szumów i zarządza napięciem wejściowym odbiornika za pośrednictwem USB Power Delivery, aby zminimalizować wewnętrzne rozpraszanie mocy. ■

Andrew Levido

Adaptacja do wydania polskiego – Andrzej Nowicki

Odnosiniki

- Nota katalogowa „STUSB4500 – Autonomiczny sterownik USB PD dla urządzeń pobierających energię – STMicroelectronics” – siliconchip.com.au/link/ab73
- Nota katalogowa Microchip „AN1953 Introduction to USB Type-C” – siliconchip.com.au/link/ab74
- Nota techniczna Texas Instruments „USB PD Power Negotiations”, 2016, 21 – siliconchip.com.au/link/ab75

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. www.siliconchip.com.au

Od lewej do prawej: Comsol COWCC30WH, XY-PDS100 i Belkin F7U060AU



Korzystanie z tanich modułów elektronicznych. Ładowarki USB

W tym artykule opisano kilka tanich modułów, które są popularne na rynku. Wykorzystują one niesamowity wzrost możliwości interfejsu USB, zwłaszcza w obszarze dostarczania energii (PD – Power Delivery). Asortyment w tym zakresie obejmuje ładowarki PD, kable i adaptory (prześciówki) kabli.

Jak wspominaliśmy w artykule pt. „Historia Uniwersalnego Złącza Szeregowego. Ekspansja USB”, znaczącym obszarem zastosowań USB, który ostatnio niesamowicie się rozwinął, jest dostarczanie zasilania. W ten trend wpisuje się również decyzja UE o przyjęciu USB-C jako standardowego złącza ładowania wszystkich telefonów komórkowych. Kiedy złącze USB pojawiło się po raz pierwszy pod koniec lat 90., mogło dostarczać przy napięciu 5 V prąd o natężeniu zaledwie do 100 mA dla urządzeń „małej mocy” lub do 500 mA dla urządzeń „dużej mocy”, takich jak dysk twardy USB.

Jednak wraz z rozszerzeniem możliwości przesyłania danych poprzez USB 2.0, USB 3.0 i wreszcie USB-C, rozszerzono również możliwości dostarczania zasilania. Szczegółowe omówienie działania procedur dostarczania zasilania po USB znajduje się w kolejnym artykule na ten temat. Zanim przejdziemy do opisu modułów, zrobimy krótkie podsumowanie.

Standard USB 3.0 zachował napięcie zasilania 5 V, ale zwiększył natężenie prądu „dużej

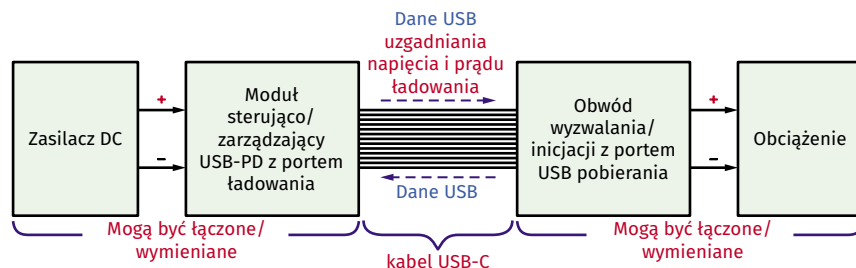
mocy” do 900 mA, umożliwiając podłączonemu urządzeniu pobieranie do 4,5 W mocy (zamiast tylko 2,5 W).

Kiedy specyfikacja USB-PD (Power Delivery) została sfinalizowana w 2012 roku, urządzenie mogło przy napięciu 5 V pobierać prąd o natężeniu do 1,5 A lub 7,5 W mocy za pośrednictwem standardowego kabla USB z wtykami typu A i typu B.

Mniejsze 24-stykowe złącza USB-C pojawiły się w 2014 roku, a kiedy

specyfikacja USB-PD została poddana kolejnym zmianom w latach 2014, 2016 i 2017, zwiększono również napięcie i natężenie dostarczanego prądu.

Teraz urządzenia mogą żądać zasilania napięciem 5 V, 9 V, 12 V, 15 V lub 20 V i mogą pobierać prąd o natężeniu do 5 A – co odpowiada 100 W przy zasilaniu napięciem 20 V. A od wersji USB-PD 3.0 z 2017 roku, urządzenia mogą również korzystać z protokołu programowalnego zasilania (PPS), który



Rysunek 1. System USB-PD składa się z pięciu elementów: głównego źródła zasilania DC, „menedżera” USB-PD z portem „wysyłania” (DFP – Downstream Facing Port), kabla USB-C, obwodu obciążenia wyposażonego w port „pobierania” (UFP – Upstream Facing Port) i wreszcie „odbiornika” zasilania. Element zarządzający USB-PD może być połączony z głównym źródłem prądu stałego, a obwód odbiorczy może być również połączony z odbiornikiem

umożliwia zmianę napięcia zasilania w krokach co 20 mV.

To znacznie rozszerza zakres możliwych zastosowań USB-PD i dlatego widzimy tak wiele tanich modułów zaprojektowanych w celu wykorzystania tych zwiększonych możliwości.

Jak działa USB-PD

Jak wspomniano wcześniej, działanie zasilania przez USB zostało szczegółowo opisane w skojarzonym artykule. Jest jednak kilka punktów, które możemy tutaj dodać, a także podsumujemy podstawy negocjacji w protokole USB-PD.

Zasadniczo, USB-PD jest możliwe dzięki niektórym dodatkowym stykom w złączu USB-C. W szczególności chodzi o styki CC1 (A5) i CC2 (B5), które są oznaczone jako piny kanału konfiguracji (Configuration Channel – CC). Układ „uzgadniania i dopasowania” został przedstawiony na **rysunku 1**.

Początkowo zasilacz obsługujący USB-PD ustawia napięcie wyjściowe VBUS na wartość 5 V. Ustawia również każdy ze styków CC swojego wyjścia (wysyłanie) złącza USB-C w stan wysokiego poziomu logicznego za pomocą rezystora polaryzującego R_p , przy czym wartość R_p dobiera się zgodnie z wydajnością prądową zasilacza.

Urządzenia zaprojektowane do zasilania ze złącza USB-C są wyposażone w rezystor polaryzujący R_d podłączony między jednym ze styków CC a masą. Wartość R_d jest wybierana w celu wskazania poziomu prądu wymaganego przez urządzenie.

W rezultacie, gdy kabel z urządzenia jest podłączony do złącza USB-C, spadek napięcia na jednej z linii CC wskazuje zasilaczowi, że:

- podłączono urządzenie obciążenia prądowego lub „przeciążenie”,
- wtyczki USB-C w złączu są odpowiednio zorientowane,

Wykorzystanie USB-PD do szybkiego ładowania

Jeszcze zanim w 2012 roku specyfikacja USB-PD została wydana, różne firmy związane z rozwijającym się rynkiem telefonów komórkowych opracowały sposoby wykorzystania gniazd USB do szybkiego ładowania akumulatorów telefonów komórkowych. Przykładami są Qualcomm, który opracował protokół Quick Charge (QC), Motorola z protokołem TurboPower i Huawei z protokołem SuperCharge (SC).

Być może ze względu na powszechne zastosowanie tych protokołów, różne wersje USB-PD stopniowo je uwzględniały. W rezultacie, gdy w 2017 roku wydano wersję 3.0 USB-PD, w tym PPS (programowalny zasilacz), zasadniczo zawierała ona prawie wszystkie wcześniejsze protokoły szybkiego ładowania.

Dlatego też specyfikacje większości modułów inicjujących USB-PD i szybkich ładowarek zapewniają zgodność z listą protokołów, takich jak PD 2.0, PD 3.0, Qualcomm QC3.0 i QC4+, Huawei SCP/FCP, Apple 2.4A, Samsung AFC, MediaTek PE2.0 i PE3.0, Oppo's VOOC i tak dalej.

- jest dostępny prąd z zasilacza o pożądanym natężeniu.

Następnie odbywa się wymiana pakietów danych między zasilaniem a obciążeniem za pośrednictwem linii CC, przy użyciu sprężonego protokołu DC BMC (Biphase Mark Code) lub różnicowego kodowania Manchester. Umożliwia to urządzeniu obciążającemu wskazanie żądanego napięcia zasilania, a następnie zasilaczowi zmianę wyjścia na żądany poziom, jeśli jest to możliwe.

Jak wspomniano powyżej, jeśli zasilacz obsługuje protokół PPS, napięcie można regulować w krokach co 20 mV.

Te negocjacje mogą mieć miejsce tylko wtedy, gdy urządzenie obciążające jest podłączone do zasilacza za pomocą złącza USB-C i odpowiedniego kabla. Nie będzie działać, jeśli używane jest złącze USB typu A, ponieważ nie ma ono żadnych styków CC ani linii wymiany danych.

Początkowa specyfikacja USB-PD Rev.1 z 2012 roku zezwalała urządzeniu podłączonemu poprzez złącza USB 2.0/3.0 typu A i typu B do jednostki nadrzędnej/zasilacza na negocjowanie wyższego napięcia niż 5 V (np. 12 V lub 20 V) za pomocą binarnego sygnału FSK na linii VBUS. Podejście to zostało jednak wycofane wraz z wydaniem USB-PD Rev.2.0 w 2014 roku

Tak więc większość zasilaczy USB-PD może przez port lub porty USB typu A dostarczać tylko 5 V (lub ewentualnie 12 V).

Należy pamiętać, że protokół negocjacji USB-PD umożliwia przesyłanie energii w dowolnym kierunku – od urządzenia głównego do urządzenia-gościa lub odwrotnie. Na przykład, laptop lub tablet PC może szybko naładować swoją baterię z zasilacza/ładowarki USB-PD, żądając, aby ładowanie odbywało się przy napięciu 9 V, 15 V lub 20 V zamiast 5 V.

Szybka ładowarka XY-PDS100

Ten pierwszy moduł to „szybka ładowarka”, którą można skonfigurować przy użyciu standardowego protokołu USB-PD tak, aby zapewniała odpowiedni zakres napięć i prądów wyjściowych.

XY-PDS100 ma wymiary 53×46×21 mm. Jest dostępna u kilku dostawców internetowych, w tym Banggood, który w chwili pisania tego tekstu ma go w cenie 15,99 USD plus wysyłka (czasami darmowa w promocji).

Jak pokazano na zdjęciach, wyjście XY-PDS100 to gniazdo USB typu A i gniazdo USB-C, a także 3-cyfrowy, 7-segmentowy wyświetlacz LED (z cyframi o wysokości 6,5 mm) i trzy diody LED wskaźnika. Jedna świeci się, gdy wyświetlane jest napięcie wyjściowe, druga, gdy moduł pokazuje prąd pobierany z gniazda USB-C, a trzecia, gdy pokazuje prąd pobierany z gniazda typu A.

Na wejściu znajdują się dwa gniazda. Jednym z nich jest małe koncentryczne gniazdo DC akceptujące zasilanie



Po lewej stronie pokazano moduł XY-PDS100, podłączony do uniwersalnego obciążenia laboratoryjnego typu XY-WPDT. Moduł obciążenia pomaga optymalizować zadany profil ładowania dla urządzenia na wejściu, które wysyła napięcie o stałym poziomie. W prawym dolnym rogu widać tylną ściankę modułu XY-PDS100; oba te zdjęcia przedstawiają moduły w przybliżeniu w naturalnej wielkości



12...28 V DC z zasilacza sieciowego, a drugim gniazdo USB-C oznaczone jako wejście zasilania („Input-PD”). To drugie wejście, na spodzie obudowy, ma napis „PD Recommended 87 W”, ale wydaje się, że jest to po prostu alternatywne wejście zasilania DC.

Zasadniczo, XY-PDS100 przekształca zwykły zasilacz z wyjściem 12...28 V DC w „inteligentną” ładowarkę USB-PD lub źródło zasilania, które może reagować na przebieg negocjacji z jednostką podrzędną, aby zapewnić jeden ze standardowych profili napięcia i prądu ładowania.

Jest to więc programowalny konwerter DC-DC obniżający napięcie, który może zapewnić do 100 W mocy przy napięciu od 5 V do 20 V z wyjścia USB-C lub do 36 W mocy przy napięciu 5 V (opcjonalnie do 12 V) z wyjścia USB typu A. Zawiera nawet trzycyfrowy odczyt LED wyświetlający aktualne napięcie i prąd wyjściowy.

Całkiem nieźle jak na bardzo kompaktowe urządzenie, które kosztuje mniej niż 20 USD.

Ponieważ XY-PDS100 jest konwerterem obniżającym napięcie, musi mieć napięcie wejściowe DC co najmniej o 2 V wyższe niż najwyższe napięcie wyjściowe, które może być wymagane. Tak więc, jeśli potrzebujesz maksymalnie tylko 12 V do ładowania przez wyjście typu A, napięcie wejściowe 14...15 V będzie w porządku. Ale dla pełnego zakresu napięć wymaganych do szybkiego ładowania USB-PD, napięcie wejściowe będzie musiało wynosić co najmniej 22...23 V.

Autor był całkiem zadowolony z wydajności XY-PDS100. Wydaje się on być w miarę kompatybilny z protokołami PD 3.0, a także z protokołem „regulacji precyzyjnego ustawiania napięcia wyjściowego” PPS.

Podczas gdy XY-PDS100 jest modulem adaptera („USB-PD Manager”), wymagającym zewnętrznego zasilania DC, pozostałe urządzenia, którym przyjrzała się Redakcja Silicon Chip-a, łączą obie funkcje, tworząc kompletne źródło zasilania USB-PD.

Autor miał jednak pewne trudności z ich zdobyciem. Zamówił kilka interesujących go modeli od chińskiego dostawcy, ale nie dotarły one do Autora (!!!) i ostatecznie odkrył on, że oferta była wirtualna i nie było ich w magazynie.

Zamiast tego Autor musiał kupić je od lokalnych dostawców, którzy dostarczyli je w ciągu kilku dni roboczych, ale kosztowały znacznie więcej niż sztuki zamówione w Chinach. Pierwszy z nich to...

Zasilacz Belkin F7U060AU 27W

To urządzenie z firmy JB Hi-Fi (www.jbhifi.com.au) kosztowało 39,95 USD.

Zachowaj ostrożność przy zakupie kabli i adapterów USB-C

Chociaż w Internecie można znaleźć wiele tanich kabli USB-C od różnych sprzedawców, należy zachować ostrożność przy zakupie wielu z nich. Na przykład, wiele z tanich kabli nadaje się tylko do zasilania i ładowania akumulatorów, a nie do przesyłania danych, a zwłaszcza do przesyłania danych z dużą prędkością.

Poza liniami związanymi z transferem zasilania (w tym liniami kanału konfiguracyjnego), mogą one nie mieć żadnych linii transferu danych, z wyjątkiem być może tych dla USB 2.0 (D+ i D-).

Dotyczy to w szczególności kabli z wtyczką typu A na jednym końcu i wtyczką USB-C na drugim. W rzeczywistości obecność wtyczki typu A jest oczywistym wskazaniem, że kabel nie nadaje się do szybkiego przesyłania danych, a całkiem możliwe, że służy tylko do przesyłania energii i ładowania. Jednocześnie przesyłanie/ładowanie będzie możliwe tylko przy napięciu 5 V, ponieważ wynegocjowanie wyższego napięcia zasilania prawie na pewno nie będzie możliwe.

Dotyczy to również wielu nominalnych przejściówek USB-C. Jeśli mają one wtyczkę lub gniazdo USB typu A na jednym końcu, oznacza to, że prawdopodobnie nadają się tylko do przesyłania energii i ładowania, chociaż mogą przysłać dane USB z niską i pełną prędkością za pośrednictwem linii D+ i D-, zakładając, że te przewody są nawet zamontowane.

Nawet jeśli tani kabel ma złącza USB-C na obu końcach, nie gwarantuje to, że nadaje się do naprawdę szybkiego przesyłania danych. To sprawia, że kupowanie takich kabli przez Internet jest nieco ryzykowne, ponieważ nie można ich przetestować przed zakupem.

W rzeczywistości, jeśli zobaczysz jeden z tych kabli za mniej niż 15 USD, możesz prawdopodobnie założyć, że nadaje się on tylko do przesyłania energii i ładowania baterii. Kable USB-C, które mogą być używane do naprawdę szybkiego przesyłania danych, prawdopodobnie będą kosztować znacznie więcej.

Zupełnie oddzielną sprawą jest jakość wykonania wtyczek USB-C. Ich niewielkie rozmiary oraz ilość zaimplementowanych styków, aż 24, powodują, że wtyczki niezbyt solidnie wykonane b. szybko ulegają zużyciu, a nawet mogą od początku nie działać. Dlatego kable zasilające USB stanowią stały element elektrośmieci.

Mierzy zaledwie 51×60×31 mm i waży 50 g. Urządzenie jest pokazane na zdjęciu po prawej stronie na początku tego artykułu; ma dwukołkową wtyczkę sieciową na jednym końcu i gniazdo USB-C na drugim końcu. To wszystko – jest to po prostu wydłużona wersja znanego zasilacza wtyczkowego USB. Napis na końcu wtyczki informuje, że została ona zaprojektowana w Kalifornii i zmontowana w Chinach.

Kiedy Autor wypróbował zasilacz z kilkoma różnymi jednostkami obciążającymi, odkrył, że chociaż zasilacz zgłasza się jako urządzenie PD 3.0, zapewnia tylko wybór jednego z trzech napięć wyjściowych: 5 V, 9 V lub 12 V. Dwa niższe ustawienia napięcia mogą dostarczyć prąd o natężeniu do 3 A, podczas gdy ustawienie 12 V może dostarczyć prąd o natężeniu do 2,25 A.

Tak więc moc znamionowa 27 W ma miejsce tylko wtedy, gdy urządzenie dostarcza 9 V lub 12 V; gdy dostarcza 5 V, jest to naprawdę źródło o mocy 15 W. Oczywiście byłoby to w porządku, gdybyś potrzebował napięcia maksymalnie 12 V i mocy zasilacza 15...27 W.

Ładowarka ścienna Comsol COWCC30WH 30W

To urządzenie w sklepie Officeworks (www.officeworks.com.au/shop/) kosztuje 39,88 USD. Mierzy 44×64×40 mm i waży 80 g. Jak widać na zdjęciu po lewej stronie na początku tego artykułu, jest ono bardzo podobne do urządzenia Belkin, z dwukołową wtyczką sieciową na jednym końcu i gniazdem USB-C na drugim końcu. Napis na końcu wtyczki mówi po prostu „Made in China”.

Kiedy Autor sprawdził to urządzenie z kilkoma różnymi jednostkami obciążenia,

zarejestrowało się ono tylko jako urządzenie PD 2.0, ale mogło zapewnić dowolne z pięciu pełnych napięć wyjściowych: 5 V, 9 V, 12 V, 15 V lub 20 V. Podobnie jak w przypadku jednostki Belkin, może ono dostarczyć do 3 A przy napięciu 5 V lub 9 V, ale przy 12 V może dostarczyć tylko do 2,5 A prądu. Następnie przy napięciu 15 V może zapewnić do 2 A, a przy 20 V do 1,5 A.

Jest to więc tylko źródło zasilania o mocy 30 W dla trzech z pięciu wybieranych napięć.

Biorąc pod uwagę, że jego cena jest praktycznie taka sama jak jednostki Belkin, fakt, że zapewnia wybór pełnych pięciu napięć PD i prawie stałą moc 30 W, sprawia, że mamy do czynienia z lepszym stosunkiem jakości do ceny.

Zakres napięć i prądów dostępnych z tego typu ładowarki oznacza, że może ona zasilć szeroką gamę urządzeń, w tym te, które można zbudować samodzielnie.

Jeśli każde z tych urządzeń zawiera obwody do negocjowania wymaganego prądu i napięcia, oznacza to, że można mieć jeden lub zaledwie kilka zasilaczy do zasilania szerokiej gamy urządzeń.

Zasadniczo, ładowarki te mogą być nowymi „wielonapięciowymi” zasilaczami wtyczkowymi, z których wszyscy będziemy korzystać w przyszłości.

Ładowarka ścienna ALOGIC WCG1X65-ANZ 65W

Trzecią ładowarką ścienną USB-PD, którą Autor kupił, była ładowarka ALOGIC WCG1X65, pokazana na zdjęciu na końcu artykułu. Ma ona bardzo podobny rozmiar jak urządzenia Belkin i Comsol. Jest nawet nieco mniejsza, mierzy 55×60×35 mm i waży blisko 95 g.



To urządzenie również pochodzi ze sklepu JB Hi-Fi i kosztuje 74 USD plus dostawa, ale jest również dostępne od TechBuy-a (www.techbuy.com.au), innego lokalnego dostawcy, za 72,70 USD plus dostawa.

Chociaż urządzenie jest prawie dwukrotnie droższe od innych ładowarek ściennych, może pochwalić się ponad dwukrotnie większą mocą, aż do 65 W. W zestawie znajduje się kabel ładujący USB-C o długości 2 m oraz niewielka (90×110 mm) czterostronowa skrócona instrukcja obsługi. Posiada również biały wskaźnik zasilania LED, tuż pod gniazdem wyjściowym USB-C.

Kiedy Autor sprawdził to urządzenie z tymi samymi modułami obciążenia, co poprzednio, zarejestrowało się ono jako urządzenie PD 3.0 i można je było łatwo zaprogramować tak, aby podawało dowolne z pięciu standardowych napięć PD: 5 V, 9 V, 12 V, 15 V lub 20 V. Może dostarczyć do 3 A przy każdym z czterech niższych napięć lub do 3,25 A przy 20 V, co jest dość imponujące, biorąc pod uwagę jego kompaktowy rozmiar i wagę.

Twórcy twierdzą, że jest to wynikiem zastosowania „najnowszej technologii ładowania GaN”. Prawdopodobnie wykorzystują oni zdolność tranzystorów i diod wykonanych z azotku galu (GaN) do pracy przy znacznie wyższych napięciach i z wyższą sprawnością.

Do 65 W mocy przy dowolnym z pięciu poziomów napięcia PD 3.0, ALOGIC WCG1X65-ANZ byłby najlepszym wyborem pomimo znacznie wyższej ceny.

Należy zauważyć, że jednym z urządzeń, które Autor próbował zakupić i nie udało się go pozyskać z Chin, była ładowarka ścienna Bakeey HC-652CA 65 W, która prawdopodobnie również byłaby dobrym wyborem, jeśli kiedyś stanie się dostępna. ■

Jim Rowe

*Adaptacja do wydania
polskiego – Andrzej Nowicki*

Przydatne linki:

- USB-C: <https://w.wiki/nto>
- USB-PD: <https://w.wiki/34dT>
- siliconchip.com.au/link/ab7l
- Szybkie ładowanie: <https://w.wiki/34dU>
- Azotek galu: <https://w.wiki/34dV>

Złącza USB-C

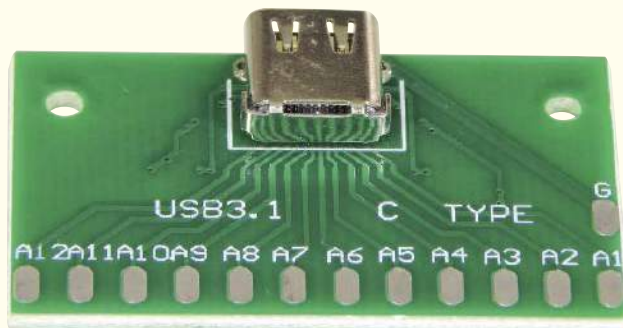
Ze względu na możliwe problemy związane z kablami USB-C, możesz być zainteresowany niedrogim modułem „prześciówki” lub płytką testową pokazaną na poniższym zdjęciu. Jest on dostępny u dostawców internetowych, takich jak Banggood, za jedyne 2,10 USD za pojedynczy moduł, 4,80 USD za paczkę pięciu modułów lub 9,00 USD za paczkę dziesięciu modułów (plus koszty wysyłki w wysokości 3,30 USD w każdym przypadku).

Płytkę drukowaną tego modułu ma wymiary zaledwie 25×40 mm i ma gniazdo USB-C zamontowane pośrodku jednego z 40-milimetrowych boków. Wszystkie 24 połączenia gniazda są wyprowadzone do dwóch rzędów po 12 pól lutowniczych na przeciwległej krawędzi płytki PCB, z jednym rzędem (A1-12) na górze i drugim (B1-12) pod spodem. Metalowa obudowa gniazda jest również wyprowadzona do kolejnego pola „G” po każdej stronie płytki drukowanej.

Para tych płytek „prześciówkowych” ułatwia przetestowanie wszystkich linii i połączeń w kablu USB-C, przez sprawdzenie obecności konkretnych linii w kablu USB-C. Autor kupił paczkę pięciu sztuk, ale nie był zachwycony jakością lutowania 24 bardzo blisko rozmieszczonych styków gniazd; jedna z „prześciówek” wydawała się mieć „zimny lut” lub nawet dwa.

Ponieważ nie byłoby łatwo naprawić te połączenia ręcznie ze względu na bardzo małe odstępstwa (około 0,5 mm), Autor zdecydował się po prostu wyrzucić płytkę do elektrośmieci. Tak więc, uwaga!

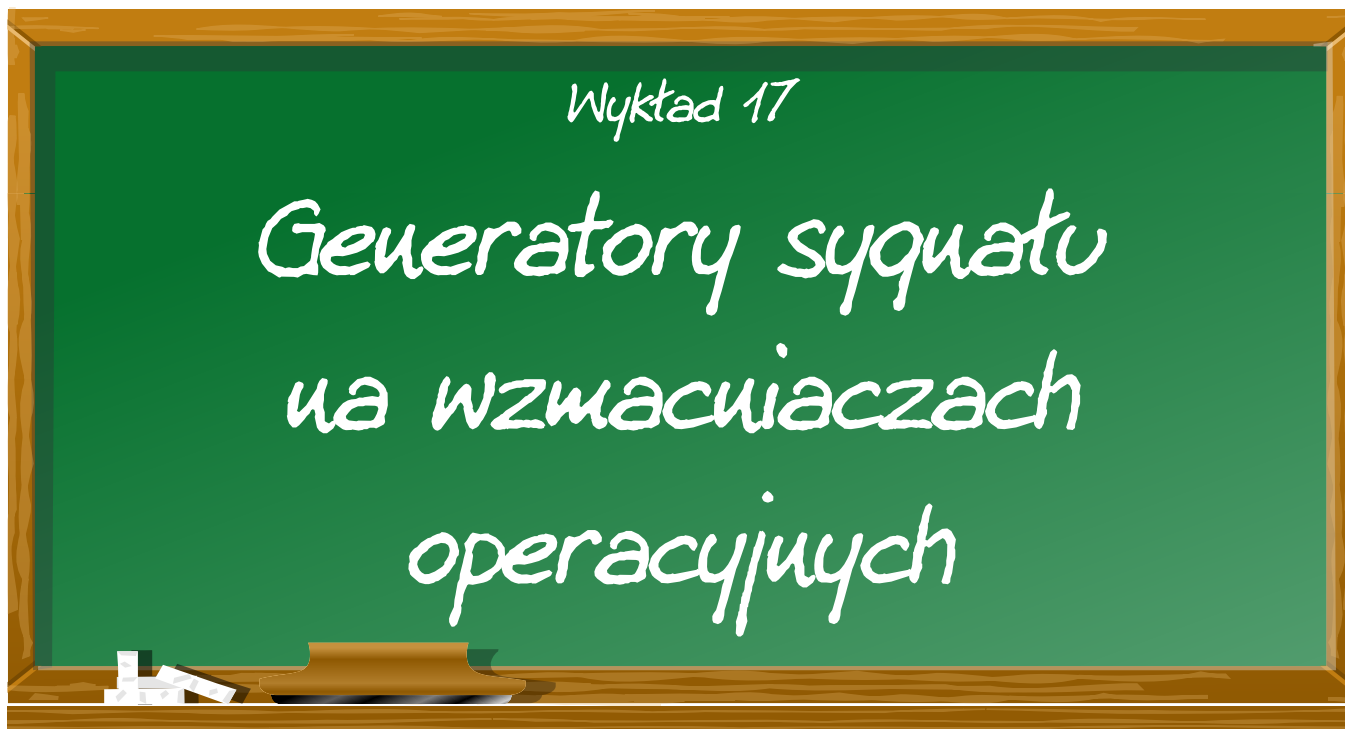
W oddzielnym artykule przyjrzmy się niektórym niedrogim modułom „inicjacji” USB-PD, które mogą być używane do ustawiania wyjściowego napięcia i prądu zasilacza USB, takich jak te opisane tutaj.



