



ELEKTRONIKA

dla wszystkich

nr 05/2026 (364) • maj • www.elportal.pl

DIY PLUS
tylko dla prenumeratorów

PROJEKTY dla elektroników

- ▶ Liniowy zasilacz laboratoryjny 0...50 V/0...2 A oraz regulowane zasilanie symetryczne
- ▶ Miernik indukcyjności, pojemności i ESR
- ▶ Wysokowydajny subwoofer aktywny do domowego sprzętu Hi-Fi, część 2

DIY dla wszystkich

- ▶ Zdalnie sterowany sygnalizator na czas pandemii koronawirusa
- ▶ Automatyczny kran umywalkowy z oświetleniem LED
- ▶ Ściemniacz taśmy LED i dzwonek

TUTORIALE

- ▶ Charakterograf diod półprzewodnikowych – przystawka do oscyloskopu
- ▶ Ekscytacje Maxa: Migające diody LED i śliniacy się inżynierowie, część 31
- ▶ Elektroniczne bloki konstrukcyjne: Wybór i stosowanie siłowników, część 2. Przekształcanie sygnałów elektronicznych w ruch fizyczny – liniowy i obrotowy
- ▶ Wynalazcy i ich wynalazki w elektronice, część 3

Liniowy zasilacz laboratoryjny 0...50 V/0...2 A



Miernik indukcyjności, pojemności i ESR



ISSN 1425-1698 Indeks 33362X

nr 05/2026 • cena: 19,90 zł (w tym 8% VAT)

Pomocna dłoń



automatykaB2B.pl

EP.com.pl

Największy portal dla elektroników konstruktorów

eprasa.pl 6a3462f436



FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

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

www.piekarz.pl



TRZECIARĘKA ZD-11P
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt z latarką, ZD11P



TRZECIARĘKA ZD-11P-1
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt z latarką i lupą, ZD11P-1



TRZECIARĘKA SN-394
Uchwyt montażowy typu „Trzecia ręka”,
pająk z lupą 50 mm, przykręcany do blatu
Proskit SN-394

BESTSELLERY sklepu AVT – sklep.avt.pl

Trzecia ręka

Rabat dla Czytelników EdW
przy zakupie podaj kod **EdW2505TR**

Kod ważny do 30.09.2025

-3%

Rabat dla Prenumeratorów EdW
przy zakupie podaj numer prenumeraty

-6%



TRZECIARĘKA ZD-11M-1
Uchwyt montażowy typu „Trzecia ręka”,
pająk – z uchwytem na szpulkę cyny, ZD11M-1



TRZECIARĘKA ZD-11M-2
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt z lupą i podświetleniem LED
ZD11M-2



TRZECIARĘKA ZD-11M-3
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt z lupą i podświetleniem LED
ZD-11M-3



TRZECIARĘKA ZD-11M
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt ZD11M



TRZECIARĘKA SN-392
Uchwyt montażowy typu „Trzecia ręka”
z lupą 90 mm, Proskit SN-392



TRZECIARĘKA
Uchwyt montażowy typu „Trzecia ręka”
z lupą 60 mm

-15%
NA START
203,00 zł

-30%
po pierwszym roku
prenumeraty
167,20 zł

-40%
po drugim roku
prenumeraty
143,30 zł

-50%
po trzecim roku
nieprzerwanej prenumeraty
119,40 zł

Odkryj korzyści z **prenumeraty drukowanej** – **większe oszczędności z każdym rokiem!**

Rozpocznij swoją przygodę z *Elektroniką dla Wszystkich*. Decydując się teraz na roczną prenumeratę drukowaną, otrzymasz nie tylko dostęp do najnowszych wydań, ale i **znakomity start dzięki zniżce 15%** na pierwsze zamówienie!

Prenumerata to nie tylko wygoda dostępu do treści, ale także sposób na znaczące oszczędności. Dołącz do grona naszych stałych czytelników i ciesz się coraz lepszymi warunkami.

Im dłużej jesteś z nami, tym więcej oszczędzasz:

- po roku nieprzerwanej prenumeraty zapewnimy Ci **30% rabatu** na kolejny rok,
- po dwóch latach wierności zaoferujemy **40% rabatu**,
- po trzech latach lojalności osiągniesz **najwyższy poziom rabatu – 50%**!

Jak otrzymać rabat za lojalność?

Zaloguj się na swoje konto prenumeratora na www.UlubionyKiosk.pl i zamów prenumeratę, korzystając z przycisku PRZEDŁUŻ w zakładce „Prenumeraty”.

Przeglądaj wcześniej, płać mniej – **postaw na e-prenumeratę!**

Wybierz prenumeratę cyfrową PDF i ciesz się dostępem do czasopisma nawet 7 dni przed oficjalną premierą w kioskach. Oszczędzaj czas i pieniądze – skorzystaj z **rabatu 30%** na roczną e-prenumeratę w cenie 133,60 zł.

Dodatkowa oferta dla prenumeratorów wersji drukowanej: jeśli już subskrybujesz wersję papierową, możesz dokupić równolegle e-wydania w cenie 38,20 zł/rok – z **niesamowitym rabatem 80%**.

Zyskaj nieograniczony dostęp do zasobów dla pasjonatów elektroniki!

Tylko prenumeratorzy mają pełny dostęp do:

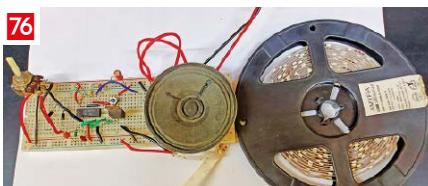
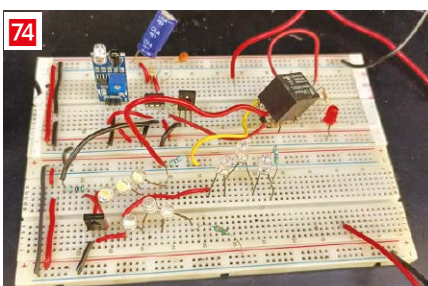
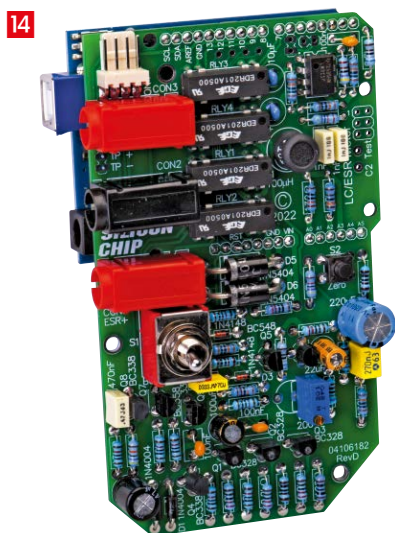
- cyfrowego archiwum *Elektroniki dla Wszystkich* na www.elportal.pl/archiwum
- projektów DIY+ na www.elportal.pl/diy

Zamów prenumeratę drukowaną lub e-prenumeratę na www.UlubionyKiosk.pl lub przez przelew na konto Wydawnictwa AVT, a po zaksięgowaniu wpłaty wyślemy Ci mailowo kod dostępu do portalu.

ARCHIWUM



Zacznij korzystać z pełnych zasobów już dziś!



Projekty dla elektroników:

Liniowy zasilacz laboratoryjny	
0...50 V/0...2 A oraz regulowane zasilanie symetryczne	8
Miernik indukcyjności, pojemności i ESR.....	14
Wysokowydajny subwoofer aktywny	
do domowego sprzętu Hi-Fi, część 2	26

Tutoriale:

Charakterograf diod półprzewodnikowych	
- przystawka do oscyloskopu.....	34
Ekscytacje Maxa:	
Migające diody LED i śliniacy się inżynierowie, część 31	37
Elektroniczne bloki konstrukcyjne:	
Wybór i stosowanie siłowników, część 2. Przekształcanie sygnałów	
elektronicznych w ruch fizyczny - liniowy i obrotowy	43
Wynalazcy i ich wynalazki w elektronice, część 3.....	48
Edukacja w EdW dla szkół i uczelni.	
Wykład 41: PLL - pętla synchronizacji fazowej.....	62

DIY dla wszystkich:

Zdalnie sterowany sygnalizator na czas pandemii koronawirusa.....	72
Automatyczny kran umywalkowy z oświetleniem LED	74
Ściemniacz taśmy LED i dzwonek	76

Junior:

Dwudzieste trzecie spotkanie z najmłodszymi pasjonatami elektroniki.....	78
Na zdjęciu na okładce Krystian - Młodzi Entuzjaści Elektroniki, Szkoła Podstawowa nr 86, Wrocław	

DIY PLUS tylko dla prenumeratorów zamawiających prenumeratę na www.UlubionyKiosk.pl	
Wszechstronny odbiornik światłowodowy Versatile Link	90
Wszechstronny nadajnik światłowodowy Versatile Link	90

Rubryki stałe:

Prenumerata	3
Od redakcji.....	5
Poczta.....	7

A za miesiąc w czerwcowym EdW



Wielokanałowy regulator głośności, część 1

Masz kilka kanałów audio i dość waleki z wieloma potencjometrami? Ten układ pozwala sterować nawet dwudziestoma kanałami jednocześnie. Ekran dotykowy, enkoder i pilot IR zapewniają wygodną obsługę, a modułowa konstrukcja ułatwia rozbudowę. W tej części: zasada działania i architektura rozwiązania.

Raspberry Pi Clock Radio, część 1

Zwykły budzik to już za mało! Ten projekt oparty na Raspberry Pi oferuje znacznie więcej: alarmy z własnymi ustawieniami oraz odtwarzanie dźwięku z sieci lokalnej, z Internetu i przez Bluetooth. Do tego sterowanie przez przeglądarkę i wbudowany wzmacniacz stereo. W tej części: koncepcja i architektura.

Przedwzmacniacz mikrofonowy

Niewielkie urządzenie, a możliwości do pracy na scenie i w studiu. Ten przedwzmacniacz zwiększa poziom sygnału z mikrofonu, oferuje wyjście symetryczne lub niesymetryczne oraz regulowane wzmocnienie w szerokim zakresie. Do tego zasilanie phantom 48 V, bardzo niski poziom zniekształceń i szumów oraz solidna konstrukcja do pracy w trudnych warunkach.

Dla studentów i uczniów: garść technicznych tutoriali

Dla szukających inspiracji: ciekawe projekty DIY

Dla najmłodszych: kolejny zestaw z serii AVTEDU

**W kioskach
od 29 maja**

Maj – od Saskiej Kępy do targowych hal

To był maj... „pachniała” Saska Kępa, szalonym, zielonym bżem” – i choć w tej historii więcej jest rozstania niż spełnienia, zostaje po niej coś znacznie cenniejszego: energia początku, momentu, kiedy na dobrą sprawę wszystko wciąż jest możliwe, może nawet bardziej niż kiedykolwiek wcześniej. Jest w tym jednocześnie nuta nostalgii i dziwnej radości, że coś się wydarzyło, nawet jeśli nie do końca tak, jak miało.

Podczas gdy słońce coraz śmielej zagląda do okien, gdzieś równolegle toczy się drugi świat – hale pełne robotów, automatyki i elektroniki użytkowej, które za chwilę trafią do naszych domów. W tym miesiącu widać to szczególnie wyraźnie: podczas Warsaw Industry Automatica 2026 królują systemy sterowania i robotyka, a kilka dni później Electronics Show 2026 pokazuje elektronikę użytkową i rozwiązania smart home w najbardziej namacalnej formie. To właśnie tam widać, jak idee znane z warsztatowych eksperymentów dojrzewają do roli gotowych produktów.

U nas trochę spokojniej, może nieco bardziej kameralnie, ale nie mniej klimatycznie. Przyglądając się galopującej technice i kibicując jej, zapraszamy na kolejny przegląd projektów z różnych zakątków świata – starannie i na bieżąco selekcjonowanych z myślą o naszych Czytelnikach.

W rytm rzeczony piosenki wypada przejść do bieżących tematów z przytupem. A zatem co czeka nas w bieżącym wydaniu?

Phil Prosser wieńczy budowę aktywnego subwoofera Hi-Fi. To moment, w którym teoria zamienia się w praktykę: montaż wzmacniacza Ultra-LD, zasilacza i całej elektroniki w obudowie, a na końcu – pierwsze odsłuchy. Czyli moment, na który czeka się od samego początku.

Z szacunkiem dla sprawdzonych rozwiązań Steve Griffin wraca do swojego projektu sprzed ponad 40 lat. Liniowy zasilacz laboratoryjny napięć stałych w zakresie 0...50 V oraz symetrycznego napięcia ± 15 V pokazuje, że dobra koncepcja nie potrzebuje rewolucji – wystarczy ją uporządkować i dostosować do współczesnych realiów.

Warsztat uzupełnia Steve Matthysen, prezentując miernik LC z pomiarem ESR. Przydaje się wszędzie tam, gdzie zwykły pomiar pojemności już nie wystarczy.

Swoją opowieść o historii elektroniki kontynuuje dr David Maddison. Czy wiedzieliście – albo czy pamiętacie – że pierwszy dysk twardy mieścił zaledwie 3,75 MB danych i ważył około tony? Że pierwszy kalkulator kieszonkowy nie miał wyświetlacza, tylko drukował wynik na papierowej taśmie? I że pierwszy komputer z graficznym interfejsem kosztował tyle, co dziś dobre mieszkanie? Trzecia część zamyka tę serię i przypomina, jak wiele musiało się wydarzyć, by elektroniczna rzeczywistość osiągnęła obecny stan.

Po tej historycznej podróży wracamy do codzienności. Julian Edgar pokazuje, jak dobrać odpowiedni siłownik do konkretnego zastosowania. To wiedza, która pozwala uniknąć błędów, zanim jeszcze powstanie pierwszy prototyp. Nieco szerszą perspektywę proponuje Max Maxfield – od prostego serwo mechanizmu aż po roboty humanoidalne. Bo niezależnie od skali projektu, wszystko zaczyna się od skrupulatnie czynionych założeń projektowych.

Jos Verstraten prowadzi przez zagadnienia PLL, pokazując, jak układy potrafią same „doganiać” sygnał odniesienia i odzyskiwać informacje.

Nie brakuje też prostych konstrukcji DIY. Suresh Chandra Dwivedi oraz Rakesh Jain proponują bezdotykowe sterowanie wodą i światłem oraz praktyczne układy oparte na NE555 i LM556 – nieskomplikowane, a przy tym potencjalnie użyteczne.

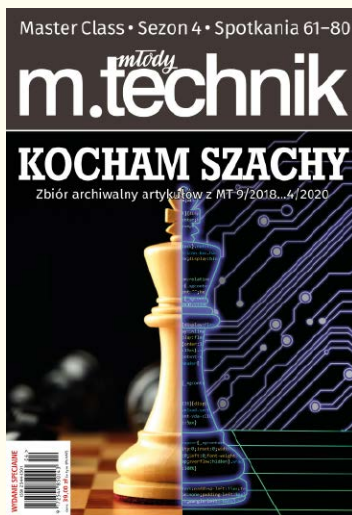
Juniorzy złożą w tym miesiącu kolejny układ reagujący na dźwięk. Diody LED reagują tu na zmiany sygnału, takie jak jego poziom i szybkość zmian. Co więcej, trzy grupy diod zachowują się nieco inaczej, dzięki czemu efekt świetlny jest bardziej zróżnicowany. W efekcie układ nie tylko miga w rytm, ale także odzwierciedla dynamikę dźwięku. To nieskomplikowany projekt, który dobrze nadaje się do pierwszych eksperymentów z elektroniką analogową i pokazuje, jak sygnał akustyczny można zamienić w efekt świetlny.

Maj ma w sobie coś szczególnego. Z jednej strony patrzymy na to, co dzieje się w wielkich halach targowych, gdzie technika pędzi do przodu. Z drugiej – wracamy do jej historii i własnych projektów, często prostszych, ale nie mniej satysfakcjonujących. I może właśnie w tym tkwi równowaga: między światem gotowych produktów a tym, co powstaje w naszych przydomowych warsztatach.

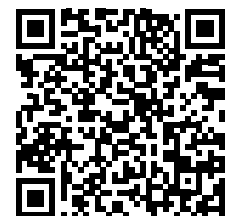
Bo ostatecznie najważniejszy moment w elektronice wciąż wygląda tak samo – to chwila, gdy do nowo zbudowanego układu po raz pierwszy podłączamy zasilanie.

A więc... do dzieła.

Mariusz Ciszewski



pakiet promocyjny
KOCHAM SZACHY
7 e-booków z rabatem
50%



Dla prenumeratorów – 30% rabatu!

Promocja internetowa – w formularzu zamówienia online zaznacz pole „Jestem prenumeratorem wydawnictwa AVT, kupuję ze zniżką” i podaj swój numer prenumeraty.

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

Centrala alarmowa AVT5252 i drugi manipulator

Szanowna Redakcjo,
piszę w sprawie centrali alarmowej AVT5252. Kupiłem ją jakiś czas temu i po zamontowaniu działa bez zarzutu.

Chciałbym zapytać, czy istnieje możliwość rozbudowy centrali o drugi wyświetlacz i klawiaturę. Zależy mi na możliwości rozbrajania alarmu z drugiego pomieszczenia – odległość przewodu komunikacyjnego wynosi około 15 m.

Alternatywnie rozważam inne rozwiązanie, na przykład system zdalnego rozbrajania drogą radiową. Czy taka modyfikacja byłaby możliwa do zrealizowania?

Jeśli temat mógłby zostać poruszony na łamach czasopisma, byłoby to dla mnie bardzo cenne. Nie jestem zaawansowanym elektronikiem, a programowanie nie jest moją mocną stroną, dlatego liczę na możliwie proste rozwiązanie, które pozwoliłoby mi – oraz innym użytkownikom – samodzielnie przeprowadzić taką modernizację.

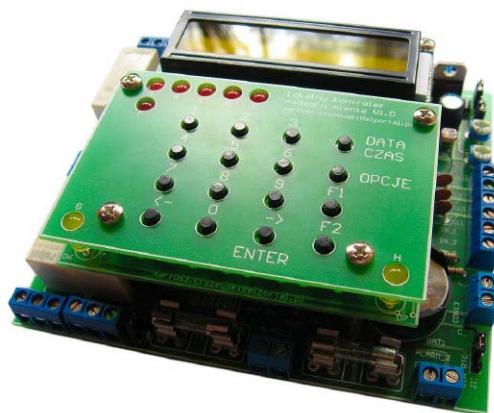
Z wyrazami szacunku,
Adam Małota

Red. Dziękujemy za list.

Wspomniana centrala została opublikowana na łamach siostrzanej EP w sierpniu 2010 roku, projekt ma zatem blisko 16 lat. Dobrze słyszeć, że jego egzemplarze są do dziś wykorzystywane i dobrze się sprawują. Trzeba przyznać, że zastosowane w centrali rozwiązanie automatycznego samouzbrajania były (i są nadal) naprawdę ciekawe, o czym można się przekonać, zaglądając do dostępnej w Internecie dokumentacji projektu: <https://serwis.avt.pl/manuals/AVT5252.pdf>.

Patrząc z czysto technicznego punktu widzenia, projekt centrali AVT-5252 nie przewiduje wprost możliwości dołączenia drugiego panelu operatorskiego (klawiatury z wyświetlaczem). Klawiatura jest co prawda zrealizowana jako osobny moduł komunikujący się przez magistralę I²C, jednak oprogramowanie obsługuje tylko jeden zestaw adresów urządzeń. Wyświetlacz natomiast jest sterowany bezpośrednio z mikrokontrolera, co praktycznie wyklucza jego proste zdublowanie. Pewną furtką do rozbudowy mógłby być interfejs RS-485 wyprowadzony na złącze CON6 i wspomniany w artykule jako element systemu sieciowego, jednak w podstawowej wersji oprogramowania nie jest on zaimplementowany. Stanowił on przygotowanie części sprzętowej pod ewentualną rozbudowę

oprogramowania, która jednak nigdy nie nastąpiła. Co ciekawe, urządzenie posiada również inne złącza, które po odpowiedniej rozbudowie oprogramowania mogłyby posłużyć do zewnętrznego sterowania centralą. Są to złącza CON2 (RESET) oraz CON3 (ON/OFF). Ze schematu wynika, że złącza te pozwalają zwierzać linie mikrokontrolera ATmega32 – odpowiednio PB0 i PB1 (domyślnie podciągnięte rezystorami R32 i R33 do VCC) – do masy. Wystarczyłoby zatem zaimplementować obsługę tych złączy w oprogramowaniu układowym (firmware mikrokontrolera), a następnie sprzęgnąć centralę, na przykład z modułem radiowym. Dla bezpieczeństwa należałoby wykorzystać zestaw zdalnego sterowania (piloty i odbiornik) pracujące z kodem zmiennym (rolling code), na przykład w standardzie KeeLoq. Warto wspomnieć, że KeeLoq to algorytm szyfrujący oraz system autoryzacji wykorzystywany w pilotach radiowych (alarmy, centralne zamki, bramy) i stosowany często w układach firmy Microchip Technology. Zamiast modułu radiowego można również rozważyć moduł Ethernet, który umożliwiłby realizację podstawowej funkcji uzbrajania i rozbrajania centrali przez sieć lokalną lub Internet. Należy jednak zaznaczyć, że chodzi o uproszczoną funkcję pozwalającą uzbroić lub rozbroić centralę, a nie zapewnić pełną dwukierunkową komunikację. W przypadku chęci uzbrajania/rozbrajania centrali lokalnie – z innego pomieszczenia w obrębie budynku lub z ogrodu mogłoby okazać się wystarczające, ponieważ stan uzbrojenia lub rozbrojenia centrali byłby sygnalizowany optycznie i dźwiękowo. Gorzej ze sterowaniem przez Internet z innej lokalizacji (brak informacji zwrotnej o uzbrojeniu bądź rozbrojeniu centrali). Ponadto komunikację dwukierunkową (klawiatura i wyświetlacz) można by w teorii zaimplementować z wykorzystaniem złącza RS-485. O ile funkcje prostego sterowania za pomocą złączy CON2 i CON3 dałoby się stosunkowo łatwo zaimplementować w oprogramowaniu układowym, o tyle dla implementacji obsługi manipulatorów po szynie RS-485 bardzo ograniczona pamięć flash zastosowanego mikrokontrolera okazałaby się tu bardzo wąskim gardłem. Puszczając wodze fantazji, można by zaprojektować moduł pozwalający w miejsce mikrokontrolera ATmega32 wpiąć popularne dziś rozwiązania, takie jak ESP32, Raspberry Pi Pico W lub Raspberry Pi Zero i opracować oprogramowanie centrali zupełnie od nowa, tym razem



w pełni korzystając z dobrodziejstw współczesnego świata elektroniki, w tym mając do dyspozycji dużą ilość pamięci na oprogramowanie układowe, łatwą modernizację oprogramowania centrali i zdalne jego aktualizacje. Ponadto można byłoby połączyć centralę ze światem zewnętrznym za pomocą dostępnych w takich modułach interfejsów Bluetooth lub Wi-Fi, już bez konieczności implementacji obsługi wspomnianych wcześniej złączy CON2, CON3 czy CON6. Warto jednak mieć na uwadze, że wraz z wygodą i wzrastającą liczbą nowych funkcjonalności centrali wzrośnie również liczba możliwych dróg i obszarów uzyskania nieautoryzowanego dostępu, co w przypadku tego typu urządzeń może uczynić je atrakcyjnym celem dla potencjalnego włamywacza. Warto zauważyć, że nawet implementacja obsługi złączy CON2 i CON3 i sterowanie centralą przez zwieranie sygnałów do masy – po zdjęciu obudowy lub przez wykonane w niej otwory zwiększa podatność systemu na potencjalny atak.