Ładowarka ścienna ALOGIC WCG1X65-ANZ 65 W, pokazana w powiększeniu dla lepszej widoczności. Rejestruje się jako urządzenie zgodne z PD 3.0, a zatem może dostarczać standardowe napięcia 5 V, 9 V, 12 V, 15 V i 20 V przy 3 A (lub 3,25 A dla 20 V). Wraz ze wzrostem mocy wyjściowej, ładowarki te mogą stać się dość kosztowne

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. www.siliconchip.com.au

Patronat EdW nad szkołami i uczelnianymi Kołami Naukowymi rozkwita i daje redakcji EdW impulsy zachęcające do wspierania edukacji szkolnej i uczelnianej. Działa sprzężenie zwrotne. Dostajemy mnóstwo wiadomości od uczniów, nauczycieli i studentów. Dla nich jest ta rubryka.



Za pomocą wzmacniaczy operacyjnych można łatwo zaprojektować układy generujące pięć podstawowych przebiegów: prostokątny, sinusoidalny, trójkątny, piłokształtny i schodkowy.

Pojęcia wprowadzające

Generatory sygnału

Generatory sygnałów to obwody, za pomocą których można wygenerować jeden lub więcej elektronicznych przebiegów napięcia przemiennego. Jednakże, aby zakwalifikować się do oznaczenia „generator sygnału”, obwód musi spełniać szereg wymagań.

- Sygnał musi być generowany całkowicie niezależnie. Oznacza to, że obwód musi uruchamiać się natychmiast po włączeniu napięcia zasilania, bez potrzeby zewnętrznego sygnału sterującego.
- Generowany sygnał musi mieć charakter okresowy, tj. powtarzać się regularnie w czasie.
- Wygenerowany sygnał musi zachować jednolity charakter, każdy okres sygnału musi wyglądać identycznie jak poprzedni i następny.
- Częstotliwość i amplituda sygnału muszą być stałe i mogą zmieniać wartość tylko poprzez świadome działanie.

Kształty przebiegów

W elektronice istnieje oczywiście niezliczona ilość różnych sygnałów, które spełniają powyższe wymagania. Istnieje jednak tylko kilka, które można określić jako sygnały standardowe. Są to:

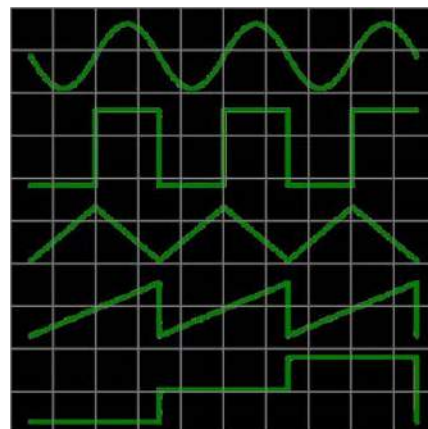
- przebiegi sinusoidalne,
- przebiegi prostokątne,
- przebiegi trójkątne,
- przebiegi piłokształtne,
- przebiegi schodkowe.

Dzięki tym pięciu podstawowym kształtom sygnału można rozwiązać niemal każdy analogowy problem elektroniczny. W kolejnych rozdziałach omówimy sposób generowania tych pięciu podstawowych przebiegów za pomocą wzmacniaczy operacyjnych.

Generator przebiegu prostokątnego

Wprowadzenie

Prostokąt to przebieg, który, jak sama nazwa wskazuje, jest zdefiniowany przez prostokątny kształt. Sygnał ma zatem tylko dwie wartości, które w większości



Omawiane pięć form sygnału
(© 2023 Jos Verstraten)

przypadków nazywane są „L” i „H” (odpowiednio stany niski i wysoki – przyp. tłum.). Sygnał będzie przeskakiwał z „L” do „H” tak szybko, jak to możliwe z określoną częstotliwością, a następnie ponownie z „H” do „L”.

Przebieg prostokątny jest jednym z najczęściej używanych sygnałów w elektronice. Cała elektronika cyfrowa, od prostego przerzutnika do najbardziej skomplikowanego komputera, działa wyłącznie dzięki istnieniu przebiegu prostokątnego. Prostokąt jest również często potrzebny w elektronice analogowej, nawet jeśli tylko do sterowania innym obwodem, takim jak generator schodkowy lub generator piłokształtny.

Podstawowy obwód ze wzmacniaczem operacyjnym

Podstawowy obwód prostokątnego generatora napięcia ze wzmacniaczem operacyjnym pokazano na poniższym rysunku.

Należy pamiętać, że schemat ten zakłada pojedyncze dodatnie zasilanie $+U_b$!

Wejście nieodwracające wzmacniacza operacyjnego jest ustawione na napięcie U_1 za pomocą dwóch rezystorów R_1 i R_2 . Wejście odwracające jest podłączone do kondensatora C_1 , który jest podłączony do masy, a także utrzymuje kontakt z tym, co dzieje się na wyjściu za pomocą rezystora sprzężenia zwrotnego R_3 . Druga pętla sprzężenia zwrotnego jest zawarta w obwodzie, a mianowicie dioda D_1 , która w pewnych warunkach podaje wejście nieodwracające z powrotem na wyjście.

Działanie układu

Działanie obwodu omówiono na podstawie wykresów po prawej stronie rysunku. Załóżmy, że obwód jest podłączony do napięcia zasilania. Wejście nieodwracające wzmacniacza operacyjnego jest dostosowane do napięcia U_1 dzięki dzielnikowi napięcia R_1/R_2 . Kondensator C_1 jest całkowicie rozładowany, napięcie na wejściu odwracającym wynosi zero. Napięcie na wejściu nieodwracającym jest bardziej dodatnie niż napięcie wejściu odwracającym, wyjście wzmacniacza operacyjnego przechodzi do napięcia zasilania. Kondensator C_1 ładuje się przez rezystor R_3 od tego wysokiego napięcia. Napięcie na wejściu odwracającym będzie zatem powoli rosnąć. Szybkość, z jaką wzrasta to napięcie, zależy od wartości rezystora R_3 i kondensatora C_1 . W czasie t_1 kondensator jest naładowany do napięcia U_1 . Nieco później napięcie na wejściu odwracającym staje się wyższe niż napięcie na wejściu nieodwracającym.

W rezultacie wzmacniacz operacyjny, który działa jako komparator, przełącza się, a na wyjściu pojawia się niskie napięcie. Dioda D_1 , która do tej pory była zatkana, będzie teraz przewodzić. W końcu katoda staje się bardziej ujemna niż anoda. Wejście nieodwracające wzmacniacza operacyjnego jest podłączone przez diodę przewodzącą do niskiego napięcia na wyjściu wzmacniacza. W rezultacie napięcie na wejściu nieodwracającym spada do napięcia U_2 . Jest ono równe napięciu wyjściowemu wzmacniacza operacyjnego plus napięcie przewodzenia diody D_1 .

Naładowany kondensator C_1 zostanie teraz rozładowany przez rezystor R_3 do niskiego napięcia wyjściowego wzmacniacza operacyjnego.

W czasie t_2 napięcie na kondensatorze staje się równe napięciu U_2 . Wejście odwracające staje się bardziej ujemne niż wejście nieodwracające, w wyniku czego wyjście wzmacniacza operacyjnego ponownie staje się równe wartości napięcia zasilania. Obwód powraca do stanu początkowego.

Podsumowując

Obwód generuje prostokątny przebieg na wyjściu, jego częstotliwość jest określona przez wartość kondensatora C_1 i rezystora R_3 .

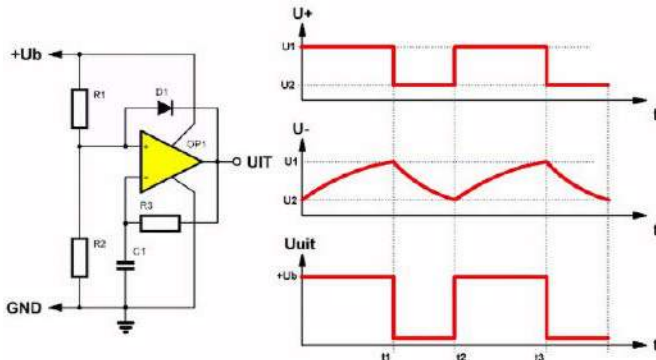
Ustawianie cyklu pracy

Omawiany obwód generuje przebieg prostokątny, ale stosunek między czasem wysokim i niskim nie jest regulowany. Ten „cykl pracy” jest określany tylko przez stałą czasową ładowania i rozładowywania kondensatora C_1 . Chociaż te stałe czasowe są takie same, nie oznacza to, że ładowanie i rozładowywanie trwa tyle samo czasu. Jest to konsekwencją różnych charakterystyk ładowania i rozładowywania kondensatora. W rezultacie na wyjściu generowany jest niesymetryczny przebieg prostokątny, napięcie, którego czas „L” nie jest równy czasowi „H”.

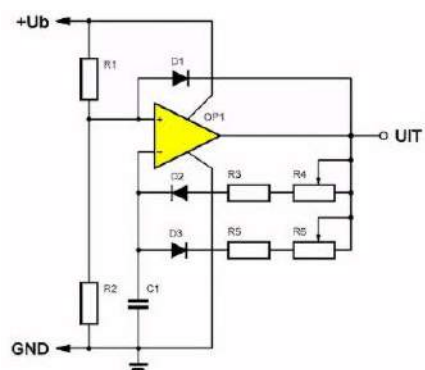
Jednak w wielu zastosowaniach konieczne jest generowanie sygnału prostokątnego o określonym stosunku włączenia do wyłączenia. Podstawowy obwód można łatwo rozszerzyć do wariantu, w którym można regulować ten stosunek czasów włączania do wyłączania.

Sposób, w jaki jest to możliwe, przedstawiono na poniższym rysunku. Kondensator C_1 jest teraz ładowany i rozładowywany przez dwa oddzielne rezystory R_3+R_4 i R_5+R_6 . Diody D_2 i D_3 zapewniają, że rezystory R_3+R_4 są używane do ładowania, a rezystory R_5+R_6 do rozładowywania. Jeśli napięcie wyjściowe wzmacniacza operacyjnego jest wysokie, przewodzić będzie tylko dioda D_2 , co spowoduje zamknięcie obwodu między C_1 i R_3+R_4 . Dioda D_3 jest zatkana, a rezystory R_5+R_6 nie są aktywne. Gdy napięcie wyjściowe staje się niskie, dioda D_2 wyłącza się, a dioda D_3 przewodzi. Następnie kondensator C_1 jest rozładowywany przez rezystor R_5+R_6 .

Projektując oba rezystory częściowo w formie potencjometrów, można indywidualnie ustawić czasy ładowania i rozładowania kondensatora C_1 . W ten sposób możliwe jest wygenerowanie impulsu wyjściowego z regulowanym cyklem pracy.



Podstawowy obwód prostokątnego generatora napięcia
(© 2023 Jos Verstraten)



Rozbudowa do obwodu z regulowanym cyklem pracy
(© 2023 Jos Verstraten)

Zalety i wady układu

Omawiany układ ma tę zaletę, że jest bardzo prosty i działa bezproblemowo. Ma jednak również kilka wad.

Wartość częstotliwości zależy w dużej mierze od wielkości napięcia zasilania. Im jest ono wyższe, tym szybciej ładuje się kondensator. Jeśli chcesz zaprojektować generator o bardzo stałej częstotliwości, absolutnie konieczne jest zasilanie wzmacniacza operacyjnego z bardzo dobrze stabilizowanego źródła zasilania.

Powszechnie wiadomo, że wzmacniacze operacyjne nie są demonami prędkości. Każdy wzmacniacz operacyjny ma określoną szybkość narastania, wyrażoną w V/ μ s. Wielkość ta wskazuje, jak szybko napięcie na wyjściu wzmacniacza operacyjnego może rosnąć lub opadać. W tym zastosowaniu wielkość ta określa, jak szybko napięcie prostokątne przełącza się z „L” na „H” i z „H” na „L”. Ponieważ ta szybkość narastania jest dość niska w większości wzmacniaczy operacyjnych, nie należy oczekiwać użycia tego obwodu do generowania napięcia prostokątnego o częstotliwości 10 MHz. Dla wielu wzmacniaczy operacyjnych, 100 kHz jest wartością graniczną, przy której można oczekiwać akceptowalnego napięcia wyjściowego.

Symetryczne zasilanie

Jeśli wzmacniacz operacyjny jest używany w układzie z zasilaniem symetrycznym, można nieco uprościć obwód. Teoretyczny schemat przedstawiono na poniższym rysunku. Obwód wokół wejścia odwracającego jest identyczny. Jednak wejście nieodwracające jest teraz również objęte rezystancyjnym sprzężeniem zwrotnym między masą a wyjściem.

Działanie układu jest następujące: po podłączeniu do symetrycznego napięcia zasilania, wejście odwracające będzie miało wartość 0V. W końcu kondensator C1 jest rozładowany. W zależności od polaryzacji napięcia przesunięcia wzmacniacza operacyjnego, wejście nieodwracające będzie nieco bardziej dodatnie lub nieco bardziej ujemne niż 0 V. Tak więc wyjście wzmacniacza operacyjnego jest zablokowane względem jednego z napięć zasilania. W zależności od polaryzacji wyjścia, wejście nieodwracające ustawi się na dodatnie lub ujemne napięcie odniesienia. Wartość tego napięcia jest określana przez stosunek między R2 i R3.

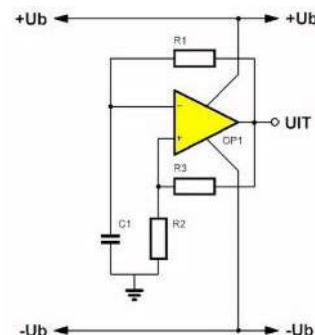
Kondensator C1 ładuje się dodatnio lub rozładowuje ujemnie, ponownie w zależności od polaryzacji wyjścia. Niezależnie od przypadku, po pewnym czasie napięcie na wejściu odwracającym staje się większe lub mniejsze niż napięcie na wejściu nieodwracającym. W rezultacie wyjście wzmacniacza operacyjnego zmienia polaryzację z + na - lub z - na +, a kondensator jest ładowany lub rozładowywany w przeciwnym kierunku.

Częstotliwość sygnału wyjściowego zależy tylko od wartości elementów pasywnych wokół wzmacniacza operacyjnego i wartości napięcia zasilania. Wyjście dostarcza niesymetryczny przebieg prostokątny, który przeskakuje tam i z powrotem pomiędzy napięciem ujemnym „L” i dodatnim „H”.

Wadą tego obwodu jest to, że napięcie na kondensatorze będzie zarówno ujemne, jak i dodatnie.

Dlatego nie jest możliwe użycie kondensatora elektrolitycznego.

Po raz kolejny można zastosować omówioną już zasadę kontrolowania stosunku włączania i wyłączania impulsu wyjściowego.



Schemat z symetrycznym zasilaniem
(© 2023 Jos Verstraten)

Generator przebiegu piłokształtnego

Wprowadzenie

Przebieg piłokształtny to sygnał charakteryzujący się stałym wzrostem lub spadkiem napięcia w czasie. Gdy napięcie wzrośnie lub spadnie do określonej wartości, powróci do poziomu wyjściowego tak szybko, jak to możliwe.

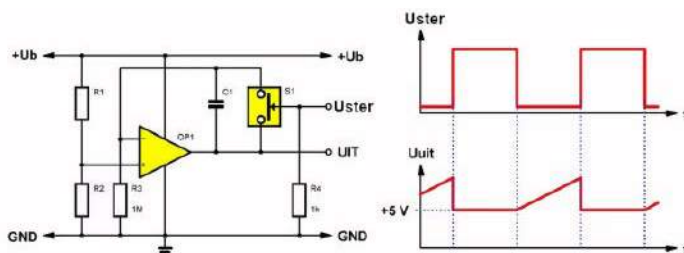
Podstawowy obwód

Podstawowy obwód generatora sygnału piłokształtnego ze wzmacniaczem operacyjnym pokazano na poniższym rysunku. Wejście nieodwracające wzmacniacza operacyjnego jest ustawione na połowę napięcia zasilania za pomocą rezystorów R1 i R2 o jednakowej wielkości. Załóżmy, że obwód jest zasilany napięciem +10 V. Wtedy wejście nieodwracające będzie miało wartość +5 V. Obwód będzie dążył do tego by na obu wejściach panowało to samo napięcie. W tym przypadku napięcie na wejściu odwracającym również będzie dążyło do osiągnięcia wartości +5 V. Jest to możliwe tylko wtedy, gdy przez rezystor R3 przepływa określony prąd. Przy podanej wartości 1 M Ω , prąd ten będzie równy 5 μ A.

Wejście odwracające ma bardzo wysoką impedancję wejściową. Prąd płynący przez rezystor R3 nie może płynąć dalej przez to wejście. Otwarta jest tylko jedna ścieżka prądowa, a mianowicie przez kondensator C1. Kondensator ten jest więc ładowany przez stały prąd o natężeniu 5 μ A. Jeśli kondensator ładowany jest stałym prądem, napięcie na nim będzie rosło liniowo. Ale górny zacisk kondensatora jest utrzymywany na stałym napięciu +5 V przez układ, w rezultacie czego napięcie na dolnym zacisku wzrasta liniowo. Liniowo rosnące napięcie powstaje na wyjściu wzmacniacza operacyjnego, co stanowi podstawową charakterystykę sygnału piłokształtnego.

Oczywiście proces ten nie może trwać wiecznie. Po pewnym czasie napięcie na wyjściu stałoby się większe niż napięcie zasilania, a to jest niemożliwe.

Dlatego kondensator C1 jest połączony równolegle z elektronicznym przełącznikiem S1 z układu scalonego CMOS CD4066B. Przełącznik ten jest sterowany napięciem zewnętrznym U_{ster} , dla którego bardzo przydatne jest



Podstawowy schemat generatora napięcia piłokształtnego (© 2023 Jos Verstraten)

napięcie prostokątne. Jeśli jest ono dodatnie, przełącznik zostanie zamknięty. Kondensator C1 jest teraz bardzo szybko rozładowywany przez bardzo niską rezystancję wewnętrzną zamkniętego przełącznika. W rezultacie napięcie na wyjściu kondensatora spada do +5 V, napięcia na obu wejściach.

Regulacja

Obwód ten wytwarza przebieg piłokształtny, którego amplituda wzrasta od napięcia ustawionego dzielnikiem z rezystorów R1 oraz R2 o równych wartościach czyli przy napięciu zasilania 10 V wzrasta od napięcia 5 V do pewnej wyższej wartości. Częstotliwość sygnału wyjściowego zależy od częstotliwości sygnału sterującego przełącznikiem elektronicznym.

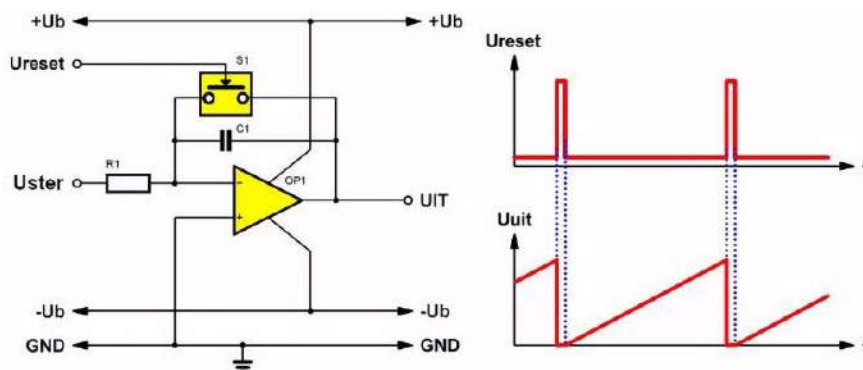
Alternatywny obwód z symetrycznym zasilaniem

Wejście nieodwracające jest podłączone do masy, w wyniku czego obwód ustawi również napięcie na wejściu odwracającym na 0V. Wejście odwracające jest podłączone do zewnętrznego napięcia sterującego U_{ster} poprzez rezystor R1. Przez rezystor R1 przepływa prąd do wejścia odwracającego. Wielkość tego prądu zależy od wartości rezystora i wielkości napięcia sterującego. Ponownie, prąd ten nie może płynąć dalej do wejścia odwracającego. Jediną ścieżką dla prądu, która pozostaje otwarta, jest kondensator C1. Przepływa przez niego stały prąd, a napięcie na kondensatorze rośnie lub opada liniowo.

Narastający lub opadający charakter napięcia na kondensatorze zależy od polaryzacji napięcia sterującego. Kondensator jest ponownie rozładowywany przez zamknięcie przełącznika elektronicznego, który jest podłączony równolegle z kondensatorem, za pomocą impulsu U_{reset} .

Mamy pełną kontrolę nad wszystkimi właściwościami generowanego przebiegu piłokształtnego. W końcu polaryzacja tego przebiegu zależy od polaryzacji napięcia sterującego. Szybkość, z jaką przebieg piłokształtny maleje lub rośnie, zależy od wielkości napięcia sterującego.

Częstotliwość przebiegu zależy od częstotliwości napięcia prostokątnego, za pomocą którego steruje się przełącznikiem elektronicznym. Jeśli przełącznik jest sterowany dodatnim napięciem prostokątnym w kształcie szpilki, jak w powyższym przykładzie, kondensator zostanie rozładowany przez ten wąski dodatni impuls, a zaraz potem rozpocznie się nowy okres przebiegu piłokształtnego.



Generator piłokształtny z symetrycznym zasilaniem (© 2023 Jos Verstraten)

Generator fali sinusoidalnej

Wprowadzenie

Przebieg sinusoidalny jest jedną z najważniejszych form sygnałów w elektronice. Cała technologia audio i wideo opiera się na wzmacnianiu i przetwarzaniu sygnałów sinusoidalnych lub ich kombinacji (z punktu widzenia matematyki przebiegi dowolnych kształtów można rozpatrywać jako kombinacje fal sinusoidalnych o różnych amplitudach, fazach i częstotliwościach, co stanowi podstawę cyfrowego przetwarzania sygnałów – przyp. tłum.). Oczywiście jest zatem, że obwody, które mogą generować napięcia sinusoidalne, są bardzo popularne.

Aby zrozumieć działanie generatora fali sinusoidalnej, należy najpierw wyjaśnić kilka ogólnych pojęć z zakresu elektroniki.

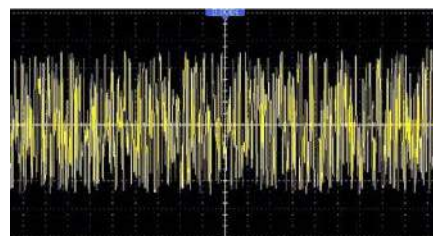
Szum

Pierwszym omawianym pojęciem jest „szum”. Szum to niepożądane napięcie zakłócające, które jest generowane w każdym obwodzie elektronicznym. Istnieją różne rodzaje szumów, ale najczęstszym jest „szum termiczny”. Pod wpływem temperatury części, wolne elektrony przeskakują z jednego atomu na drugi. Ten ruch elektronów powoduje bardzo mały prąd, który generuje napięcie szumu na komponentach. Szum jest zatem zjawiskiem całkowicie statystycznym. Nie można przewidzieć momentu, w którym wolny elektron zdecyduje się na migrację z jednego atomu do drugiego. W rezultacie szum składa się z sygnałów o bardzo różnych częstotliwościach. Gdybyś miał przeanalizować sygnał szumu pod kątem zawartości częstotliwości, odkryłbyś, że obecne są prawie wszystkie możliwe częstotliwości. Poniższy oscylogram przedstawia typowy przebieg szumu.

Filtr

Drugą podstawową zasadą, którą należy wyjaśnić, jest zasada działania filtra. Filtr to obwód, który pozwala na przepuszczanie lub tłumienie sygnałów o jednych określonych częstotliwościach w jak największym stopniu oraz na tłumienie lub przepuszczanie sygnałów o innych częstotliwościach. Istnieją niezliczone obwody filtrów. Generatory sinusoidalne są prawie zawsze zbudowane wokół filtru pasmowego, który nosi nazwę „mostka Wiena”. Typowy układ takiego filtra pokazano na poniższym rysunku. Wejście jest połączone z wyjściem poprzez rezystor R1 i kondensator C1. Pomiedzy wyjściem a masą znajduje się równoległe połączenie rezystora R2 i kondensatora C2. W większości przypadków $R1=R2$ i $C1=C2$, ale nie jest to wymóg bezwzględny.

Jeśli do wejścia filtra przyłożymy sygnał sinusoidalny o zmiennej częstotliwości i zmierzmy, ile tego sygnału można znaleźć na wyjściu, powstanie typowy obraz naszkicowany po prawej stronie na rysunku. Dla bardzo niskich częstotliwości na wyjściu nie znajdziemy nic. Jest to całkiem logiczne, ponieważ kondensator szeregowy C1 ma



Sygnał szumu na oscyloskopie (© 2023 Jos Verstraten)

bardzo wysoką rezystancję dla tych niskich częstotliwości. Wraz ze wzrostem częstotliwości sygnału wejściowego, będzie można zmierzyć coraz więcej sygnału wyjściowego. Istnieje pewna częstotliwość f_0 , przy której sygnał wyjściowy jest maksymalny. Jeśli częstotliwość sygnału wejściowego wzrośnie jeszcze bardziej, okaże się, że na wyjściu jest mniej sygnału. Ma to również sens, ponieważ wraz ze wzrostem częstotliwości sygnału kondensator C2 będzie miał coraz niższą rezystancję i będzie zwracał coraz więcej sygnału do masy.

Częstotliwość f_0 jest nazywana częstotliwością drgań własnych filtra. Przy tej częstotliwości mostek Wienera tłumia sygnał dokładnie trzykrotnie. Napięcie wyjściowe jest zatem równe $1/3$ napięcia wejściowego.

Przesunięcie fazowe mostka Wienera

Jest jeszcze jedna ważna cecha mostka Wienera, o której nie można nie wspomnieć. Jeśli zmierzmy przesunięcie fazowe między przebiegiem wejściowym i wyjściowym, okaże się, że przy naturalnej częstotliwości f_0 to przesunięcie fazowe wynosi dokładnie 0° . Zjawisko to zostało przedstawione graficznie na poniższym wykresie. Dla wszystkich innych częstotliwości sygnału występuje dodatnia lub ujemna różnica faz między sygnałem wejściowym i wyjściowym.

Szum na wejściu mostka Wienera

Jeśli podłączysz źródło szumu do wejścia filtra, składowa sygnału w szumie o częstotliwości równej f_0 pojawi się na wyjściu mniej słyszalna niż wszystkie inne składowe szumu. Co więcej, będzie to jedyny sygnał z szumu, w którym nie wystąpi przesunięcie fazowe.

Podstawowy generator fali sinusoidalnej

Podstawowy schemat generatora fali sinusoidalnej ze wzmacniaczem operacyjnym został pokazany na rysunku obok. Należy zauważyć, że ten obwód nie będzie działał. Ale omówienie tego schematu jest niezbędnym krokiem w kierunku ostatecznej wersji obwodu generatora.

Mostek Wienera włączony jest między wyjściem wzmacniacza operacyjnego i jego wejściem nieodwracającym, dzieląc z nim masę. Wejście mostka jest połączone z wyjściem wzmacniacza, zaś wyjście mostka jest połączone z wejściem nieodwracającym wzmacniacza.

Wejście odwracające wzmacniacza operacyjnego jest włączone w rezystorową pętlę sprzężenia zwrotnego. Rezystory R3, R4 i R5 zapewniają określony współczynnik wzmocnienia obwodu. Wartość współczynnika wzmocnienia można regulować potencjometrem regulacyjnym R5. Załóżmy, że potencjometr zostanie ustawiony na wartość minimalną. W tym momencie wzmocnienie wzmacniacza operacyjnego jest określone przez stosunek rezystorów R3 i R4.

Założmy teraz, że obwód jest podłączony do symetrycznych napięć zasilania, a do wyjścia podłączony jest oscyloskop.

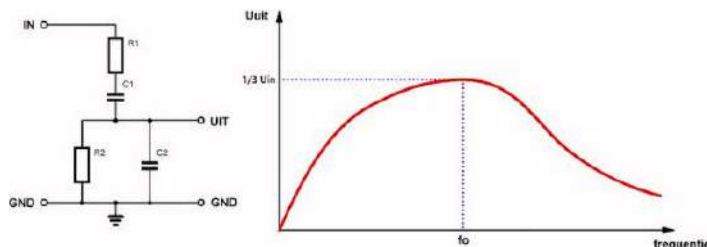
Nic się nie stanie, wyjście obwodu pozostanie na poziomie 0 V. Następnie należy bardzo powoli obracać R5, tak aby wzmocnienie wzmacniacza operacyjnego powoli wzrastało. W pewnym momencie zauważysz, że na wyjściu obwodu pojawiają się bardzo małe oscylacje sinusoidalne. W tym momencie nie można już obracać potencjometru R5.

To, co dzieje się później, przedstawiono na poniższym rysunku. Amplituda sygnału wyjściowego będzie powoli wzrastać i początkowo sygnał będzie sinusoidalny. Jednak po pewnym czasie amplituda staje się tak duża, że sygnał wyjściowy wzmacniacza operacyjnego osiągnie wartości bliskie napięciu zasilania. Fala sinusoidalna zostanie najpierw zniekształcona, a po pewnym czasie zniekształcenie to stanie się tak duże, że sygnał wyjściowy będzie przypominał przebieg prostokątny. Jest to stan stabilny obwodu.

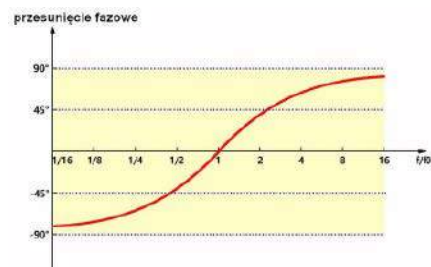
Wyjaśnienie działania

Niezwykle ważne jest zrozumienie szczególnego zachowania obwodu. Co się dzieje? Gdy zasilacz jest włączony, szum pojawi się na wszystkich rezystorach (a także wewnątrz krzemowej struktury wzmacniacza operacyjnego – przyp. tłum.). Te szumy oczywiście pojawiają się również na wyjściu wzmacniacza operacyjnego. Są one podawane z powrotem na wejście nieodwracające przez mostek Wienera. Filtr ten zapewnia, że sygnały szumu o częstotliwości równej f_0 są podawane z powrotem na wejście nieodwracające z najmniejszym tłumieniem. Co więcej, tylko te sygnały pojawiają się w fazie na wyjściu filtra. Jednak mostek Wienera również tłumia te sygnały trzykrotnie. Ponieważ wzmocnienie wzmacniacza operacyjnego jest obecnie mniejsze niż trzy, sygnał o częstotliwości f_0 pojawi się na wyjściu wzmocniony, ale nie na tyle, aby przewyższyć tłumienie filtra Wienera. Obwód pozostaje w spoczynku.

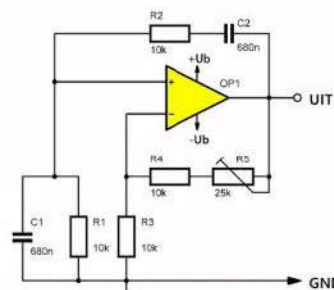
Jeśli jednak, obracając potencjometrem R5, zwiększysz wzmocnienie obwodu do dokładnie trzech, sygnał podawany z powrotem i tłumiony trzykrotnie przy częstotliwości f_0 zostanie następnie trzykrotnie wzmocniony przez wzmacniacz operacyjny. Sygnał ten jest



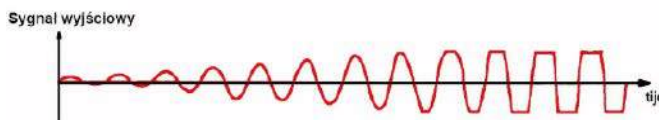
Strojony filtr znany jako „mostek Wienera” (© 2023 Jos Verstraten)



Przesunięcie fazowe mostka Wienera (© 2023 Jos Verstraten)



Podstawowy obwód generatora fali sinusoidalnej (© 2023 Jos Verstraten)



Profil napięcia na wyjściu obwodu (© 2023 Jos Verstraten)

podawany z powrotem na wejście mostka, pojawia się trzykrotnie stłumiony, ale w fazie na wejściu nieodwracającym, a następnie jest trzykrotnie wzmacniany przez wzmacniacz operacyjny.

W rezultacie tylko ten jeden sygnał może nadal krążyć w obwodzie, a obwód zasadniczo prezentowałby na wyjściu ładną, nieznkształconą falę sinusoidalną o częstotliwości równej f_o . Jednak całe działanie układu zależy od dokładnego skompensowania tłumienia mostka Wiena przez wzmacniacz operacyjny. Jeśli wzmacnienie wzmacniacza operacyjnego stanie się nieco mniejsze niż trzy, sygnał ponownie zaniknie. Jeśli wzmacnienie wzrośnie nieco powyżej trzech, amplituda sygnału będzie rosła, aż do napięcia zasilania.

Automatyczna regulacja wzmacnienia AVR

Jedyną możliwością przekształcenia opisanego wyżej obwodu w praktycznie użyteczny generator fali sinusoidalnej jest wprowadzenie automatycznej regulacji wzmacnienia (ARW), która zapewnia automatyczne dostosowanie wzmacnienia do tłumienia filtra. Innymi słowy, należy wprowadzić obwód, który automatycznie dostosowuje wzmacnienie wzmacniacza operacyjnego w taki sposób, że amplituda napięcia wyjściowego pozostaje stała.

Przez lata opracowano niezliczone układy i zasady tej absolutnie niezbędnej automatycznej kontroli wzmacnienia. Niektóre z nich są bardzo proste, inne zawierają znacznie więcej komponentów niż sam generator fali sinusoidalnej. Nie jest to błędem, ponieważ jakość generatora fali sinusoidalnej jest całkowicie zdeterminowana przez właściwości automatycznej regulacji wzmacnienia.

Wpływ na zniekształcenie sinusoidy

Problem polega na tym, że układ ten wprowadza do obwodu nieliniowe sprzężenie zwrotne. Specyfikacje tego sprzężenia zwrotnego określają, między innymi, zniekształcenia występujące w fali sinusoidalnej na wyjściu. Istnieją generatory fali sinusoidalnej o zniekształceniu 1,0% i 0,01%. Wartość ta zależy przede wszystkim od dokładności komponentów w mostku Wiena, ale w drugiej kolejności, prawie równie ważnej, od jakości automatycznej regulacji wzmacnienia.

W poniższych akapitach pod lupę wzięto kilka powszechnie używanych rozwiązań układowych.

Termistor jako element sprzężenia zwrotnego

Termistor, NTC lub rezystor o ujemnym współczynniku temperaturowym to rezystor, którego rezystancja maleje wraz z nagrzewaniem się rezystora. Charakterystykę niektórych termistorów przedstawiono na poniższym rysunku. Na osi poziomej widoczna jest temperatura, a na osi pionowej rezystancja termistora NTC.

Takiego termistora można użyć zgodnie ze schematem po prawej stronie do automatycznej stabilizacji wzmacnienia generatora fali sinusoidalnej. Termistor jest włączony w obwód sprzężenia zwrotnego określający wzmacnienie między wyjściem, a wejściem odwracającym. Po podłączeniu obwodu do zasilania termistor będzie zimny, a jego rezystancja będzie dość wysoka. Współczynnik wzmacnienia wzmacniacza operacyjnego jest określany przez stosunek rezystancji termistora do wartości rezystora R_3 . Współczynnik ten jest duży i w każdym przypadku obwód będzie wzmacniał ponad trzykrotnie. Wzmacniacz kompensuje osłabienie mostek Wiena i obwód zaczyna oscylować. Napięcie wyjściowe jest początkowo dość wysokie, istnieje nawet szansa, że amplituda sygnału na wyjściu będzie blisko napięcia zasilania. Tak duża amplituda wyjściowa wymusza przepływ prądu przez termistor, powodując jego nieznaczne nagrzanie i zmniejszenie rezystancji.

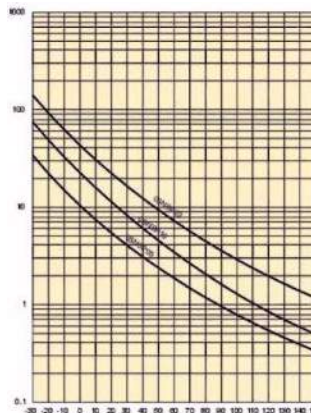
W rezultacie zmniejsza się również współczynnik wzmacnienia wzmacniacza operacyjnego. Sztuczka polega teraz na znalezieniu odpowiedniego termistora i połączeniu go z prawidłowo dobraną wartością rezystora R_3 . Tylko wtedy obwód ustabilizuje się w stanie, w którym na wyjściu powstaje stabilna i w miarę wolna od zniekształceń fala sinusoidalna.

Wadą tego układu jest to, że temperatura otoczenia ma duży wpływ na jego działanie. Dlatego generator sinusoidalny stabilizowany termistorem można znaleźć tylko w miejscach, w których nie ma wysokich wymagań dotyczących jakości fali sinusoidalnej.

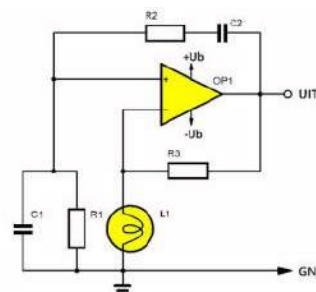
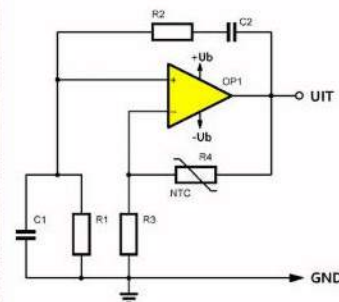
Żarówka jako element sprzężenia zwrotnego

Żarówka jest zasadniczo termistorem PTC lub rezystorem o dodatnim współczynniku temperaturowym. Zimny żarnik żarówki ma niewielką rezystancję, która wzrasta w miarę nagrzewania się żarnika. Ta właściwość umożliwia stabilizację generatora fali sinusoidalnej poprzez włączenie małej żarówki do pętli sprzężenia zwrotnego.

Podstawowy schemat przedstawiono na poniższym rysunku. Lampka L_1 jest podłączona między wejściem odwracającym wzmacniacza operacyjnego a masą. Gdy jest zimna, rezystancja L_1 jest niska. Wzmacnienie wzmacniacza operacyjnego zależy od stosunku rezystancji tej lampki do wartości rezystora R_3 . Jeśli stosunek ten jest duży, wzmacniacz operacyjny wzmacnia znacznie więcej niż trzy razy i obwód zaczyna oscylować. Duża amplituda wyjściowa wymusza duży prąd przez R_3 i L_1 , powodując nagrzewanie się żarnika i wzrost rezystancji. Współczynnik wzmacnienia maleje i ponownie sztuczka polega na wybraniu L_1 i R_3 w taki sposób, aby obwód ustabilizował się przy współczynniku wzmacnienia równym trzy.



Termistor używany do regulacji wzmacnienia (© 2023 Jos Verstraten)



Żarówka przypominająca AVR w generatorze sinusoidalnym (© 2023 Jos Verstraten)

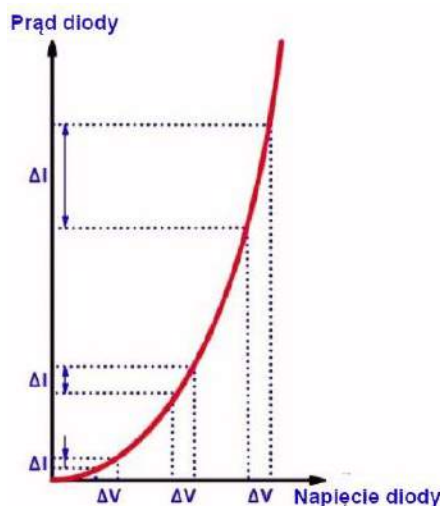
Diody krzemowe jako element sprzężenia zwrotnego

Dioda krzemowa ma charakterystykę prądowo-napięciową narysowaną po lewej stronie na poniższym rysunku. Jeśli napięcie na diodzie jest niskie, prąd jest również bardzo mały. Jednak wraz ze wzrostem napięcia prąd będzie wzrastał bardziej niż liniowo. Wartość rezystancji wewnętrznej można wyprowadzić z charakterystyki prądowo-napięciowej elementu. Zostało to zrobione w trzech różnych miejscach wykresu. Dla trzech równych zmian napięcia ΔU wykreślono odpowiadające im zmiany prądu ΔI . Pokazuje to, że ΔI_1 jest znacznie mniejsze niż ΔI_3 . Logiczną konsekwencją jest to, że rezystancja wewnętrzna diody jest znacznie większa przy najniższym ΔU niż przy najwyższym ΔU . W końcu zmiana prądu przy tej samej zmianie napięcia jest całkowicie zdeterminowana przez rezystancję!

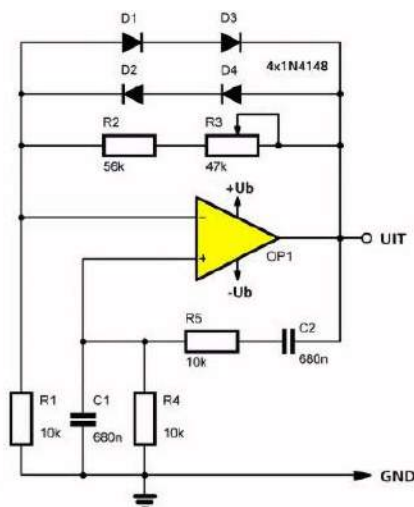
Można zatem stwierdzić, że rezystancja wewnętrzna diody krzemowej nie jest stała, ale maleje wraz ze wzrostem napięcia na diodzie. Będzie więc jasne, że można użyć takiej części do ustabilizowania wzmocnienia generatora fali sinusoidalnej.

Podstawowy schemat został narysowany po prawej stronie na poniższym rysunku. Cztery antyrównoległe połączone diody D1 do D4 są zawarte w pętli sprzężenia zwrotnego wzmacniacza operacyjnego. Gdy zasilanie jest włączone, na diodach nie ma napięcia. Mają one bardzo

wysoką rezystancję. Współczynnik wzmocnienia wzmacniacza operacyjnego jest bardzo duży, obwód natychmiast zaczyna oscylować. Wraz ze wzrostem napięcia wyjściowego do diod przykładane jest większe napięcie. Ich rezystancja wewnętrzna maleje, a tym samym maleje wzmocnienie wzmacniacza operacyjnego. Ponownie, zadaniem projektanta jest upewnienie się, że R1, R2 i R3 są dobrane w taki sposób, aby obwód ostatecznie ustabilizował się przy współczynniku wzmocnienia równym trzy. Za pomocą potencjometru R3 można wyregulować obwód tak, aby na wyjściu uzyskać ładną falę sinusoidalną.



Cztery diody krzemowe jako AVR w generatorze fali sinusoidalnej (© 2023 Jos Verstraten)



Tranzystor FET jako element sprzężenia zwrotnego

Ostatnim przykładem użytecznego elementu sprzężenia zwrotnego jest tranzystor FET. FET to tranzystor, który ma bardzo liniową część swojej charakterystyki prądowo-napięciowej, patrz rysunek poniżej. Jeśli bramka jest spolaryzowana ujemnie w stosunku do źródła, powstaje obszar, w którym istnieje całkowicie liniowa zależność między napięciem bramka-źródło a prądem drenu. W tym obszarze FET zachowuje się jak rezystor sterowany, z rezystancją między drenem a źródłem wprost proporcjonalną do napięcia między bramką a źródłem.

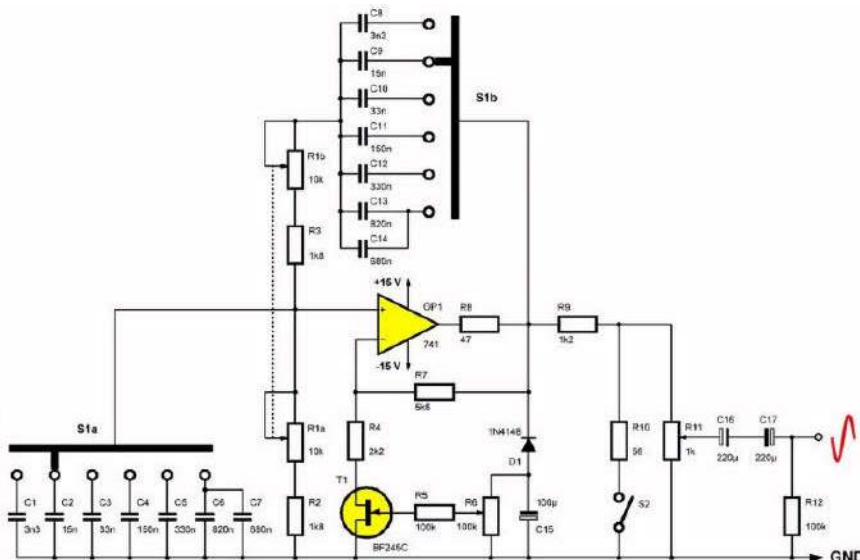
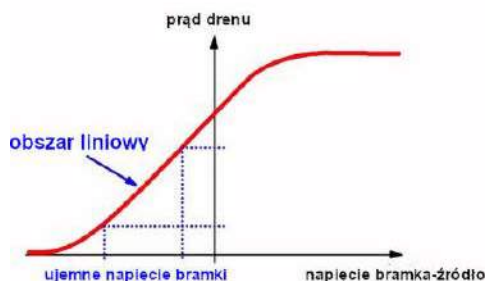
Dzięki tej właściwości, tranzystor FET może być wykorzystywany do stabilizacji wzmocnienia generatora fali sinusoidalnej. A dzięki całkowicie liniowemu zachowaniu jest to wręcz najlepsze rozwiązanie, ponieważ nie ma potrzeby stosowania nieliniowego sprzężenia zwrotnego i związanych z nim zniekształceń harmonicznych sygnału wyjściowego. W większości praktycznych generatorów fali sinusoidalnej można znaleźć tranzystor FET jako element sprzężenia zwrotnego.

Bardzo prosty przykładowy obwód wykorzystujący zasadę FET pokazano po prawej stronie na poniższym rysunku. FET T1 jest zawarty w sprzężeniu zwrotnym określającym wzmocnienie wzmacniacza operacyjnego OP1. To sprzężenie zwrotne składa się również z elementów R4 i R7. Problemem w takich układach jest uzyskanie sygnału bramki dla tranzystora FET. W końcu bramka musi być sterowana napięciem stałym, a generator dostarcza napięcie zmienne. W tym przypadku rozwiązano ten problem w najprostszym sposób. Napięcie wyjściowe generatora sinusoidalnego jest prostowane za pomocą diody D1. Wyprostowany sygnał jest wygładzany za pomocą kondensatora C15. Wygładzone napięcie podawane jest na potencjometr regulacyjny R6, a z niego trafia do bramki tranzystora FET. Za pomocą potencjometru R6 można określić najkorzystniejszy punkt pracy tranzystora FET.

Działanie układu nie różni się zasadniczo od układów omówionych do tej pory. Gdy obwód jest podłączony do zasilania, napięcie na kondensatorze C15 będzie wynosić 0 V. W rezultacie prąd drenu tranzystora FET jest dość wysoki, a jego rezystancja wewnętrzna jest niska. Sprzężenie zwrotne składa się z tej niskiej rezystancji w szeregu z również dość niskim R4 i rezystorem 5,6 kΩ R7. Obwód ma współczynnik wzmocnienia większy niż trzy i generator może oscylować. Rosnące napięcie wyjściowe jest prostowane i tworzy ujemne napięcie na kondensatorze C15. Napięcie to przesuwając ujemnie punkt pracy tranzystora FET, powodując spadek prądu drenu i wyższą rezystancję tranzystora. Wzmocnienie obwodu zmniejszy się. Jeśli przyłożysz zbyt duże napięcie przez R5 do bramki FET, wzmocnienie spadnie tak bardzo, że oscylacja wygaśnie.

Jednak poprzez regulację R6 można ustawić obwód na stabilny, sinusoidalny przebieg wyjściowy.

W tym przykładzie należy również zwrócić uwagę na sposób ustawiania zakresu częstotliwości. Przełącznik zakresu S1a+S1b przełącza dwa identyczne kondensatory w sprzężeniu zwrotnym wzmacniacza operacyjnego. Za pomocą potencjometru stereo R1a+R1b można precyzyjnie dostroić częstotliwość w dowolnym zakresie.



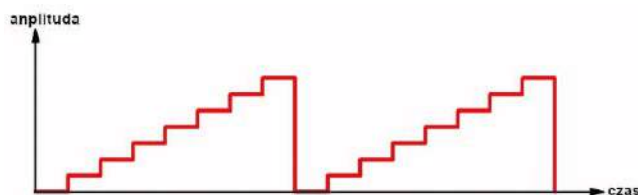
FET jako AVR w generatorze fali sinusoidalnej (© 2023 Jos Verstraten)

Obwód wyjściowy jest również wart dalszego przestudiowania, ponieważ można go użyć w dowolnym obwodzie generatora. Wyjście wzmacniacza operacyjnego trafia do dzielnika rezystorowego R9/R10/R11. Gdy przełącznik S2 jest otwarty, R10 nie odgrywa żadnej roli i można ustawić napięcie wyjściowe na żadaną wartość za pomocą potencjometru R11. Jeśli chcesz uzyskać bardzo małe napięcie sinusoidalne, zamykasz przełącznik S2, który umieszcza bardzo mały rezystor R10 równolegle do R11 i zmniejsza zakres regulacji napięcia wyjściowego dziesięciokrotnie.

Generator sygnału schodkowego

Wprowadzenie

Przebiegi schodkowe są niezbędne, jeśli chcesz zaprojektować narzędzie do śledzenia krzywej w celu wyświetlenia charakterystyki prądowo-napięciowej komponentu na ekranie oscyloskopu. Typowy przebieg schodkowy przedstawiono na poniższym rysunku. Przebieg charakteryzuje się tym, że amplituda wzrasta lub maleje do pewnej stałej wartości w regularnych odstępach czasu. W pierwszym przypadku mamy do czynienia z dodatnim przebiegiem schodkowym, a w drugim z ujemnym. Wielkość każdego stopnia i liczba stopni na okres zależy od sposobu zaprojektowania generatora.



Typowy kształt przebiegu schodkowego (© 2023 Jos Verstraten)

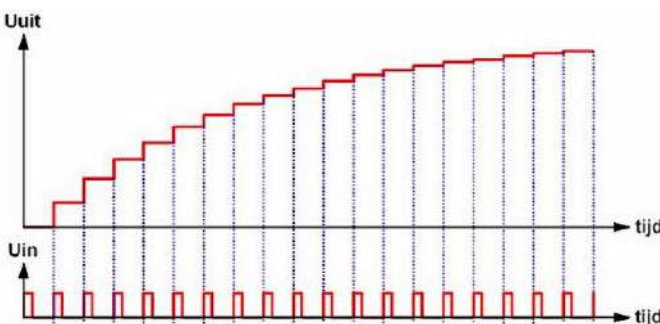
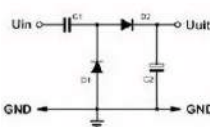
Pompa diodowa

Podstawową częścią generatora schodkowego jest pompa diodowa. Dlatego należy najpierw zrozumieć działanie takiego obwodu. Teoretyczny schemat pompy diodowej pokazano na poniższym rysunku.

Obwód jest sterowany z generatora sygnału prostokątnego. Napięcie wyjściowe tego generatora przechodzi przez kondensator C1 do diody D1. Te dwie części tworzą tak zwany obwód ograniczający napięcie, który zapewnia, że przejścia od niskiego do wysokiego napięcia prostokątnego są przekształcane w wąskie dodatnie impulsy szpilkowe. Impulsy te powodują przewodzenie diody D2. W rezultacie kondensator C2 jest krótko ładowany przy każdym impulsie i pozostaje pod stałym napięciem między dwoma kolejnymi impulsami. Zależność między napięciem wejściowym i wyjściowym tego obwodu jest pokazana po prawej stronie rysunku.

Oczywiście prąd ładowania płynący przez diodę D2 do kondensatora C2 będzie malał wraz ze wzrostem napięcia na kondensatorze. W końcu prąd ten zależy od różnicy napięć na diodzie. Lewa strona zawsze pokazuje te same dodatnie impulsy, prawa strona pokazuje rosnące napięcie kondensatora. Będzie zatem jasne, że kolejne schodki napięcia wyjściowego nie są tej samej wielkości, ale powoli zmniejszają amplitudę.

Pompa diodowa generuje napięcie schodkowe, ale jest ono całkowicie bezużyteczne w większości zastosowań. Jednak dzięki wprowadzeniu wzmacniacza operacyjnego można łatwo pokonać tę przeszkodę.



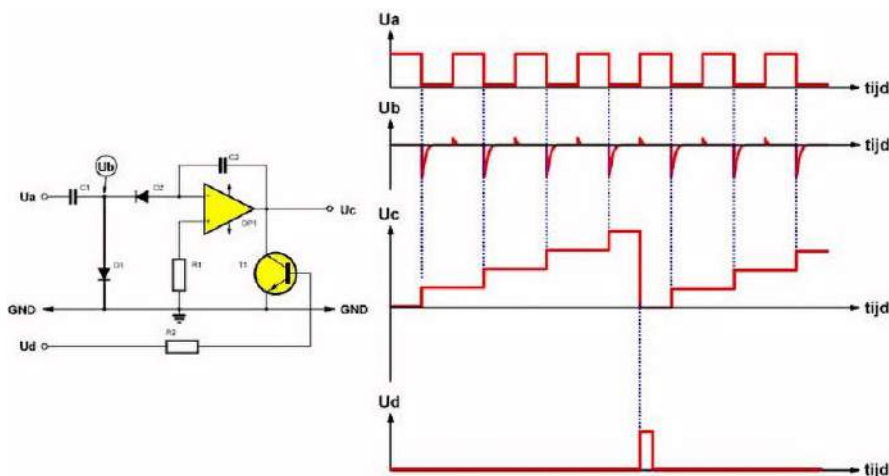
Działanie pompy diodowej (© 2023 Jos Verstraten)

Obwód ze wzmacniaczem operacyjnym

Poniższy rysunek przedstawia aktywną pompę diodową, w której wada niestalego rozmiaru schodka została przezwyciężona poprzez włączenie wzmacniacza operacyjnego do obwodu. Działanie obwodu wyjaśniono za pomocą wykresów na rysunku.

Kondensator C1 wraz z diodą D1 tworzy obwód ograniczający napięcie. Ale teraz dioda została odwrócona, w wyniku czego prostokątne napięcie U_a na wejściu jest przekształcane w ujemne impulsy szpilkowe U_b . Wzmacniacz operacyjny jest podłączony jako integrator (układ całkujący – przyp. tłum.). Z każdym ujemnym impulsem na diodzie D1, dioda D2 na krótko przewodzi. Prąd płynący przez diodę może pochodzić tylko z kondensatora C2. Wejście odwracające wzmacniacza operacyjnego ma bardzo wysoką impedancję. Przez kondensator C2 zatem płynie prąd przy każdej szpilce, w wyniku czego kondensator się ładuje. Ponieważ wejście nieodwracające wzmacniacza operacyjnego jest podłączone do masy, wejście odwracające również będzie na tym potencjale. W rezultacie wyjście wzmacniacza operacyjnego wzrasta o stałe napięcie przy każdym impulsie. Wielkość impulsów prądowych jest określana przez różnicę napięć na diodzie D2. W przeciwieństwie do poprzedniego obwodu, ta różnica napięć jest teraz stała.

Na wyjściu wzmacniacza operacyjnego powstaje zatem stopniowane napięcie dodatnie. Oczywiście po wykonaniu określonej liczby kroków należy wszystko zresetować. Odbywa się to za pośrednictwem tranzystora T1. Jest on sterowany przez wąski dodatni impuls U_d na bazie. Impuls ten powoduje przewodzenie tranzystora, a napięcie na prawej płytce kondensatora C2 jest zwierane do masy. Napięcie wyjściowe U_c spada do zera i układ jest gotowy do wygenerowania drugiego okresu napięcia schodkowego.



Aktywna pompa diodowa jako generator napięcia schodkowego (© 2023 Jos Verstraten)

Generator przebiegu trójkątnego

Wprowadzenie

Spośród wszystkich podstawowych form sygnałów w elektronice, trójkąt ma prawdopodobnie najmniej praktycznych zastosowań. Jednak dzięki wszechobecnym w dzisiejszych czasach generatorom funkcyjnym, przebieg trójkątny jest obecnie jednym ze standardowych sygnałów w laboratorium elektronicznym.