Możliwości, jak widać, jest wiele, ale żadna z nich nie jest dostępna „od ręki”. Gdyby zainteresowanie było naprawdę spore, moglibyśmy rozważyć możliwość implementacji obsługi złączy CON2 i CON3, co należy jednak poważnie przemyśleć z uwagi na wszystkie tematy wspomniane powyżej.

Młodszych Czytelników, dysponujących większą ilością wolnego czasu oraz chęcią – a może potrzebą – stworzenia własnego elektronicznego portfolio, zachęcamy do odtworzenia centrali AVT5252 z wykorzystaniem współczesnych dobrodziejstw techniki, na przykład modułów opartych na Raspberry Pi. Chętnie opublikujemy projekt stanowiący kontynuację tej dość ciekawej i na swój sposób unikatowej centrali.

Redakcja EdW



Linowy zasilacz laboratoryjny

0...50 V/0...2 A oraz regulowane zasilanie symetryczne

Zajrzałem niedawno do wnętrza zasilacza laboratoryjnego, zbudowanego przeze mnie 40 lat temu, i stwierdziłem, że jego konstrukcja nie spełnia już aktualnych standardów. Postanowiłem zrobić wszystko od nowa, używając nowoczesnych elementów i upraszczając okablowanie. Tak powstał przedstawiony tutaj projekt.

W 1980 roku, przeglądając letnie wydanie Elektor Summer Circuits, natknąłem się na to, czego wtedy szukałem: przyzwoity zasilacz do mojej pracowni elektronicznej [1]. Miał to być mój pierwszy poważny projekt (tzn. taki z obudową). Po dodaniu prostego, niskoprądowego zasilacza symetrycznego przeznaczonego dla wzmacniaczy operacyjnych, jedna skrzynka zaspokajałaby wszystkie moje potrzeby w zakresie zasilania. Po kilku tygodniach tworzenia płytki drukowanej, dodaniu miernika z ruchomą cewką, montażu i poprowadzeniu okablowania, wszystko było gotowe. Działał jak marzenie – przez ponad 40 lat.

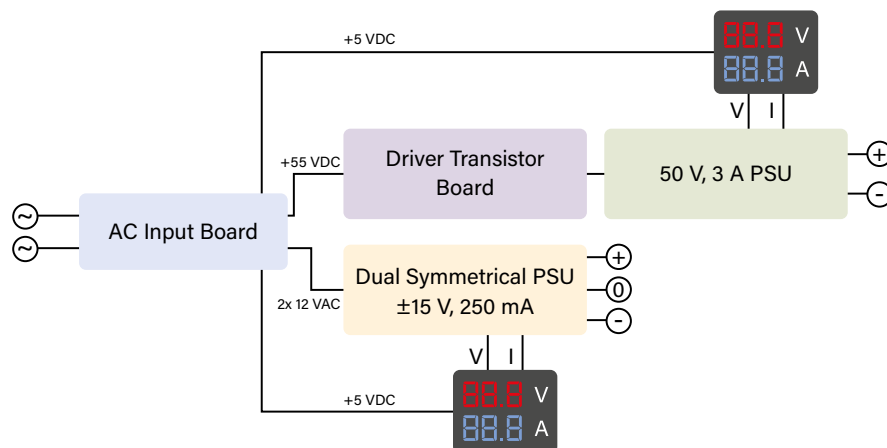
40 lat później...

... ząb czasu wyraźnie dał o sobie znać. Przede wszystkim wytarły się opisy wykonane wyklejkami. Dlatego uznałem, że obudowa wymaga odświeżenia. Zdjąłem przedni panel z zamiarem wykonania nowego (z użyciem wspomagania komputerowego). Kiedy jednak zajrzałem do środka zasilacza, byłem przerażony tym, co zobaczyłem. „Czy ja naprawdę tak to podłączyłem?” Zasilacza z pewnością się

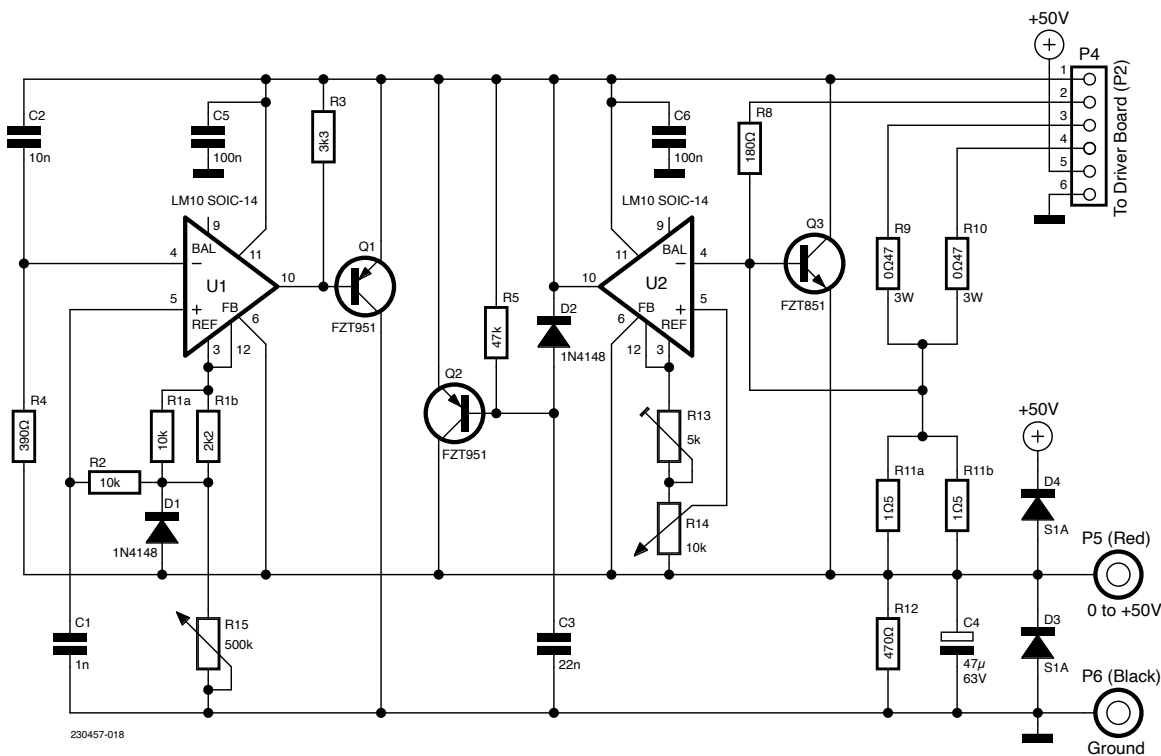
nie dałoby się rozebrać i ponownie złożyć w ten sposób. Pojawił się pomysł całkowitej przebudowy. Jeśli miałem to zrobić, musiałem użyć elementów nowoczesnych, a więc podzespołów do montażu powierzchniowego (SMD) i współczesnych złączy. Schemat blokowy mojego nowego zasilacza pokazuje **rysunek 1**. Zaczniemy od bloku zasilacza 0...50 V.

Nowa wersja zasilacza 0...50 V

Jednym z moich głównych założeń było uporządkowanie okablowania. Doprowadziło to do podziału tego bloku na trzy części: główną płytkę stabilizatora, płytkę sterownika oraz radiator tranzystorów mocy. Wszystkie te elementy zostały ze sobą połączone poprzez uporządkowane okablowanie i złącza wielostykowe.



Rysunek 1. Schemat blokowy zasilacza laboratoryjnego



Rysunek 2. Układ stabilizatora 0...50 V. Wartości R1a i R1b dobrałem doświadczalnie, używając posiadanych rezystorów. Dobór tych rezystorów ustala maksymalne napięcie wyjściowe na poziomie 50 V i zapewnia przepływ prądu 100 μ A przez potencjometr P1, pod warunkiem że jego rezystancja wynosi dokładnie 500 k Ω

Mogłem oczywiście kupić gotowy zasilacz impulsowy (SMPSU) za rozsądną kwotę pieniędzy. Po co więc się męczyłem? Odpowiedź jest prosta: elektronika była moim hobby przez całe życie, a tworzenie przydatnych urządzeń elektronicznych jest czymś, co zawsze sprawia mi ogromną satysfakcję.

UWAGA: Projekt wymaga użycia transformatorów zasilanych z sieci. Osoby nie mające doświadczenia w pracy z napięciem sieciowym nie powinny samodzielnie realizować tego projektu; w tej części montażu powinny skorzystać z pomocy kogoś bardziej doświadczonego.

Zmiany w projekcie

Nowy układ (rysunek 2 – płytką stabilizatora i rysunek 3 – płytką układu wykonawczego) dość wiernie naśladuje projekt z 1980 roku. Zachowałem nawet oryginalną numerację elementów. Oczywiście różnicą jest zastosowanie podzespołów SMD. Spowodowało to pewne problemy, ponieważ większość półprzewodników użytych w oryginalnym projekcie nie jest dostępna w formie SMD, co wymagało znalezienia odpowiedników. Jedynymi elementami, które przetrwały aktualizację, były: transformator, kondensator za prostownikiem, niezawodne tranzystory mocy 2N3055 oraz obudowa.

Potencjometry

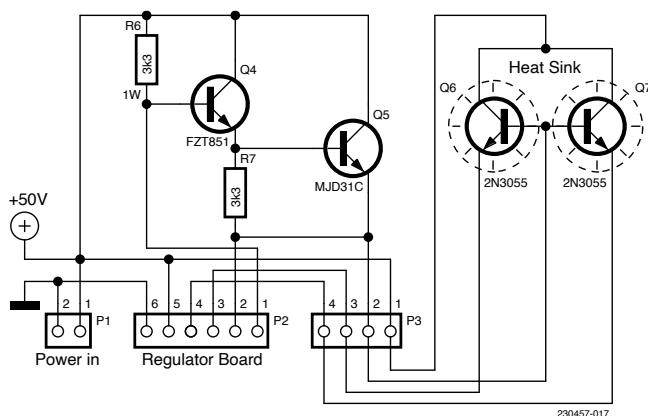
Podstawowych ustawień w zasilaczu dokonuje się potencjometrami. R14 nastawia prąd wyjściowy, a R15 – napięcie wyjściowe. Zakres działania każdego z nich trzeba przed użyciem wyregulować. Maksymalny prąd wyjściowy należy ustawić potencjometrem montażowym R13. W przypadku R15 zastosowałem dodatkowe rezystory, które ustalają maksymalne napięcie wyjściowe na poziomie 50 V; przy tej wartości przez R15 płynie prąd 100 μ A.

Oryginalny układ nie został zmieniony, z wyjątkiem przejścia na elementy SMD i dodania kondensatorów odsprężających przy U1 i U2. Inną zmianą jest użycie jako R11 dwóch rezystorów, aby lepiej rozproszyć ciepło i ułatwić skonfigurowanie obwodu z Q3 ograniczającego prąd.

Bardziej szczegółowe informacje na temat działania układu można znaleźć w oryginalnym artykule [1] oraz w notach aplikacyjnych National Semiconductor dotyczących wzmacniacza operacyjnego LM10.

Uproszczenie okablowania

Jedną ze zmian w stosunku do oryginalnego projektu jest montaż potencjometrów bezpośrednio na płytce stabilizatora (rysunek 4). Eliminuje to pięć przewodów i daje więcej możliwości montażu mechanicznego płytki drukowanej. Główną



Rysunek 3. Płytką układu wykonawczego z elementami mocy, w tym tranzystorami montowanymi na radiatorze. Tranzystory Q6 i Q7 są łączone przewodami do P3. Tranzystor Q5 wymaga niewielkiego radiatora; wystarczy radiator przyklejany o rezystancji termicznej około 5 $^{\circ}$ C/W. Tranzystory Q6 i Q7 należy zamontować na dużym radiatorze lub przymocować je do obudowy, stosując elementy izolujące elektrycznie

różnicą jest umieszczenie tranzystorów wykonawczych na oddzielnej płytce. Znów zmniejsza to liczbę przewodów w obudowie. Wynika to z faktu, że zmontowałem płytkę układu wykonawczego i tranzystory mocy razem na tylnym panelu, a jedyne połączenie kablowe do płytki stabilizatora zapewnia złącze P4. Kolejną kwestią, o której należy wspomnieć, jest to, że złącze należy zamontować z tyłu głównej płytki drukowanej, przewidzianej do montażu pionowego na panelu przednim. Można to jednak zmienić, jeśli płytka ma być zamontowana poziomo.

Płytką wejściową (AC)

Dla stabilizatora 50 V płytka ta zapewnia jedynie ochronę bezpiecznikiem, prostowanie oraz wygodne połączenie z kondensatorem C1 montowanym na obudowie (**rysunek 5**).

Dla zasilania symetrycznego (patrz poniżej) płytka doprowadza napięcie sieciowe (**rysunek 6**) do transformatora z dwoma uzwojeniami wtórnymi 12 V, który jest mocowany do obudowy i podłączony do płytki. Pozwala to na pewną elastyczność w wyborze tego podzespołu. Transformator montowany na płytce drukowanej musiałby być ściśle określonego typu.

Dwie linie zasilania są zabezpieczone bezpiecznikami polimerowymi (PTC) F2 i F3 o dopuszczalnym prądzie ciągłym 100 mA każdy. Każde uzwojenie jest dołączone bezpośrednio do płytki stabilizatora poprzez złącze P5.

Wprowadziłem kilka zmian w stosunku do oryginalnego projektu. Główne elementy wejściowe są mocowane do obudowy. Transformator sieciowy, którego użyłem, ma dwa uzwojenia wtórne 20 V, które, aby zapewnić 40 V RMS, połączyłem szeregowo. Po wyprostowaniu i wygładzeniu uzyskuje się około 55...56 V napięcia stałego. Oryginalny projekt przewidywał transformator o mocy 80...100 VA, ale wybrałem wersję 120 VA, która zapewnia wystarczający zapas mocy dla zakładanego zakresu pracy zasilacza i pozostaje w granicach dopuszczalnych parametrów zastosowanych tranzystorów wyjściowych. Warto zauważyć, że transformatory do montażu na obudowę są elementami drogimi, więc przed zakupem należy zdecydować, jaki maksymalny prąd wyjściowy jest naprawdę wymagany.

Transformator zasilą mostek prostowniczy BR1 o prądzie 4 A, na którym występuje łączny spadek napięcia rzędu około 1...2 V (na całym mostku), zależnie od obciążenia, dzięki czemu na kondensatorze elektrolitycznym C1 o pojemności 4700 µF i napięciu znamionowym

63 V uzyskuje się około 55 V napięcia stałego. Ten element również jest kosztowny i jego wybór wymagał przemyślenia. Napięcie znamionowe 63 V jest w tym zastosowaniu wystarczające, ale zapas nie jest duży, dlatego napięcie wtórne transformatora nie powinno przekraczać 40 V RMS, zwłaszcza przy pracy bez obciążenia i przy podwyższonym napięciu sieci.

Prosty zasilacz symetryczny

Układ ten, zbudowany z łatwo dostępnych elementów SMD, jest łatwy do wykonania. Schemat wywodzi się głównie z not aplikacyjnych producenta, z kilkoma dodatkowymi ulepszeniami. Stabilizator został przewidziany do pracy w zakresie od 0 V do ± 15 V (wyjście symetryczne) z maksymalnym prądem wyjściowym 250 mA. Oba te parametry można w razie potrzeby stosunkowo łatwo zwiększyć. Dla moich celów są one wystarczające. Taki uniwersalny zasilacz małej mocy jest dokładnie tym, czego potrzebuję na przykład do zasilania wzmacniaczy operacyjnych.

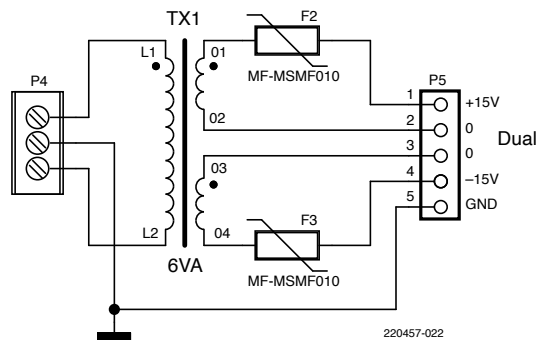
Jak widać na schemacie (**rysunek 7**), układ jest bardzo prosty. Opiera się na liniowych regulowanych stabilizatorach napięcia dodatniego (LM317) i ujemnego (LM337). Napięcia są regulowane podwójnym potencjometrem. Układ zawiera diody zabezpieczające przed stanami nieustalonymi. Zawiera również obwody kompensujące, które ograniczają wpływ napięcia odniesienia 1,25 V właściwego dla tych stabilizatorów. Zwykle w układach tego typu napięcia wyjściowe



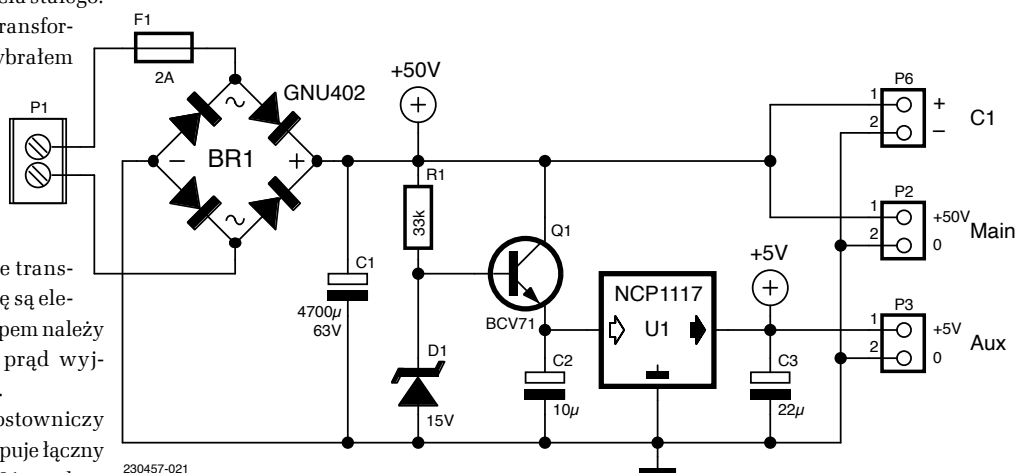
Rysunek 4. Trójwymiarowy widok projektu płytki stabilizatora 50 V

nie mogą być regulowane od zera. Tutaj, dzięki wykorzystaniu w każdym ze stabilizatorów spadku napięcia na dwóch diodach krzemowych, udało się uzyskać regulację od napięcia bardzo bliskiego 0 V.

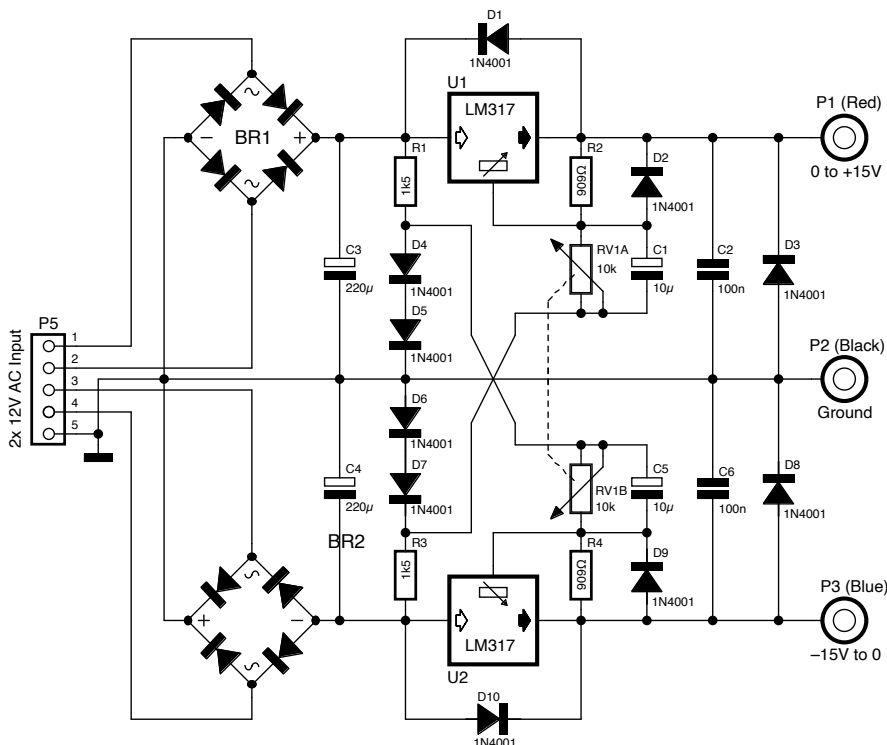
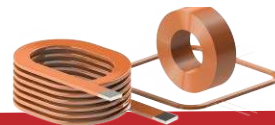
Aby uprościć końcowy montaż, potencjometr i gniazda wyjściowe zasilacza zostały umieszczone na płytce drukowanej



Rysunek 6. Na płytce AC znajdują się również układy zasilania transformatorowego z bezpiecznikami dla płytki podwójnego zasilania



Rysunek 5. Układ na płytce AC. Zasilą blok regulacji i zapewnia pomocnicze zasilanie 5 V dla mierników panelowych. C1 jest zbyt duży, aby znajdować się na płytce, dlatego jest montowany na obudowie i podłączony poprzez złącze P6



Rysunek 7. Układ symetrycznego zasilacza regulowanego. Dla uzyskania większego dodatniego prądu wyjściowego układ U1 można zastąpić stabilizatorem LT1083, LT1084 lub LT1085, stosując odpowiedni radiator. Dla wyższych prądów ujemnych U2 można zastąpić LT1033, również z radiatorem

przewidzianej do montowania na panelu przednim (rysunek 8). Złącze, tak jak w bloku stabilizatora, jest umieszczone z tyłu płytki.

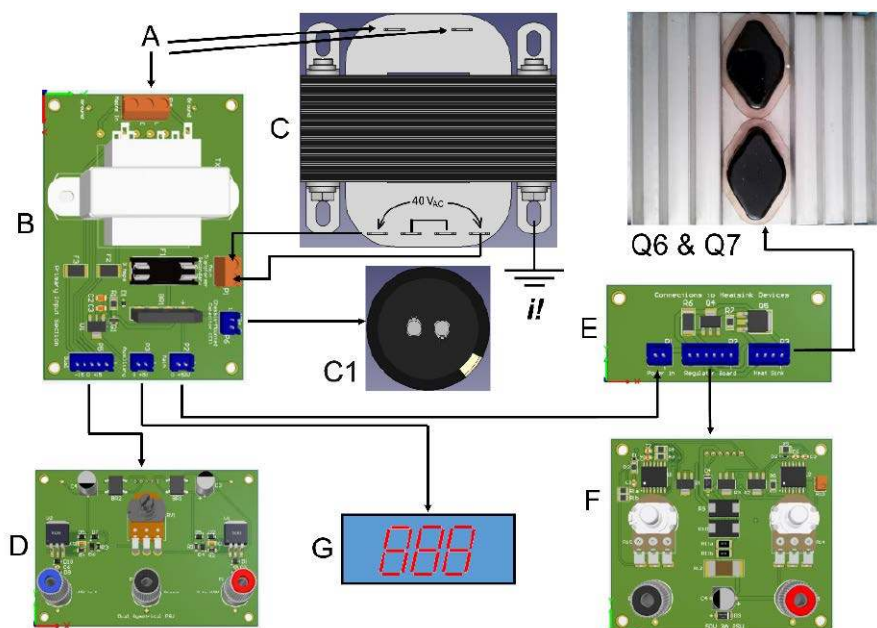
Płytki zasilacza symetrycznego

Płytki wymaga doprowadzenia tylko dwóch napięć zmiennych 12 V RMS, ponieważ znajdują się na niej prostowniki mostkowe i kondensatory wygładzające. Kondensatory muszą być na napięcie znamionowe co najmniej 1,5 razy większe od szczytowego napięcia z transformatora, czyli $12 \text{ V} \cdot 1,41 \cdot 1,5 \approx 25 \text{ V}$. Następnie w kolejności są stabilizatory liniowe. Szczegółowe informacje na temat działania



Rysunek 8. Trójwymiarowy widok projektu płytki symetrycznego zasilacza regulowanego

stabilizatorów można znaleźć w notach aplikacyjnych producenta. Swoje obliczenia zamieściłem jednak w ramce, zatytułowanej – jakżeby inaczej – Obliczenia.



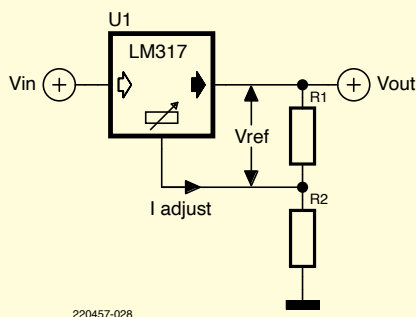
Rysunek 9. Schemat okablowania systemu. „A” to zasilanie prądem przemiennym z sieci 230 V przez bezpiecznik na panelu. „B” to płytka AC (wejściowa). „C” to transformator sieciowy – powinien dostarczać napięcie 40 V RMS. „C1” to kondensator C1 z rysunku 4. „D” to płytka symetrycznego zasilania regulowanego. „E” to płytka układu wykonawczego z tranzystorami. Tranzystory mocy Q6 i Q7 są zamontowane na wspólnym radiatorze z użyciem zestawów izolacyjnych TO-3. Płytki „F” realizuje regulowane zasilanie 50 V. Wreszcie „G” to mierniki panelowe. „!!” oznacza, że wszystkie odsoniowane elementy metalowe powinny być połączone z przewodem ochronnym

Przy określaniu dolnej wartości napięć wyjściowych należy pamiętać, że napięcie przewodzenia diod krzemowych zależy od temperatury złącza i przepływającego przez nie prądu. Po zapoznaniu się z charakterystyką napięciowo-prądową diody 1N4001 użyłem rezystorów szeregowych R1 i R3 o wartościach 1,5 kΩ. W notcie aplikacyjnej stabilizatora LM317 można znaleźć inne rozwiązanie problemu nastawiania napięcia wyjściowego od zera, ale odkryłem, że opisana tam metoda ma wpływ na maksymalne napięcie wyjściowe.

Wyższe napięcia wyjściowe

Układ stabilizatora napięć symetrycznych, z transformatorem o podwójnym uzwojeniu wyjściowym, może być używany jako zasilacz samodzielny. Zastosowane układy scalone mogą, przy odpowiednim chłodzeniu, dostarczyć prąd wyjściowy do 1,5 A i pracować przy maksymalnym napięciu wejściowym 37 V. Nadają się więc do uzyskania napięć wyższych niż przyjęte w tym projekcie. Jeśli napięcia te miałyby być znacznie wyższe, transformator, kondensatory i stabilizatory będą wymagały zmiany. Producenci oferują stabilizatory low-dropout (LDO) o lepszych parametrach, takie jak LM1084, który może dostarczyć prąd do 5 A, oraz LM1085, którego maksymalny prąd wyjściowy wynosi 3 A. Podczas ewentualnej modyfikacji należy jednak pamiętać o obciążalności prądowej innych elementów układu.

Obliczenia



W projekcie maksymalne osiągalne napięcie wyjściowe zostało obliczone na podstawie wzoru:

$$V_{out} = V_{ref} \cdot (1 + R_2 / R_1)$$

gdzie (dotyczy każdego z obu wyjść):

V_{out} = żądane napięcie wyjściowe (15 V)

$V_{ref} = 1,25$ V

R_1 = obliczana wartość rezystancji (w rzeczywistym układzie są to rezystory R_2 i R_4)

R_2 = maksymalna wartość rezystancji potencjometru (w rzeczywistym układzie:

$R_{V1} = 10$ kΩ)

Aby znaleźć R_1 , wzór musi zostać przekształcony:

$$R_1 = R_2 \cdot (V_{out} / V_{ref} - 1)$$

W praktyce dokładna wartość napięcia maksymalnego nie jest aż tak istotna, więc R_1 nie musi być dokładnie równy obliczonej wartości. Jeśli jednak zależy nam na dobrej dokładności, będziemy musieli przed rozpoczęciem obliczeń zmierzyć i uwzględnić maksymalną wartość rezystancji obu sekcji potencjometru.

Przypis redaktora: Autor „z rozpędu” użył wzoru podawanego w notach aplikacyjnych LM317/LM337, zapominając o fakcie, że w jego układzie rezystancja R_2 nie jest dołączona do masy, lecz do potencjału równego $-V_{ref}$, a więc odpowiednio około $-1,25$ V lub $+1,25$ V. Jeśli ten fakt uwzględnimy, to otrzymamy poprawny wzór, obowiązujący dla przedstawionego układu:

$$R_1 = R_2 \cdot V_{ref} / V_{out}$$

Wskaźniki pomiarowe

Do monitorowania dodatnich wyjść obu stabilizatorów użyłem prostych modułów woltomierzy i amperomierzy, które można kupić bardzo tanio przez Internet. Na płytce drukowanej znajdują się dodatkowe styki w złączu wyjściowym, które umożliwiają dołączenie rezystorów pomiarowych do pomiaru prądu przed gniazdami bananowymi 4 mm. Dodanie rezystora pomiarowego będzie wymagało takiego odseparowania elektrycznego gniazd wyjściowych, aby cały prąd obciążenia przepływał przez

ten rezystor. W tym celu można użyć zestawu izolacyjnego do tranzystorów w obudowie TO-3, choć wystarczą również tulejka izolacyjna i plastikowe podkładki.

Zasilanie pomocnicze 5 V

Układem najbardziej zaawansowanym na płytce AC jest zasilanie pomocnicze 5 V, potrzebne do modułów wskaźników (rysunek 5). Napięcie doprowadzone do scalonego stabilizatora LDO 5 V musi zostać zmniejszone, by nie przekraczało dopuszczalnej górnej granicy. Realizuje to dioda Zenera

(D1), ograniczająca napięcie bazy tranzystora NPN (Q1) do około 15 V. Rozwiązanie to obniża napięcie linii zasilającej 50...55 V do bezpiecznej wartości około 14,3 V, odpowiedniej do zasilania stabilizatora U1.

Montaż całości

Kompletny zasilacz składa się z czterech płytek drukowanych, dużego kondensatora (C1), transformatora sieciowego, radiatora z Q6 i Q7 oraz dwóch mierników panelowych. **Rysunek 9** pokazuje, jak wszystkie te części powinny być ze sobą połączone. Wszystko zmontowałem w metalowej obudowie. Rezultat, jak pokazuje **rysunek 10**, jest całkiem satysfakcjonujący.

Płytki drukowane w panelu

Ostatni szczegół konstrukcyjny, o którym warto wspomnieć, jest związany z produkcją płytek drukowanych. Aby zaoszczędzić na kosztach, płytki zostały przewidziane do produkcji w parach na jednym panelu. Wymaga to wykonania nacięć technologicznych pomiędzy płytkami, umożliwiających ich łatwe rozdzielenie po montażu.

Wszystkie pliki projektowe i pełną listę materiałów można znaleźć na stronie [2]. ■

Steve Griffin (Wielka Brytania)

Odnosiniki

- [1] „Variable Power Supply 0–50 V/0–2 A”, Elektor 7/1980: <https://elektormagazine.com/magazine/elektor-197999/44508>.
- [2] Niniejszy projekt w Elektor Labs: <https://elektormagazine.com/labs/bench-power-supply-ensemble>
- [3] Część 1 na Elektor Labs: <https://elektormagazine.com/labs/variable-0-50v-2a-supply-refresh>
- [4] Część 2 na Elektor Labs: <https://elektormagazine.com/labs/simple-dual-voltage-bench-power-supply>



Rysunek 10. Zmontowany zasilacz, przygotowany na następne 40 lat





FN-SWM10

Zgrzewarka do ogniw – spawarka punktowa
z kolorowym wyświetlaczem
i funkcją powerbank FNIRSI SWM10



FN-DPOS-350P

Dwukanałowy oscyloskop 350 MHz,
FNIRSI DPOS350P



FN-2C53T

Dwukanałowy oscyloskop z multimetrem
i generatorem 50 MHz FNIRSI 2C53T

BESTSELLERY sklepu AVT – sklep.avt.pl

Mierniki Testery FNIRSI

Rabat dla Czytelników EdW
przy zakupie podaj kod **EdW2505FN**

Kod ważny do 30.09.2025

Rabat dla Prenumeratorów EdW
przy zakupie podaj numer prenumeraty

-3%

-6%



FN-LCR-ST1

Miernik pojemności, tester elementów
FNIRSI LCR-ST1



FN-LCR-P1

Tester elementów FNIRSI LCR-P1



FN-HRM10

Tester rezystancji wewnętrznej
akumulatorów FNIRSI HRM-10



FN-G1200

Mikroskop cyfrowy G1200 z wyświetlaczem
7 cali, powiększenie x1200, tryb foto/video



FN-DWS200-F245

Stacja lutownicza 200 W z kolbą F245,
FNIRSI DWS200



FN-1014D

Oscyloskop dwukanałowy 100 MHz,
Generator sygnału DDS, FNIRSI 1014D

Miernik indukcyjności, pojemności i ESR

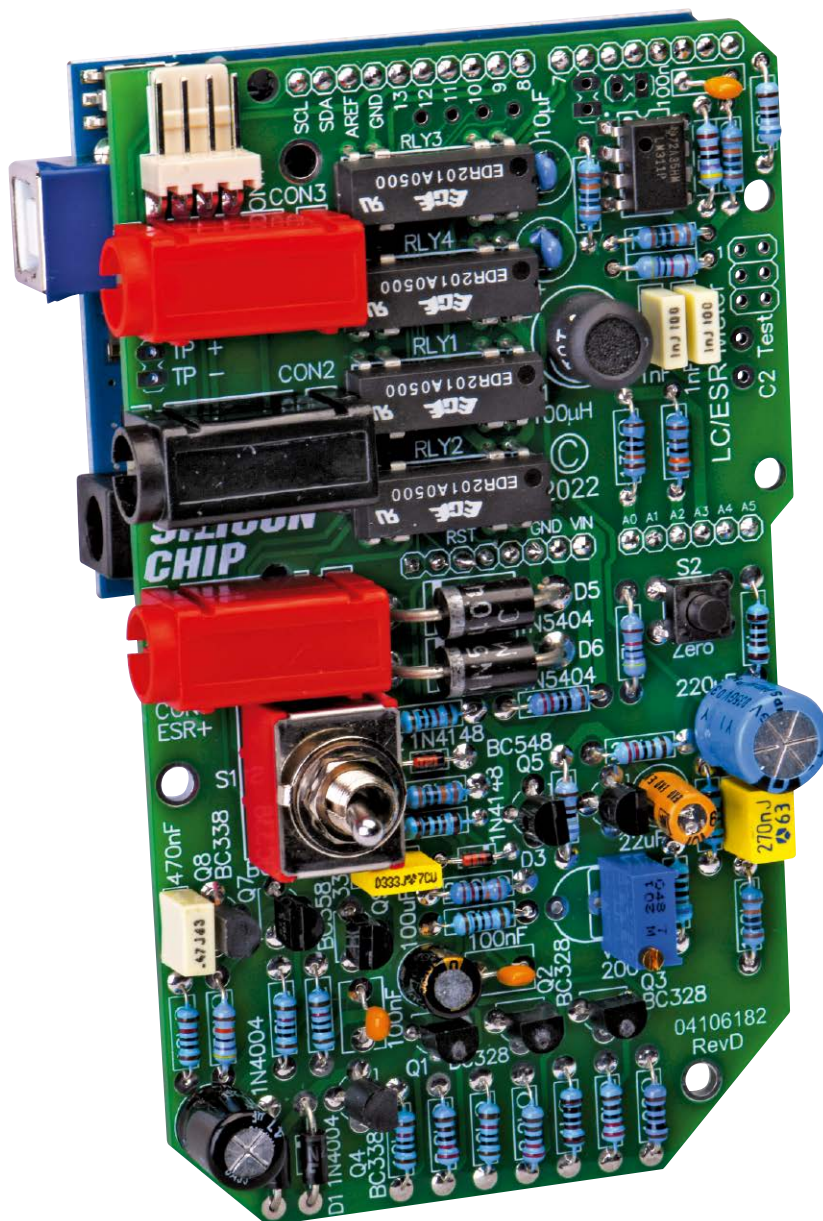
Jest to udoskonalona wersja szerokokresowego cyfrowego miernika indukcyjności i pojemności (LC), opublikowanego w Silicon Chip w czerwcu 2018 r. (www.silicon-chip.au/Article/11099). Przyrząd opisany tutaj umożliwia dodatkowo pomiar ESR kondensatorów. Jest to funkcja niezwykle przydatna przy diagnozowaniu wadliwego sprzętu, ponieważ wzrost ESR z biegiem czasu jest jednym z najczęstszych objawów awarii kondensatorów elektrolitycznych.

Przypis redaktora: ESR (ang. *equivalent series resistance*) to „zastępcza rezystancja szeregową” – jeden z parametrów rzeczywistego kondensatora. Każdy kondensator jest w przybliżeniu szeregowym połączeniem kondensatora idealnego i pewnej pasożytniczej rezystancji, oznaczanej właśnie skrótem ESR.

W wydaniu Silicon Chip z czerwca 2018 roku Tim Blythman przedstawił szerokokresowy miernik LC o doskonałej dokładności. Miernik, oparty na specjalnie zaprojektowanej nakładce (shieldzie) do Arduino, był prosty w budowie. Jego dokładność została zoptymalizowana dzięki funkcji automatycznej kalibracji oraz kompensacji nieodłącznej pojemności pasożytniczej wnoszonej przez sondy pomiarowe czy nawet piny Arduino.

Przyrząd jest niewątpliwie przydatny do sprawdzania nieznanymi elementami. W przypadku kondensatorów elektrolitycznych istotne jest również to, czy mają one małą impedancję dla prądów przemiennych, co wymaga niskiej wartości zastępczej rezystancji szeregowej (ESR).

Ostatnim kompletnym miernikiem ESR opublikowanym w Silicon Chip był miernik „Mk.2” autorstwa Boba Parkera (numery z marca i kwietnia 2004 roku; www.silicon-chip.au/Series/99). Pierwsza wersja tego projektu powstała jeszcze w latach 90. Artykuł zawiera szereg informacji na temat wykorzystania kondensatorów w projektach elektronicznych oraz znaczenia znajomości ich wartości ESR.



Pomyślałem sobie, że warto byłoby połączyć funkcje pomiaru LC i ESR w jednym urządzeniu.

Dlaczego ESR jest tak istotny?

Kondensatory elektrolityczne są stosowane tam, gdzie wymagane jest przechowywanie dużych ładunków elektrycznych. Gdy kondensator jest ładowany i rozładowywany, musi do niego wpływać i wypływać prąd. Za każdym razem, gdy ma miejsce przepływ prądu, rezystancja ESR, dołączona niejako w szereg z kondensatorem, powoduje straty energii.

ESR utrudnia również kondensatorowi prawidłowe wykonywanie jego standardowego zadania – wygładzania napięcia. Załóżmy, że kondensator jest ładowany prądem 1 A, a następnie nagle zacznie się rozładowywać prądem również 1 A. Jeśli jego ESR wynosi 1 Ω , to napięcie widziane przez resztę obwodu zmieni się skokowo o $(1 \text{ A} + 1 \text{ A}) \cdot 1 \Omega$, czyli o 2 V. Wpływ ESR będzie się na przykład objawiał w zasilaczu zwiększeniem tętnień na kondensatorze za prostownikami.

Duża wartość ESR prowadzi również do nagrzewania się kondensatora elektrolitycznego,

co może wpływać na jego pojemność i przyspieszać degradację elektrolitu.

Jednym z najczęstszych objawów uszkodzenia kondensatora elektrolitycznego jest nagły lub stopniowy wzrost jego wartości ESR. Zwiększona ESR kondensatorów może powodować w układzie różnego rodzaju awarie, często trudne do zidentyfikowania. W zasilaczach impulsowych może to być obniżenie lub brak napięcia stabilizowanego, niesprawność filtru, zwiększenie poziomu tętnień albo brak startu. Dlatego podczas testowania kondensatorów elektrolitycznych warto sprawdzać nie tylko ich pojemność, ale również to, czy wartość ESR mieści się w odpowiednim zakresie.

Nowa wersja projektu

Miernik ESR „Mk.2” z marca 2004 roku opiera się na mikrokontrolerze Z86E0412 sterującym dwoma wyświetlaczami siedmio-segmentowymi. Miernik ten był niezwykle popularny, podobnie jak jego poprzednik.

W niniejszej wersji nie staramy się ponownie wynaleźć koła. Stosujemy ten sam układ pomiarowy co w mierniku Mk.2, ale jego wyjście doprowadzamy do Arduino Uno sterującego wyświetlaczem LCD. Zaleta jest taka, że układ pomiarowy można zbudować na stosunkowo małej płytce drukowanej i dołączyć do miernika LC przedstawionego w numerze z czerwca 2018 r. (www.siliconchip.au/Article/110999).