Podstawowy obwód

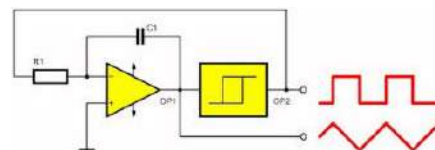
W zasadzie bardzo łatwo jest wyprowadzić przebieg trójkątny z sygnału prostokątnego. Wystarczy dostarczyć prostokąt do integratora. Jeśli napięcie prostokątne jest dodatnie, prąd ładowania przepływający przez kondensator całkujący zapewni liniowe rozładowanie tego komponentu. Napięcie wyjściowe obwodu spadnie. Jeśli napięcie prostokątne jest ujemne, prąd płynący przez kondensator odwraca się, a integrator dostarcza rosnące napięcie.

Wadą tego układu jest to, że amplituda trójkąta na wyjściu jest całkowicie zależna od częstotliwości przebiegu prostokątnego. Jest to tak poważna wada, że w praktyce nigdy nie znajdziemy tak prostego układu.

Analogowy generator funkcyjny

Obecnie przebiegi trójkątne są zawsze generowane w analogowym generatorze funkcyjnym. Analogowy generator funkcyjny jest podstawowym układem elektronicznym składającym się z dwóch wzmacniaczy operacyjnych, które są włączone w pętlę sprzężenia zwrotnego. Podstawowy schemat takiego generatora funkcyjnego przedstawiono na poniższym rysunku. Integrator składa się z lewego wzmacniacza operacyjnego OP1, rezystora ładowania R1 i kondensatora integratora C1. Wyjście tego integratora steruje drugim wzmacniaczem operacyjnym OP2, podłączonym jako komparator z histerezą. Oznacza to, że obwód ma dwa zdefiniowane punkty przełączania. Wyjście komparatora z kolei steruje wejściem integratora.

Na wyjściu integratora powstaje trójkąt, a na wyjściu komparatora prostokąt. Oba sygnały mają tę samą częstotliwość, która jest określona przez wartości rezystora R1 i kondensatora C1.



Zasada działania generatora funkcyjnego (© 2023 Jos Verstraten)

Praktyczny obwód

Działanie generatora funkcyjnego można najlepiej wyjaśnić za pomocą praktycznego schematu. Został on przedstawiony na poniższym rysunku.

Lewy wzmacniacz operacyjny OP1 jest komparatorem. Wejście odwracające jest bezpośrednio podłączone do masy, więc będzie jasne, że napięcie wyjściowe wzmacniacza operacyjnego zmienia się, gdy napięcie na wejściu nieodwracającym przechodzi przez zero.

Wejście nieodwracające jest podłączone do punktu między dwoma rezystorami połączonymi szeregowo. Jedną gałąź tego dzielnika (R5) idzie do wyjścia integratora OP2, druga gałąź (R1+R2) idzie do wyjścia komparatora. Oczywiście jest, że napięcie na wejściu nieodwracającym jest określone przez napięcia wyjściowe obu wzmacniaczy operacyjnych.

Aby zrozumieć działanie obwodu, konieczne jest założenie, że w chwili t_0 napięcie wyjściowe integratora jest dodatnie (założmy +5 V), a wyjście komparatora jest również dodatnie (założmy +10 V). Oczywiście jest, że wejście nieodwracające komparatora jest zatem dodatnie. Dokładna wartość tego dodatniego napięcia zależy od stosunku między rezystorami R_5 i R_1+R_2 .

Wykresy zakładają, że w opisywanym momencie napięcie na wejściu nieodwracającym jest równe +6,7 V. Ponieważ wejście nieodwracające komparatora jest bardziej dodatnie niż wejście odwracające, obwód ten wygeneruje wysokie napięcie wyjściowe. Jedno z założeń jest już prawidłowe! Wejście integratora jest sterowane napięciem dodatnim. Prąd płynie przez rezystor R_3+R_4 do wejścia odwracającego integratora. Jednak wejście to jest praktycznie połączone z masą, prąd może płynąć tylko przez kondensator C_1 integratora do wyjścia prawego wzmacniacza operacyjnego. Ponieważ wzmacniacz operacyjny zapewnia, że lewe połączenie kondensatora pozostaje na potencjale masy, napięcie wytworzone przez prąd na kondensatorze integratora odłoży się na prawym połączeniu. Kierunek prądu oznacza, że napięcie w tym punkcie będzie spadać.

Po czasie t_0 sytuacja wygląda następująco. Wyjście komparatora ma stałe napięcie dodatnie, wyjście integratora ma malejące napięcie dodatnie. Zmniejszenie tego ostatniego napięcia oznacza, że napięcie na nieodwracającym wejściu komparatora również spadnie.

W pewnym momencie t_1 , napięcie na wyjściu integratora spadnie tak bardzo, że napięcie na nieodwracającym wejściu komparatora osiągnie wartość 0 V. Komparator przełączy się i jego napięcie wyjściowe stanie się ujemne.

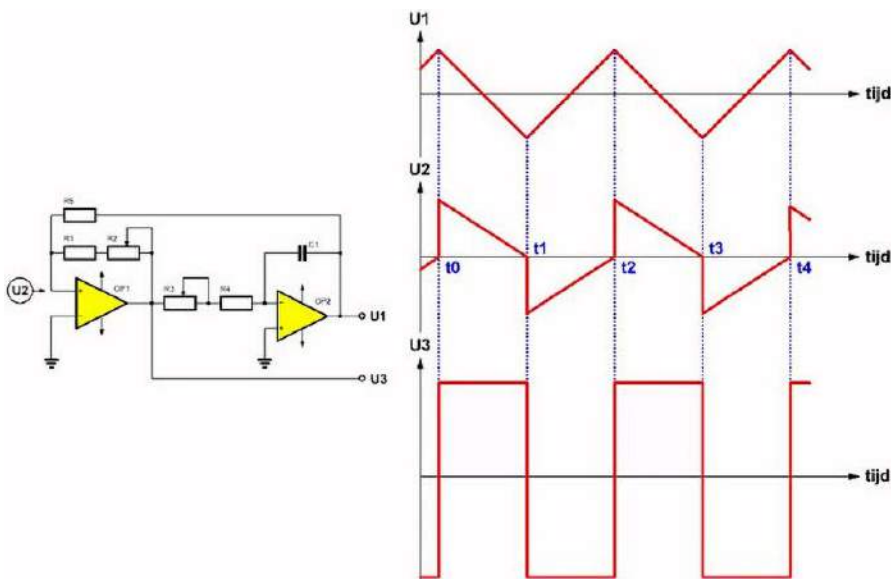
Zjawisko to ma daleko idące konsekwencje dla napięcia na nieodwracającym wejściu komparatora. Punkt ten jest teraz zasilany z dwóch ujemnych napięć wyjściowych, w wyniku czego napięcie nagle staje się dość ujemne.

To, co stanie się później, będzie jasne. Integrator jest teraz sterowany napięciem ujemnym, prąd płynący przez kondensator integratora zmienia polaryzację, napięcie wyjściowe integratora zaczyna rosnąć. Aż do czasu t_2 , czyli momentu, w którym napięcie na nieodwracającym wejściu komparatora ponownie przechodzi przez zero i obwód ponownie się odwraca.

Regulacja

Generator funkcyjny dostarcza synchroniczne przebiegi prostokątne i trójkątne, których częstotliwość może być regulowana w bardzo prosty sposób poprzez zmianę prądu ładowania integratora. ■

Jos Verstraten



Praktyczny obwód generatora funkcyjnego (© 2023 Jos Verstraten)



Generatory

Rozwiązanie znajdziesz na www.elportal.pl/quizy

1. Pompa diodowa jest częścią generatora przebiegu:

- ☐ a. trójkątnego;
- ☐ b. piłkowszczytowego;
- ☐ c. schodkowego.

2. Generator przebiegu sinusoidalnego jest zbudowany wokół mostka:

- ☐ a. wiedeńskiego;
- ☐ b. Wiena;
- ☐ c. Wheatstone'a.

3. Generator przebiegu trójkątnego zbudowany jest z obwodów:

- ☐ a. integratora i komparatora;
- ☐ b. integratora i generatora przebiegu piłkowszczytowego;
- ☐ c. dwóch integratorów.

4. Główną wadą podstawowego generatora przebiegu prostokątnego jest:

- ☐ a. niesymetryczny czas trwania stanu wysokiego i niskiego;
- ☐ b. duża wrażliwość na zmianę temperatury.
- ☐ c. niska amplituda.

5. Prędkość narastania napięcia wyrażana w woltach na mikrosekundę to istotny parametr:

- ☐ a. generatorów przebiegu prostokątnego;
- ☐ b. mostków Wiena;
- ☐ c. wzmacniaczy operacyjnych w aplikacji generatora przebiegu prostokątnego.

6. Praktyczna aplikacja generatora przebiegu trójkątnego generuje też:

- ☐ a. przebieg piłkowszczytowy;
- ☐ b. przebieg prostokątny;
- ☐ c. szpilki dla innych układów.

7. Żarówka w generatorze sinusoidalnym służy do:

- ☐ a. automatycznej regulacji wzmocnienia;
- ☐ b. ograniczenia prądu w obwodzie sprzężenia zwrotnego;
- ☐ c. sygnalizacji obecności oscylacji.

8. Termistor w generatorze sinusoidalnym służy do:

- ☐ a. jego zachowanie jako rezystor sterowany napięciem w liniowym obszarze charakterystyki;
- ☐ b. automatycznej regulacji wzmocnienia;
- ☐ c. tworzenia skompensowanego termicznie obwodu rezonansowego RC.

9. W bardziej rozbudowanych generatorach stosuje się tranzystor FET ze względu na:

- ☐ a. jego zachowanie jako rezystor sterowany napięciem w liniowym obszarze charakterystyki;
- ☐ b. jego absolutnie liniową charakterystykę U-I;
- ☐ c. jego ekstremalnie wysoką impedancję wejściową.

10. W praktycznym generatorze symetrycznego sygnału piłkowszczytowego amplituda sygnału sterującego określa:

- ☐ a. częstotliwość;
- ☐ b. amplitudę sygnału wyjściowego;
- ☐ c. nachylenie zbocza przebiegu.



Migające diody LED i śliniący się inżynierowie (7)

Witaj! Jak miło znów Was widzieć. Prawdopodobnie nie będzie zaskoczeniem, jeśli usłyszycie, że sprawy w Maxfield Manor posuwają się naprzód w kontekście mojego zestawu piłeczek pingpongowych 12×12. Być może powinienem ostrzec, że w tej części cyklu będziemy przeskakiwać z tematu na temat ze zwinnością młodych kozic górskich, więc poświęć chwilę, aby upewnić się, że jesteś na to gotowy, zanim przejdziemy dalej.

Oślepienie przez światło

Zacznijmy od przypomnienia sobie, że każdy z naszych pikseli (piłeczek pingpongowych) zawiera trójkolorową diodę LED zwaną NeoPikselem. Co więcej, każdy NeoPixel zawiera czerwone, zielone i niebieskie subpiksele, z których każdy może być sterowany 256 różnymi poziomami od 0 (0x00 w systemie szesnastkowym = 0% = całkowicie wyłączony) do 255 (0xFF w systemie szesnastkowym = 100% = całkowicie włączony).

Nie wiem jak wy, ale ja zawsze, gdy pracuję nad takim projektem, to w mojej głowie pojawiają się pytania i zdaję sobie sprawę, że jest wiele rzeczy, których nie wiem. Na przykład, założmy, że chcemy podświetlić jeden z naszych pikseli w 100% na czerwono. Założmy teraz, że sąsiedni piksel zostanie podświetlony na 100% zielono i 100% niebiesko, co da kolor cyjanowy. Czy oznacza to, że cyjanowy piksel jest dwa razy jaśniejszy od czerwonego? Czy to z kolei oznacza, że tak naprawdę powinniśmy podświetlić zielone i niebieskie subpiksele na 50% (128 w systemie dziesiętnym = 0x80 w systemie szesnastkowym), aby uzyskać taką samą jasność jak czerwony piksel?

Podobnie, jeśli weźmiemy inny piksel, który sąsiaduje z naszym oryginalnym czerwonym pikselem i zaświecimy wszystkie trzy subpiksele, aby wygenerować biel, czy oznacza to, że naprawdę powinniśmy ustawić czerwone, zielone i niebieskie subpiksele na 33% (85 w systemie dziesiętnym = 0x55 w systemie szesnastkowym), aby uzyskać taką samą jasność jak czerwony piksel?

Pomyślałem, że dobrym pomysłem będzie przetestowanie tego rozwiązania. Jak być może pamiętasz z poprzedniego odcinka, nasze NeoPiksele są połączone ze sobą jako pojedynczy ciąg o numerach od 1 do 144 (jest też „ofiarny” piksel 0, którego używamy jako konwertera poziomu napięcia, ale nie odgrywa on żadnej roli w widocznej części wyświetlacza).

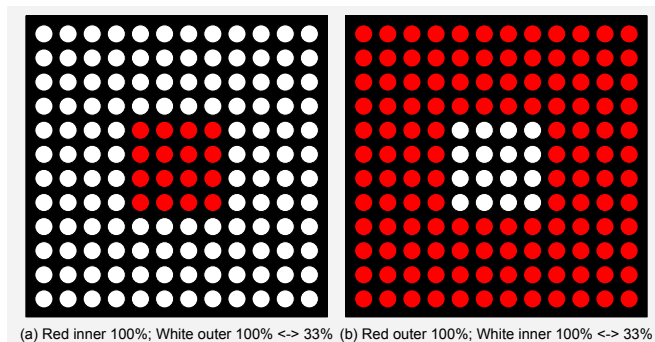


To właśnie nazywam tablicą pikseli... a koszula też jest fajna!

Uważamy, że nasza tablica jest macierzą (X,Y) z pozycjami X (poziowymi) i Y (pionowymi) ponumerowanymi od 0 do 11 oraz z pikselem (0,0) w lewym dolnym rogu. Stworzyliśmy również prostą funkcję o nazwie `GetNeoNum()`, która przyjmuje wartości X i Y jako argumenty i określa numer odpowiedniego NeoPikseli w łańcuchu.

Jako pierwszy test dla wyjaśnienia zagadki jasności/kontrastu stworzyłem prosty szkic (program), który wypełnia zewnętrzne piksele 100% bielą, a wewnętrzny kwadrat 4×4 pikseli 100% czerwienią (rysunek 1a). Następnie zmniejszamy otaczającą biel do 33%. Następnie szkic wypełnia zewnętrzne piksele w 100% kolorem czerwonym, a wewnętrzny kwadrat w 100% kolorem białym (rysunek 1b). Ponownie zmniejszamy poziom bieli do 33%. Następnie wykonujemy całą sekwencję w kółko. Jeśli chcesz, możesz pobrać ten szkic i zapoznać się z nim w wolnym czasie (jest to plik CB-Aug20-01.txt – dostępny do pobrania, wraz z pozostałymi trzema plikami omówionymi w tej kolumnie, na stronie PE we wrześniu 2020 r.). Dla waszej przyjemności stworzyłem również krótki film pokazujący to wszystko w akcji (<https://bit.ly/2OG9lgg>).

Na marginesie, być może pamiętasz, że twórcy biblioteki NeoPixel firmy Adafruit udostępnili nam dwa sposoby określania koloru pikseli. Założmy, że nasz ciąg nazywa się `Neos` i że chcemy ustawić kolor powiązany z pikselem 42, aby składowe koloru czerwonego, zielonego i niebieskiego wynosiły odpowiednio 128, 0 i 255 (byłyby



Rysunek 1. Pierwszy test kontrastu

to 0x80, 0x00 i 0xFF w systemie szesnastkowym). Jeśli chcemy, możemy określić kolory indywidualnie za pomocą instrukcji takiej jak:

```
Neos.setPixelColor(42, 128, 0, 255)
```

lub

```
Neos.setPixelColor(42, 0x80, 0x00, 0xFF)
```

Alternatywnie możemy określić wszystkie trzy kolory razem jako pojedynczą wartość liczbową, używając instrukcji takiej jak:

```
Neos.setPixelColor(42, 8388863)
```

gdzie liczba jest określona w systemie dziesiętnym lub:

```
Neos.setPixelColor(42, 0x8000FF)
```

z liczbą określoną w systemie szesnastkowym. Używanie pojedynczej wartości liczbowej w systemie szesnastkowym znacznie ułatwia nam życie, więc tak właśnie zrobimy. To wyjaśnia, dlaczego definicje kolorów w moich szkicach wyglądają mniej więcej tak:

```
#define COLOR_WHITE 0xFFFFFFF
```

```
#define COLOR_BLACK 0x0000000
```

```
#define COLOR_RED 0xFF00000
```

```
#define COLOR_GREEN 0x00FF000
```

```
#define COLOR_BLUE 0x0000FF0
```

Należy pamiętać, że dodanie znaków „U” (lub „u”) na końcu tych wartości informuje kompilator, aby traktował je jako unsigned (bez znaku) (więcej informacji na temat różnicy między wartościami signed i unsigned można znaleźć dalej, w rubryce „Sprytne porady i sztuczki...”). W większości przypadków nie będzie to konieczne, ale od czasu do czasu...

Następnie zacząłem się zastanawiać, czy mój środkowy kwadrat 4×4 nie jest zbyt mały (16 pikseli) w porównaniu do otaczającego go obszaru (128 pikseli), więc postanowiłem zwiększyć jego rozmiar do 6×6. Na szczęście, ze względu na sposób, w jaki stworzyłem swój oryginalny program, wszystko, co musiałem zrobić, aby wygenerować nową wersję, to zmienić dwie definicje: `XY_CENTER_MIN` z 4 na 3, a `XY_CENTER_MAX` z 7 na 8. Można pobrać ten szkic (CB-Aug20-02.txt). Stworzyłem też krótki film pokazujący to wszystko w akcji (<https://bit.ly/2OG9lgg>).

Osobiście uważam, że 33% bieli wygląda na nieco wyblakłe. Mój wniosek jest taki, że gdybyśmy próbowali użyć naszej tablicy do reprezentowania obrazów lub odtwarzania filmów w bardzo niskiej rozdzielczości, to prawdopodobnie zmienilibyśmy intensywność każdego subpiksela, aby wszystko zrównoważyć. Dla porównania, w przypadku naszych prostych efektów, prawdopodobnie chcemy, aby wszystko było tak jasne, jak to tylko możliwe, co – na szczęście – i tak jest moim punktem wyjścia.

Czujesz się przypadkowo?

Aby nadać naszym efektom „organiczny” charakter, użyjemy funkcji C/C++ `random()`, która przyjmuje dwa argumenty i określa losowy wynik w następujący sposób:

```
result = random(min, max);
```

Argument `min` jest opcjonalny. Jeśli ta wartość nie zostanie określona, zostanie ona domyślnie ustawiona na 0. Ta wartość jest również włączona, co oznacza, że może pojawić się w wyniku. Dla porównania, argument `max` jest wymagany i (z różnych powodów, których nie będziemy tutaj omawiać) jest również wyłączony, co oznacza, że nie pojawi się w wyniku. Innymi słowy, gdy ta funkcja zostanie wywołana, określi wynik między `min` a `(max - 1)`.

Trochę kropli

W przypadku naszej pierwszej serii efektów, wyobraźmy sobie, że nasza tablica leży płasko na podłożu, a od czasu do czasu z nieba spadają krople wody – wirtualne krople deszczu. Za każdym razem,

gdy wirtualna kropla wylądowuje na jednym z naszych pikseli (piłeczek pingpongowych), piksel ten zaświeci się na biało przez krótki czas.

Korzystne może być pobranie szkicu związanego z tym efektem (plik CB-Aug20-03.txt) i odniesienie się do niego w kontekście naszych rozważań.

Jak widać, deklarujemy dwie całkowite zmienne globalne o nazwach `OldX` i `OldY`, których będziemy używać do śledzenia współrzędnych (X,Y) ostatniego podświetlonego piksela. Inicjalizujemy obie te zmienne losowymi wartościami od 0 do 11 w naszej funkcji `setup()`.

Większość tego szkicu nie wymaga wyjaśnień. Interesująca część to ta, w której generujemy współrzędne (X,Y) związane z nową kroplą w następujący sposób (+1 celem osiągnięcia (domyślnie wyłączonej z przedziału losowania) wartości maksymalnej przez funkcję `random` (patrz wcześniejszy opis funkcji `random`):

```
do
{
    newX = random(MIN_X, (MAX_X + 1));
    newY = random(MIN_Y, (MAX_Y + 1));
} while (condition);
```

Pętla ta zostanie wykonana co najmniej jeden raz i będzie wykonywana tak długo, aż jej warunek (`condition`) zostanie spełniony (tj. zwróci wartość `true`). Ale jakiego warunku powinniśmy użyć? Jedną z alternatyw byłaby fraza:

```
((OldX == newX) && (OldY == newY))
```

W tym przypadku pętla będzie wykonywać się do momentu, gdy nowy piksel będzie różnił się od starego piksela, ale nic nie stoi na przeszkodzie, aby nowy piksel przylegał poziomo lub pionowo do starego piksela.

Osobiście wolę nieco większe zróżnicowanie, więc zdecydowałem się zamienić operator `&&` (logiczny AND) na operator `||` (logiczny OR) i użyć go:

```
((OldX == newX) || (OldY == newY))
```

Oznacza to, że nowy piksel nie może znajdować się w tym samym wierszu lub kolumnie co stary piksel.

Jedyną rzeczą, na którą należy zwrócić uwagę w tym szkicu, jest to, że używamy również funkcji `random()` do zmiany czasu trwania błysków związanych z każdą z kroplą i przerw między kroplami. Ponadto ostatnią rzeczą, którą robimy na końcu każdego cyklu, jest ustawienie wartości `OldX` i `OldY` odpowiednio na wartości `newX` i `newY`. Oznacza to, że piksel, który właśnie błysnął, będzie uważany za nasz stary piksel przy następnym przejściu przez pętlę. I tym razem przygotowałem wideo pokazujące to wszystko w akcji (patrz: <https://bit.ly/2OG9lgg>).

REKLAMA

OBWODY DRUKOWANE Produkcja, Projektowanie, Montaż			
Certyfikat Underwriters Laboratories UL 94V-0 E480140 TYPE 1	Platki jednostronne	Serie dowolne	Dokumentacja technologiczna
Zakład produkcyjny: 05-660 Warka ul. M. Ropielewskiej 17 tel. 22 781 63 95 22 761 95 40 fax. 22 781 63 95 w.13 www.elmax.waw.pl elmax@elmax.waw.pl	Platki dwustronne	Prototypy	Montaż elektroniczny
	Platki na podłożu aluminium	Maksymalny wymiar płytek 630 mm	Dokumentacja konstrukcyjna
			Ilości modelowe produkcyjne
	Aktywny kalkulator prototypów na stronie internetowej	Pokrycie Sn lub SnPb linie na życzenie	Płyty czolowe FR4
		Maski, opisy montażowe w różnych kolorach	Trawione szablony SMD
			Krótkie terminy
			Wykonania super ekspresowe

Czy wierzysz w kolorowe wróżki?

Kiedy byłem małym dzieckiem, miałem książkę, która opisywała, jak czerwone, żółte i niebieskie wróżki spędzały noc na przemalowywaniu wszystkiego, co wyjaśniało, dlaczego wszystkie kolory wyglądały tak jasno i świeżo każdego ranka.

To wszystko miało dla mnie sens, zwłaszcza że wiedziałem ze szkoły, że wróżki mogą generować inne kolory, pracując razem: czerwony i żółty tworzą pomarańczowy, czerwony i niebieski tworzą fioletowy, żółty i niebieski tworzą zielony i tak dalej.

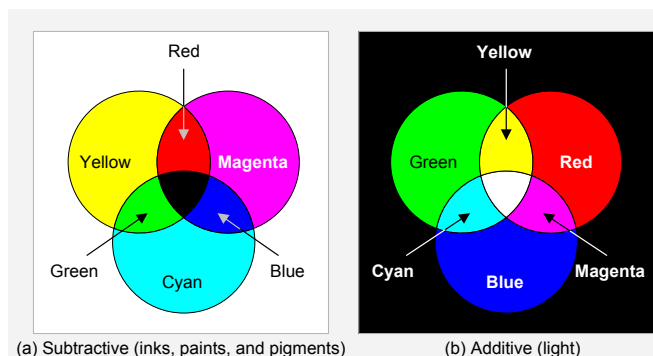
Później moja nauczycielka wyjaśniła, że czerwony (R), żółty (Y) i niebieski (B) to kolory podstawowe. Nie powiedziała jednak, że jest to tylko jeden z możliwych zestawów kolorów podstawowych, ponieważ termin „kolory podstawowe” odnosi się po prostu do niewielkiego zbioru trzech lub więcej kolorów, które można łączyć w celu utworzenia szeregu dodatkowych barw, odcieni i tonów. Na przykład drukarki używają atramentów cyjan (C), magenta (M) i żółty (Y) jako kolorów podstawowych. Co więcej, chociaż możliwe jest mieszanie cyjanu, magenty i żółtego w celu uzyskania czerni, drukarki zazwyczaj używają specjalnego czarnego pigmentu. Używamy litery K do reprezentowania tego pigmentu, ponieważ B można pomylić z niebieskim. To wyjaśnia, dlaczego używamy terminu CMYK, gdy mówimy o tuszu do drukarek.

Atramenty, farby i pigmenty mają charakter subtraktywny. Dla zobrazowania tej sytuacji można sobie wyobrazić białe tło oświetlone białym światłem. Białe tło odbija wszystkie częstotliwości białego światła. Kiedy zaczynamy dodawać nasze kolory, każdy kolor pochłania (odejmuje) niektóre częstotliwości z białego światła (**rysunek 2a**). Co więcej, gdy atramenty, farby lub pigmenty są ze sobą mieszane, każdy z nich pochłania inne częstotliwości. Zasadniczo widzimy to, co pozostało, czyli to, co nie zostało pochłonięte. Gdy wszystkie kolory są zmieszane razem, pochłaniają wszystkie częstotliwości, nie pozostawiając nam nic do zobaczenia (tj. czerni).

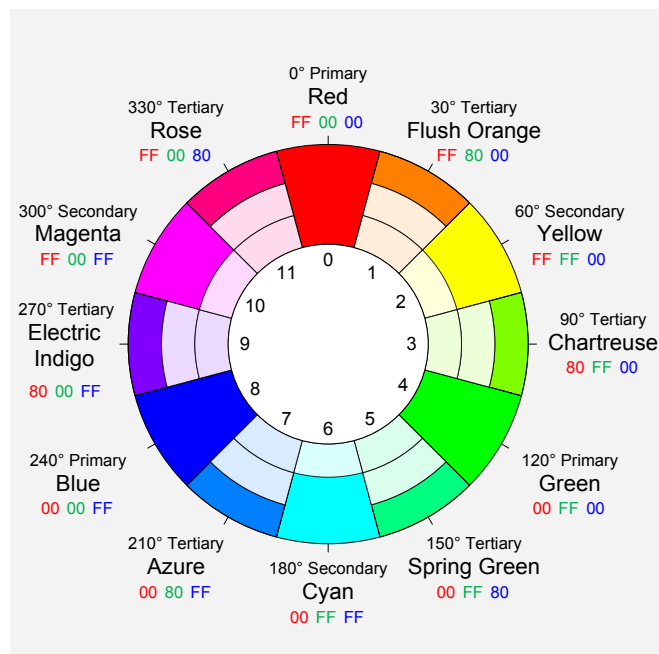
W przeciwieństwie do powyższych, światło jest addytywne. Zestaw podstawowych kolorów światła, z którymi jesteśmy najbardziej zaznajomieni, to czerwony (R), zielony (G) i niebieski (B). Dla zobrazowania tej sytuacji można sobie wyobrazić, że znajdujemy się w ciemnym pokoju i załączamy nasze podstawowe źródła światła. Czerwony i zielony tworzą żółty, czerwony i niebieski tworzą magentę, a zielony i niebieski tworzą cyjan. Po ich połączeniu czerwony, zielony i niebieski są postrzegane jako białe.

Krople Technicolor

Powodem naszego zainteresowania kolorami podstawowymi jest to, że następna wersja naszego programu kroplującego będzie oświetlanie naszych pikseli losowo wybranymi kolorami. Ponieważ mamy czerwone, zielone i niebieskie subpiksele, z których



Rysunek 2. Subtraktywne i addytywne kolory podstawowe



Rysunek 3. Koło kolorów, którego będziemy używać w naszych eksperymentach

każdy może mieć przypisane wartości od 0 do 255, mamy potencjał $256 \times 256 \times 256 = 16\,777\,216$ (2^{24}) różnych barw, odcieni i tonów.