Pozwoli to otrzymać doskonały przyrząd ogólnego przeznaczenia, który może nie

Budowa kondensatora elektrolitycznego

Kondensatory elektrolityczne, w swojej najbardziej podstawowej postaci, zawierają dwie płytki przewodzące (anodę i katodę) rozdzielone materiałem izolacyjnym zwanym dielektrykiem. Istnieją trzy główne typy kondensatorów elektrolitycznych w zależności od materiału anody i rodzaju dielektryka: aluminiowe, z tlenku niobu i tantalowe.

Pojemność kondensatora jest wprost proporcjonalna do całkowitej powierzchni płytek i odwrotnie proporcjonalna do odległości między nimi. Dlatego im cieńszy dielektryk, tym większa pojemność kondensatora przy tych samych wymiarach.

Dielektryki mają dużą rezystancję. W przypadku kondensatorów o małej pojemności, jako materiał dielektryka stosuje się różne polimery, mikę, ceramikę, a nawet wybrane ciecze i gazy, w tym powietrze.

We wszystkich trzech rodzajach kondensatorów elektrolitycznych anoda składa się z materiału pierwotnego (aluminium, tlenku niobu czy tantalu), a dielektrykiem jest bardzo cienka warstwa odpowiedniego tlenku (np. pięciotlenku niobu), osadzona na powierzchni anody. Taki dielektryk jest bardzo cienki i znajduje się w bliskiej odległości od katody, co jest zaletą użycia elektrolitu. Elektrolit stanowi zasadniczo katodę. Wymagane jest jednak jego elektryczne połączenie z wyprowadzeniem kondensatora. Aby zapewnić pewne połączenie o małej rezystancji, elektrolit jest wysoce przewodzącą cieczą, żelam lub ciałem stałym.

W aluminiowych kondensatorach elektrolitycznych skutecznym sposobem na zapewnienie wysokiej jakości połączenia jest umieszczenie w kondensatorze cienkiego, nasączonego elektrolitem arkusza papieru. W kondensatorach niobowych i tantalowych, do łączenia katody z dielektrykiem stosuje się zwykle jako elektrolit dwutlenek manganu. Więcej szczegółów można znaleźć w artykule w Silicon Chip „All About Capacitors” („wszystko o kondensatorach”) w wydaniu z marca 2021 r. (www.siliconchip.au/Article/14786).

tylko sprawdzać ESR kondensatorów, ale także ich pojemność (do pewnej granicy), ponadto może być używany do pomiaru cewek i nie tylko.

Inna możliwość to uzupełnienie układu pomiarowego modulem Arduino Uno (lub jego klonem) z 4-wierszowym wyświetlaczem alfanumerycznym LCD, dołączonym przez magistralę I²C, i w ten sposób otrzymać samodzielną miernik ESR. Program, zarówno

dla wersji zintegrowanej z miernikiem LC, jak i samodzielnej, jest na stronie www.siliconchip.com.au/Shop/6/234.

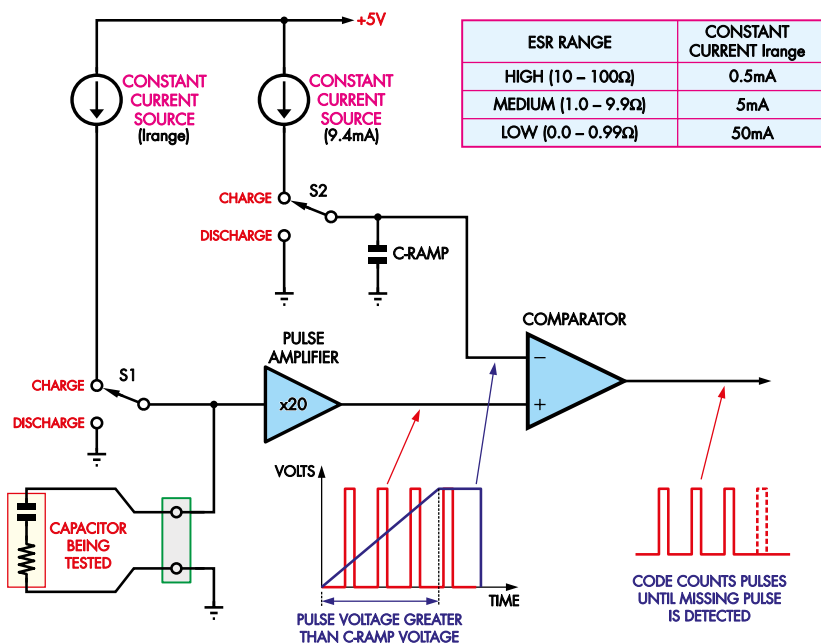
Jeśli już zbudowaliście miernik LC, to możecie dołączyć do niego moduł ESR, chociaż prawdopodobnie łatwiejsze będzie zbudowanie całości od zera.

Pomiar ESR

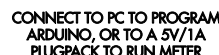
Rysunek 1 pokazuje uproszczony schemat ilustrujący sposób pomiaru. S1 i S2 to przełączniki elektroniczne sterowane przez Arduino. Gdy nie jest wykonywany żaden pomiar, to zarówno S1 jak i S2 znajdują się w pozycji DISCHARGE („rozładowanie”). Testowany kondensator oraz kondensator C-RAMP są utrzymywane w stanie rozładowanym.

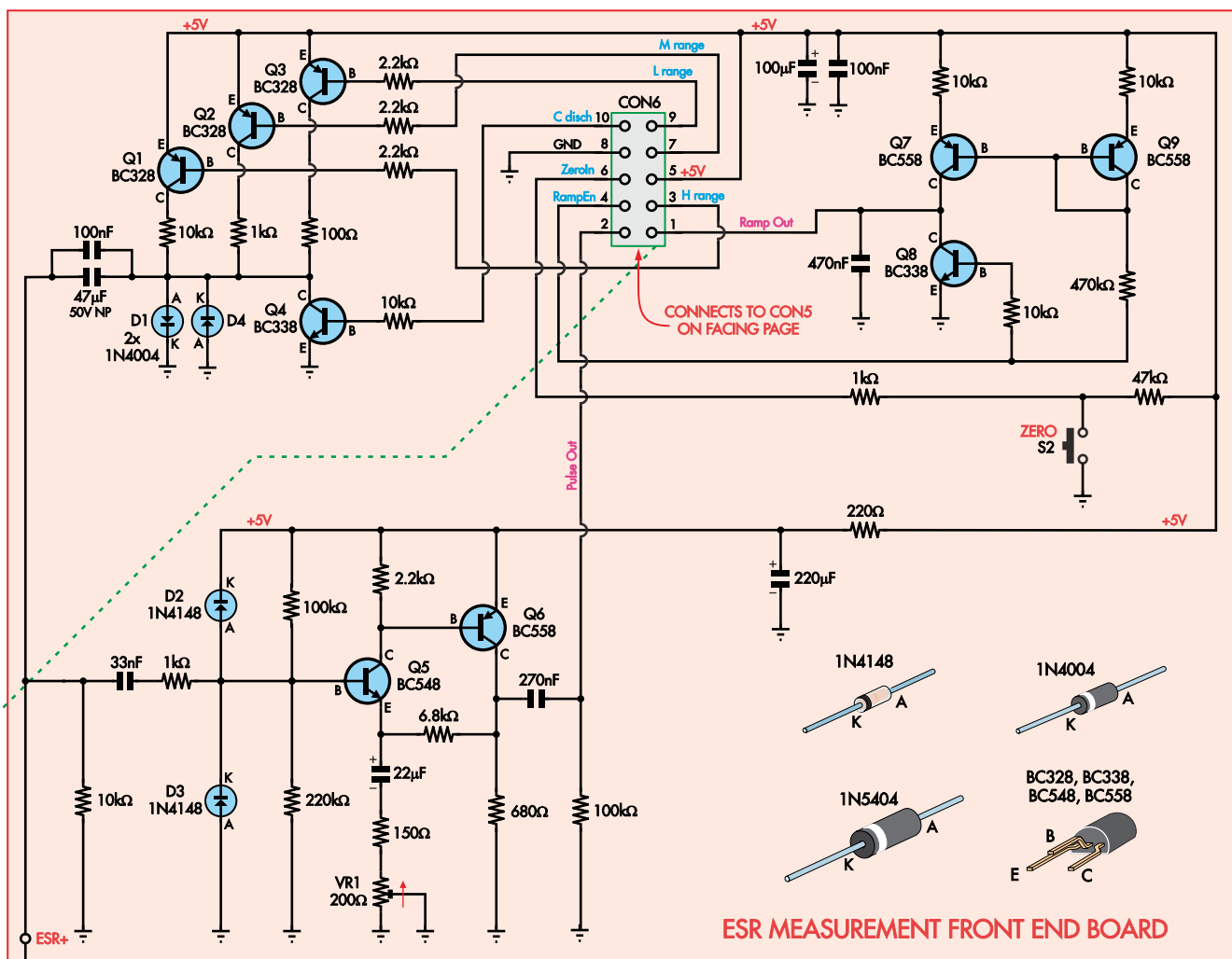
Na początku cyklu pomiarowego program w Arduino ustawia S2 w pozycję CHARGE („ładowanie”) i ładuje kondensator C-RAMP stałym prądem 9,4 μ A. Wskutek tego napięcie na wejściu odwracającym komparatora narasta liniowo z szybkością 20 mV/ms, czyli 20 V/s. **Przypis redaktora:** takie liniowo rosnące lub opadające napięcie jest w elektronice zwane „rampą”.

Po czasie 480 μ s przełącznik S1 jest na 20 μ s przełączany w pozycję CHARGE („ładowanie”), co do testowanego kondensatora dołącza źródło stałego prądu. Prąd, w zależności od zakresu, wynosi 0,5 mA, 5 mA lub 50 mA. Impuls prądu jest bardzo krótki, aby dostarczyć jak najmniej ładunku do kondensatora. Chodzi tu jedynie o pomiar impulsu napięcia, który



Rysunek 1. Klucz S1 rozładowuje kondensator, a następnie wielokrotnie przykład do niego krótki impuls prądowy. Impulsy są na tyle krótkie, że wpływ ładowania kondensatora jest pomijalny, dlatego napięcie na nim jest w przybliżeniu proporcjonalne do ESR. Wzmacniacz doprowadza wzmacnione impulsy do komparatora, który porównuje ich amplitudę z napięciem liniowo narastającym. Zliczając impulsy na wyjściu komparatora możemy dokładnie określić ESR





Rysunek 2. Oryginalny układ miernika LC (z kilkoma dodatkami) jest po lewej stronie (na białym tle). Dodany układ do pomiaru ESR znajduje się po prawej stronie (tło pomarańczowe). Na płytce połączonej nie ma złączy CON5 i CON6 – dziesięć połączeń jest poprowadzonych na płytce ścieżkami. W przypadku oddzielnych płytek styki złączy są połączone kabelkiem taśmowym

gdy odpowiednie wyjście jest w stanie niskim. Arduino załącza jeden z tranzystorów w zależności od wybranego zakresu pomiarowego. Rezystory w kolektorach – 10 k Ω , 1 k Ω i 100 Ω – ustawiają wartość impulsu prądowego odpowiednio na 0,5 mA, 5 mA lub 50 mA.

Ścisłe rzecz biorąc nie są to klasyczne źródła prądowe. Opieramy się na fakcie, że zasilanie 5 V jest stabilizowane, a mierzony kondensator jest na początku rozładowany. Na załączonym rezystorze wystąpi zatem napięcie bliskie 5 V, a wartość prądu wynika z prawa Ohma.

Impuls prądowy jest do testowanego kondensatora przykładany poprzez równoległe połączone kondensatory 100 nF i 47 μ F, blokujące wszelkie składowe stałoprądowe. Wpływ ESR tej kombinacji kondensatorów można pominąć, biorąc pod uwagę

stosunkowo duże wartości rezystorów źródeł prądu. A co istotne – pomiar jest dokonywany bezpośrednio na zaciskach mierzonego kondensatora, więc ESR tych dwóch kondensatorów nie wpływa na wynik.

Kondensator 100 nF zapewnia małą impedancję dla wysokich częstotliwości, co jest korzystne z uwagi na krótkie czasy trwania impulsów prądowych.

Za każdym razem, gdy Q1, Q2 i Q3 są wyłączane, wyjście cyfrowe D13 Arduino załącza Q4, podając prąd do jego bazy. Oba kondensatory sprzęgające są wówczas rozładowywane i gotowe do następnego impulsu prądowego.

Diody D1 i D4, włączone antyrównolegle, chronią Q4 przed ewentualnym silnym impulsem prądu w przypadku dołączenia do zacisków testowych naładowanego kondensatora. Nawet przy ESR równym 100 Ω spadek napięcia na nim podczas pomiaru wyniesie

tylko około 500 mV, więc diody D1 i D4 będą miały znikomy wpływ na to napięcie.

Wzmacniacz impulsów napięciowych

Napięcie impulsowe wytworzone na testowanym kondensatorze jest poprzez kondensator 33 nF i rezystor szeregowy 1 k Ω podawane do wzmacniacza. Impuls jest wzmacniany w dwustopniowym wzmacniaczu tranzystorowym na tranzystorach Q5 i Q6. Stosunek rezystancji sprzężenia zwrotnego 6,8 k Ω do sumy rezystancji 150 Ω i VR1 (ustawionego na około 200 Ω) wyznacza wzmocnienie układu równe $1 + 6,8 \text{ k}\Omega / (150 \Omega + 200 \Omega)$ czyli około 20. **Przypis redaktora:** jako VR1 zalecany byłby raczej potencjometr 470 Ω lub 500 Ω . Diody D2 i D3 chronią Q5 w przypadku dołączenia do zacisków pomiarowych naładowanego kondensatora.

Wzmocnione napięcie impulsu trafia poprzez kondensator 270 nF do nieodwracającego wejścia komparatora analogowego w mikrokontrolerze na module Arduino Uno. Kondensator blokuje napięcie stałe występujące na kolektorze Q6 i rezystorze 680 Ω . Rezystor ten utrzymuje kondensator 270 nF w stanie rozładowanym podczas przerwy między impulsami.

Generator napięcia rampy

Tranzystory PNP Q7 i Q9 stanowią lustro prądowe służące do ładowania kondensatora C-RAMP o pojemności 470 nF ze stałym prądem. Gdy Arduino przełączy napięcie na styku 4 złącza CON6 do stanu niskiego, załącza się Q9, powodując przepływ przez rezystor 470 k Ω prądu około 9,4 μ A. W tym czasie Q8 jest wyłączony, dzięki czemu kondensator

może się ładować. Q7 powtarza prąd płynący przez Q9, więc kondensator zaczyna być ładowany prądem 9,4 μ A począwszy od napięcia 0 V. Napięcie rosnące na kondensatorze jest dołączone do wewnętrznego komparatora w mikrokontrolerze Arduino Uno (wejście odwracające) poprzez styk 1 złącza CON6. Na zakończenie cyklu pomiarowego Arduino Uno wyłącza generator rampy, ustawiając styk 4 złącza CON6 w stan wysoki, co wyłącza prąd ładowania płynący przez Q7 i jednocześnie w celu rozładowania kondensatora C-RAMP załącza Q8.

Połączenie z miernikiem LC

Miernik LC wykorzystywał wejścia komparatora analogowego Arduino (D6 i D7) jako wyjścia cyfrowe do sterowania cewkami przekazywników RLY1 i RLY2. Konieczna była

modyfikacja programu miernika LC w celu przeniesienia funkcji sterowania przekaźnikami do pinów D3 i D4 i umożliwienia blokowi pomiarowemu ESR korzystania z komparatora. Większa płyta drukowana uwzględniła to przekierowanie.

D3 i D4 jako cyfrowe wejścia/wyjścia są współdzielone z miernikiem ESR za pośrednictwem przełącznika S1, który wybiera między trybami LC i ESR. Jest to niezbędne, ponieważ w Arduino Uno nie ma wystarczającej ilości wolnych wejść/wyjść. Oryginalny miernik LC nie ma złącza CON5, zatem w moim prototypie do pinów Arduino i przełącznika idą linie ze złącza CON6.

Dodatkowe zabezpieczenie wejścia

Gdybyśmy przypadkiem dołączyli do miernika ESR naładowany kondensator, energia jego rozładowania mogłaby uszkodzić układ pomimo istnienia zabezpieczeń opisywanych powyżej. Zatem, tak samo jak w oryginalnym projekcie opublikowanym w 2004 roku, bezpośrednio do gniazd wejściowych są dołączone dwie wysokoprądowe diody 1N5404 (D5 i D6). Nie zmienia to faktu, że należy pamiętać o rozładowaniu kondensatorów przed ich testowaniem!

Oprogramowanie

Program pomiaru ESR jest oparty na algorytmie Boba Parkera opublikowanym w Silicon Chip w numerze z marca 2004 roku, z niewielkimi zmianami w czasach impulsów, aby się lepiej dopasować do możliwości Arduino Uno. W oryginalnym projekcie szerokość impulsu wynosiła 8 μ s, a czas przerwy 492 μ s. W Arduino te oryginalne ustawienia powodowały niewielkie wahania w odczytach.

Na przykład, przy rezystancji ESR wynoszącej 0,6 Ω odczyt wahałby się między 0,59 a 0,61. Zmiana szerokości impulsu na około 20 μ s poprawiła stabilność odczytów bez wpływu na dokładność. Aby zachować okres 500 μ s, przyjęto czas wyłączenia równy 480 μ s.

Program rozpoczyna pomiary w górnym zakresie ESR, ustawiając wyjście D12 w stan niski, a D10 i D11 w stan wysoki. Powoduje to ustawienie prądu impulsu 0,5 mA. W tym samym czasie inicjowany jest generator napięcia odniesienia poprzez ustawienie wyjścia D3 w stan niski.

Gdy napięcie impulsu na nieodwracającym wejściu komparatora w Arduino przekroczy napięcie odniesienia, komparator zgłasza przerwanie. W podprogramie jego obsługi ustawiany jest znacznik wskazujący,

Jaka wartość ESR jest prawidłowa?

Kondensatory elektrolityczne zawierają elementy reakcyjne, więc wartość ESR jest różna w zależności od częstotliwości przykadanego napięcia. Ujmuje to inny parametr: „zastępcza indukcyjność szeregową” (ESL). Na oba parametry mają wpływ procesy produkcyjne. Wartości parametrów zmieniają się również przy zmianach temperatury. Karty katalogowe producentów podają wartości ESR lub impedancji przy określonej częstotliwości, najczęściej 100 kHz, a niekiedy także współczynnik strat przy 120 Hz. W wielu kartach nie ma jednak takich informacji lub są one podawane w inny sposób, np. jako współczynnik strat.

W związku z tym podawanie oczekiwanych definitywnych wartości ESR dla każdego typu kondensatora elektrolitycznego jest niemożliwe. Ogólnie jednak nie spodziewamy się wartości przekraczających kilku omów. Kondensatory o dużej pojemności powinny generalnie mieć niskie wartości ESR. Kondensatory przewidziane do użytku w zasilaczach impulsowych (określane zwykle jako „Low ESR”) powinny mieć wartość ESR wynoszącą ułamek oma lub mniej.

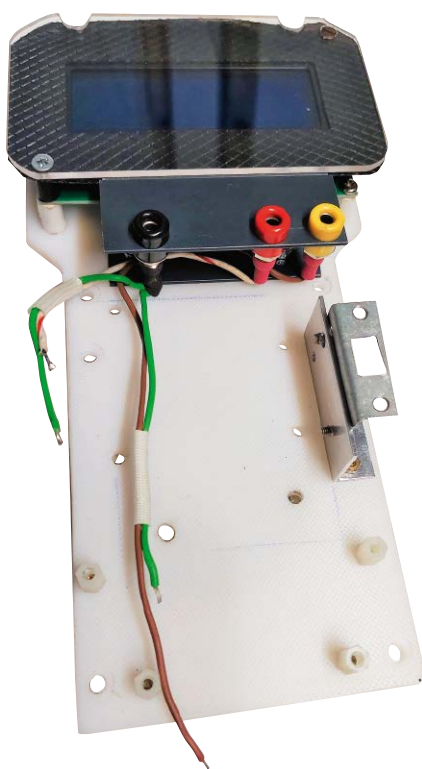
Przykładowo – karta katalogowa kondensatorów aluminiowych Panasonic z serii FM-A podaje wartości ESR od 0,012 Ω do 0,34 Ω w zależności od napięcia znamionowego kondensatora (od 6,3 V do 50 V) i pojemności (od 22 μ F do 6800 μ F). Karta serii RubyCon YXF dla podobnych pojemności i zakresów napięcia wymienia maksymalne wartości ESR od 0,025 Ω do 1,3 Ω .

W tabeli 1 przedstawiono typowe wartości ESR wzięte z projektu miernika „Mk.2”. Są to bardzo ogólne wartości oczekiwane, więc jako odniesienie powinny być raczej używane karty katalogowe producentów. Powinno być jednak jasne, że jeśli zmierzona wartość ESR przekracza dziesiątki a nawet setki omów, to kondensator jest wadliwy!

Wersję tabeli 1 w formacie PDF można pobrać ze strony www.siliconchip.com.au/Shop/11/238.

Tabela 1. Typowe wartości ESR sprawnych kondensatorów

	10 V	16 V	25 V	35 V	63 V	160 V	250 V
1 μ F			5	4	6	10	20
2,2 μ F			2,5	3	4	9	14
4,7 μ F			6	3	2	6	5
10 μ F		1,6	1,5	1,7	2	3	6
22 μ F	3	0,8	2	1	0,8	1,6	3
47 μ F	1	2	1	1	0,6	1	2
100 μ F	0,6	0,9	0,5	0,5	0,3	0,5	1
220 μ F	0,3	0,4	0,4	0,2	0,15	0,25	0,5
470 μ F	0,15	0,2	0,25	0,1	0,1	0,2	0,3
1000 μ F	0,1	0,1	0,1	0,04	0,04	0,15	
4700 μ F	0,06	0,05	0,05	0,05	0,05		
10 mF	0,04	0,03	0,03	0,03			



W prototypie miernika zastosowano specjalną obudowę dopasowaną do wyświetlacza. Można, jak na zdjęciu, użyć dwóch płytek, aczkolwiek płytka pojedyncza (połączona), pokazana na zdjęciu tytułowym, będzie wymagać znacznie mniej okablowania

że wykryto impuls. W konsekwencji zwiększony jest licznik, znacznik przerwania jest kasowany, a do testowanego kondensatora jest przykładany kolejny impuls.

Proces ten powtarza się, dopóki program nie wykryje, że po wystąpieniu impulsu prądowego znacznik nie został ustawiony. Oznacza to, że napięcie odniesienia (rampa) osiągnęło poziom wyższy niż napięcie impulsu, a więc zliczanie zostało zakończone.

Jeśli całkowita liczba impulsów po zakończeniu cyklu zliczania była mniejsza niż 10, wybierany jest następny niższy zakres pomiarowy. Cykle są powtarzane tak długo, aż uzyskana liczba impulsów wyniesie od 10 do 100. W niskim zakresie pomiarowym zliczenie między 10 a 100 odpowiada wartości ESR między 0,1 Ω a 1 Ω ; w zakresie średnim reprezentuje wartość od 1 Ω do 10 Ω , natomiast w zakresie wysokim – od 10 Ω do 100 Ω .

Jeśli liczba zliczeń przekroczy 100, program automatycznie spróbuje przejść do następnego wyższego zakresu, aż liczba nie znajdzie się w przedziale od 10 do 100. Jeśli w najwyższym zakresie liczba zliczeń nadal przekracza 100, na wyświetlaczu pojawi się komunikat „Over range!” („Przekroczenie zakresu!”).

Rezystancja przewodów pomiarowych

Naszym celem jest pomiar wartości ESR leżących często znacznie poniżej

1 Ω . Rezystancja przewodów pomiarowych i styków w gniazdach bananowych może wprowadzać znaczące błędy.

Naciśnięcie przycisku ZERO (S2) wymusza stan niski na wejściu D4, wykrywany przez Arduino. Zostaje wyświetlony komunikat „Short test leads and press zero...” („Zewrzyj przewody testowe i naciśnij ZERO”). Gdy przewody zostaną zwarte, Arduino wielokrotnie mierzy ich rezystancję i wyświetla ją w omach w czwartym wierszu wyświetlacza. Program oczekuje na ponowne naciśnięcie przycisku ZERO i zapisuje wówczas wartość rezystancji przewodów do pamięci EEPROM Arduino. Rezystancja ta będzie odejmowana od wyników kolejnych pomiarów ESR kondensatora.

Oprócz wyświetlania wyników pomiaru ESR na wyświetlaczu, Arduino wysyła również dane pomiarowe poprzez port USB.

Po każdym impulsie prądowym jest wyświetlany stan zliczania. Po zakończeniu cyklu pomiarowego pokazywany jest końcowy stan zliczania, wybrany zakres pomiarowy i liczba zmian zakresu dokonanych podczas pomiaru. Na stan zliczania ma wpływ rezystancja przewodów pomiarowych.

Miernik połączony LC/ESR

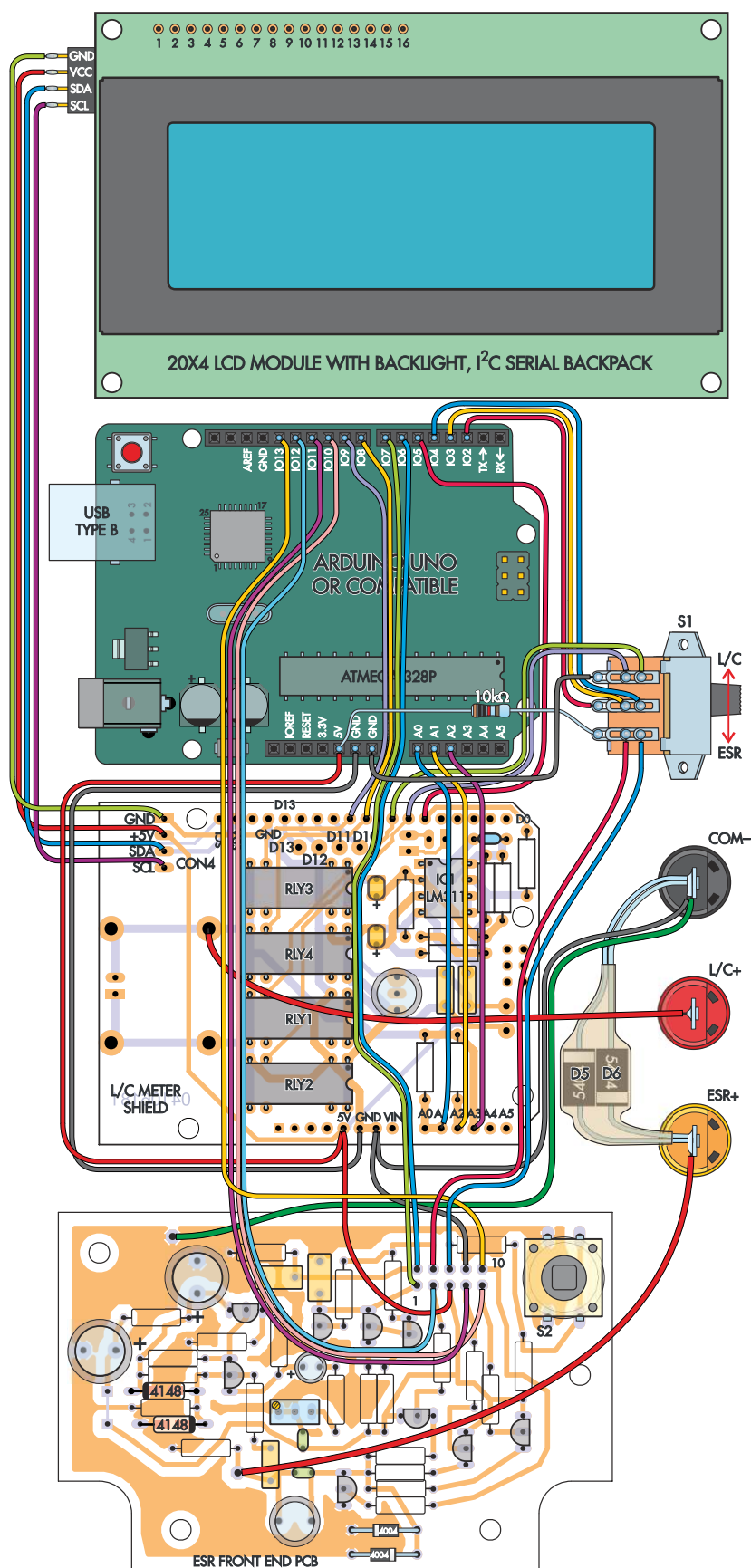
W połączonym mierniku LC/ESR jedna z sekcji przełącznika LC/ESR (S1) informuje mikrokontroler przez wejście cyfrowe

D2, który z trybów jest wybrany. Jeśli poziom logiczny wejścia D2 jest niski, miernik jest w trybie ESR. Przejście z jednego trybu na drugi następuje po zakończeniu przez program bieżącego cyklu pomiarowego.

Jak wspominaliśmy wcześniej, piny Arduino D3 i D4 są współdzielone między blokami LC i ESR. D3 jest w obu trybach wyjściem cyfrowym. D4 jest wejściem cyfrowym dla bloku ESR (wejście od przełącznika ZERO), ale dla bloku LC jest wyjściem (sterującym RLY1). Podczas przełączania



Jeśli budujecie miernik z oddzielnych płytek, może być potrzebny układ montażowy do gniazd bananowych, przedstawiony tutaj i na rysunku 8. W przypadku użycia go, diody D5 i D6 są umieszczone w białej rurce termokurczliwej



Rysunek 3. Okablowanie niezbędne do dodania funkcji pomiaru ESR do istniejącego miernika LC poprzez uzupełnienie go specjalną płytką (na dole). Większość konstruktorów wybierze zapewne sposób znacznie łatwiejszy: zbudowanie całości na pojedynczej „płytkie połączonej”

trybów pracy D4 jest odpowiednio rekonfigurowany przez program.

Wybór obudowy

Obudowę użytą w prototypie można nabyć u dystrybutorów Mouser (563-HH-3421) i Digi-Key (HH-3421-ND), aczkolwiek zapas tych obudów jest ograniczony. Opcjonalną uchylną podstawkę można w Digi-Key kupić osobno (377-1171-ND).

Płytką połączoną jest znacznie węższa niż płytką z samym miernikiem ESR, więc również powinna zmieścić się w tej obudowie. Głębokość wewnętrzna wynosi 37 mm (nie licząc elementów takich jak tuleje montażowe, które można usunąć), zatem Arduino, płytki pomiarowe oraz przełącznik trybu S1 powinny się tam zmieścić, choć będzie trochę ciasno.

Można jednak użyć dowolnej prostokątnej obudowy. Jej długość wewnętrzna musi wynosić co najmniej 175 mm, aby zmieścić się w niej wyświetlacz LCD 20x4, połączona płytką i moduł Arduino. Wyświetlacz ma około 87 mm szerokości, co wyznacza minimalną szerokość wewnętrzną. Głębokość musi wynosić co najmniej 30 mm, aby pomieścić Arduino, wyświetlacz i przełącznik S1.

W liście elementów jest wymieniona obudowa pulpituowa Altronics H0401. Powinna zapewnić wystarczająco dużo miejsca. Ze względu na pochylą pokrywę konieczne będzie zamontowanie wyświetlacza i płytek drukowanych od wewnętrznej strony pokrywy. Ten sposób montażu uwzględniają śruby i elementy dystansowe podane w liście elementów. Do przytrzymania płytek może być użyta nakrętka przełącznika S1. Pamiętajmy, by umieścić płytkę pomiarową tak, aby był dostęp do gniazd bananowych. Gniazda można także zamontować poza płytką.

Konstrukcja

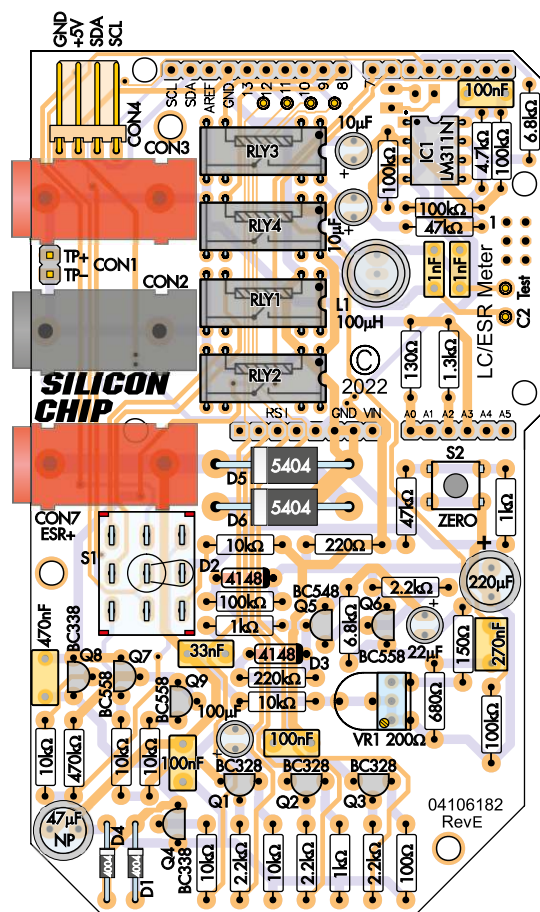
Należy przede wszystkim podjąć decyzję, czy budujemy oryginalny miernik LC i dołączamy do niego kabelkami dodatkowym moduł ESR, czy też robimy łączoną płytkę drukowaną. Naszym zdaniem to drugie rozwiązanie jest o wiele prostsze.

Rysunek 3 przedstawia okablowanie wykonywane w przypadku płytek oddzielnych, natomiast **rysunek 4** – płytkę połączoną. W przypadku wersji połączonej, jedynym elementem, który należy dodać do płytki, jest wyświetlacz, dołączany poprzez złącze CON4.

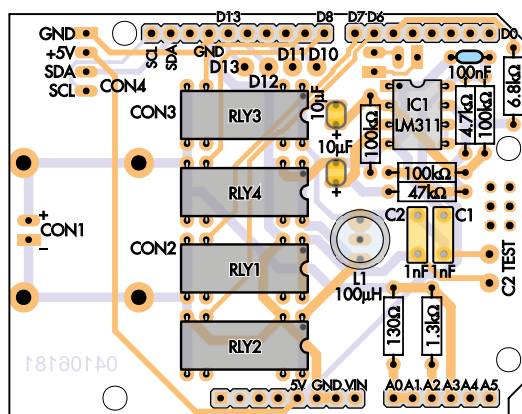
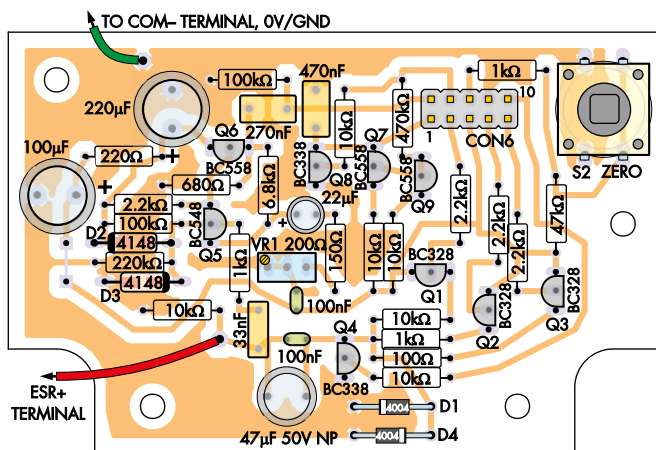
Płytką dodatkową do pomiaru ESR jest pokazana na **rysunku 5**, natomiast płytką miernika LC – na **rysunku 6** (bez gniazd, ponieważ znajdują się one poza płytką).

Następnym krokiem jest montaż kondensatorów. Zaczynamy od typów niepolaryzowanych – MKT i ceramicznych. Wartości wydrukowane na nich będą zapewne miały postać kodów, np. 102=1 nF, 104=100 nF itd. Dalej montujemy kondensatory elektrolityczne. Są one polaryzowane. Wkładamy dodatnie (dłuższe) wyprowadzenia do padów oznaczonych + (pasek na kondensatorze oznacza biegun ujemny). Jest jeden typ niepolaryzowany – 47 μ F. Znajduje się on w lewym dolnym rogu płytki. Jeśli nie

Najpierw na górną część gniazda USB nakładamy jakąś izolację, np. odcinek taśmy izolacyjnej lub taśmy Kapton. Następnie wkładamy do płytki od spodu złącza modułu Arduino. Kładziemy kawałek płytki np. uniwersalnej na wystające piny złącz



w górnej części płytki i płaskim przedmiotem popychamy złącza w dół, tak aby wierzchołki pinów pokrywały się z powierzchnią płytki uniwersalnej. Płytkę ostrożnie zdejmujemy, tak aby nie poruszyć złączy, a następnie



21

lutujemy końcowe piny złączy. W ten sposób powstanie szczelina między spodem płytki drukowanej a plastikową obudową złączy. Piny złączy będą wystawać nieco dalej i wejdą w gniazda Arduino mimo obecności gniazda USB.

I na koniec na górnej stronie płytki montujemy przełącznik 3PDT (S1). Umieszczamy go w otworach pasujących do prostokątnych wyprowadzeń przełącznika. Dzięki temu nie ma konieczności prowadzenia dziewięciu przewodów. Przewodów użyjemy tylko wtedy,

jeśli zdecydujemy się umieścić ten przełącznik gdzie indziej. W takim przypadku najlepiej będzie użyć krótkiego kabla taśmowego.

Testowanie

Dokonyjemy ostatecznej kontroli połączeń lutowanych. Upewniamy się, że nie ma zwarcia między ścieżkami i że wszystkie elementy są umieszczone na właściwych miejscach i prawidłowo zorientowane.

Jeśli zbudowaliście oddzielną płytkę ESR, możecie wykonać kilka testów przed

jej podłączeniem. Podłączamy pin 5 CON6 do zasilania +5 V, a pin 8 do 0 V. Mierzymy pobór prądu. Powinien wynosić około 1 mA. Jeśli prąd jest znacznie wyższy (albo wynosi zero), odłączamy zasilanie i szukamy błędów montażowych.

Do mocowania płytki do Arduino zalecamy użycie gwintowanych tulei dystansowych 12 mm i krótkich wkrętów maszynowych, ponieważ złącza nie obejmują w pełni gniazda Arduino i nie zapewniają wystarczającej siły łączenia. Cztery tuleje dystansowe mocujemy w otworach montażowych na Arduino. Tylko jedna musi być przykręcona do płytki. Pozostałe po prostu zapewniają odpowiedni odstęp.

Jeśli na płytce od strony gniazda USB na Arduino znajduje się punkt lutowniczy, który uniemożliwia dokręcenie śrub, należy go przyciąć na ile to możliwe.

Okablowanie

Jeśli robimy płytkę połączoną, to nie ma zbyt wiele do okablowania. Wystarczy jeden kabel 4-żyłowy od złącza CON4 do wyświetlacza. Upewniamy się, że połączenia są wykonane zgodnie z oznaczeniami na obu płytkach drukowanych, tj. GND do GND, SDA do SDA itd. W razie wątpliwości należy odnieść się do rysunku 3. Zadanie ułatwi użycie kabelka taśmowego.

Jeśli jeszcze nie przyłutowaliście adaptera I²C do wyświetlacza, to zróbcie to teraz, ponieważ łączy się z nim 4-żyłowy kabel biegnący z płytki.

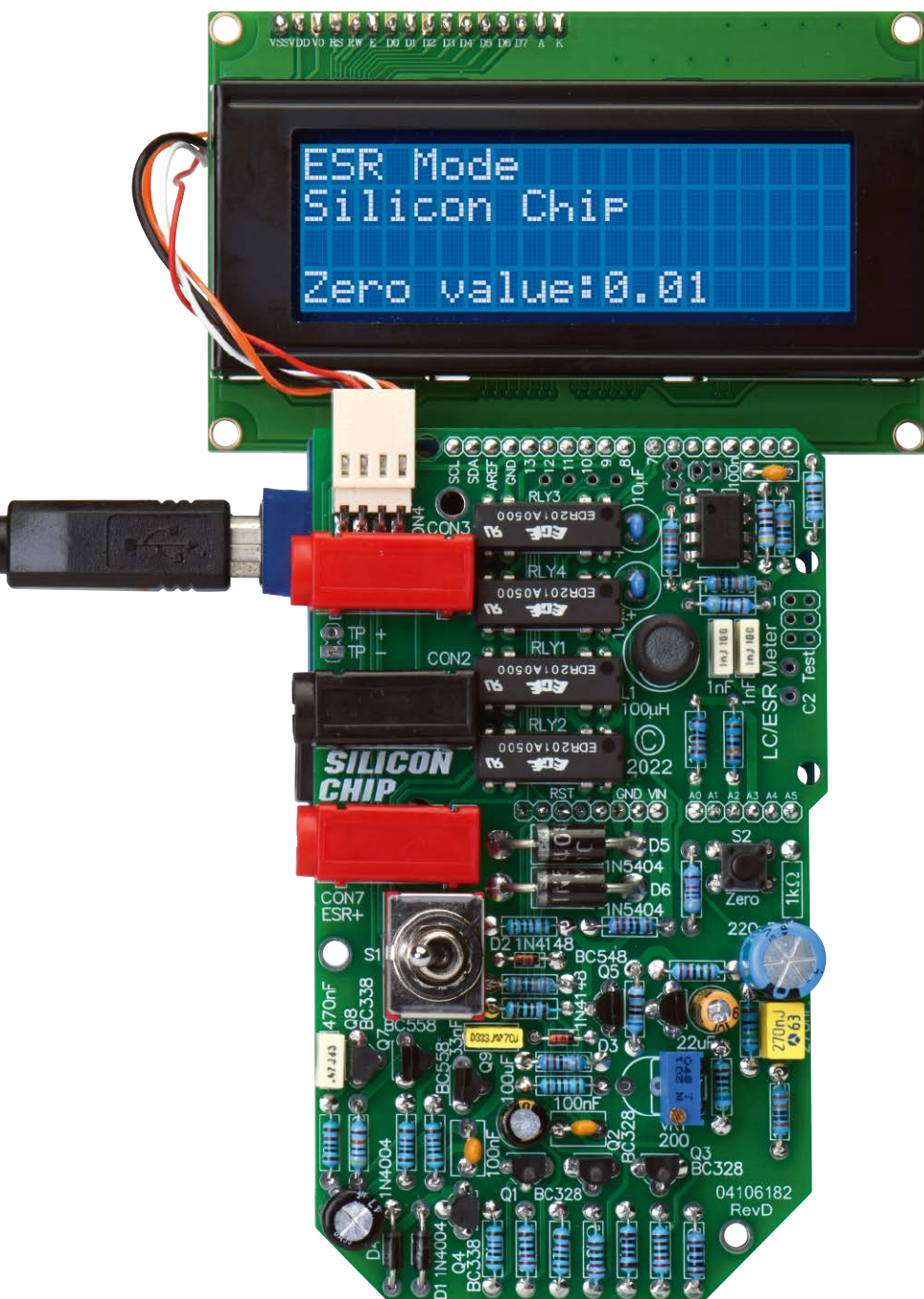
Jeśli dodajemy płytkę ESR do istniejącego miernika LC lub oddzielne płytki stosujemy z innego powodu, podłączamy je zgodnie z rysunkiem 3. Dziesięć żył idących z CON6 dla przejrzystości pokazano oddzielnie, ale również tu najlepiej jest użyć 10-żyłowego kabla taśmowego i rozdzielić poszczególne przewody tylko w takim stopniu, na ile to konieczne.