Oczywiście niektóre z nich będą przyćmione i nijakie, ale – z naszych eksperymentów z kontrastem na początku tej kolumny – wiemy, że chcemy, aby nasze kolory były tak jasne i błyszczące, jak to tylko możliwe. Jednym ze sposobów przedstawienia kolorów jest użycie koła kolorów (**rysunek 3**).

Jak widzimy, nasze trzy kolory podstawowe (czerwony, zielony i niebieski) można łączyć, tworząc trzy kolory drugorzędne (żółty, cyjan i magenta) oraz sześć kolorów trzeciorzędnych (o bardziej egzotycznych nazwach: rumieniec pomarańczowy, chartreuse, wiosenna zieleń, lazur, elektryczne indygo i róż).

Tych 12 kolorów użyjemy w naszym szkicu, który możesz pobrać, jeśli chcesz (CB-Aug20-04.txt). Jeśli porównasz ten szkic z naszą poprzednią wersją, zauważysz, że jedyną różnicą jest to, że – w przeciwieństwie do oświetlania naszych pikseli kolorem białym – dla każdej wirtualnej kropli używamy teraz funkcji `random()`, aby wybrać jedną z 12 opcji kolorów. Po raz kolejny stworzyłem film pokazujący to wszystko w akcji (<https://bit.ly/2OG9lgg>).

Następnym razem

Muszę przyznać, że podobają mi się kolorowe krople przedstawione w naszym ostatnim szkicu, ale zwykle włączanie i wyłączanie naszych pikseli nie jest zbyt subtelne i brakuje mu czegoś. Jednym ze sposobów, w jaki możemy dodać odrobinę wyrafinowania, jest rozjaśnianie wybranego koloru, utrzymywanie go na stałym poziomie, a następnie stopniowe acz równomiernie jego wygaszenie i to właśnie zrobimy w następnej części tego cyklu. Do tego czasu życzę miłej zabawy! ■



Komentarze lub pytania?
Napisz do Maxa na adres:
max@CliveMaxfield.com

Sprytne porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania



W poprzednim odcinku z moimi poradami i wskazówkami zauważyliśmy, że kiedy ludzie po raz pierwszy zapoznają się z językami programowania C/C++, jednym z pierwszych typów danych, z którymi się zapoznają, jest `int`, co oznacza „liczbę całkowitą”. Zauważyliśmy również, że możemy zadeklarować nasze liczby całkowite jako ze znakiem lub bez znaku:

```
signed    int John;  
unsigned  int Jane;
```

Liczba `int` ze znakiem może być używana do reprezentowania zarówno wartości dodatnich, jak i ujemnych, podczas gdy liczba `int` bez znaku może być używana do reprezentowania tylko wartości dodatnich. Co ciekawe, standardy C/C++ nie definiują jednoznacznie rozmiaru `int`. Mówią jedynie, że `int` powinien mieć rozmiar co najmniej dwóch bajtów (16 bitów). W przypadku Arduino Uno z 8-bitową magistralą danych, rozmiar `int` rzeczywiście wynosi dwa bajty; w innych komputerach mogą to być dwa bajty, cztery bajty lub więcej.

Zmienna 2-bajtowa (16-bitowa) może być używana do reprezentowania $2^{16}=65\,536$ różnych ciągów zer i jedynek. W przypadku `int` ze znakiem ciągi te mogą być używane do reprezentowania zarówno wartości ujemnych, jak i dodatnich w zakresie od $-32\,768$ do $+32\,767$. Dla porównania, w przypadku `int` bez znaku, ciągi te mogą być używane do reprezentowania tylko wartości dodatnich w zakresie od 0 do $65\,535$.

Hej, krótkie rzeczy!

Istnieją również krótkie (`short`) i długie (`long`) odmiany `int`:

```
signed    short int Fred;  
signed    long  int Bert;  
unsigned  short int Mary;  
unsigned  long  int Beth;
```

W ten sam sposób, w jaki zakładamy, że liczba taka jak 42 jest dodatnia bez wyraźnego zapisywania jej jako +42, kompilator założy, że niewykwalifikowany `int` jest podpisany, co oznacza, że możemy pominąć podpisaną (`signed`) część deklaracji, jeśli chcemy:

```
int        John;  
short      int Fred;  
long       int Bert;
```

Podobnie, jeśli zadeklarujemy coś jako typ `short` lub `long`, kompilator z radością założy, że mówimy o wartościach całkowitych, więc możemy również pominąć część `int`, jeśli chcemy:

```
int        John;  
short      Fred;  
long       Bert;
```

Specyfikacje C/C++ określają, że minimalny rozmiar `short int` (krótkiego) to dwa bajty (16 bitów), podczas gdy minimalny rozmiar `long int` (długiego) to cztery bajty (32 bity).

Ustalanie szerokości

Może to być nieco niepokojące, gdy po raz pierwszy odkrywa się, że rozmiar różnych typów liczb całkowitych może się różnić w zależności od maszyny. Powodem tego jest umożliwienie kompilatorowi wykorzystania wiedzy o wewnętrznej architekturze docelowego procesora w celu użycia optymalnych rozmiarów danych dla tej konkretnej maszyny.

Sytuacja nie wygląda źle, jeśli jesteś hobbystą piszącym kod tylko dla jednego procesora. Jednak w zależności od rozmiaru liczb, które próbujesz reprezentować, mogą pojawić się problemy, jeśli spróbujesz przenieść swój kod z maszyny obsługującej czterobajtowe liczby całkowite na taką, której liczby całkowite mają na przykład tylko dwa bajty. Aby obejść ten problem, możemy użyć typów liczb całkowitych o stałej szerokości, które są dostępne od wersji v11 języka C++; na przykład:

```
int8_t Cuthbert; // Jednoznacznie 8bitów (jeden bajt)  
uint8_t Clarence; // Niewątpliwie 8bitów (jeden bajt)
```

Istnieją również 16-bitowe (`int16_t` i `uint16_t`), 32-bitowe (`int32_t` i `uint32_t`) i 64-bitowe (`int64_t` i `uint64_t`) odpowiedniki tych typów o stałej szerokości.

Niebezpieczna rzecz

„Niewielka wiedza jest niebezpieczna”, jak mówi stare powiedzenie. Załóżmy, że piszesz program dla Arduino Uno i wiesz, że licznik używany do kontrolowania pętli `for()` będzie liczył tylko do 100. Czy w takim przypadku lepiej byłoby użyć `int8_t` (8 bitów) zamiast `int` (16 bitów), oszczędzając w ten sposób bajt (8 bitów) w pamięci? Na przykład:

```
for (int8_t i = 0; i < 100; i++)  
{  
    // Zrób tutaj coś sprytnego  
}
```

Ogólnie rzecz biorąc, odpowiedź brzmi „nie”. Problem polega na tym, że użycie mniejszej zmiennej w celu zaoszczędzenia pamięci może negatywnie wpłynąć na wydajność programu. O ile nie jesteś dogłębnie zaznajomiony z działaniem kompilatora i wewnętrzną architekturą procesora, najlepiej jest użyć `int` i pozwolić kompilatorowi wykorzystać swoją szczegółową wiedzę na temat architektury procesora do podjęcia optymalnych decyzji.

Tak czy inaczej, jeśli masz specyficzne wymagania, to jak najbardziej używaj odpowiednich typów o stałej szerokości. Na przykład w przypadku programów dla mojego projektu tablicy piłeczek ping-pongowych 12×12 można zauważyć, że deklaruję wszystkie wartości kolorów jako typu `uint32_t` (tj. 32-bitowe liczby całkowite, bez znaku). W przypadku Arduino Uno odpowiednikiem byłby `unsigned long`, ale w przypadku tych deklaracji wolę być jednoznaczny.

Ale czekaj, to nie wszystko!

Pochodzę z czasów, gdy każda lokalizacja pamięci i każdy cykl zegara były uważane za cenne towary. W czasopiśmie komputerowych można było znaleźć kolumny poświęcone sprytnym sztuczkom, takim jak zamiana górnych i dolnych czterech bitów w bajcie w jak najmniejszej liczbie cykli zegara, użycie jak najmniejszej liczby instrukcji lub zużycie jak najmniejszej ilości pamięci.

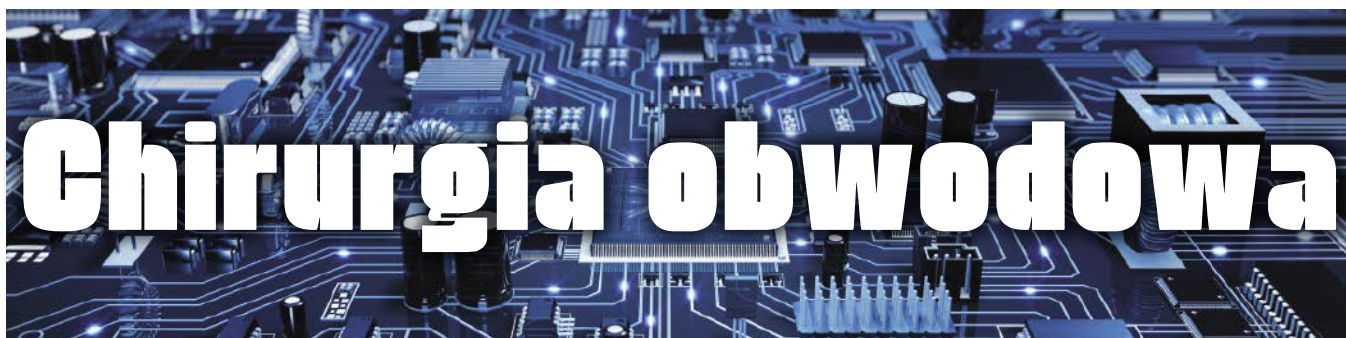
Niektórzy projektanci systemów wbudowanych przesuwają granice procesorów, których mają używać. W takim przypadku mogą zdecydować się na użycie typów o minimalnej szerokości i najszybszej minimalnej szerokości, które są obsługiwane przez C od 1999 roku (znane jako kompilatory zgodne z C99); na przykład:

```
int8_t Bob; // Podpisana stała szerokość  
           // 8-bitowa liczba całkowita  
int_least8_t Dan; // Minimalna podpisana szerokość  
              // 8-bitowa liczba całkowita  
int_fast8_t Sam; // Najszybsza minimalna szerokość  
                // 8-bitowa liczba całkowita  
                // ze znakiem
```

Oczywiście istnieją niepodpisane wersje tych typów, wraz z wariantami 16-, 32- i 64-bitowymi, ale radziłbym wstrzymać się z używaniem tych małych sztuczek, dopóki nie będziesz w 100% pewien, że jesteś świadomy wszelkich implikacji. ■

Clive „Max” Maxfield

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, wrzesień 2020 (www.epemag3.com)



Chirurgia obwodowa

Transformatory i LTspice, część 2

Kontynuujemy tematykę z zeszłego miesiąca, przyglądając się podstawom transformatorów i niektórym aspektom symulowania obwodów z ich użyciem w LTspice. Zaczniemy od szybkiego podsumowania, a następnie przyjrzymy się transformatorom z wieloma uzwojeniami wtórnymi. Następnie przyjrzymy się, w jaki sposób transformatory są używane z diodami prostowniczymi w celu uzyskania niskiego napięcia prądu stałego z napięcia sieci 230 V AC.

Przegląd transformatorów

W zeszłym miesiącu wyjaśniliśmy, że transformator składa się z dwóch lub więcej cewek lub uzwojeń znajdujących się blisko siebie, co umożliwia przesyłanie energii elektrycznej w postaci prądu przemiennego z jednego obwodu do drugiego, bez konieczności stosowania połączenia przewodzącego prąd elektryczny. Kluczowe właściwości transformatorów polegają na tym, że zapewniają one izolację galwaniczną między obwodami, mogą zmieniać poziomy napięcia i zmieniać efektywną impedancję obciążenia podłączonego za pośrednictwem transformatora.

Zależność między napięciami i prądami w transformatorze zależy od stosunku liczby zwojów między uzwojeniami wejściowym (pierwotnym) i wyjściowym (wtórnym). Napięcie wtórne jest równe napięciu pierwotnemu pomnożonemu przez stosunek zwojów strony wtórnej do pierwotnej. Prąd wtórny jest dzielony przez ten sam stosunek. W zależności od przekładni napięcie wtórne może zostać podniesione (wydajność prądowa uzwojenia wtórnego będzie mniejsza w porównaniu z prądem płynącym przez uzwojenie pierwotne) lub obniżone (wydajność prądowa uzwojenia wtórnego będzie większa w porównaniu z prądem płynącym przez uzwojenie pierwotne). Moc wyjściowa

i wyjściowa są równe dla idealnego transformatora, ale w rzeczywistych podzespołach występują straty i nie osiąga się 100% przeniesienia mocy. Charakter źródeł tych strat jest złożony (np. wielkość sprzężenia strumienia, właściwości materiałów i struktury rdzenia), zatem chociaż podstawowa zasada działania transformatora jako elementu obwodu jest prosta, bardzo dokładny i szczegółowy projekt, analiza i symulacja obwodu transformatora mogą stanowić wyzwanie.

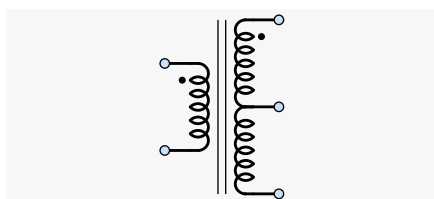
W zeszłym miesiącu przyjrzelśmy się także podstawowym symulacjom transformatorów w LTspice. Transformatory są prawdopodobnie najtrudniejszym z podstawowych elementów pasywnych stosowanych w SPICE – nie ma elementu transformatorowego – trzeba go utworzyć z wielu cewek i użyć wspólnego współczynnika indukcji, aby połączyć je ze sobą. Współczynnik sprzężenia należy umieścić w LTspice jako element tekstowy SPICE (np. K1 L1 L2 1 w celu połączenia cewek L1 i L2 z idealnym współczynnikiem sprzężenia wynoszącym 1). Wartości cewek są ustawione w taki sposób, że pierwiastek kwadratowy ich stosunku $\sqrt{L_1/L_2}$ jest równy przekładni uzwojeń. Rzeczywistą wartość indukcyjności można uzyskać z noty katalogowej elementu, zmierzyć lub oszacować tak, aby była zgodna z prądem roboczym przy stosowanej częstotliwości.

Wiele uzwojeń

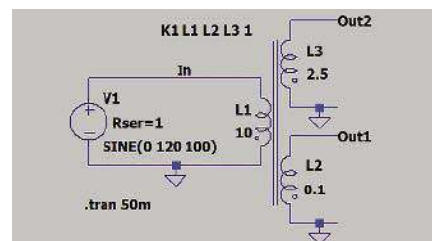
Do tej pory przyglądaliśmy się tylko transformatorom z jednym uzwojeniem pierwotnym i jednym wtórnym. Tymczasem często występuje w nich wiele uzwojeń wtórnych, a niektóre transformatory mają wiele uzwojeń

pierwotnych; zazwyczaj można je przystosować do pracy w różnych warunkach poprzez wybór pierwotnego źródła – chociaż należy przy tym zachować szczególną ostrożność, szczególnie w przypadku wyższych napięć, np. napięcia sieci 230 V AC. Jak zawsze należy zapoznać się z danymi producenta. Często zdarza się, że dwa uzwojenia pierwotne lub wtórne są połączone ze sobą w ramach konstrukcji transformatora, a nie są całkowicie oddzielne. Nazywa się to „uzwojeniem z odczepem centralnym”, a połączenie środkowe to odczep środkowy. Oczywiście odczepy nie muszą być realizowane w połowie długości uzwojenia, a niektóre transformatory mają wiele odczepów. Symbol transformatora z uzwojeniem wtórnym z odczepem centralnym pokazano na **rysunku 1**.

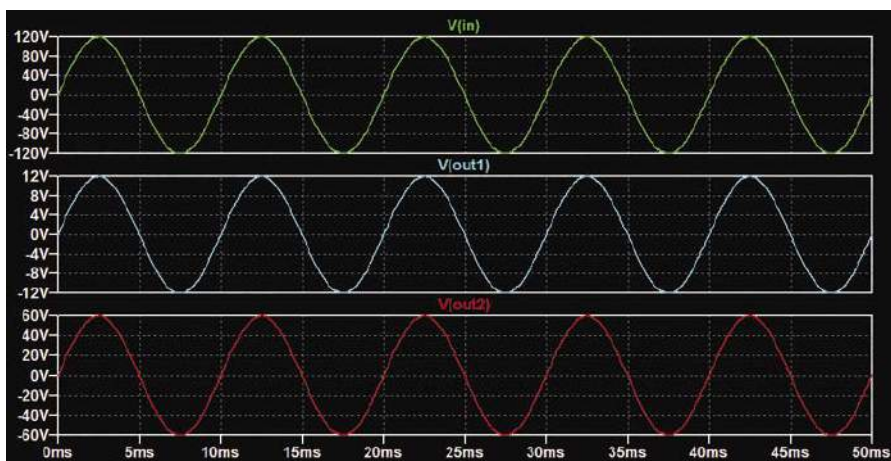
Wiele niezależnych uzwojeń wtórnych można używać oddzielnie lub łączyć szeregowo w celu połączenia napięć wtórnych. Jeżeli są one połączone, należy wziąć pod uwagę fazę połączeń – napięcia w fazie sumują się, a napięcia w przeciwfazie odejmują. Równoległe połączenie uzwojeń wtórnych może być również możliwe, ale tylko wtedy, gdy ich napięcia są równe, a transformator został określony jako odpowiedni dla tego rodzaju



Rysunek 1. Symbol transformatora z uzwojeniem wtórnym z odczepem w środku



Rysunek 2. Schemat LTspice transformatora z dwoma uzwojeniami wtórnymi



Rysunek 3. Wyniki symulacji dla obwodu z rysunku 2

konfiguracji. Nawet niewielka różnica napięcia może spowodować przepływ szkodliwie wysokich prądów w uzwojeniach wtórnych o niskim oporze uzwojeń.

Rysunek 2 pokazuje przykład transformatora z dwoma uzwojeniami wtórnymi w LTspice. Pamiętaj, że same cewki indukcyjne (tutaj L1, L2 i L3) nie tworzą transformatora, niezależnie od tego, w jaki sposób zostaną narysowane na schemacie. Również linie rdzenia pomiędzy uzwojeniami są tworzone za pomocą funkcji rysowania i mają charakter wyłącznie graficzny – nie mają wpływu na symulację. Jak wspomniano wcześniej, współczynnik sprzężenia (element K1) łączy cewki indukcyjne, tworząc transformator. Dla uproszczenia stosuje się współczynnik sprzężenia równy 1. Zwoje, a co za tym idzie stosunek napięcia, są równe pierwiastkowi kwadratowemu stosunku indukcyjności. W tym przykładzie stosunek indukcyjności uzwojenia pierwotnego (L1) do wtórnego (L2) wynosi 10:0,1 lub 100:1, dlatego stosunek napięcia (zwojów) jest pierwiastkiem kwadratowym z tego (10:1), więc przy napięciu 120 V wejście pokazane na Out1 będzie wynosić 12 V. Podobnie Out2 (z wtórnego L3) będzie wynosić 60 V (10:2,5 = współczynnik indukcyjności 4:1, napięcie i współczynnik zwojów 2:1). Na rysunku 3 przedstawiono wyniki symulacji dla obwodu z rysunku 2 – napięcia są zgodne z obliczeniami. Wszystkie przebiegi są w fazie, ponieważ wszystkie uzwojenia są w fazie – w tym przypadku z uziemionym końcem uzwojenia z „kropką fazową”.

Na rysunku 4 przedstawiono schemat LTspice z dwiema kopiami transformatora z rysunku 2. W jednym obwodzie uzwojenia wtórne są połączone w fazie (z wyjściem OutA), a w drugim są one przeciwfazowe (z wyjściem OutS). Napięcia z obu uzwojeń sumują się w pierwszym przypadku – OutA wynosi 60+12=72 V, a w drugim

odejmują – OutS wynosi 60–12=48 V. Kropki fazowe na schemacie pokazują fazę każdego uzwojenia. Wyniki symulacji dla obwodu z rysunku 4 pokazano na rysunku 5, potwierdzając właśnie obliczone napięcia. OutS jest w fazie niezgodnej z OutA, ponieważ uzwojenie 60 V jest podłączone w przeciwnej fazie do przewodów OutA. Odejmowanie napięć jest ważnym aspektem teorii transformatora, ale prawdopodobnie nie będzie zbyt często stosowane w rzeczywistych obwodach – lepiej wybrać transformator z uzwojeniem, które bezpośrednio zapewnia odpowiednie parametry elektryczne w tym oczekiwanym napięciu na uzwojeniu wtórnym. Dodawanie napięć następuje w uzwojeniach z odczepem centralnym i wieloodcepowym (lub w tak połączonych oddzielnych uzwojeniach).

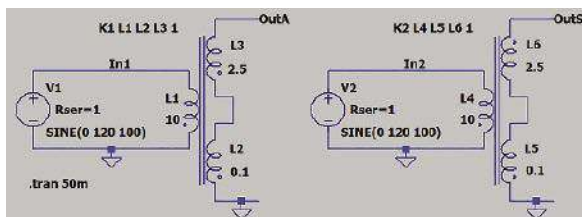
Transformatory do zasilaczy

Transformatory są szeroko stosowane bezpośrednio do obniżania napięcia sieciowego do niższych napięć celem

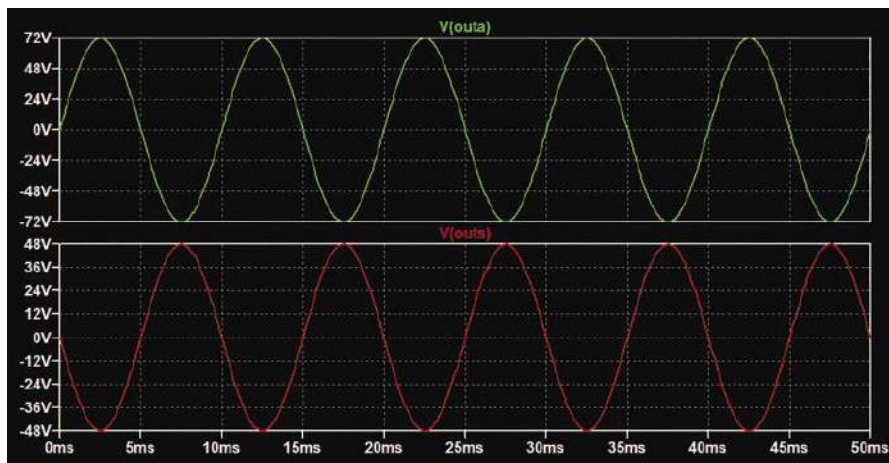
ich wykorzystania w układach elektronicznych – jednak są już mniej powszechne niż kiedyś, ponieważ lepszą wydajność oraz mniejsze gabaryty, wagę i koszt zapewniają zasilacze impulsowe typu direct-off-line (SMPS). SMPS mogą nadal wykorzystywać transformatory, ale takie, które pracują przy znacznie wyższych częstotliwościach niż częstotliwość sieci, co ułatwia użycie transformatorów o znacznie mniejszych gabarytach (lub po prostu cewki indukcyjnej). Jednak w niektórych sytuacjach transformatory pracujące na częstotliwości sieciowej w połączeniu ze stabilizatorami liniowymi mogą okazać się lepsze – szczególnie jeśli nie można uniknąć sytuacji, w której szum przełączania SMPS znajduje się w tym samym zakresie częstotliwości, co przetwarzane w danym układzie elektronicznym sygnały.

Rysunek 6 przedstawia początek typowego obwodu zasilacza sieciowego – transformator i mostek prostowniczy. W zasilaczu byłby on zwykle podłączony do kondensatora wygładzającego i stabilizatora napięcia, ale nie patrzmy tutaj na pełne obwody zasilania – tylko na sposób wykorzystania transformatora.

Źródło napięcia V1 reprezentuje sieć 230 V AC. Napięcie wynosi 325 V i odpowiada (szczytowemu) napięciu sieciowemu w Wielkiej Brytanii. Napięcie zwykle podawane dla sieci elektrycznej to napięcie skuteczne (RMS), które w Wielkiej Brytanii jest bardziej znanym napięciem 230 V. Wartość 230 V jest równa napięciu prądu stałego, które spowodowałoby



Rysunek 4. Przykłady połączonych ze sobą uzwojeń wtórnych – zwróć uwagę na kropki fazowe. Uzwojenia są w fazie dla OutA i przeciwfazie dla OutS



Rysunek 5. Wyniki symulacji dla obwodu z rysunku 4

wydzielenie takiej samej mocy na rezystorze, jak sygnał prądu przemiennego. Musimy jednak określić napięcie szczytowe dla źródła fali sinusoidalnej w SPICE. Napięcie szczytowe sygnału prądu przemiennego jest większe niż wartość skuteczna, w szczególności w przypadku czystej fali sinusoidalnej napięcie szczytowe wynosi $\sqrt{2}$ (1,4142) razy wartość skuteczna (dla sieci brytyjskiej $230 V_{sk} \times 1,4142 = 325,3 V_{szczyt}$). Przełożenie zwojów transformatora można zastosować zarówno przy napięciu skutecznym, jak i napięciu szczytowym. Na przykład stosunek zwojów 10:1 da $23 V_{sk}$ na wyjściu dla $230 V_{sk}$ na wejściu lub $32,5 V_{szczyt}$ na wyjściu dla $325 V_{szczyt}$ na wejściu.

Transformator w obwodzie jest skonfigurowany tak, aby wytwarzać napięcie szczytowe 12 V na uzwojeniu wtórnym – stosunek napięcia (zwojów) wynoszący $325:12=27,08:1$ wymaga współczynnika indukcyjności 733,5:1. Zastosowano tutaj stosunek 10:0,01363, tak że uzwojenie pierwotne (przy 10 H) mieści się w zakresie właściwym dla transformatora sieciowego, chociaż nie reprezentuje konkretnego rzeczywistego elementu. Rzeczywista indukcyjność transformatora nie ma większego znaczenia dla podstawowej symulacji obwodów zasilania (w przeciwieństwie do innych sytuacji, takich jak obwody w.cz.), ale musimy wybrać wartość, aby skonfigurować symulację LTSpice.

Testowanie

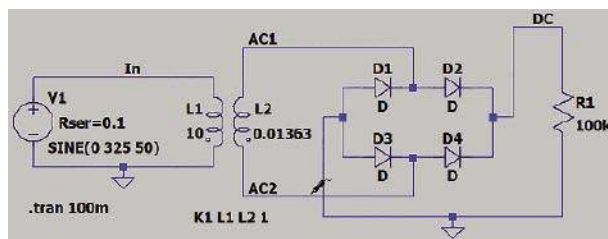
Wyniki symulacji obwodu z rysunku 6 pokazano na **rysunku 7**. Aby wyświetlić napięcie prądu przemiennego na uzwojeniu wtórnym transformatora, musimy określić jeden z zacisków transformatora jako punkt odniesienia do pomiaru napięcia na drugim. W przeciwnym razie LTSpice używa masy jako odniesienia do definiowania napięć wyświetlanych przebiegów. Punkt odniesienia można ustawić, klikając prawym przyciskiem myszy na przewód, który ma być użyty jako odniesienie, i wybierając „Zaznacz odniesienie” z menu podręcznego. Spowoduje to umieszczenie na schemacie czarnego symbolu sondy (rysunek 6). Następnie należy kliknąć na drugi przewód, dla którego ma zostać wyświetlony przebieg – po najechaniu myszką na przewód wyświetli się czerwona sonda.