Zwróćmy uwagę, że, jak widać na rysunku 3, wyświetlacz nie jest podłączany bezpośrednio do Arduino, ponieważ sporo jego pinów zostało przekierowanych. Zauważmy też, że w tym wariancie diody D5 i D6 są zamontowane poza płytką.

Załadowanie oprogramowania

Aby wgrać do płytki Arduino Uno program, musimy mieć zainstalowane na komputerze oprogramowanie Arduino IDE („Integrated Development Environment” – „scalone środowisko programistyczne”). Jeśli go jeszcze nie mamy, pobieramy je ze strony www.arduino.cc/en/main/software.

Program na Arduino wymaga zewnętrznej biblioteki dla wyświetlacza I²C. Otwieramy



Po przełączeniu miernika w tryb ESR pojawia się na krótko ekran powitalny, pokazujący „wartość zerowania” („Zero value”). Jest to dodatkowa rezystancja przewodów pomiarowych i innych elementów włączonych w szereg z obwodem pomiarowym

IDE i wybieramy Sketch → Include Library → Manage Libraries... Następnie znajdujemy bibliotekę „liquidcrystal_pcf8574” autorstwa Matthiasa Hertela i instalujemy ją.

Otwieramy teraz plik szkicu – „ESR.ino” dla wersji samodzielnej albo „LC_ESR_Meter.ino” dla wersji połączonej. Wybieramy typ płytki „Arduino Uno” (Tools → Board Type → Arduino AVR Boards). Następnie w pozycji menu Tools → Port wybieramy port szeregowy, do którego podłączone jest Arduino. W większości wersji Uno będzie w menu rozwijanym wyświetlane jako COMx: (Arduino Uno lub podobne).

Jeśli stosujemy wyświetlacz 16×4 zamiast zalecanego 20×4, zmieniamy wiersz „lcd.begin(20,4)” na „lcd.begin(16,4)”. Szkic kompilujemy i przesyłamy do Arduino, naciskając Ctrl-U. Jeśli w dolnej części okna pojawi się komunikat „Done Uploading”, to wszystko jest w porządku. Jeśli pojawi się komunikat o błędzie, to sprawdzamy, czy biblioteka wyświetlacza I²C jest poprawnie zainstalowana i czy wybrano prawidłowy port szeregowy.

Regulacja wyświetlacza

Jeśli podświetlenie wyświetlacza nie świeci, sprawdzamy, czy na płytce adaptera I²C jest włożona zworka podświetlenia. Jeśli podświetlenie działa, ale nie widać tekstu, regulujemy potencjometr kontrastu z tyłu płytki adaptera I²C.

Zerowanie przewodów pomiarowych

Program najpierw sprawdza, czy do pamięci EEPROM została wpisana rezystancja przewodów pomiarowych. Jeśli nie, to zostanie wyświetlone żądanie przeprowadzenia procesu zerowania. Postępujemy zgodnie z instrukcjami. Wymagane będzie zwarcie przewodów pomiarowych. Gdy odczyty wartości rezystancji będą stabilne, naciskamy przycisk ZERO (S2).

Na wyświetlaczu powinna się na krótko pojawić informacja, że proces zerowania został zakończony, po czym powinno nastąpić przejście do zwykłego trybu pomiarowego. Program oczekuje, że całkowita rezystancja przewodów będzie mniejsza niż 1 Ω. W przeciwnym razie wynik pomiaru nie jest akceptowany i na chwilę pojawia się komunikat „Invalid reading or bad leads” („nieprawidłowy odczyt lub wadliwe przewody”), a proces zerowania zostaje przerwany.

W zwykłym trybie pomiarowym, gdy sondy pomiarowe są rozłączone, wyświetlacz powinien wskazywać „Over range” („przekroczenie zakresu”).

Wykaz elementów:

- 1 obudowa [Altronics H0401]
- 1 moduł mikrokontrolera Arduino Uno lub inny równoważny
- 1 alfanumeryczny wyświetlacz LCD 20×4, podświetlany na niebiesko, z interfejsem I²C [SC4203]
- 1 dwustronna płytka drukowana, 68,5 × 115,5 mm; kod Silicon Chip 04106182
- 1 cewka 100 μH, np. wysokoprądowa cewka osiowa w.cz. (L1)
- 4 przełączniki kontaktowe DIL z cewką 5 V DC (RLY1...RLY4) [Altronics S4100, Jaycar SY4030]
- 1 wieloobrotowy potencjometr montażowy 200 Ω (lepiej 470 Ω lub 500 Ω; przyp. red.), regulowany od góry (VR1)
- 1 przełącznik 3PDT (S1) [Jaycar ST0505]
- 1 przełącznik przyciskowy (S2)
- 3 gniazda bananowe 4 mm do montażu na płytce drukowanej (prostokątne);
- 1 czarne, 2 czerwone (CON2, CON3, CON7) [Silicon Chip SC4983]
- ALBO
- 3 gniazda bananowe do montażu na panel; 2 czarne, 1 czerwone (CON2, CON3, CON7)
- 1 złącze polaryzowane 1×4, kątowe, z dopasowanym wtykiem (CON4)
- 1 zestaw standardowych złączy do Arduino (jedno 10-pinowe, dwa 8-pinowe, jedno 6-pinowe)
- 1 kabel taśmowy 4-żyłowy, długość 100 mm, z gniazdami DuPont na jednym końcu
- 8 tulei dystansowych 12 mm z gwintem M3
- 9 śrub z łbem walcowym M3 × 6 mm
- 4 śruby maszynowe M3 × 6 mm z łbem stożkowym, czernione

Półprzewodniki

- 1 komparator LM311, obudowa DIP-8 (IC1) [Altronics Z2516, Jaycar ZL3311]
- 3 tranzystory PNP BC327 lub BC328, 500 mA (Q1...Q3)
- 2 tranzystory NPN BC337 lub BC338, 500 mA (Q4, Q8)
- 1 tranzystor NPN BC548 lub BC547, 100 mA (Q5)
- 3 tranzystory PNP BC558 lub BC557, 100 mA (Q6, Q7, Q9)
- 2 diody 1N4004, 400 V/1 A (D1, D4)
- 2 diody 1N4148, 75 V/200 mA (D2, D3)
- 2 diody 1N5404, 400 V/3 A (D5, D6)

Kondensatory

- 1 szt. elektrolityczny 220 μF/16 V
- 1 szt. elektrolityczny 100 μF/16 V
- 1 szt. elektrolityczny niepolaryzowany 47 μF/16 V
- 1 szt. elektrolityczny 22 μF/16 V
- 2 szt. tantalowy lub ceramiczny 10 μF/6,3 V
- 1 szt. MKT 470 nF/63 V
- 1 szt. MKT 270 nF/63 V
- 3 szt. ceramiczny wielowarstwowy lub MKT 100 nF/50 V
- 1 szt. MKT 33 nF/63 V
- 2 szt. ceramiczny NP0/C0G, MKP lub polistyrenowy 1 nF 1% [Silicon Chip SC4273]

Rezystory (wszystkie osiowe, 1/4 W, 1%)

- 1 szt. 470 kΩ
- 1 szt. 220 kΩ
- 5 szt. 100 kΩ
- 2 szt. 47 kΩ
- 7 szt. 10 kΩ
- 2 szt. 6,8 kΩ
- 1 szt. 4,7 kΩ
- 4 szt. 2,2 kΩ
- 1 szt. 1,3 kΩ
- 3 szt. 1 kΩ
- 1 szt. 680 Ω
- 1 szt. 220 Ω
- 1 szt. 150 Ω
- 1 szt. 130 Ω
- 1 szt. 100 Ω

Elementy dodatkowe w przypadku miernika z oddzielnymi płytkami drukowanymi

- 1 dwustronna płytka drukowana, 68,5×53 mm; kod Silicon Chip 04106181
- 1 przełącznik suwakowy 3PDT (S1) [Mouser 502-50209LX].
- 1 wtyk IDC 2×5 z dopasowanym gniazdem (CON6)
- Kabel taśmowy, rurka termokurcząca

Kalibracja

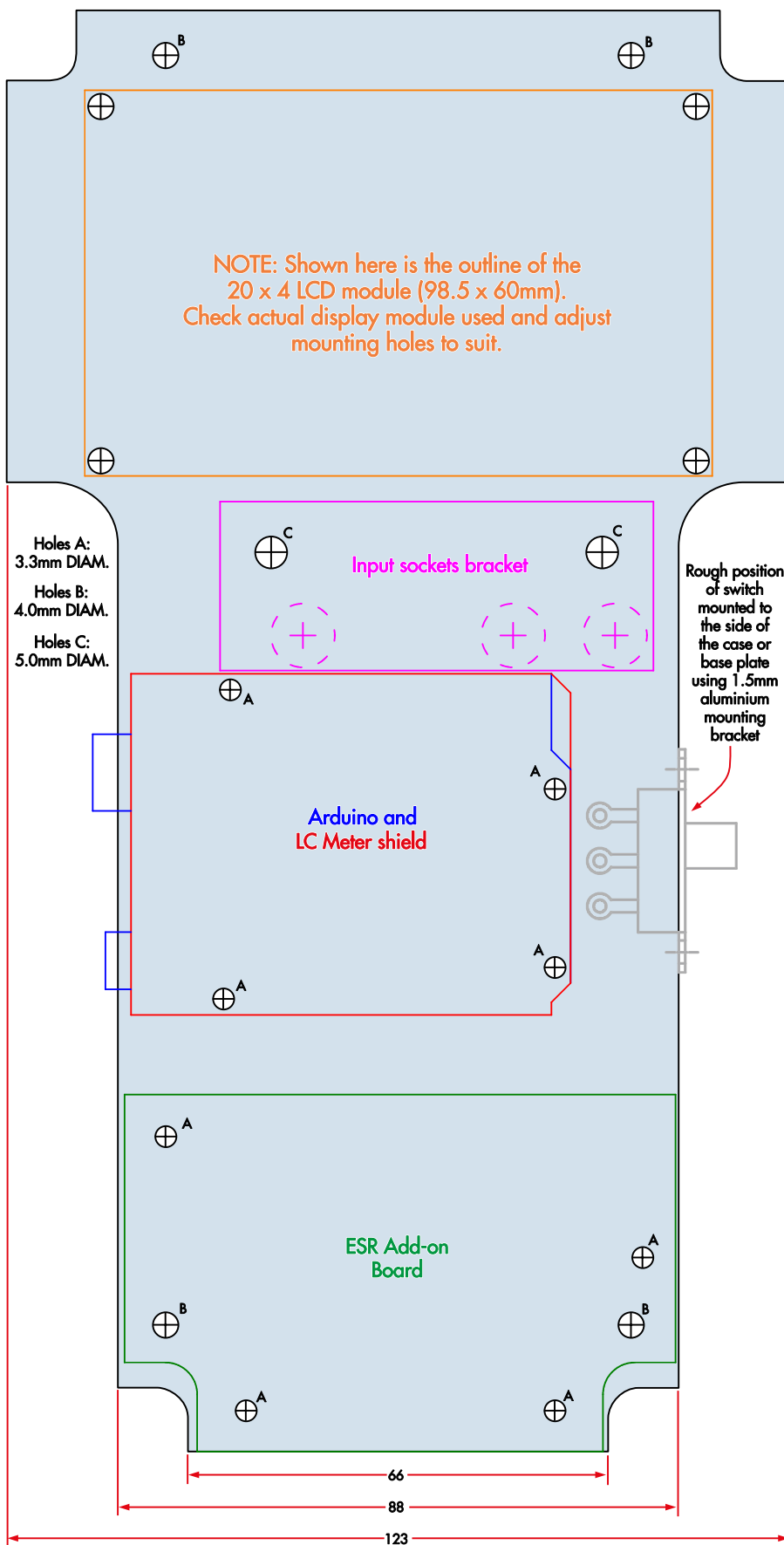
Proces kalibracji jest prosty i wymaga użycia rezystancji o znanej wartości – około 68 Ω lub 82 Ω. Wartość rezystora należy wcześniej zweryfikować multimetrem, poza tym trzeba odjąć od niej rezystancję przewodów multimetru zmierzoną przy zwarciu przewodów. Przełącznik S1, jeśli jest, przełączamy w tryb ESR. Po dołączeniu rezystora sondami, na wyświetlaczu powinna pojawić się w przybliżeniu jego wartość. Regulujemy potencjometr VR1 aż odczyt będzie równy wartości rezystora. Następnie przewody dołączamy do rezystora 1...9,9 Ω i sprawdzamy w średnim zakresie pomiarowym, czy odczyt jest zbliżony do oczekiwanego. Analogicznie sprawdzamy

również rezystor 0,1...0,9 Ω. Wynik pomiaru powinien być bardzo zbliżony do rzeczywistej wartości rezystora.

Aby zapoznać się z działaniem miernika, dobrze jest po zakończeniu kalibracji przetestować kilka wybranych kondensatorów elektrolitycznych.

Wyświetlacz pokazuje w pierwszym wierszu zmierzoną wartość ESR, w trzecim wierszu zakres (górny, środkowy lub dolny), a w czwartym – zapamiętaną rezystancję przewodu. Jest ona automatycznie odejmowana od zmierzonej wartości ESR przed wyświetleniem.

Jeśli zbudowaliście połączony miernik LC/ESR (a naszym zdaniem robi tak większość



Rysunek 7. Tu widzimy, jak w prototypie oddzielne płytki miernika zamontowano na podstawie akrylowej

konstruktorów), to teraz przyszła pora, aby przełączyć się na tryb LC i sprawdzić, czy układ przełącza się między trybami po naciśnięciu przełącznika oraz czy pomiar cewek i kondensatorów jest prawidłowy.

Montaż końcowy

Jeśli budujecie samodzielny miernik ESR, pozostanie tylko umieścić Arduino i płytkę pomiarową ESR w odpowiedniej obudowie, z odsłoniętym wyświetlaczem, dostępnymi zaciskami testowymi i ew. przyciskiem S2.

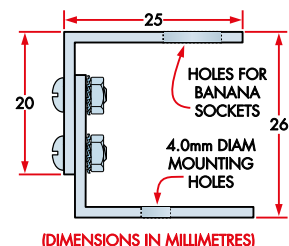
Dla tego przypadku nie pokazujemy okablowania, ale jest ono zbliżone do pokazanego na rysunku 3. Różnicą jest brak płytki pomiarowej LC, poza tym połączenia z pinów D3 i D4 Arduino przez S1 powinny być poprowadzone bezpośrednio do płytki ESR, a pin D2 – dołączony do masy. Połączenia zasilania 5 V i masy idą z Arduino do płytki ESR bezpośrednio, a nie przez płytkę pomiarową LC.

Jeśli zbudowaliście płytkę połączoną, to zamontowanie jej w obudowie jest nawet nieco prostsze. Będzie potrzebne wycięcie na wyświetlacz (chyba że obudowa ma przezroczystą pokrywę) i ew. jakieś dojdzie do przycisku S2 (np. mały otwór w obudowie). Płytkę powinna być zamontowana przy lewej krawędzi obudowy, aby gniazda bananowe przeszły przez otwory z boku, chyba że postanowiliście zamontować je gdzieś indziej i podłączyć do płytki przewodami.

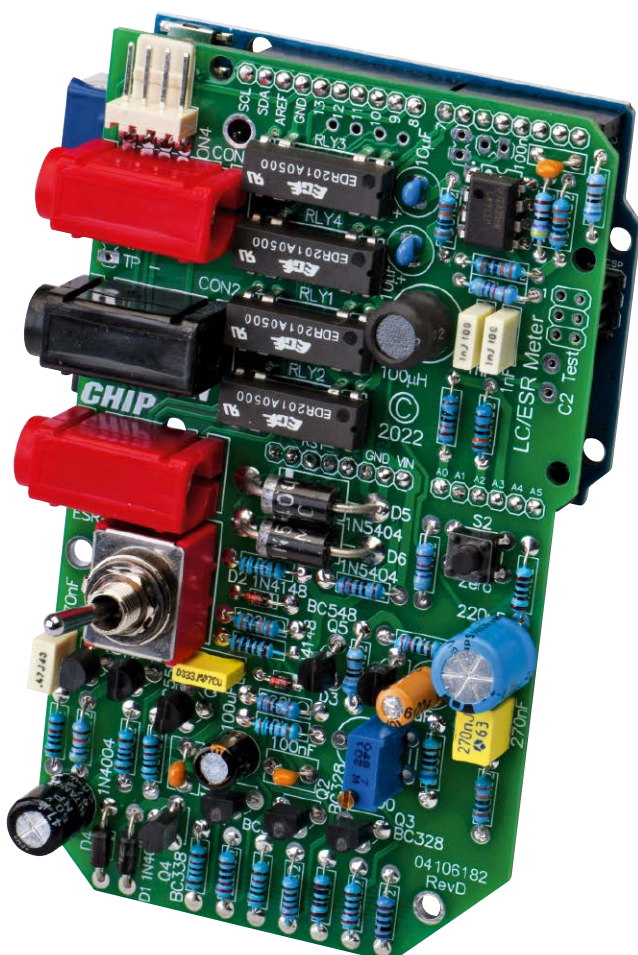
Wyświetlacz montujemy zwykle w górnej części obudowy, a główną płytkę drukowaną – poniżej, tak jak w prototypie.

W prototypie zasilanie jest doprowadzane przez złącze zasilania modułu Uno, a wtyczka przechodzi przez wycięcie po lewej stronie obudowy. Można użyć takiego sposobu lub zamontować na obudowie własne gniazdo zasilania i dołączyć je przewodami do pinów VIN i GND na module Uno.

Jeśli obudowa nie ma przezroczystej pokrywki, robimy dla wyświetlacza okienko z kawałka przezroczystego akrylu lub innego tworzywa. Można je przykleić do spodu pokrywki obudowy lub przymocować na górze wyświetlacza.



Rysunek 8. Wspornik montażowy użyty w prototypie do zamocowania gniazd bananowych. Nie jest konieczny w wersji z płytką połączoną



Jeszcze jedno zdjęcie w rzeczywistym rozmiarze zaprojektowanej przez nas płytki dołączonej do Arduino Uno. Zwróćmy uwagę, że z lewej strony wystają nie tylko gniazda bananowe, ale również złącza USB i DC zasilania modułu Uno. Dla wszystkich tych złączy trzeba wykonać otwory z boku obudowy

```
ESR: Over range!
Range: High
Leads: 0.01 ohms
```

Ekran 1. Taki komunikat pojawia się, gdy nie jest podłączony żaden element lub gdy zostanie wykryta rezystancja większa niż 100 Ω . W dolnym wierszu jest cały czas wyświetlana rezystancja przewodów. Jeśli taki ekran pojawi się po podłączeniu kondensatora, to raczej jest on już niesprawny!

```
Short test leads
and press ZERO..
0.01
```

Ekran 2. Taki ekran pojawia się po naciśnięciu i przytrzymaniu przycisku ZERO. Należy zewrzeć ze sobą sondy pomiarowe, upewnić się, że wyświetlana jest rezystancja o małej wartości, po czym potwierdzić wynik, ponownie naciskając ZERO

Szczegóły montażu prototypu

Wiele osób będzie chciało zmontować układ w sposób podobny jak w prototypie. Jest to jednak zalecane tylko gdy stosujemy oddzielne płytki, natomiast nieodpowiednie, gdy mamy płytkę połączoną. W prototypie obie płytki, wsporniki montażowe i wyświetlacz zostały zamontowane na akrylowej płytce o grubości 4 mm, pokazanej na **rysunku 7**.

Na **rysunku 8** pokazano wspornik wzmacniający, używany do przymocowania gniazd bananowych do obudowy. W zależności od rodzaju obudowy może on nie być potrzebny, a gniazda można zamontować bezpośrednio na obudowie lub użyć gniazd montowanych na płytce. Podobny wspornik należy wykonać dla przełącznika S1, jeśli zamontowanie go bezpośrednio do obudowy nie jest możliwe.

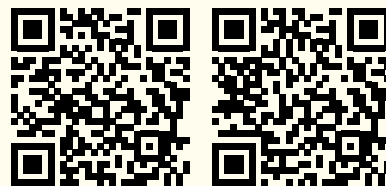
Należy wziąć pod uwagę, że w obudowie powinien być otwór o średnicy około 3 mm, umożliwiający dostęp do przełącznika ZERO (S2). Dostęp przyda się, jeśli będzie trzeba ponownie skompensować rezystancję przewodów pomiarowych. Bez otworu konieczne będzie uciążliwe otwieranie obudowy przed dokonaniem kalibracji.

Tabela 1 może zostać wydrukowana na papierze samoprzylepnym – albo na zwykłym i zostać zaalaminowana – i przyklejona do obudowy urządzenia jako informacja. Tak właśnie zrobiłem w przypadku prototypu. Jeśli używacie płytki połączonej, to pamiętajcie, że przełącznik trybu pracy znajduje się na płytce. Jeśli gniazda bananowe również będą zmontowane na płytce, to trzeba będzie starannie wymierzyć miejsca otworów w obudowie na przełącznik i gniazda.

Uwagi

Nie zapominajcie rozładowywać kondensatorów przed ich testowaniem! ■

Steve Matthysen



Materiały dodatkowe dostępne są na stronie Silicon Chip:
<https://www.siliconchip.com.au/Shop/8/4618>
<https://www.siliconchip.com.au/Shop/8/4634>

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

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

Wysokowydajny subwoofer aktywny do domowego sprzętu Hi-Fi, część 2

W poprzednim odcinku przedstawiliśmy parametry nowego subwoofera o wysokiej jakości oraz szczegółowo omówiliśmy konstrukcję jego obudowy. W tym artykule, kończącym opis projektu, dokończymy budowę aktywnego subwoofera: opiszemy wykonanie i montaż wewnętrznego wzmacniacza o mocy 180 W, wykonamy okablowanie, zamontujemy głośnik oraz zainstalujemy nożyki.

Po zbudowaniu wzmacniacza Ultra-LD Mk.3 lub Mk.4 należy wykonać niestandardowy metalowy wspornik, wywiercić otwory w radiatorze oraz połączyć wspornik, radiator, wzmacniacz i zasilacz w kompaktowy moduł wzmacniający. Gotowy moduł można następnie umieścić w prostokątnym wycięciu o wymiarach 220 mm × 170 mm, przygotowanym wcześniej w tylnej ścianie subwoofera.

Jeśli moduł wzmacniacza nie został jeszcze zbudowany, warto sięgnąć do oryginalnych artykułów opisujących jego konstrukcję. W przypadku wersji Ultra-LD Mk.3 szczególnie przedstawiono w wydaniu magazynu Silicon Chip z sierpnia 2011 roku (siliconchip.au/Article/1129), natomiast konstrukcję wersji Ultra-LD Mk.4 opisano w numerze z września 2015 roku (siliconchip.au/Article/8959).

Aby uzyskać najlepsze działanie subwoofera, należy zwrócić uwagę na kilka istotnych szczegółów konstrukcyjnych, takich jak nawijanie i montaż cewki filtra wyjściowego. Dlatego zdecydowanie zalecamy zapoznanie się z odpowiednim artykułem przed rozpoczęciem lub w trakcie budowy modułu wzmacniacza Ultra-LD. Przed montażem tranzystorów wyjściowych warto jednak przeczytać poniższe wskazówki dotyczące budowy wzmacniacza.

Konieczne będzie również wykonanie modułu zabezpieczenia głośników w wersji wielokanałowej, jednak z wykorzystaniem tylko jednego przełącznika. Elementy związane

z drugim przełącznikiem można pominąć. Na przykład można zamontować przełącznik RLY2, pomijając wszystkie elementy znajdujące się na lewo od diody D2 i rezystora 100 kΩ nad nią.

Po zmontowaniu obu modułów i przygotowaniu pozostałych elementów można przystąpić do montażu całości.

Wykonanie wspornika

Jako główną płytę montażową zastosowałem aluminiowy panel o grubości 3 mm. Do niego zamocowałem gięty wspornik aluminiowy o grubości 1,5 mm do montażu transformatora oraz kątownik z blachy aluminiowej przeznaczony do instalacji modułu zabezpieczenia głośników.

Zmontowane panele pokazano na fotografii 11 (zwróć uwagę na różnice w wycięciu względem wersji ostatecznej). Wszystkie elementy modułu wzmacniacza należy zamontować na tych panelach, głównie na wsporniku centralnym.

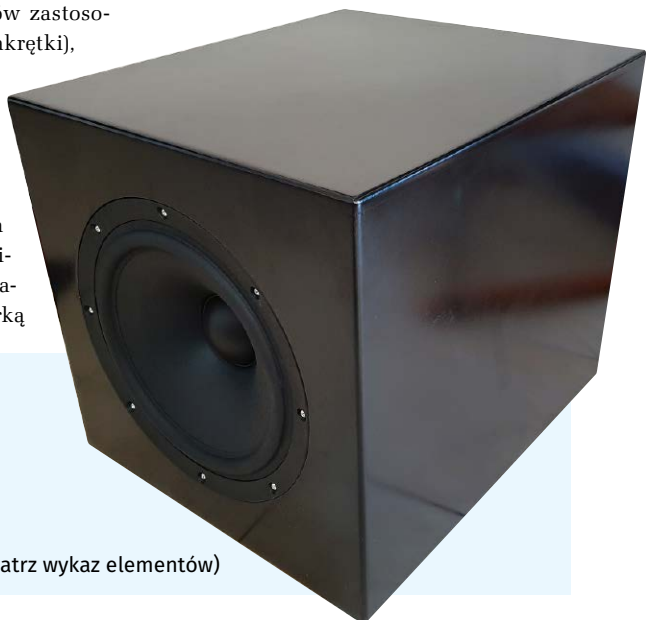
Do połączenia tych elementów zastosowałem nakrętki nitowane (nitonakrętki), które zapewniają estetyczny wygląd i wygodę montażu. Zamiast nich można jednak użyć zwykłych śrub z nakrętkami.

Wspornik w kształcie litery L dla modułu zabezpieczenia głośników można wykonać, zginając blachę aluminiową w imadło. Jeśli nie dysponujesz giętarką

do blachy, wykonanie większego wspornika dla zasilacza może być trudniejsze. Odpowiednie wsporniki można znaleźć w sklepach z narzędziami. Należy przy tym pamiętać, że transformator jest ciężki, a konstrukcja musi uwzględniać obciążenia uderzeniowe, na przykład podczas upuszczenia urządzenia.

Zasilacz jest prosty. Jego schemat ideowy jest doprowadzane przez złącze CON1 i podawane na transformator T1 poprzez bezpiecznik F1 oraz wyłącznik zasilania S1. W zależności od typu transformatora może mieć pojedyncze uzwojenie pierwotne 230 V lub podwójne 115 V. Dwa uzwojenia wtórne 40 V AC są połączone z mostkiem prostowniczym BR1 oraz zespołem kondensatorów filtrujących, tworząc symetryczne szyny zasilania ± 57 V DC.

Wzmacniacz subwoofera musi dostarczać dużą moc przez dłuższy czas, dlatego zastosowano zespół kondensatorów o łącznej



Co jest potrzebne do zbudowania aktywnego subwoofera

Wzmacniacz Ultra-LD Mk.3 lub Mk.4

Mk.3 – lipiec-wrzesień 2011, siliconchip.au/Series/286

Mk.4 – sierpień-październik 2015, siliconchip.au/Series/289

Wielokanałowe zabezpieczenie głośnika (4-CH)

Styczeń 2022, siliconchip.au/Article/15171

Materiały i elementy mechaniczne:

płyty na obudowę, materiał tłumiący, radiator, przewody oraz inne elementy (patrz wykaz elementów)

pojemności 16 mF w każdej szynie zasilania. Takie rozwiązanie zapewnia duży zapas energii, jednak nawet przy zastosowaniu tylko dwóch kondensatorów 8000 μ F urządzenie będzie działać poprawnie.

Rezystor 270 Ω o mocy 10 W obniża napięcie do poziomu odpowiedniego do zasilania modułu zabezpieczenia głośników oraz zmniejsza straty mocy w jego stabilizatorze.

Konstrukcja wzmacniacza płytkowego

Wzmacniacz Ultra-LD zamontowałem do płyty montażowej oraz radiatora. Płyta o grubości 3 mm znajduje się pomiędzy tranzystorami wyjściowymi a radiatorem, co pokazano na fotografii 12.

Ważne jest, aby płyta montażowa była wolna od wgnieceń i zarysowań, a radiator został do niej zamocowany z użyciem odpowiedniej warstwy pasty termicznej. Ułatwi to montaż i poprawi skuteczność chłodzenia.

Aby zapewnić prawidłowe dopasowanie otworów montażowych radiatora i płyty montażowej do tranzystorów, radiator i płytę montażową nawierciłem oraz skrzyłem jeszcze przed zmontowaniem wzmacniacza. Następnie zamontowałem tranzystory, a dopiero potem przylutowałem je do płytki drukowanej. Dzięki temu ich położenie względem otworów montażowych i płytki drukowanej było prawidłowe.

Na tym etapie nie należy stosować izolatorów – zostaną dodane później. Po przylutowaniu tranzystorów w ten sposób można rozmontować układ, mając pewność, że podczas ponownego montażu wszystkie elementy będą prawidłowo dopasowane.

Wiercenie radiatora

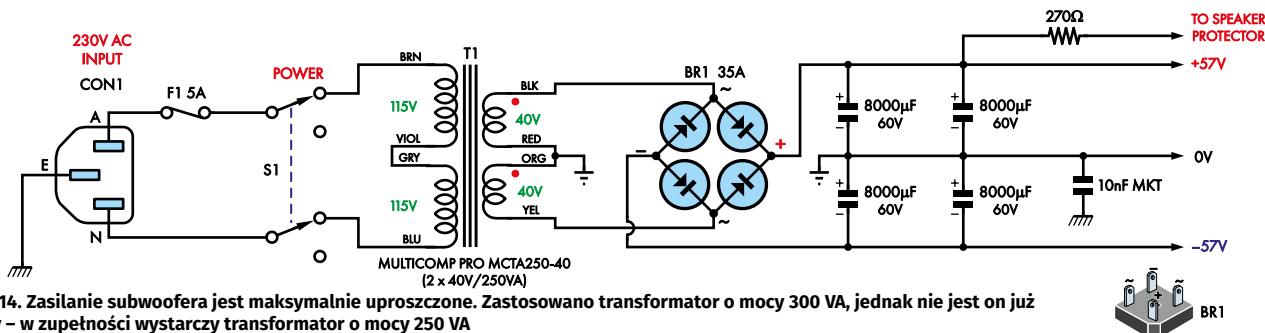
Na rysunku 15 pokazano miejsca, w których należy wywiercić otwory w radiatorze. Otwory najpierw zaznaczono na płycie montażowej (rysunek 16), a następnie w radiatorze wywiercono i nagwintowano cztery narożne otwory montażowe, po czym przykręcono go do płyty montażowej śrubami M3. Kolejno wywiercono otwory o średnicy



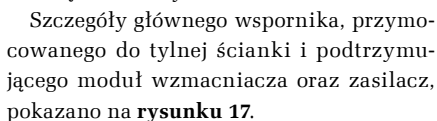
Fotografia 11. Na tym wsporniku zamontowana jest większość elementów modułu wzmacniacza



Fotografia 12. Wzmacniacz Ultra-LD Mk.4 zamontowany na wsporniku, gotowy do podłączenia



Rysunek 14. Zasilanie subwoofera jest maksymalnie uproszczone. Zastosowano transformator o mocy 300 VA, jednak nie jest on już dostępny – w zupełności wystarczy transformator o mocy 250 VA



Rysunek 15. Szczegóły wiercenia otworów w radiatorze. Zastosowany radiator jest taki sam jak w oryginalnych artykułach Ultra-LD Mk.3/4, jednak sposób jego montażu jest inny

Rysunek 16. Tylna ścianka obudowy została wykonana z aluminium o grubości 3 mm, wyciętego i nawierconego zgodnie z rysunkiem. Po zakończeniu montażu warto ją pomalować na czarno. Należy upewnić się, że prostokątny otwór na przetwornik kołtyskowy jest możliwie najmniejszy, ale wystarczający do jego zatrzasknięcia

Wspornik modułu zabezpieczenia głośników jest mocowany za pomocą dwóch śrub radiatora. Wykonano go z blachy aluminiowej o grubości 1,5 mm, kątem 90° – szczegóły pokazano **18.**

Do zamocowania rezystora drutowego $270\ \Omega$ o mocy $10\ \text{W}$, obniżającego napięcie dodatniej szyny zasilania $+57\ \text{V}$ o około $15\ \text{V}$, zastosowano niewielki zacisk. Rezystor ten jest włączony szeregowo w torze zasilania modułu zabezpieczenia głośników.

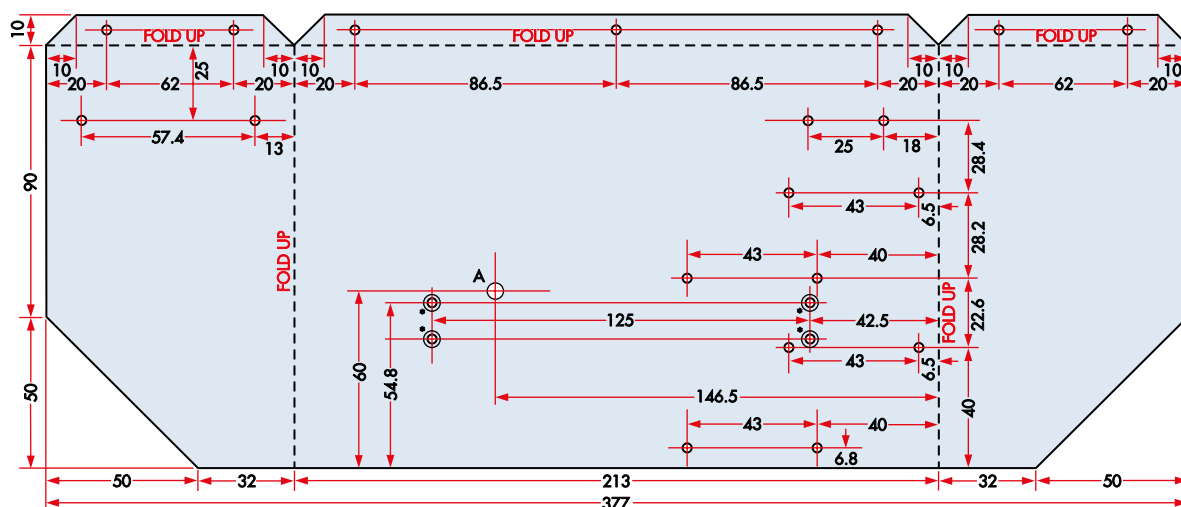
Po przygotowaniu elementów metalowych należy najpierw przeprowadzić montaż „na sucho” i zaplanować kolejność prac. Gdy wszystkie elementy są prawidłowo

dopasowane, układ powinien wyglądać jak na **rysunkach 19...21** i **fotografiach 12...14**.

Na tym etapie należy tymczasowo zamontować płytkę wzmacniacza, przykręcić tranzystory wyjściowe w miejscach montażu (bez izolatorów) i przyłutować je do płytki drukowanej. Zapewni to prawidłowe dopasowanie wszystkich otworów.

Montaż końcowy należy rozpocząć od zamocowania listwy zaciskowej, transformatora, śruby uziemiającej oraz mostka prostowniczego. Pod mostkiem należy zastosować niewielką ilość pasty termicznej.

Następnie należy zamontować wsporniki o wysokości 15 mm dla modułu wzmacniacza (tylko w dwóch narożnikach najbardziej oddalonych od radiatora), upewniając się, że otwór pod śrubę przechodząca



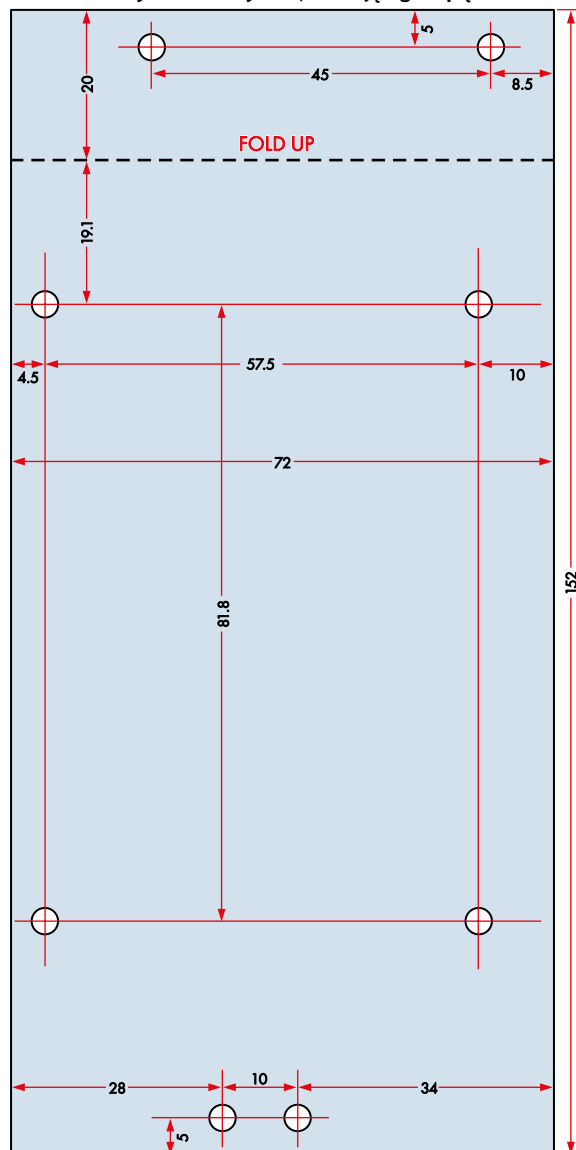
HOLE A IS 6mm DIAMETER; ALL OTHER HOLES 3mm DIAMETER
(* HOLES MARKED WITH AN ASTERISK ARE COUNTERSUNK)

(SHOWN HERE AT 40% OF FULL SIZE)

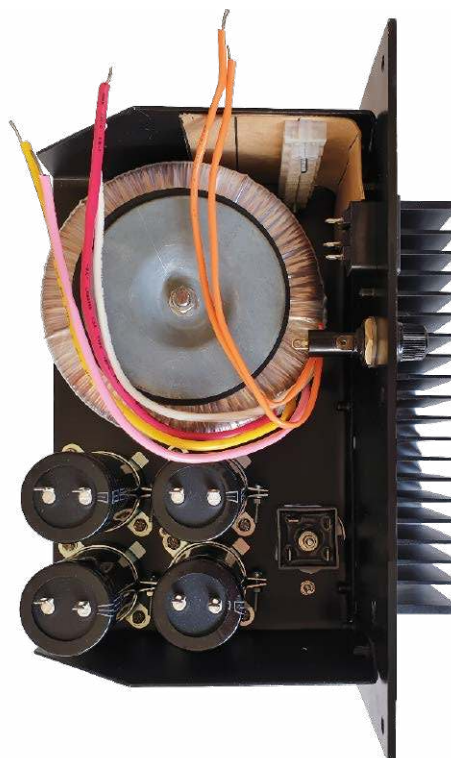
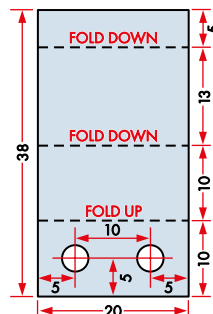
ALL DIMENSIONS IN MILLIMETRES

Rysunek 17. Wspornik należy wykonać z blachy aluminiowej o grubości 1,5 mm, zagiętej pod odpowiednim kątem, i pomalować na czarno. Mocuje się go prostopadłe do tylnej ścianki

Rysunek 18. Większy wspornik umożliwia zamontowanie modułu zabezpieczenia głośników w przestrzeni obok wzmacniacza. Mniejszy wspornik służy do zamocowania rezystora o mocy 10 W, obniżającego napięcie zasilania tego modułu



ALL HOLES IN THESE
BRACKETS ARE 3.0mm IN DIAMETER
ALL DIMENSIONS IN MILLIMETRES



Fotografia 13. Spód całkowicie zmontowanego modułu wzmacniacza, jeszcze niepodłączanego

pod transformatorem został pogłębiony oraz że zastosowano śrubę z łbem stożkowym.

Należy wyciąć kawałek preszpanu lub podobnego materiału izolacyjnego i umieścić go pod listwą zaciskową, aby zapewnić izolację w przypadku połuzowania lub odłączenia przewodów. Listwę należy zamontować w odpowiedniej odległości od tylnej ścianki, tak aby nie kolidowała z okablowaniem gniazda sieciowego (IEC).

Wystarczy zastosować trójtrową listwę zaciskową do podłączenia uzwojenia pierwotnego transformatora, w tym do jego połączenia szeregowego.

Następnie należy zamontować kondensatory. Dla zachowania porządku wszystkie ich zaciski ujemne powinny być skierowane w tę samą stronę.

Moduł wzmacniacza montuje się po odwróceniu go i zamocowaniu z użyciem tulei izolacyjnych oraz podkładek, zgodnie z opisem w artykułach magazynu Silicon Chip z sierpnia 2011 lub września 2015 roku. Moduł należy przykręcić do wcześniej zamontowanych wsporników o wysokości 15 mm, stosując podkładki sprężynujące pod śrubami M3.