Przebieg napięcia pomiędzy węzłami X i Y (X mierzony względem Y) będzie w oknie wykresu zatytułowany V(X,Y), a nie V(x) w celu wyświetlenia przebiegu napięcia w węzle X względem masy (patrz V(AC1,AC2) na rysunku 7). Czerwony wskaźnik i opcja „Zaznacz odniesienie” pojawiają się dopiero po przeprowadzeniu symulacji i otwartym oknie wykresu. Naciśnięcie klawisza Esc usuwa czarną sondę referencyjną ze schematu i wznawia normalne sondowanie

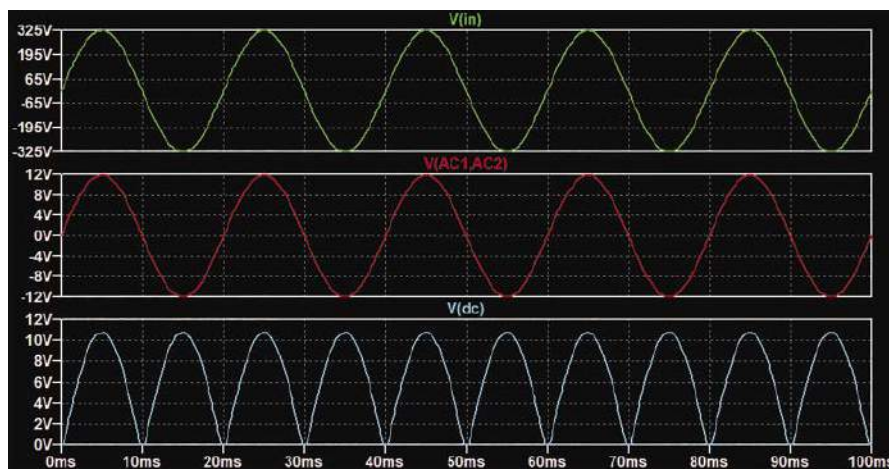
względem masy w celu wyświetlenia przebiegów.

Wyniki na rysunku 7 pokazują, że napięcie wejściowe o wartości szczytowej 325 V skutkuje wartością szczytową 12 V w obwodzie wtórnym, zgodnie z oczekiwaniami na podstawie omówionej powyżej przekładni cewek indukcyjnych. Sygnał wyjściowy prostownika pełnookresowego (D1 do D4), V(dc), to seria jednokierunkowych impulsów odpowiadających

przewodzenia diod prostowniczych (w tym przypadku około 0,6 V). W tym przykładzie diody są po prostu domyślnymi diodami



Rysunek 6. Transformator i prostownik mostkowy do stosowania w zasilaczu sieciowym

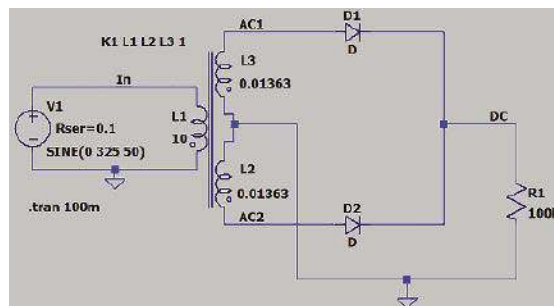


Rysunek 7. Wyniki symulacji z obwodu z rysunku 6

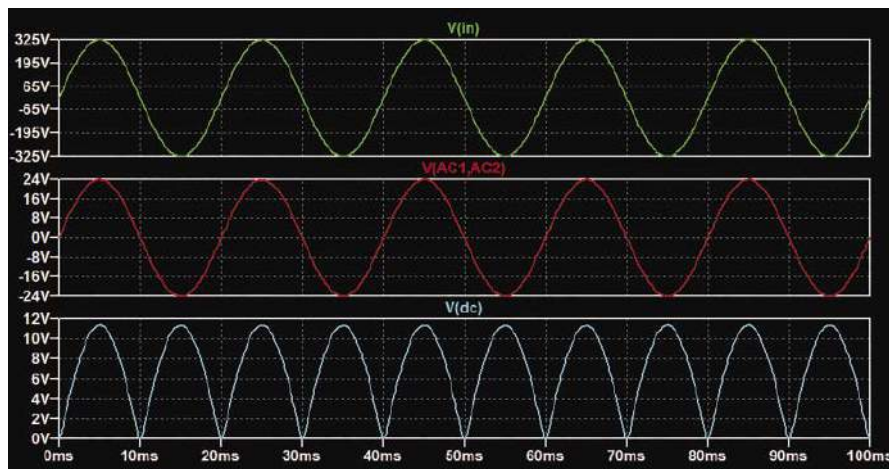
półcykłem przebiegu prądu przemiennego. Należy zauważyć, że jest to prostowanie pełnookresowe, ponieważ oba półcykle wejścia prądu przemiennego przyczyniają się do wyjścia prądu stałego.

Spadki napięcia

Napięcie szczytowe wynosi około 10,8 V, a nie 12 V szczytowe z transformatora. Różnica wynika ze spadku napięcia w kierunku



Rysunek 8. Prostownik pełnookresowy wykorzystujący transformator z odczepem środkowym



Rysunek 9. Wyniki symulacji z obwodu z rysunku 8

LTspice, a nie prawdziwymi modelami urządzeń. Szczyt impulsu prądu stałego wynosi $12 - (2 \times 0,6) = 10,8$ V. W konstrukcji z pełnym zasilaniem kondensator wygładzający byłby podłączony do wyjścia prostownika w celu wytworzenia niemal ciągłego napięcia prądu stałego, które zwykle byłoby następnie wprowadzane do obwodu stabilizatora. Nie uwzględniliśmy tutaj kondensatora, aby było jaśniejsze, w jaki sposób przebieg prądu przemiennego z transformatora jest kierowany przez diody prostownicze do wyjścia prądu stałego.

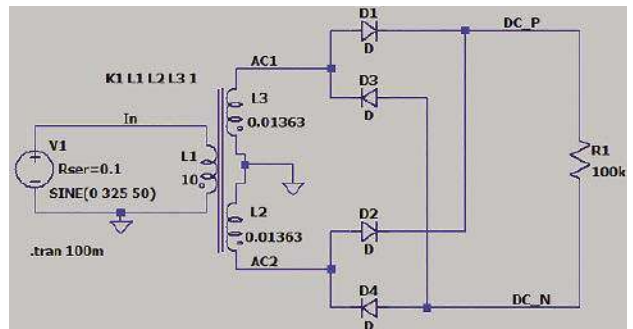
Obwód pokazany na rysunku 6 wykorzystuje rezystor o dużej wartości rezystancji jako obciążenie, aby zakończyć obwód i zapewnić pewien prąd diody. Jest to wyidealizowana i uproszczona sytuacja w porównaniu z rzeczywistym zasilaczem z kondensatorem wygładzającym nie uwzględniamy tu również charakteru obciążenia i wpływu ewentualnego zastosowania obciążenia reaktancyjnego i w miejsce rezystancyjnego. Obwód nadaje się do pokazania podstawowych zasad działania transformatora i prostownika, ale nie obejmuje wszystkich szczegółów pełnego zasilania. U_{gr} i wygładzone napięcie będzie niższe niż jednokierunkowy szczyt prądu stałego (10,8 V w powyższym przykładzie), a stabilizator będzie potrzebował pewnej „nadwyżki” (napięcia spadku), więc maksymalne regulowane napięcie wyjściowe będzie niższe niż obserwowany szczyt w kształcie fali.

Odczep środkowy

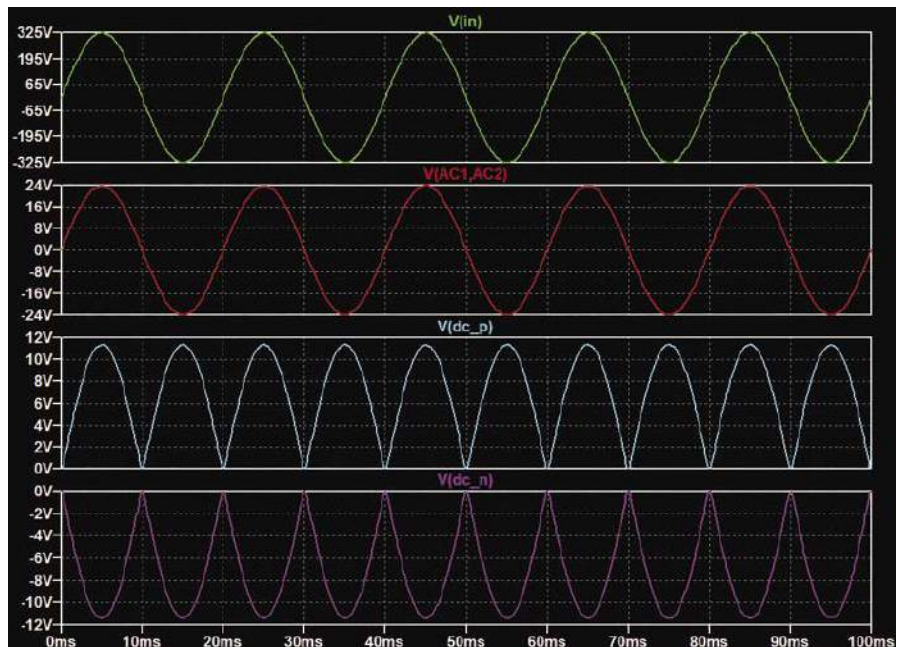
Rysunek 8 przedstawia alternatywny układ transformatora i prostownika. Wykorzystuje się w nim uzwojenie wtórne z odczepem środkowym i prostownik dwudiodowy. Obydwa uzwojenia są takie same jak na rysunku 6, zatem całkowite napięcie na uzwojeniu z odczepem środkowym wynosi 24 V. Każda połówka dostarcza napięcie 12 V do jednej z diod prostowniczych, więc napięcie wyjściowe prądu stałego, podobnie jak w obwodzie na rysunku 6, wynosi około 12 V. Dokładniej, szczyt impulsów jednokierunkowych wynosi 11,4 V (12 V – 0,6 V), czyli o jeden spadek na diodzie więcej niż napięcie z obwodu na rysunku 6. Jednakże obwód ten wymaga dwóch uzwojeń i podwójnego napięcia wyjściowego prądu przemiennego, aby uzyskać to samo podobne napięcie prądu stałego jak obwód na rysunku 6. Obwód był popularny w przeszłości, gdy diody prostownicze były stosunkowo drogie (szczególnie w epoce prostowników próżniowych), ale obecnie jest rzadziej używany ze względu na konieczność stosowania dodatkowego uzwojenia transformatora.

Dzielone zasilanie

Chociaż poprzedni obwód nie jest szczególnie przydatny, istnieje nieco inny układ z odczepem środkowym. Jeśli połączymy transformator z odczepem środkowym z prostownikiem mostkowym, możemy wytworzyć



Rysunek 10. Prostownik pełnookresowy z transformatorem z odczepem centralnym zapewniającym rozdzielone zasilanie prądem stałym



Rysunek 11. Wyniki symulacji z obwodu z rysunku 10

zarówno dodatni, jak i ujemny sygnał wyjściowy prądu stałego w odniesieniu do centralnego odczepu jako punktu odniesienia (masy). Obwód ten wraz z dwoma kondensatorami wygładzającymi i dwoma stabilizatorami stanowi podstawę rozdzielonego zasilania (dodatni i ujemny prąd stały), które jest bardzo powszechnie stosowane w obwodach przetwarzania sygnału analogowego, takich jak obwody zasilania wzmacniaczy operacyjnych.

Rysunek 10 przedstawia schemat LTspice pełnookresowego prostownika mostkowego z transformatorem z odczepem środkowym. Wyniki symulacji pokazano na rysunku 11. Uzwojenia transformatora są takie same jak w poprzednim przykładzie, więc oba zapewniają szczytowe napięcie prądu przemiennego 12 V. Podobnie jak w obwodzie na rysunku 8, mamy szczytowe napięcie prądu przemiennego o wartości 24 V na dwóch uzwojeniach (V(AC1, AC2)). Odczep środkowy jest podłączony do masy, a diody prostownicze dostarczają w stosunku

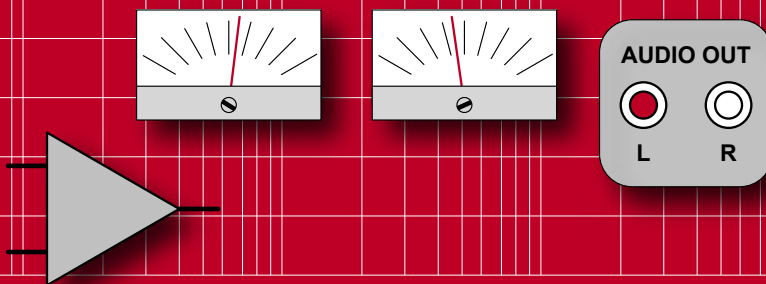
do tego dwa zestawy jednokierunkowych impulsów – dodatni w DC_P i ujemny w DC_N (rysunek 11). Podobnie jak w obwodzie na rysunku 8, wartości szczytowe wynoszą spadek napięcia przewodzenia o jedną diodę poniżej szczytowego napięcia prądu przemiennego – w tym przykładzie ponownie 11,4 V.

W celu podłączenia do sieci należy wybrać transformator przeznaczony do stosowania przy napięciu i częstotliwości sieci elektrycznej w kraju, w którym jest używany. Podłączenie do sieci niewłaściwego transformatora może skutkować jego zniszczeniem. Napięcie sieciowe jest potencjalnie śmiertelne i należy podjąć odpowiednie środki bezpieczeństwa w przypadku konstruowania obwodów przeznaczonych do podłączenia do sieci 230 V AC. ■

Ian Bell

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, lipiec 2021 (www.epemag3.com)

AUDIO OUT



Przedwzmacniacz mikrofonowy (do wokodera), część 1

Zamierzam rozpocząć serię najwyższej jakości modułów do pracy studyjnej. Wkrótce zrobię analogowy wokoder, aby upamiętnić dźwięk piosenek „Mr Blue Sky” zespołu ELO oraz „Radio Ga Ga” zespołu Queen. Pierwszą rzeczą, jakiej potrzebujemy do tego projektu, jest przedwzmacniacz mikrofonowy, który będzie tematem tego i następnego miesiąca. Muszę znaleźć czas na zaprojektowanie płytek drukowanych wokodera, ponieważ w oryginalnym projekcie wykorzystano przestarzały układ wzmacniacza transkonduktancyjnego CA3080. Jego produkcja została przerwana w 2005 roku, ku konsternacji branży technologii muzycznej. Podstawowa maska tego chipa przetrwała w układzie LM13700, który na szczęście jest jeszcze wciąż produkowany. Aby posłuchać wokodera w użyciu, odwiedź Soundcloud mojego syna i posłuchaj: <http://bit.ly/pe-may21-voco>.

Praca samodzielna

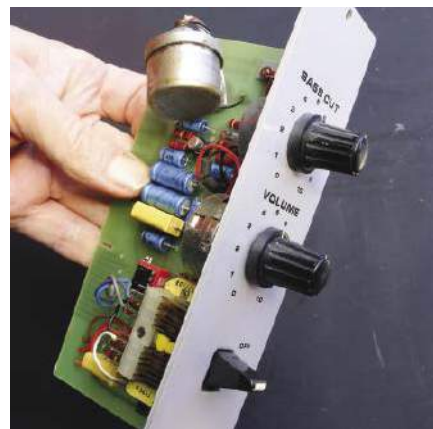
Wokoder mógłby wykorzystywać tylko jeden stopień odwracającego wzmacniacza operacyjnego jako przedwzmacniacza mikrofonowego, tak jak w konstrukcjach Korga i Rolanda. Odkryłem jednak, że użycie najwyższej jakości zewnętrznych przedwzmacniaczy i korektorów z wokoderami dało lepsze rezultaty. Jest zatem szansa, że mój projekt można przekształcić w pełnoprawną, samodzielną jednostkę porównywalną z tymi kosztującymi setki funtów.

Impedancja wejściowa i szum

Mikrofony wymagają dedykowanych wzmacniaczy, ponieważ ich sygnał wyjściowy często wynosi tylko kilka mV przy impedancji wyjściowej od około 25 Ω do 600 Ω . To znaczy, że standardowe obwody wzmacniaczy są narażone na dodatkowe szumy, ponieważ są dostosowane do optymalnej impedancji źródła (OSI) zapewniającej minimalny poziom szumu, która jest zazwyczaj znacznie wyższa niż impedancja wyjściowa mikrofonu. Dla wzmacniacza operacyjnego 5534, OSI wynosi około 5 k Ω , co oznacza, że należy zastosować specjalne sztuczki, aby zmniejszyć efektywne OSI. Analiza pełnego szumu układów wzmacniających jest złożona. Występuje również szum migotania, szum śrutowy, szum wybuchowy, szum kontaktowy, efekty zanieczyszczeń, termiczne prądy upływowe i defekty kryształów. Nie mamy tu miejsca, aby omówić je wszystkie, ale zobacz: <http://bit.ly/pe-may21-noise>.

Najstarszą metodą dopasowania niskiej impedancji do wyższej jest zastosowanie transformatora podwyższającego napięcie między mikrofonem a wzmacniaczem, jak pokazano na **rysunku 1**. Wiele inżynierów studyjnych nadal uważa to rozwiązanie za najlepiej brzmiące podejście. Naturalnie istnieją problemy związane z transformatorami: mają one opadającą charakterystykę częstotliwościową w skrajnych obszarach, odbierają przydźwięki, wprowadzają zniekształcenia na niskich częstotliwościach i kosztują około 60 funtów.

Najbardziej opłacalną metodą (tj. bez użycia transformatora) zmniejszenia efektywnego OSI jest zastosowanie pary tranzystorów bipolarnych o niskiej rezystancji rozproszonej (R_{bb}), pracujących przy stosunkowo wysokim prądzie kolektora, po których następują wzmacniacze operacyjne. Wymagana jest niska rezystancja, ponieważ wszystkie rezystancje powodują szum Johnsona w wyniku termicznego przemieszczania się elektronów. Im wyższa rezystancja, tym większy szum (można to zmniejszyć, chłodząc elektronikę ciekłym helem!). Zakładając jednak, że nie interesujemy się elektroniką astronomiczną i jej astronomicznymi kosztami, ezoteryczną kriogeniczną pozostawimy Jodrell Bank. Najlepiej jest po prostu przestrzegać podstawowych zasad audio dotyczących projektowania niskoszumowego, „utrzymywać wszystkie rezystory i impedancje na jak najniższych wartościach” oraz „w pierwszym stopniu wzmacniać tyle, ile jest możliwe”.



Rysunek 1. We wczesnych przedwzmacniaczach mikrofonowych stosowano na wejściu transformator podwyższający napięcie – tutaj na górze znajduje się puszka ekranująca z Mu-metalu (stopu metali magnetycznych, który charakteryzuje się wysoką przewodnością magnetyczną i zdolnością do efektywnego osłabiania pól magnetycznych). Ten moduł Astronic z lat 70. XX wieku wykorzystywał tranzystor wejściowy BC109, a następnie wzmacniacz operacyjny 741. Transformator, OEP X187B, kosztuje 64,00 GBP z VAT od RS (nr katalogowy 210-6352). Obrotowy montaż pozwala uzyskać minimalny poziom wychwytywania przydźwięków był świetnym pomysłem

Dzięki temu powstający szum nie jest późniejszy wzmacniany.

Mikrofon (lub dowolne źródło) dodaje własny szum. Mikrofony charakteryzują się szumem Johnsona oraz szumem wytwarzanym przez ruchy Browna cząsteczek powietrza na membranie. Celem konstruktora jest, aby szum wzmacniacza był co najmniej równy



Rysunek 2. Electro Voice RE20 to mikrofon wyposażony w przedwzmacniacz mikrofonowy o niskim poziomie szumów i niskiej impedancji. To „dźwięk amerykańskiego radia FM” (zdjęcie: Harvey Rothman)

szumowi własnemu źródła. Wtedy szum wzrósł tylko o 3 dB.

Wszystkie mikrofony pojemnościowe mają wewnętrzny wzmacniacz, który często jest dominującym źródłem szumu elektrycznego w tego typu urządzeniach. Niektóre mikrofony dynamiczne (elektromagnetyczne), np. wybrany przez amerykańskiego nadawcę mikrofon Electro Voice RE20 z impedancją wyjściową 150 Ω pokazany na **rysunku 2** wymaga szczególnie dobrego przedwzmacniacza. Mój projekt ma optymalną impedancję źródła około 200 Ω , ale rzeczywista impedancja wejściowa wynosi około 10 k Ω , aby zmniejszyć obciążenie. Dzięki temu mój przedwzmacniacz jest bardzo tani (5 funtów) i nawet stereofoniczny mikrofon dynamiczny Sony F-99A brzmiał dobrze, pomimo tego, że zawsze był stłumiony i zaszumiony (**rysunek 3**). Musiałem odciąć kabel z wtyczką mono jack F-99A i podłączyć go ponownie, aby zapewnić pracę symetryczną złącza XLR.

Elementy

Aby zachować spójność, tranzystory wejściowe przedwzmacniacza mikrofonowego muszą być wyselekcjonowane pod kątem niskiego poziomu szumów. Większość firm audio używa specjalnie wykonanych ekranowanych komór, przeznaczonych do testowania tranzystorów i wzmacniaczy operacyjnych. Kiedy w latach 80. pracowałem w branży konsol mikerskich, sortowano elementy pod kątem poziomu generowanego szumu. Bardziej szumiące wykorzystywano w korekcji basów i sterownikach LED. Te naprawdę złe

były wykorzystywane w syntezatorowych generatorach szumu. Rozstrzał poziomu szumów w segregowanych elementach był bardzo szeroki. Istnieje możliwość zakupu selekcjonowanych, komponentów o niskim poziomie szumów, ale za to płacisz. W nowoczesnych komponentach poziom szumów jest z reguły bardziej powtarzalny.

Tranzystory

Tranzystory o niskim R_{bb} są trudne do znalezienia, ponieważ jest to parametr prawie nigdy nie wyszczególniany w kartach katalogowych i trudny do zmierzenia. Przewidziano go tylko dla takich elementów, jak przestarzałe tranzystory produkcji Rohm i Toshiba, typu 2SB737 z $R_{bb}=2 \Omega$, wykonane specjalnie dla przedwzmacniaczy do gramofonów z wkładką z ruchomą cewką. Obecnie kwitnie branża sprzedająca podróbki z Hongkongu w serwisie eBay i tak, mam szufladę pełną tych bezużytecznych tranzystorów. Wiele inżynierów wybiera swoje ulubione tranzystory o niskim poziomie szumów spośród elementów ogólnodostępnych. Przez lata jednym z popularnych typów był 2N4403, który miał R_{bb} wynoszący 40 Ω . Pierwszy w Europie tranzystor o niskim poziomie szumów, BC109, ma rezystancję 400 Ω . Horowitz i Hill, autorzy *The Art of Electronics* zmierzili mnóstwo tranzystorów i ich wyborem do przedwzmacniacza mikrofonu wstęgowego (urządzenie o najniższej impedancji w audio) był Ferranti/Zetex/Little Diode ZTX751, który miał 1,7 Ω .

Inną metodą jest połączenie równolegle wielu małych tranzystorów, co jest sztuczką stosowaną w przedwzmacniaczach Ortofon. Inne popularne typy to tranzystory mocy BD139/40 z R_{bb} 30 Ω i typy TO5, takie jak BC461 i BC143 z R_{bb} 20 Ω . Używam Philipsa BFW16A, tranzystora RF średniej mocy, który ze względu na swoją splecioną strukturę faktycznie składa się z wielu małych tranzystorów połączonych równolegle, co zmniejsza rezystancję R_{bb} . Zdjęcie chipa pokazano na **rysunku 4**, a R_{bb} ma tu wartość około 5 Ω . Znalazienie nadwyżki w ilości 200 sztuk również miało wpływ na decyzję o ich wykorzystaniu.

Tranzystory PNP są czasami nieco cichsze (również mają niższe R_{bb}) niż komplementarne typy NPN. John Linsley-Hood powiedział, że było to spowodowane niższym szumem generowanym na powierzchni kryształu w wyniku rekombinacji dziur i elektronów. Nie zamierzam zagłębiać się w fizykę, która za tym stoi, ale stwierdziłem, że różnica jest raczej subtelna w praktyce. Na płycie drukowanej zapewnię możliwość odwrócenia polaryzacji, w razie potrzeby poprzez zworki do różnych szyn zasilających. Tranzystory o najniższym

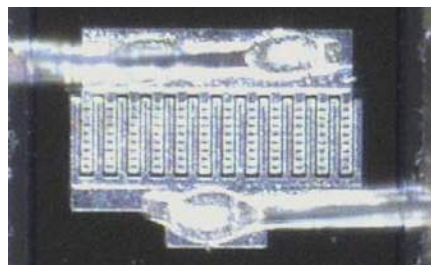


Rysunek 3. Tani mikrofon Sony F-99A 200 Ω (po lewej) nie nadawał się wcześniej do użytku ze standardowymi przedwzmacniaczami mikrofonowymi, a nawet dobry mikrofon wokalny do występów na żywo AKG D7 LTD 600 Ω wymagał dodatkowej charakterystyki wysokich częstotliwości. Obydwa dobrze działają z użyciem odpowiedniego przedwzmacniacza, takiego jak ten, który będziemy budować w przyszłym miesiącu

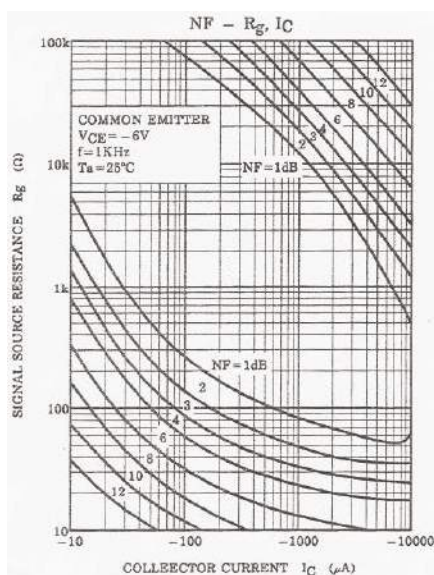
poziomie szumów to tranzystory JFET, które nie mają szumów podziału (gdzie następuje oddzielenie prądu bazowego od prądu kolektora). Klasycznym przykładem jest 2SK170, prawdziwie niskoszumny tranzystor FET, który oczywiście już nie jest produkowany. Istnieją zamienniki firmy InterFET, które kosztują fortunę. Ogólnie rzecz biorąc, im większy prąd tym elementy działają przy niższym OSI i niższym poziomie szumu aż do optymalnego punktu, specyficznego dla danego typu. Często jest to zaznaczone na specjalnym wykresie współczynnika szumu (**rysunek 5**). Współczynnik szumu to różnica w dB pomiędzy teoretycznym szumem Johnsona a szumem rzeczywistym.

Wzmacniacze operacyjne

Mój ulubiony wzmacniacz operacyjny audio to 5532/4. Co dziwne, w podwójnym wzmacniaczu operacyjnym 5532 otrzymujesz dwa wzmacniacze operacyjne 5534 za mniej niż cena jednego! Produkuje się



Rysunek 4. Tranzystor BFW16A składający się z przeplatających się ze sobą pasków małych tranzystorów. Zmniejsza to rezystancję rozproszoną bazy R_{bb} do około 5 Ω . Idealny wybór dla wzmacniaczy o niskiej impedancji (zdjęcie: dr Joe Botting)



Rysunek 5. Mapa współczynnika szumów dla Toshiba 2SA1312, nowoczesnego niskoszumowego tranzystora audio SMT PNP. Można zauważyć, że aby uzyskać najniższy współczynnik szumu wynoszący 1 dB przy najniższej rezystancji źródła wynoszącej 40 Ω, potrzebny jest prąd kolektora wynoszący 7 mA. Należy pamiętać, że specyfikacja tego współczynnika szumu może się zmieniać pomiędzy 0,2 dB i 3 dB, więc dobór jest nadal wymagany

znacznie więcej modeli podwójnych 5532 niż pojedynczych, co znajduje odzwierciedlenie w niższych kosztach. Istnieją niewielkie różnice: pojedynczy ma zewnętrzny kondensator kompensacyjny, który zapewnia lepszą charakterystykę wysokich częstotliwości przy wysokich wzmocnieniach. Ponadto model 5534 ma nieco niższy poziom szumów, zwłaszcza model 5534A. Szum wzmacniacza operacyjnego jest zwykle podawany w nanowoltach na pierwiastek z herca (nV/√Hz), w przypadku modelu 5534A z 3,5 nV/√Hz i modelu 5532 z 5 nV/√Hz. Dobre tranzystory dyskretnie mają zwykle napięcie od 0,5 do 1 nV/√Hz. Rzeczywiste napięcie zależy od prądu szumu (pA), a także napięcia szumu i konfiguracji obwodu. Ważne jest także pasmo częstotliwości. W przypadku układów audio zwykle przyjmujemy, że jest to 20 kHz, więc pamiętaj o pierwiastku kwadratowym



Rysunek 6. Specjalny potencjometr firmy Blore Edwards do regulacji wzmocnienia (VR1). To jest seria CTS 45, wykładniczy, podwójny, z wyłącznikiem

z 20 000 (~141). Zatem nasze 5 nV/√Hz wynosi $5 \times 141 = 705$ nV.

Istnieje bardzo niewiele wzmacniaczy operacyjnych, które stanowią przydatne ulepszenie w stosunku do 5534/2 w pracy audio. Jedyne jakie znalazłem to LM4562. Jest on dostępny tylko w wersji podwójnej, więc w przypadku przyszłych aktualizacji musimy zaprojektować płytkę drukowaną dla wzmacniaczy podwójnych. Zauważ też, że maksymalne zasilanie dla LM4562 wynosi ± 18 V, a nie ± 22 V jak dla 5532/4.

Rezystory

Oprócz szumu Johnsona, rezystory generują również dodatkowy szum powodowany przez mikrołuki pomiędzy cząsteczkami przewodzącymi. Szum ten jest proporcjonalny do prądu przepływającego i nazywany jest „szumem nadmiarowym”. Im bardziej jednorodny jest materiał oporowy, tym lepiej. Najlepsze są rezystory z folii tantalowej Charcroft Vishay, a najgorsze to stare węglowe rezystory kompozytowe, w których zasadniczo związane są cząsteczki sadzy razem z klejem. Dla ekonomicznego profesjonalisty prac audio powszechnie stosowane są osiowe 0,25 W pełnowymiarowe, naporowywane rezystory nichromowe z powłoką metaliczną MR25. Uważaj na tanie typy SMT z metalową powłoką, które są tak naprawdę grubowarstwowymi typami cermetowymi, które generują większe szumy. Zwiększanie objętości materiału oporowego zwiększa skuteczną liczbę cząstek, więc typy o wyższej mocy generują na ogół mniejsze szumy.

Potencjometry

Potencjometr regulacji wzmocnienia jest elementem trudnym do uzyskania, ponieważ musi być wykładniczy (typ C), w tym sensie, że szybkość zmiany rezystancji maleje wraz z obrotem potencjometru w kierunku zgodnym z ruchem wskazówek zegara. Jest to przeciwieństwo normalnej logarytmicznej kontroli głośności (typ A). Jest niezbędne aby w takich sytuacjach mieć odpowiednie prawo regulacji lub typ, aby mieć pewność, że efekt jest równomiernie rozłożony w całym zakresie obrotowym. Wykładniczy typ C nadal nie jest wystarczający, aby uniknąć dużego skoku wzmocnienia, na końcu obrotu. Alps produkuje specjalny typ zwany „RD”, zaprojektowany specjalnie do tego zastosowania, ale trzeba je kupować w tysiącach sztuk. Wszystkie potencjometry generują szumy, większość wykorzystuje warstwę z kompozytu węglowego. Istnieje rodzaj toru formowanego, który jest grubszy, ale nadal stanowi urządzenie kompozytowe lub zawiera bardziej nowoczesny przewodzący plastik, którym jest węgiel związany epoksydowo.

Te typy generują mniejszy poziom szumów. Jedną z rzeczy, które często robię, aby



Rysunek 7. Kondensatory tantalowe oferujące wysokie wartości pojemności przy niskich prądach upływu (10...20 μA) – chociaż producenci tego nie gwarantują (powinieneś sprawdzić swoje komponenty!). Uzyskanie wystarczająco wysokich wartości (ponad 100 μF) może być trudne. Kondensator metalowy to montowany śrubowo kondensator typu mokrego Castanet, a kondensator osiowy to standardowy, solidny europejski TAA w metalowej obudowie, odpowiednik amerykańskiego CTS13

zmniejszyć szum potencjometru, jest użycie podwójnego podłączonego równolegle.

Zmniejsza to o połowę rezystancję, ale zmniejszenie wartości potencjometru do 1...2,5 kΩ daje mniej gwałtowny skok przy dużych wzmocnieniach kosztem wyższego minimalnego wzmocnienia. Używam podwójnego potencjometru 5 kΩ z przełącznikiem i całkowicie go wyłączam, aby uzyskać minimalne wzmocnienie (rysunek 6). Jeśli jesteś bardzo wybredny, możesz pójść drogą API lub Neve i mieć połączany przełącznik Elma z 1% rezystorami z folii metalowej, skalibrowanymi w równych odstępach dB.

Kondensatory

Wracając do kondensatorów i szumów, rodzaje folii z tworzywa sztucznego generują najmniejsze szumy, ponieważ mają bardzo niski poziom prądu upływu.

Tam, gdzie wymagana jest duża pojemność i trzeba zastosować elektrolity, typy tantalowe (rysunek 7) charakteryzują się mniejszymi upływnościami niż standardowe mokre aluminiowe, szczególnie tam, gdzie nie były używane przez jakiś czas i muszą się uformować. Wielowarstwowe typy ceramiki X7R są bardzo słabe; w rzeczywistości same zachowują się jak mikrofony piezoelektryczne. Dotknij ich, a zaczną szumieć.

W następnym miesiącu

To wszystko na ten miesiąc. W następnym zaprojektujemy i zbudujemy obwód przedwzmacniacza. ■

Jake Rthman

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, maj 2021 (www.epemag3.com)

Wydłużenie żywotności pompy wody wraz z oszczędnością energii

Układ proponowany w bieżącym projekcie jest przewidziany do prostych systemów klimatyzacji i nawilzaczy powietrza. Systemy te wykorzystują zjawisko pochłaniania ciepła podczas parowania. Zawierają też pompę wodną, która nie powinna pracować na sucho.

Proponowany układ zabezpieczy przed taką sytuacją, przyczyniając się równocześnie do oszczędności energii elektrycznej. Proste systemy schładzania pomieszczeń są powszechnie stosowane w wielu domach, biurach, w szkołach i w przemyśle. W okresach letnich są często włączane na wiele godzin i pozostawione bez kontroli. Bieżący projekt wykonano z myślą o takiej kontroli w zakresie obecności wody w zbiorniku. Pompa nie powinna pracować na sucho. Drugą funkcją układu jest włączanie pompy (gdy woda jest w zbiorniku) w kilkuminutowych interwałach czasowych z wypełnieniem 30-to procentowym. Na **rysunku 1** pokazano prototyp poddany testom, zaś na **rysunku 2** jest schemat ideowy zaproponowanego tu układu.

W układzie wykorzystano następujące elementy i podzespoły: transformator obniżający napięcie sieciowe (X1), mostek prostowniczy BR1, stabilizator 5-cio voltowy LM7805 (IC1), licznik dekadowy 4017 (IC2), diodę prostowniczą 1N4007 (D1), kilka diod sygnałowych

1N4148 (D2 do D5), tranzystor NPN 2N2219 (T1), pięciowoltowy przekaźnik (RL1) oraz niewiele pasywnych elementów.

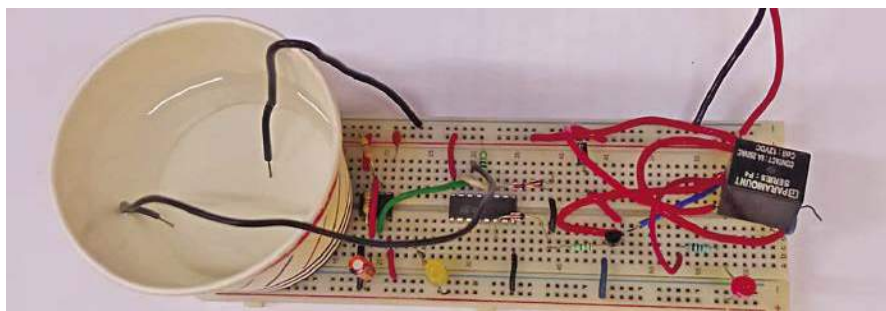
Układ jest zasilany napięciem 5 V, które pozyskano stabilizatorem liniowym. Zasilanie czerpane jest z sieci 230 VAC za pośrednictwem transformatora obniżającego napięcie do 9 VAC o znamionowej obciążalności do 0,5 A. Do podłączenia sieci energetycznej przewidziano złącze CON1 w obwodzie uzwojenia pierwotnego transformatora. Uzwojenie wtórne wytwarza napięcie o wartości skutecznej ok. 9 V. Podano je na wejście mostka Graetza, w celu prostowania dwupołkowego, na wyjściu którego zamontowano kondensator dużej pojemności dla celu wstępnej filtracji. Takie zasilanie podane jest na wejście stabilizatora liniowego, który wypracowuje stabilne +5 VDC. Na wyjściu IC1 wstawiono także diodę LED1, której świecenie potwierdza obecność napięcia 5 VDC.

W układzie wykorzystano też popularny timer 555, który połączono w konfigurację

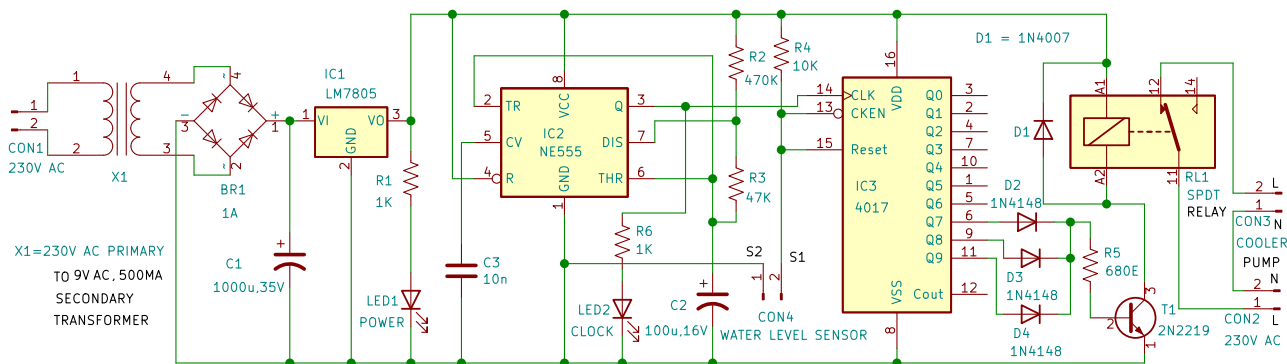
przerzutnika astabilnego. Częstotliwość multiwibratora uwarunkowana jest rezystorami R2 i R3 oraz pojemnością kondensatora C2. To długie stałe czasowe dające okres oscylacji na poziomie 40 sekund. Częstotliwość multiwibratora można zmieniać w szerokim zakresie poprzez dobór elementów RC. W miejsce rezystora R3 można także zastosować potencjometr, co pozwoli na płynną regulację częstotliwości (**wiekszy wpływ na timing multiwibratora ma rezystor R2, gdyż jest on o rząd wielkości większy od R3 – przypis redakcji**). Na wyjściu timera 555 wstawiono także diodę LED. Jej miganie świadczy o poprawnej pracy przerzutnika astabilnego.

Kolejnym kluczowym elementem jest licznik dekadowy IC3. Wejście zegarowe licznika jest wyzwalone z wyjścia Q (pin 3) multiwibratora. Zasilanie licznika podane jest na jego nóżkę 16, zaś masę GND stanowi wyprowadzenie 8. Istotne znaczenie mają wejścia pin 13 i 15. Na n.15 jest Reset aktywny w stanie wysokim, a na n.13 jest Clock-Enable aktywny w stanie niskim. Oba wejścia połączono razem i przez pull-up rezystor R4 do zasilania 5 V. Wysoki stan w tym miejscu blokuje pracę licznika, gdyż równocześnie nie pozwala na takt zegarowy, a aktywny Reset ustawia licznik dekadowy do pozycji, gdy jedynka logiczna jest tylko na wyjściu Q0.

Kontrola poziomu wody w zbiorniku dokonywana jest właśnie na wejściu Reset/Clock-Enable licznika. W wodzie zanurzane są dwie elektrody, z których jedna jest na masie a druga na tymże wejściu licznika. Zakłada się,



Rysunek 1. Prototyp poddany testom



Rysunek 2. Schemat ideowy układu

że przewodność elektryczna wody jest dużo większa od rezystancji R4, i wejście n.13 i n.15 zostanie ściągnięte do stanu niskiego. Jedynie obniżenie się poziomu wody poniżej miejsca zamontowania elektrod spowoduje, iż R4 podniesie potencjał Resetu do poziomu bliskiego zasilania. Dla podłączenia kontrolnych elektrod przewidziano dwu-pinowe złącze CON4.

Normalna praca licznika jest wtedy, gdy woda w zbiorniku jest obecna, czyli elektrody są w wodzie zanurzone. Wtedy pompa wody jest włączana, jednak nie cały czas, ale w interwałach czasowych z wypełnieniem 30%. Jeśli poziom wody jest niewystarczający, pompa nie jest uruchamiana.

Licznik dekadowy 4017 ma dziesięć wyjść: od Q0 do Q9. Pierwsze siedem wyjść (Q0 do Q6) jest niewykorzystanych i pozostają one permanentnie „w powietrzu”. Wyjścia Q7, Q8 i Q9 połączono w obwód sumy logicznej za pośrednictwem diod D2, D3 i D4. Katody tych diod połączono razem i sterują one bazą tranzystora T1. Wysoki stan na dowolnym z wyjść (Q7, Q8 lub Q9) włącza ten tranzystor i za jego pośrednictwem przełącznik. Właśnie wykorzystanie trzech spośród dziesięciu wyjść licznika dekadowego, narzuca 30-procentowe wypełnienie interwałów czasowych włączających pompę.

Wykaz elementów, kupuj w sklepie [avt.pl](http://www.avt.pl) (W-wa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl);

Półprzewodniki:

IC1: LM7805 – stabilizator linii 5 V
IC2: timer NE555
IC3: licznik dekadowy 4017
BR1: mostek prostowniczy 1 A
LED1, LED2: dioda LED 5 mm
D1: 1N4007 – dioda prostownicza
D2-D5: 1N4148 – dioda sygnałowa
T1: 2N2219 – tranzystor NPN

Rezystory: (wszystkie 0,25 W/±5%)

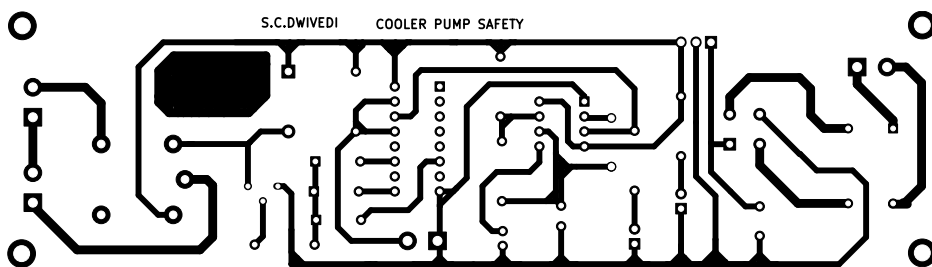
R1, R6: 1 kilo-ohm
R2: 470 kilo-ohm
R3: 47 kilo-ohm
R4: 10 kilo-ohm
R5: 680 ohm

Kondensatory:

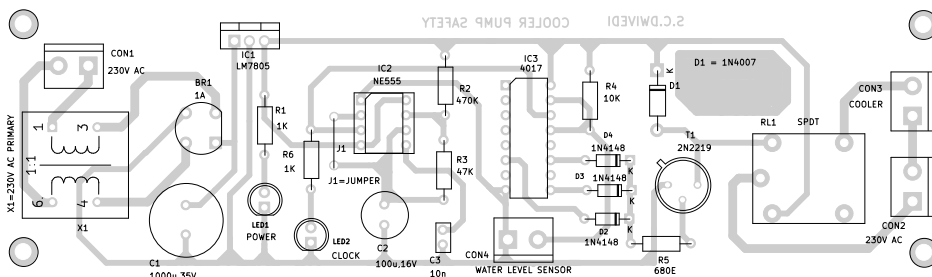
C1: 1000 uF/35 V elektrolityczny
C2: 100 uF/16 V elektrolityczny
C3: 10 nF ceramiczny

Inne:

CON1-CON4: złącze 2-pinowe
S1: elektrody zanurzone w wodzie – odcinki ze stali nierdzewnej
RL1: przełącznik 5 V z pojedynczymi stykami NO/NC
X1: transformator sieciowy 230 VAC – uzwojenie pierwotne, 9 VAC/500 mA – wtórne
pompa wodna



Rysunek 3. Projekt druku PCB



Rysunek 4. Schemat montażowy układu elementów na PCB

Konstrukcja i testowanie pracy układu

Po zrealizowaniu wszystkich połączeń, należy podłączyć zasilanie 230 VAC. Przekaznik RL1 włączany jest driverem w postaci tranzystora T1 wtedy gdy stan wysoki panuje na jednym z trzech wyjść licznika. Styki wykonawcze przełącznika włączają pompę. Zatem, gdy woda w zbiorniku jest obecna, pompa pracuje z interwałami narzuconymi częstotliwością pracy multiwibratora oraz logiką licznika dziesiętnego. Gdy wody brak, praca pompy jest zatrzymana.

Czasy włączenia i wyłączenia pompy uwarunkowane są następującymi zależnościami: stała czasowa $(R2+R3) \times C2$ to ok. 51,7 sekundy; aby oszacować czas stanu wysokiego na wyjściu Q multiwibratora, należy τ przemnożyć przez $\ln 2$ (ok. 0,69), co daje czas ok. 36 sekund; aby ocenić okres multiwibratora trzeba ten czas powiększyć o 10% (bo $R3=1/10 \times R2$), to daje czas 40 sekund; licznik jest dekadowy, więc pełny okres należy przemnożyć przez 10; 400 sekund to 6,7 minuty; z tego czas włączenia pompy stanowi 30% czyli 2 minuty i przerwa 4 minuty i 40 sekund – przypis redakcji EdW.

Dla niniejszego projektu przygotowano jednostronną płytkę drukowaną PCB, którą można odwzorować (w naturalnej skali) z rysunku 3. Na rysunku 4 widoczny jest schemat ułatwiający montaż elementów na PCB.

Po zmontowaniu elementów na PCB, układ należy umieścić w odpowiednio przygotowanej obudowie. We frontowej części obudowy należy zamontować diody LED1 i LED2 oraz złącze dla elektrod zanurzanych w zbiorniku wody. Złącza dla pompy i jej zasilania (CON2 i CON3) najlepiej umieścić na tyle obudowy.

Zasadniczą częścią prac montażowych jest poprawne zamontowanie elektrod mających kontakt z wodą. Powinny być wykonane ze stali nierdzewnej i ulokowane w odległości od 30 do 40 mm od siebie. Gdy poziom wody podniesie się na założoną wysokość, układ powinien zareagować jak na zwarcie między punktami S1 i S2.

Rezystancja widziana między elektrodami mającymi kontakt z wodą musi być wielokrotnie mniejsza od rezystora R4, który tu ma oporność 10 kilo-ohmów; przewodność elektryczna wody silnie zależy od jej zasolenia; zatem wymagana odległość montażu elektrod powinna być eksperymentalnie sprawdzona w konkretnej aplikacji; i przykładowo, przewodność właściwa wody destylowanej wynosi ok. 10^{-4} S/m (Siemensa na metr); woda rzeczna ma przewodność właściwą ok. 0,005 S/m; zaś w przypadku wody morskiej jest to już ok. 3 do 5 S/m; Siemens/metr to jednostka mało przemawiająca do wyobraźni, więc przeliczmy: odcinek 3 cm o przekroju 1 cm² z materiału o przewodności σ będzie miał w przypadku wody destylowanej rezystancję kilku mega-ohmów, dla wody rzecznej będzie to kilkadziesiąt kilo-ohmów, a dla zasolonej wody morskiej jedynie ok. kilkadziesiąt ohmów – przypis redakcji EdW.

Bonus

Działanie układu wg bieżącego projektu można obejrzeć pod adresem: <https://www.electronicsforu.com/videos-slideshows/live-diy-cooler-pump-safety>. ■

Suresh Dwivedi

Ten artykuł był wcześniej opublikowany na łamach „EFY”, lipiec 2023 (efymag.com)

Alarm kuchenny z wykorzystaniem czujnika gazu MQ2

Z gazem „nie ma żartów”. Najczęstszą przyczyną pożaru w kuchni jest ulatnianie się gazu. Ten projekt powinien zabezpieczyć przed tak poważnym zdarzeniem. Zastosowany czujnik powinien wykryć niebezpieczne stężenie gazu LPG, jak również powinien zareagować na zadymienie. W takiej sytuacji powinien rozleć się alarm w celu natychmiastowej interwencji i stwierdzenia przyczyny ew. ulatniania się gazu.

Działanie układu wg bieżącego projektu opiera się na wykorzystaniu niedrogiego i niewielkich rozmiarów sensoru czułego zarówno na obecność gazu jak i dymu. MQ2 jest sensorem półprzewodnikowym typu MOS (Metal Oxide Semiconductor). Jest on czuły na obecność gazów palnych jak: LPG, propan-butan, metan, wodór, także tlenek węgla. MQ2 umieszczony w środowisku o dużym stężeniu tego typu gazów, jak również w sytuacji zadymienia, zmienia rezystancję między elektrodami czynnymi. Sytuację taką jest łatwo wykryć włączając czujnik w obwód prostego dzielnika rezystancyjnego z potencjometrem lub zwykłym rezystorem. MQ2 zawiera dwie elektrody wykonane z tlenku aluminium Al_2O_3 oraz niewielkiej mocy grzejnik. Warstwą „czynną” czułą na obecność gazów jest SnO_2 .

Schemat układu i jego działanie

Schemat alarmu z wykorzystaniem czujnika MQ2 pokazano na **rysunku 1**.

W układzie wykorzystano transformator obniżający napięcie sieciowe AC (X1), mostek prostowniczy (BR1), 5-cio voltowy stabilizator LM7805 (IC1), czujnik MQ2 (Sensor1), wzmacniacz operacyjny LM358 (IC2), timer NE555 (IC3), pięcio-voltowy przełącznik RL1, piezoelektryczny buzzer (BZ1), oraz niewielką ilość prostych elementów pasywnych. Za mostkiem prostowniczym wstawiono kondensator elektrolityczny C1 w celu odzyskania i filtrowania składowej stałej napięcia. Układ alarmu zasilany jest napięciem stałym 5 VDC. Napięcie to jest stabilizowane popularnym LM7805 ze zmiennego 9 VAC pozyskanego transformatorem X1 o obciążalności uzwojenia wtórnego na poziomie 0,5 A. Napięcie sieciowe 230 VAC podłączono do uzwojenia pierwotnego transformatora za pośrednictwem złącza CON1. 9 VAC z uzwojenia wtórnego prostowane jest mostkiem BR1. Wstępnie filtrowane napięcie stałe podane jest na stabilizator liniowy LM7805 na wyjściu

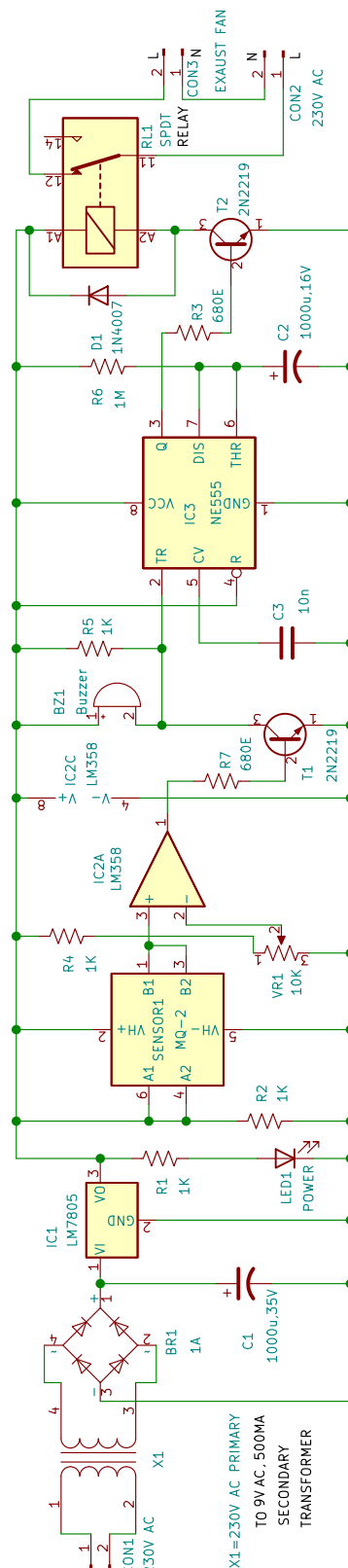
którego dostępne jest 5 VDC. Świecenie diody LED1 informuje o obecności tego napięcia.

Wyjście czujnika MQ2 podłączono na wejście nieodwracające wzmacniacza operacyjnego, który pracuje tu w roli komparatora. Napięcie referencyjne podane jest na wejście odwracające WO (pin 2 LM358). Tu zastosowano potencjometr, dzięki któremu można ustawić próg działania, co przekłada się na czułość alarmu względem stężenia gazu w powietrzu. Wysoki stan wyjścia wzmacniacza operacyjnego oznacza alarm. Za pośrednictwem tranzystora T1 włączany jest sygnał alarmowy z buzzera.

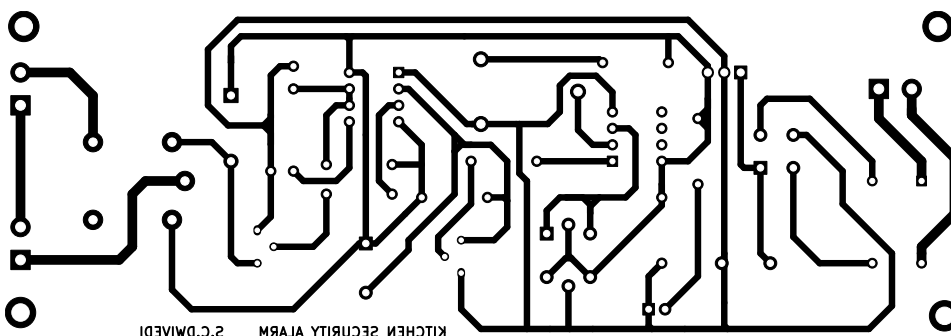
Zatem działanie całości układu jest bardzo proste.

Po podłączeniu zasilania 230 VAC do złącza CON1, powinna zaświecić się dioda LED1 informując o gotowości działania alarmu. W przypadku zadymienia lub niebezpiecznego stężeniu gazu z powietrza, napięcie na wyjściu czujnika MQ2 podnosi się. Po przekroczeniu wartości ustawionej potencjometrem VR1, następuje zmiana stanu wyjścia WO z niskiego na bliski zasilania +5 V. Wówczas zostanie wysterowana baza T1. Tranzystor ten przejdzie w stan nasycenia, konsekwencją czego jest włączenie sygnału dźwiękowego z buzzera, co powinno wzbudzić czujność w celu sprawdzenia przyczyny takiego stanu rzeczy.

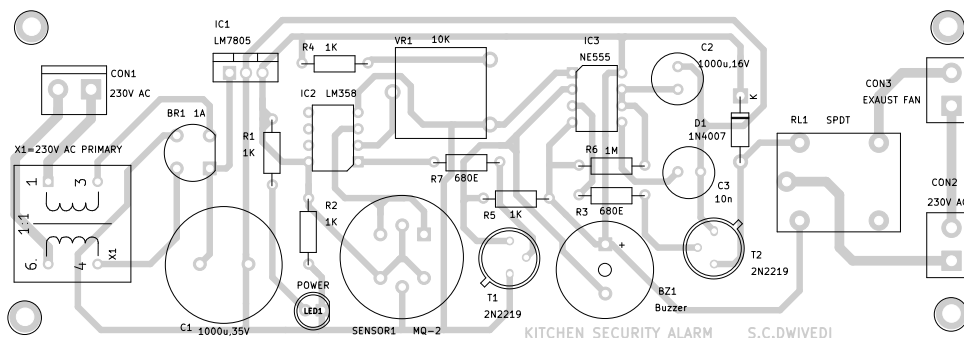
Oprócz sygnału dźwiękowego, przewidziano też uruchomienie wentylatora w celu przewietrzenia pomieszczenia. Na timerze NE555 wykonano przerzutnik monostabilny. Wyzwalany jest on tym samym sygnałem który włącza buzzer. Czas działania tego przerzutnika ustalony jest iloczynem rezystora R6 i kondensatora C2. W tym przypadku jest to ekstremalnie duża wartość $1000 \mu F \times 1 M\Omega = 1000 s$. Stałą czasową należy przemnożyć przez $1,1 (\ln 3)$, co daje wartość 1100 s, czyli ok. 18 minut. Wysoki stan wyjścia Q timera uruchamia przełącznik RL1 za pośrednictwem tranzystora T2. W obwodzie styków przełącznika zamontowano dwa złącza CON2 i CON3. Do jednego należy podłączyć wentylator, a do drugiego zasilanie



Rysunek 1. Schemat ideowy alarmu z czujnikiem MQ2



Rysunek 2. Płytkę PCB „alarmu kuchennego”



Rysunek 3. Schemat montażowy układu

wentylatora (w naszym przypadku przewidywano 230 VAC).

W celu skalibrowania układu, proponujemy następującą procedurę. Potencjometrem VR1 na wejściu odwracającym IC2 ustaw napięcie ok. 0,5 V. Teraz należy wytworzyć dym lub w okolicy czujnika MQ2 odkręcić lekko kurek z gazem LPG. Po krótkim czasie napięcie na pinie 3 LM358 powinno się podnosić i stan wyjścia powinien zmienić się z niskiego na wysoki. To powinno uruchomić sygnał alarmu oraz zwarcie styków przekaźnika RL1. Podniesienie napięcia potencjometrem VR1 zmniejsza czułość, a obniżenie tego napięcia zwiększa czułość systemu.

Należy ustawić bezpieczną wartość możliwie niską, jednak na tyle wysoką, aby alarm nie uruchamiał się bez przyczyny.

Konstrukcja i testowanie pracy układu

Dla naszego układu przewidziano płytkę drukowaną PCB. Projekt jednostronnego druku pokazano na **rysunku 2**, który w papierowej wersji czasopisma powinien odpowiadać skali 1:1. **Rysunek 3** pokazuje ułożenie elementów na PCB.

Po zmontowaniu układu, należy go umieścić w odpowiednio przygotowanej obudowie. Diodę LED1 i złącze CON2 najlepiej przewidzieć z przodu obudowy. Od tyłu należy podłączyć wydajny wentylator, oraz zamontować go tak, aby skutecznie przewietrzył pomieszczenie. Całość układu lub sam czujnik MQ2 należy ułożyć możliwie blisko pieca gazowego. Tak, aby ew. ulatnianie się gazu lub stężenie dymu w okolicy pieca uruchamiało alarm.

Bonus

Działanie modelu wykonanego przez autora można obejrzeć pod adresem: <https://www.electronicsforu.com/videos-slideshows/live-diy-kitchen-security-alarm-using-mq2-sensor>. ■

Suresh Dwivedi

Od Redakcji EdW: Bieżący projekt jest bardzo prosty. Mimo to na schemacie z rysunku 1 wkradł się błąd. Czujnik MQ2 działa na zasadzie zmiennej rezystancji. Aby oddać informację

o ew. obecności gazu lub zadymienia, musi współpracować z jakimś rezystorem tworząc dzielnik rezystancyjny. Zmienna rezystancja widziana jest między elektrodami A1-A2 i B1-B2. Impedancja wejściowa wzmacniacza operacyjnego LM358 jest bardzo duża i układ wg zamieszczonego schematu działać nie będzie. Powtarzalność charakterystyki czujnika gazu jest bardzo kiepska. Dlatego często zamiast rezystora na wyjściu jest potencjometr. Oczywiście może być i tak jak na schemacie, gdzie potencjometr jest na drugim wejściu komparatora. Ale na wyjściu MQ2 trzeba wstawić rezystor o wartości zbliżonej do rezystancji czujnika w okolicy „centrum” jego charakterystyki. Przyglądając się schematowi stwierdzimy, iż na wejściu jest rezystor R2, który tutaj nie ma sensu. Należy go zatem przenieść na wyjście, a wartość 1 kilo-ohm powinna być odpowiednia, aby potencjometrem VR1 ustawić właściwy próg przy, którym powinien uruchomić się alarm.

Ponadto na warstwie opisowej (silkscreen) PCB znajdują się błędy: wprowadzające w błąd oznakowanie uzwojeń transformatora (brak separacji galwanicznej/podanie napięcia sieci 230 V AC na elektronikę) oraz brak zaznaczonej polaryzacji dla kondensatorów elektrolitycznych C1, C2 a także brak oznakowania wyprowadzenia o numerze jeden dla sensora MQ2.

Ten artykuł był wcześniej opublikowany na łamach „EFY”, czerwiec 2023 (efymag.com)

Wykaz elementów:

Półprzewodniki:

IC1: LM7805 – stabilizator 5 V
IC2: LM358 – wzmacniacz operacyjny
IC3: timer NE555
T1,T2: tranzystor npn 2N2219
BR1: 1 A mostek prostowniczy
LED1: dioda LED 5 mm
D1: 1N4007 – dioda prostownicza

Rezystory:

(wszystkie 0,25 W/±5%)
R1,R2,R4,R5: 1 kΩ
R3,R7: 680 Ω
R6: 1 MΩ

Kondensatory:

C1: 1000 µF/35 V elektrolityczny
C2: 1000 µF/16 V elektrolityczny
C3: 10 nF ceramiczny

Inne:

Wentylator: wydajny wentylator zasilany napięciem 230 VAC
RL1: przekaźnik SPDT z cewką 5 V
CON1-CON3: złącze dwu-pinowe
Sensor1: czujnik gazu typu MQ2
X1: transformator sieciowy 230 VAC; 9V AC/500 mA – uzwojenie wtórne

Wyłącznik czasowy

Często zapominamy o wyłączeniu np. światła lub wentylatora w łazience. Przyczynia się to do marnowania energii elektrycznej i niepotrzebnego wzrostu rachunków za energię. Ten projekt jest bardzo prosty, ale efekt ekonomiczny może być całkiem pokaźny. Wyłącznik czasowy „nie zapomni” i wyłączy urządzenie elektryczne po ściśle zaprogramowanym interwale czasowym. Taka „przystawka” jest pożądana nie tylko w mieszkaniu, ale tym bardziej w biurach, gdzie wiele anonimowych osób korzysta np. z toalety.

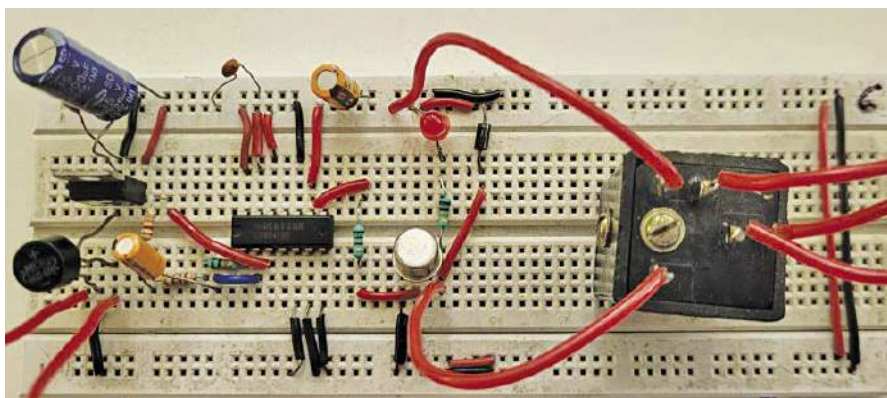
Urządzenie będące tematem bieżącego projektu nie jest adresowane dla pomieszczeń w których przebywamy przez długi czas. Ale przede wszystkim tam, gdzie światło (lub inny odbiornik energii) potrzebne jest na stosunkowo krótki odcinek czasu.

Proponowany układ należy włączyć między źródło zasilania i odbiornik. Obciążenie zostanie automatycznie odcięte po ściśle określonym czasie (w prototypie ustawiono ten czas na ok. 15 minut). Czas ten można ustawić z rozsądnym marginesem względem potrzeb, a oszczędności rachunków za energię elektryczną powinny być znaczące. Pewną niedogodnością wynikającą z konstrukcji jest fakt, iż aby ponownie załączyć światło, trzeba na chwilę wyłączyć zasilanie i z powrotem je włączyć. Na **rysunku 1** widzimy prototyp który pomyślnie przetestowano w laboratorium redakcji „Electronics For You”.

Budowa układu i jego działanie

Schemat ideowy proponowanego wyłącznika czasowego pokazuje **rysunek 2**.

Układ wykonano z wykorzystaniem transformatora sieciowego (X1) obniżającego napięcie z 230 VAC do 12 VAC (o nominalnej obciążalności 1 A). Kluczowym elementem jest timer/licznik CD4541BE (IC1).



Rysunek 1. Prototyp układu poddany testom

Ponadto wykorzystano 12-to voltowy stabilizator napięcia LM7812 (IC2), 5 diod 1N4007 (D1 do D5), tranzystor npn typu 2N2219 (T1), dwunasto-voltowy przekaźnik (RL1) i niewielką ilość rezystorów i kilka kondensatorów.

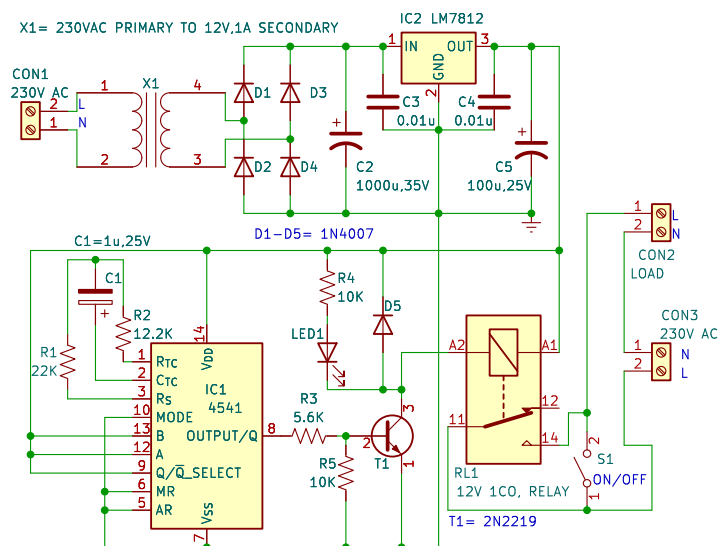
Wyłącznik czasowy zostanie załączony z chwilą podania napięcia 230 VAC na wejście (złącze CON1 na rysunku 2). Na wejściu musi być umiejscowiony switch ON/OFF, którego na schemacie ideowym nie pokazano. Styki wykonawcze timera zostaną wyłączone po zadanym czasie mimo dalszej obecności napięcia na złączu CON1. Układ scalony CMOS CD4541BE (IC1) umożliwia łatwe odmierzenie długich czasów.

Działanie polega bowiem na oscylatorze i wielobitowym dzielniku częstotliwości. Częstotliwość oscylatora ustalona jest ze względu na elementy RC zgodnie ze wzorem: $f=1/(2,3 \times R2 \times C1)$. Dzielnik częstotliwości działa w oparciu o liczniki binarne, których „długość” można programować. Służą

do tego piny 12 i 13. Poniższa tabelka pokazuje jakie współczynniki podziału są dostępne. Wyprowadzenia 12 i 13 (oznaczone A i B) można zwierzać do masy lub plusa zasilania, co oznacza stan zera lub jedynki logicznej.

W naszym przypadku chcemy odmierzyć czas ok. 15-tu minut. W tym celu ustawiamy częstotliwość oscylatora 35,6 Hz. Współczynnik podziału ustawiamy maksymalny dostępny, tj. 16 bitów. To daje zwielokrotnienie okresu zegara 2^{16} czyli 65536.

$1/35,6 \text{ Hz} \times 65536 = 1840 \text{ sekund} = \text{ok. 30 minut}$; nie oznacza to jednak, że autor popełnił błąd; w istocie dzielnik częstotliwości odpowiada krotności 2^{N-1} , gdyż najstarszy bit licznika oznacza już przepełnienie i to ostatni bit zeruje (lub ustawia) przerzutnik na wyjściu (jeśli timer pracuje w trybie „jednego impulsu”; zatem odmierzony timerem czas powinien faktycznie wynieść ok. 15 minut; powtórzmy, iż Set czy Reset jak również co jest zerem a co jedynką logiczną, jest kwestią czysto umowną; tutaj Reset ustawia wyjście Q do stanu wysokiego, a jedynka logiczna na najstarszym bicie licznika skutkuje zmianą wyjścia Q do stanu



Rysunek 2. Schemat ideowy układu

Tabela 1.

Pin-12	Pin-13	2 ^N	współczynnik podziału
0	0	2 ¹³	8192
0	1	2 ¹⁰	1024
1	0	2 ⁸	256
1	1	2 ¹⁶	65536

niskiego; i w trybie „jednego impulsu” ten stan pozostaje stabilny – przypis redakcji EdW.

Elastyczność wykorzystania timera CD4541 sprawia obecność dodatkowych wyprowadzeń: MODE, Auto Reset, Master Reset oraz wybór aktywności stanu logicznego wyjścia (aktywne zero lub jedynka logiczna). MODE (pin 10) ustawia cykl pracy ciągły lub „pojedynczy impuls”. W naszym przypadku zachodzi potrzeba odmierzenia jednego (stosunkowo długiego) impulsu, dlatego pin 10 należy połączyć z potencjałem masy układu. Na wyjściu potrzeba aby stan aktywny był stanem wysokim, dlatego wyprowadzenie 9 łączymy z plusem zasilania. Układ timera ma się resetować po załączeniu zasilania, dlatego pożądanym trybem jest Auto Reset. Dlatego pin 5 AR łączymy z potencjałem masy. Pełne informacje na temat konfiguracji timera CD4541 zawarte są w karcie katalogowej tego układu scalonego.

Działanie układu wyłącznika czasowego

Układ zasilany jest napięciem 12VDC ze stabilizatora 12-to woltowego, na którego wejście podane jest wyprostowane napięcie z uzwojenia wtórnego transformatora. Podanie napięcia zasilania uruchamia timer. Stan wysoki wyjścia poprzez driver na tranzystorze T1

Wykaz elementów:

Półprzewodniki:

IC1: LM7812 – stabilizator napięcia 12 V
IC2: CD4541BE – układ scalony timera
D1-D5: 1N4007 – diody prostownicze
T1: 2N2219 – tranzystor NPN

Rezystory: (wszystkie 0,25 W/±5%)

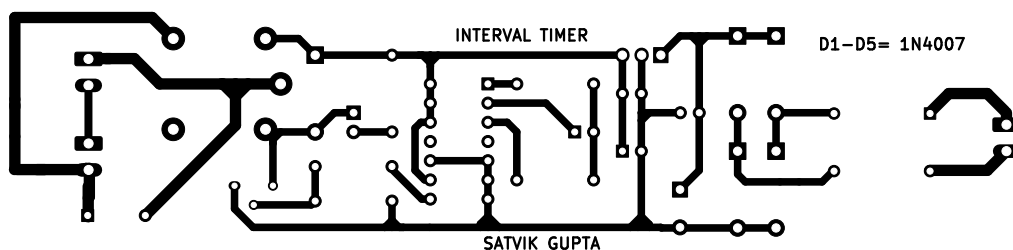
R1: 22 kΩ
R2: 12,2 kΩ
R3: 5,6 kΩ
R4, R5: 10 kΩ

Kondensatory:

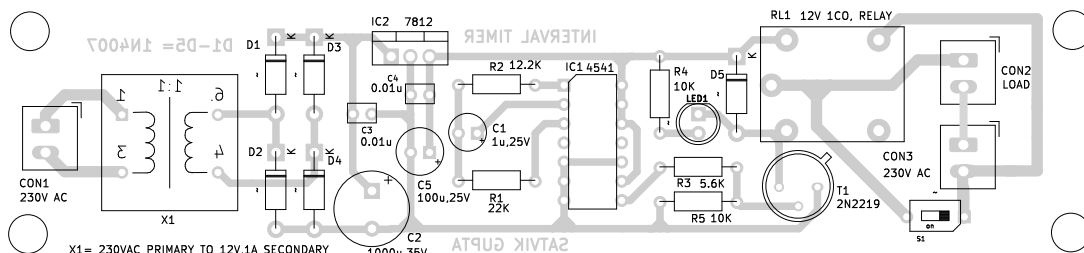
C1: 1 μF/16 V elektrolityczny
C2: 1000 μF/25 V elektrolityczny
C3, C4: 0,01 μF ceramiczny
C5: 100 μF/25 V elektrolityczny

Inne:

CON1-CON3: złącze 2-piniowe
S1: ON/OFF switch, obciążenie do 6 A/230 VAC; ze złączem 2-pinowym
RL1: Przełącznik 12 V ze stykami NO
X1: Transformator sieciowy 230 VAC – uzwojenie pierwotne, 12 VAC/1 A wtórne



Rysunek 3. Projekt druku płytki PCB wyłącznika czasowego



Rysunek 4. Schemat montażowy elementów na PCB

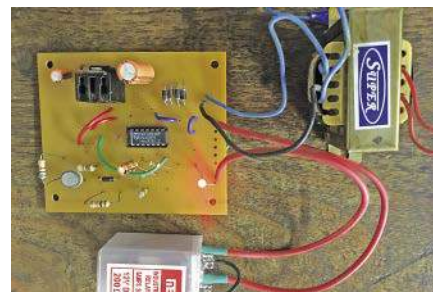
uruchamia przełącznik. Urządzenie wykonawcze (zwykle będzie to światło lub wentylator) załączane jest poprzez styki NO przełącznika i odrębne zasilanie podane na złącze CON3. Po zadanim czasie (ok. 15-tu minut) wyjście pin 8 IC1 przejdzie do stanu niskiego i RL1 zostanie wyłączony. Tym samym zostanie wyłączone obciążenie (obwód wykonawczy). Ponowne włączenie „światła” wymaga chwilowego odłączenia zasilania z CON1 oraz ponowne jego załączenie.

Ta procedura nie zawsze musi być wygodna i pożądana. Równolegle do styków przełącznika zamontowano „normalny” przełącznik – switch S1. Chcąc korzystać z timera powinien być on w pozycji wyłączzonej. Natomiast jego obecność umożliwia załączenie obciążenia na dowolnie długi czas. Równolegle do cewki przełącznika wstawiono także diodę LED1, której świecenie informuje o pracy timera.

Konstrukcja i testowanie pracy układu

Dla naszego wyłącznika czasowego przygotowano płytkę drukowaną PCB, którą w skali 1:1 pokazuje rysunek 3. Na rysunku 4 widzimy schemat montażowy ułożenia elementów na PCB.

Po zmontowaniu układu, należy go zamknąć w odpowiednio przygotowanej obudowie. Można także do montażu wykorzystać płytkę uniwersalną. Tak uczynił autor projektu i efekt widoczny jest na rysunku 5. Wszystkie komponenty widoczne na tym rysunku należy umieścić w obudowie. Napięcie 230 VAC należy doprowadzić oddzielnie do złącza zasilającego CON1 i wykonawczego CON3. W widocznym miejscu należy umieścić diodę LED1 oraz dodatkowy przełącznik „bypassu”



Rysunek 5. Prototyp wyłącznika czasowego zmontowany przez autora projektu

S1. Po zmontowaniu i przetestowaniu całości należy dokończyć prace montażowe i uzbroić wszystkie złącza tak, iż fizycznie wyłącznik czasowy ułożony jest między zasilaniem i odbiornikiem energii.

Uwaga od Redakcji „Electronics For You”:

Styki przełącznika RL1 mają dopuszczalną obciążalność prądową do 6 A AC. Dlatego także obciążenie (oświetlenie bądź moc wentylatora) nie powinno przekroczyć tej wartości lub należy zastosować przełącznik o większej obciążalności styków wykonawczych.

Od Redakcji EdW: Resetowanie timera w istocie uruchamia odmierzenie czasu. Wbieżącym projekcie reset realizowany jest przez chwilowy zanik zasilania. Wydaje się, iż taki sposób może nie być wygodny. Tymczasem bez większych trudności można układ zmodyfikować tak, aby był wyzwalany stanem logicznym z dodatkowego switcha. W układzie scalonym CD4541 można w tym celu wykorzystać np. wejście Master Reset, a dezaktywować Auto Reset stanem wysokim na pinie 5 IC1. ■

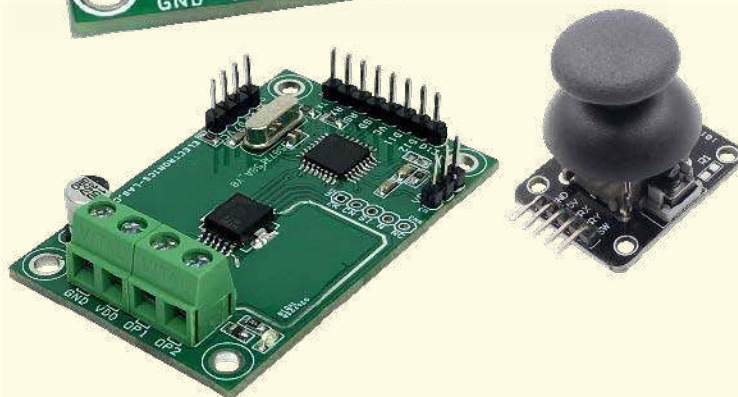
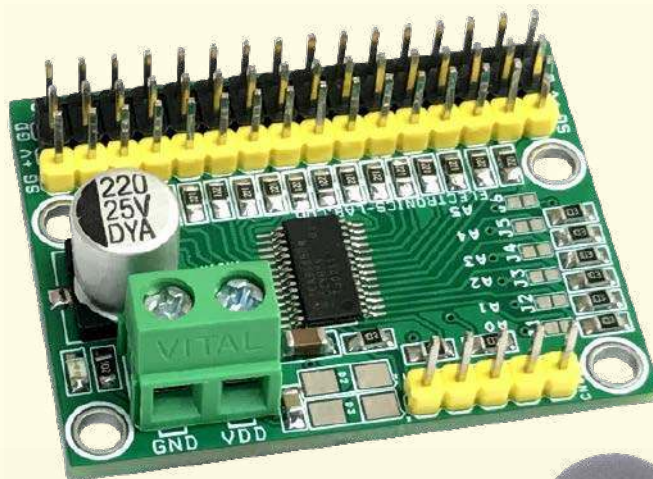
P. Balasubramanian

Ten artykuł był wcześniej opublikowany na łamach „EFY”, sierpień 2023 (efymag.com)

Przedstawiamy początkowe fragmenty dwóch projektów ze zbioru kilkudziesięciu projektów dostępnych wyłącznie dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl. W rubryce **DIY PLUS** zamieszczamy aktualnie najciekawsze projekty publikowane w Internecie w formacie open source. Prenumeratorów EdW zapraszamy do zapoznania się na www.elportal.pl z niezwykle inspirującymi zasobami rubryki **DIY PLUS**.

16-kanalowy sterownik serwomechanizmów RC z interfejsem I²C

Jest to 16-kanalowy sterownik serwomechanizmów, który może sterować 16 serwomechanizmami RC za pośrednictwem interfejsu I²C. Projekt został zbudowany przy użyciu układu PCA9685, 16-kanalowego generatora PWM, który może jednocześnie sterować 16-kanalowymi serwomechanizmami. Płytkę można podłączyć do Arduino lub innego mikrokontrolera. Sterowanie robotami i lalkami animatronicznymi jest łatwe dzięki tej płytce. Płytkę działa z napięciem 5 VDC, a oddzielne złącze zasilania służy do zasilania serwomechanizmu RC napięciem 6 VDC. 6× adres Zworka lutownicza służy do ustawiania adresu płytki, co pomaga użytkownikom podłączyć 62 płytki na jednej linii I²C i obsługiwać do 992 serwomechanizmów. Płytkę można również używać do innych zastosowań wymagających 16-kanalowych sygnałów PWM. Częstotliwość można regulować w zakresie od 24 Hz do 1526 Hz, a cykl pracy wynosi od 0 do 100%.



Sterowanie silnikiem DC za pomocą joysticka

Jest to kompaktowy, niskoprofilowy i wysokowydajny sterownik silnika DC, który zapewnia łatwe sterowanie silnikiem DC za pomocą joysticka. Ta niewielka płytka może sterować małymi i średnimi szczotkowymi silnikami prądu stałego o natężeniu ciągłym do 2,5 A i szczytowym do 6 A. Jest to projekt kompatybilny z Arduino, który pomaga użytkownikom opracować własny kod i sterować silnikiem szczotkowym DC zgodnie z wymaganiami aplikacji. Sprzęt kompatybilny z Arduino składa się z mikrokontrolera ATMEGA328 i układu IFX9201 H-Bridge.

Niektóre projekty aktualnie dostępne tylko dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl:

1. Programowalny kondycjoner sygnału z czujnika rezystancyjnego mostkowego
2. 20-segmentowy wyświetlacz słupkowy w rozmiarze jumbo
3. Stacja pogodowa lilygo ttgo t5-4.7 z wyświetlaczem typu e-papier
4. Półprzewodnikowy przekaźnik mocy DC z prądowym sprzężeniem zwrotnym
5. Wyłącznik nadprądowy – przekaźnik wyłączający nadprądowy
6. Uniwersalny konwerter napięcia AC – wyjście 18 V DC z wejścia 85...265 V AC
7. Moduł procesora echa głosu – urządzenie opóźniające do efektów dźwiękowych, echo, reverb
8. Najlepszy sposób na próbkowanie dźwięku za pomocą ESP32
9. Sterownik silnika krokowego z joystickiem
10. Choinka z Arduino i pikselowymi diodami
11. RPI – stacja pogodowa IoT
12. Niskobudżetowy monitor jakości powietrza IoT oparty o RaspberryPi 4
13. Automatyczny system ogrodnictwa z NodeMCU i Blynk, ArduFarmBot 2
14. TinyML – Rozpoznawanie ruchu przy pomocy Raspberry Pi Pico
15. Wzmacniacz piezoelektryczny do gitary i skrzypiec
16. Wysokowydajny i niezawodny sterownik bipolarnego silnika krokowego
17. Sterownik silnika prądu stałego z wykorzystaniem przekaźnika i mosfetu – interfejs Arduino
18. Przedwzmacniacz do mikrofonu MEMS
19. Super prosty czuły wykrywacz metali
20. Stymulator czaszkowy Arduino (Bio-BrainTuner)
21. Generator sygnałów AD9833
22. Obserwacja charakterystyk tranzystora
23. Wyświetlacz EKG z użyciem Arduino
24. Łatwy do zbudowania robot krocący
25. Sonarowy theremin MIDI
26. Zamek elektroniczny na kod
27. Prosty tester tranzystorów
28. Zegar binarny z użyciem Microbit
29. Przetwornik częstotliwości na napięcie (tachometr) – przetwornik częstotliwości na napięcie z czujnikiem magnetycznym o zmiennej reluktancji
30. Izolowany obwód wykrywania napięcia 250 V AC z pojedynczym wyjściem (wejście 250 V prądu przemiennego, wyjście 5 V)

Miesięcznik „Elektronika dla Wszystkich” (12 numerów w roku) jest wydawany we współpracy z kilkoma redakcjami zagranicznymi



Wydawnictwo:
AVT-Korporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Adres redakcji:
03-197 Warszawa, ul. Leszczyńska 11
e-mail: edw@elportal.pl, www.elportal.pl

Redaktor merytoryczny:
Mariusz Ciszewski, Paweł Sujko

Dział Reklamy:
Katarzyna Gugała
katarzyna.gugala@elportal.pl, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobański
jakub.sobanski@elportal.pl

Copyright AVT-Korporacja Sp. z o.o., Warszawa, ul. Leszczyńska 11. Projekty publikowane w „Elektronice dla Wszystkich” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki dla Wszystkich”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice dla Wszystkich” jest dozwolone wyłącznie po uzyskaniu pisemnej zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice dla Wszystkich”.

DTP, okładka, redakcja strony internetowej www.elportal.pl:
MAD Sp. z o.o.

Prenumerata:
W Wydawnictwie AVT, e-mail: prenumerata@avt.pl
tel. 22 257 84 22, (godz. 10:00–14:00)

W RUCH S.A., e-mail: prenumerata@ruch.com.pl
tel. 801 800 803, 22 717 59 59, www.prenumerata.ruch.com.pl



O PCBWAY:

PCBWay to dostawca usług w zakresie produkcji płytek drukowanych, montażu PCB oraz integracji produktów końcowych.

W czasie ponad 10 lat rynkowej obecności firma osiągnęła czołową pozycję w branży.



 Adres URL:
www.pcbway.com

 Poczta:
service@pcbway.com

DOSTĘPNE

- KRÓTKIE CZASY REALIZACJI ◀
- NAJWYŻSZEJ JAKOŚCI ◀
- OBWODY DRUKOWANE
- WSPARCIE DLA KLIENTÓW 24/7 ◀
- PROSTY I BEZPROBLEMOWY ◀
- PROCES ZAMÓWIEŃ