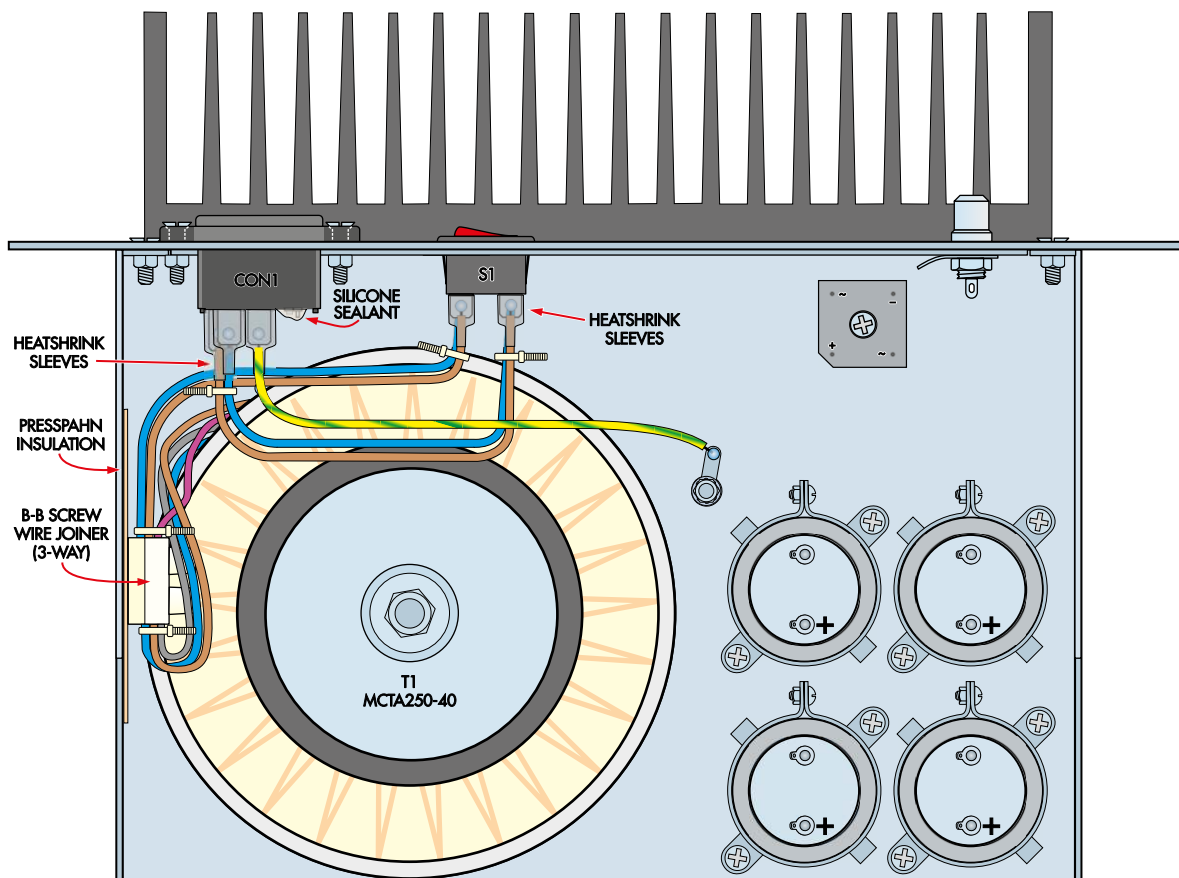
Fotografia 14. Widok spodu ukończonego modułu wzmacniacza, pokazujący całe okablowanie. Należy jednak pamiętać, że w tej wersji zastosowano oddzielną oprawkę bezpiecznika oraz przełącznik. Włastry układ należy wykonać zgodnie z poprawionym projektem, z oprawką bezpiecznika zintegrowaną w gnieździe zasilania sieciowego IEC.



Następnie należy zamontować moduł zabezpieczenia głośników na jego wspornikach. Należy przy tym pamiętać o podłączeniu do wejścia zasilania cienkiego przewodu o długości około 20 cm, ponieważ później dostęp do tego złącza będzie utrudniony.

Należy pamiętać, aby rezystor 270 Ω o mocy 10 W włączyć szeregowo w torze zasilania modułu zabezpieczenia głośników. Zmniejsza to straty mocy w stabilizatorze tego układu.

W przypadku zastosowania tylko jednego przekaźnika nie jest to konieczne,



Rysunek 19. Widok spodu modułu wzmacniacza, pokazujący okablowanie sieciowe. Przewody należy prowadzić możliwie krótko, związając opaskami kablowymi i zaizolować wszystkie odstłonięte połączenia. Podczas montażu transformatora należy upewnić się, że nie znajduje się on zbyt blisko narożnika, aby nie kolidował z okablowaniem gniazda zasilania sieciowego IEC. Należy stosować tę konfigurację, a nie tę pokazaną na zdjęciach prototypu, ponieważ zapewnia ona oddzielenie obwodów sieciowych od części niskonapięciowej



Fotografia 15. Widok od strony wzmacniacza po zakończeniu montażu i okablowania

pod warunkiem użycia radiatora Altronics H0655 w module zabezpieczenia głośników, choć jego zastosowanie nie zaszkodzi.

Po zamontowaniu wszystkich elementów większość pozostałych prac sprowadza się do wykonania okablowania, zgodnie z **rysunkami 19** (okablowanie sieciowe), **20**

(okablowanie zasilania niskiego napięcia) oraz **21** (okablowanie modułu wzmacniacza).

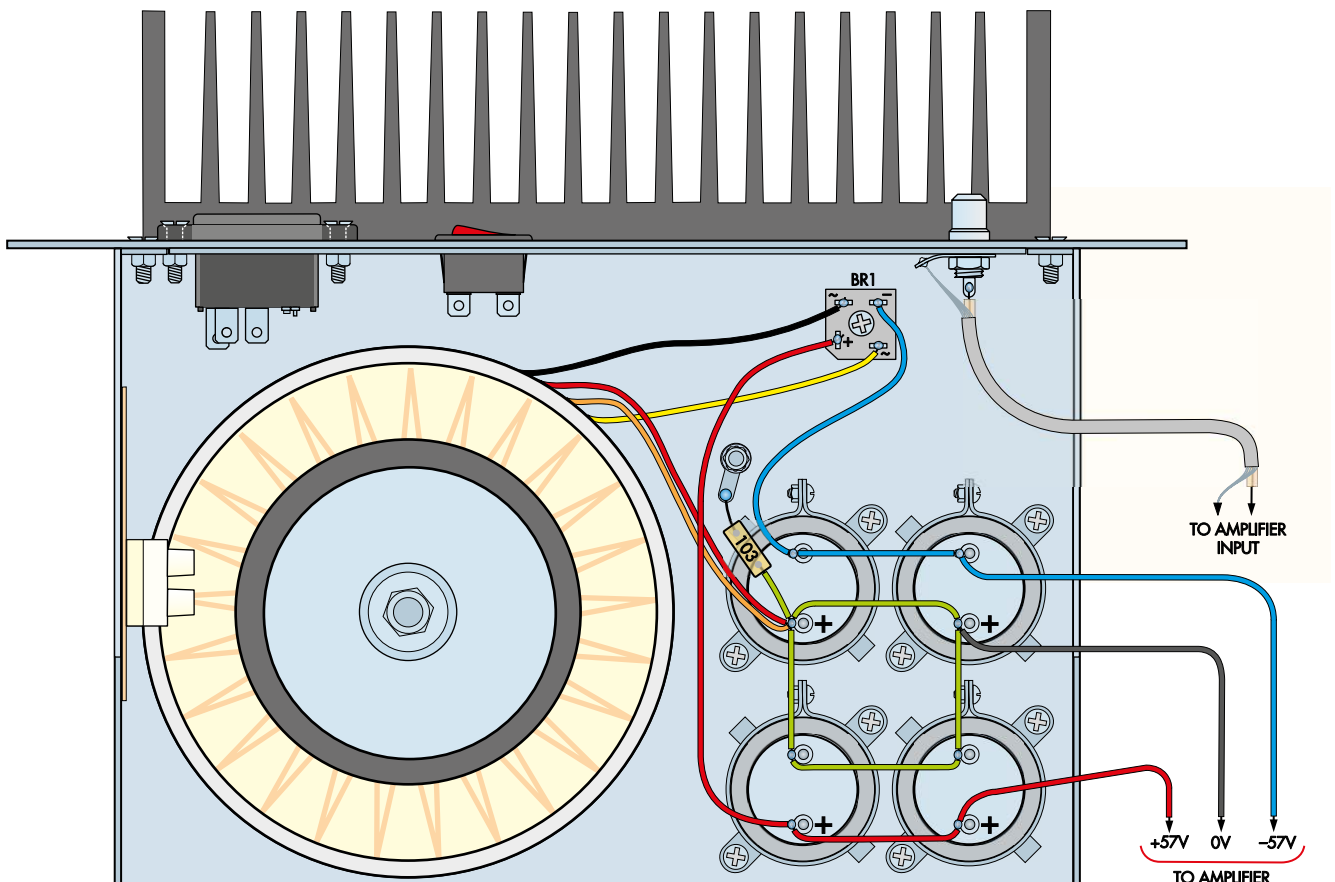
Do wszystkich obwodów zasilania sieciowego należy stosować przewody o odpowiednim przekroju, przeznaczone do prądu co najmniej 7,5 A. Wszystkie połączenia należy starannie zaizolować, aby zapobiec

przypadkowemu dotknięciu elementów znajdujących się pod napięciem.

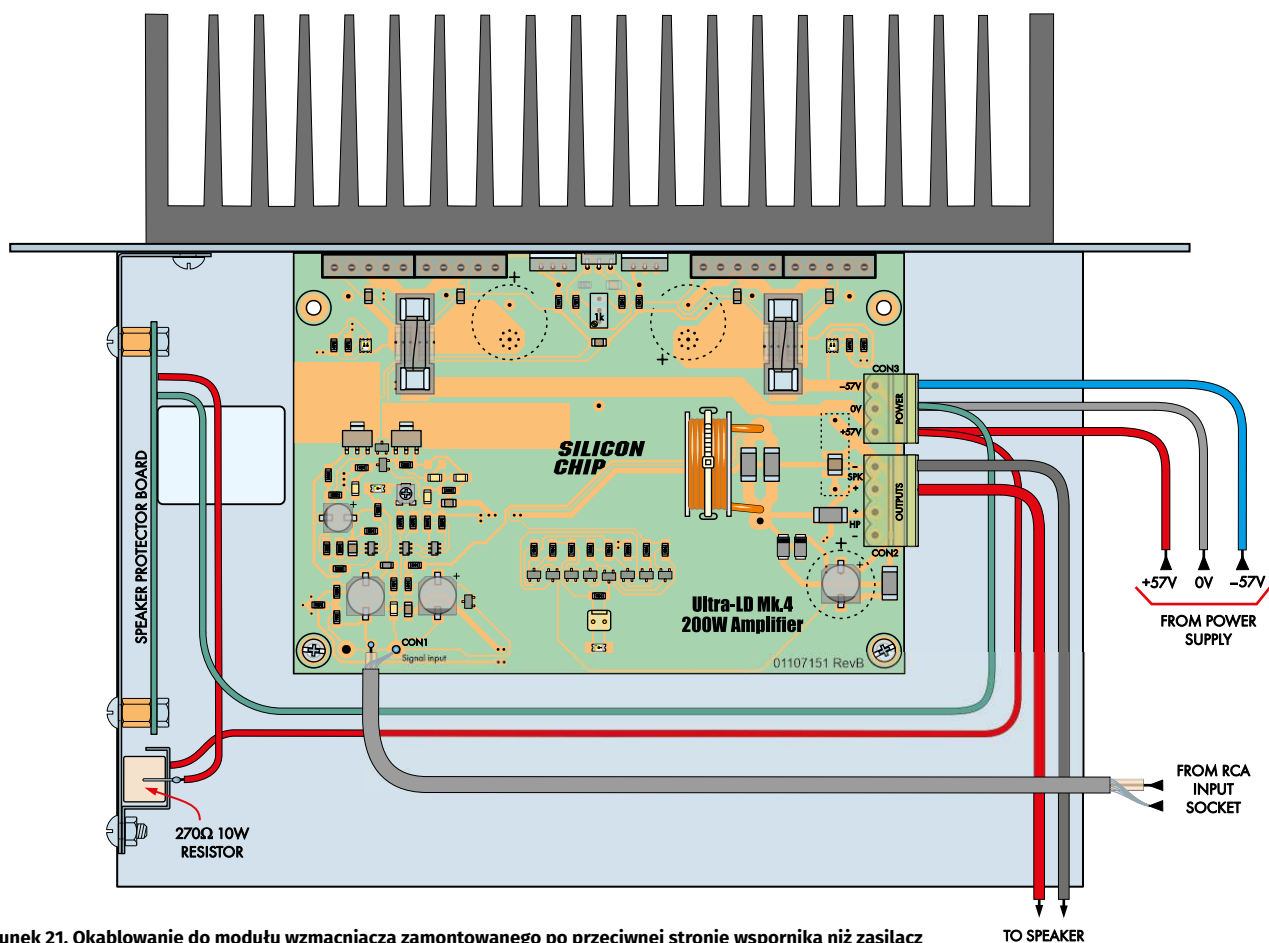
Należy zauważyć, że ostateczna wersja urządzenia różni się nieco od tej pokazanej na zdjęciach. Zamiast oddzielnej oprawki bezpiecznika zastosowano gniazdo zasilania sieciowego IEC ze zintegrowaną oprawką bezpiecznika, a przełącznik zasilania zastąpiono przełącznikiem kołyskowym. Rozwiązanie to upraszcza okablowanie i pozwala oddzielić obwody sieciowe od części niskonapięciowej. W tym zakresie należy kierować się schematami, a nie fotografiami.

Okablowanie można wykonywać zgodnie z poniższymi krokami.

Należy zamontować zacisk uziemienia i połączyć przewód ochronny (zielono-żółty) z punktem uziemienia gniazda zasilania sieciowego IEC. *Zalecane jest stosowanie zaciskanych końcówek oczkowych, które przy prawidłowym wykonaniu są bardziej wytrzymałe niż połączenia lutowane – przyp. red. wydania oryginalnego.* Uziemienie można i należy umieścić możliwie blisko gniazda IEC; na schemacie pokazano je dalej jedynie dla zachowania czytelności. Śruba uziemienia powinna łączyć zacisk



Rysunek 20. Przedstawione okablowanie jest podobne do tego z rysunku 19, jednak pokazano tu wyłącznie połączenia zasilacza o niższym napięciu (około 114 V DC). Aby zminimalizować tętnienia napięcia zasilającego moduł wzmacniacza, należy ściśle przestrzegać tego schematu



Rysunek 21. Okablowanie do modułu wzmacniacza zamontowanego po przeciwnej stronie wspornika niż zasilacz

przewodu ochronnego wyłącznie z obudową i z żadnym innym elementem. Należy upewnić się, że w miejscu styku nie ma farby ani innych powłok uniemożliwiających dobry kontakt elektryczny; w razie potrzeby należy je usunąć. Drugi zacisk należy połączyć z kondensatorem 10 nF, a następnie krótkim



Fotografia 16. Materiał tłumiący zamocowany zszywaczem, stosując podwójną warstwę włókna poliestrowego. Jest to konieczne do tłumienia fal akustycznych emitowanych przez tylną stronę głośnika oraz ograniczenia rezonansów

zielonym przewodem połączyć kondensator z punktem 0 V zespołu kondensatorów.

Wyprowadzenia uzwojenia wtórnego transformatora należy przyciąć do odpowiedniej długości, tak aby sięgały do wejść mostka prostowniczego. Następnie należy je zaciśnąć i podłączyć lub przyłutować do mostka prostowniczego.

Dodatnie i ujemne wyjścia mostka prostowniczego należy połączyć z zespołem kondensatorów, stosując przewody o odpowiednim przekroju (np. czerwony i biały). Opcjonalnie można zastosować złącza zaciskane przy mostku prostowniczym.

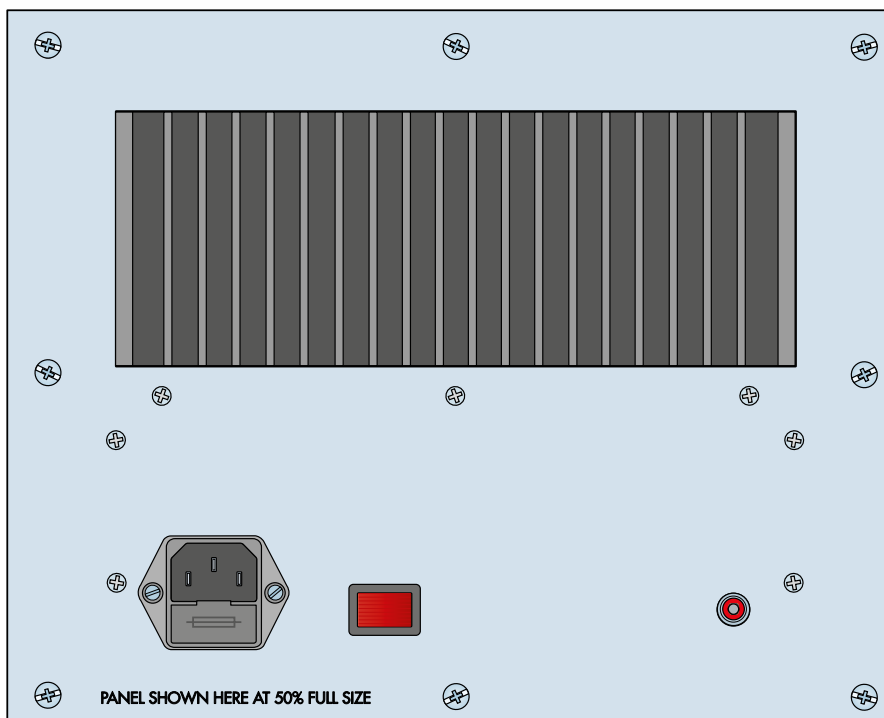
Odsłonięte metalowe elementy gniazda zasilania sieciowego IEC należy pokryć silikonem o neutralnym utwardzaniu.

Przewód fazowy (brązowy) należy przyłutować do gniazda zasilania sieciowego IEC, następnie poprowadzić do jednego styku przełącznika i dalej do listwy zaciskowej. W analogiczny sposób należy podłączyć przewód neutralny (niebieski). Wszystkie odsłonięte połączenia należy zabezpieczyć koszulką termokurczliwą. Przewody należy skrócić razem i unieruchomić opaskami kablowymi, aby zapobiec ich poluzowaniu w przypadku uszkodzenia połączenia.

Nie zaleca się stosowania końcówek widelkowych do podłączeń przy gnieździe zasilania sieciowego IEC (z wyjątkiem przewodu ochronnego), ponieważ dostępna przestrzeń jest ograniczona ze względu na bliskość transformatora. Zaleca się lutowanie przewodów w taki sposób, aby były prowadzone nad korpusem transformatora w kierunku przełącznika. Dla uzyskania większego prześwitu nie ma potrzeby wyginania wyprowadzeń gniazda IEC, choć w razie potrzeby można to zrobić. Do podłączenia przełącznika można zastosować zaciskane końcówki widelkowe, ponieważ znajduje się on tuż nad transformatorem.

Uzwojenie pierwotne transformatora należy podłączyć do przewodu fazowego prowadzonego przez przełącznik, za pośrednictwem listwy zaciskowej. Połączenia należy starannie zaizolować. Jeśli transformator ma dwa uzwojenia pierwotne, należy je połączyć szeregowo, wykorzystując dodatkowy zacisk na listwie zaciskowej (najlepiej pomiędzy zaciskami używanymi do pozostałych połączeń uzwojeń pierwotnych).

Kondensatory należy połączyć grubymi przewodami: czerwonym dla dodatniej szyny zasilania i czarnym dla ujemnej. Zaciski



Rysunek 22. Widok tylnej strony modułu wzmacniacza po zakończeniu montażu

ujemne kondensatorów należy połączyć razem grubym przewodem (np. zielonym) i dołączyć do odczepu środkowego transformatora.

Następnie należy skrócić razem odcinki przewodów o długości około 40 cm: czerwony, czarny i zielony, o odpowiednim przekroju. Przewody należy podłączyć odpowiednio do zacisków +57 V, -57 V oraz punktu 0 V zespołu kondensatorów. Następnie należy je

poprowadzić do wejścia zasilania wzmacniacza mocy i przyciąć do odpowiedniej długości.

Odsłonięte połączenia ± 57 V na kondensatorach należy zabezpieczyć kawałkami koszulki izolacyjnej, przymocowanymi silikonem o neutralnym utwardzaniu. Zapobiega to przypadkowemu dotknięciu punktów znajdujących się pod napięciem stałym około 114 V.



Fotografia 17. Widok z tyłu subwoofera, nieco różniący się od wersji ostatecznej

Szynę +57 V ze wzmacniacza należy podłączyć do rezystora 270 Ω (jeśli jest stosowany), a następnie z jego drugiego końca poprowadzić połączenie do dodatniego wejścia modułu zabezpieczenia głośników. Do tego połączenia można użyć cienkiego przewodu.

Punkt 0 V wzmacniacza należy podłączyć do wejścia GND modułu zabezpieczenia głośników.

Wyjście wzmacniacza należy podłączyć do wejścia „AMP” modułu zabezpieczenia głośników. Zacisk „SPKR” należy podłączyć do dodatniego zacisku głośnika.

Punkt 0 V (masa) wzmacniacza należy połączyć z ujemnym zaciskiem głośnika.

Montaż końcowy

Montaż aktywnego subwoofera jest prosty, ponieważ sprowadza się głównie do wykonania obudowy i modułu wzmacniacza. Materiał tłumiący należy umieścić po bokach, na górze i na dole obudowy, jak pokazano na **fotografii 16**. Nie należy zasłaniać otworu bass-reflex, ponieważ podczas pracy przepływa przez niego duża ilość powietrza.

Wyjście wzmacniacza należy podłączyć do głośnika za pomocą grubego przewodu głośnikowego, zwracając uwagę na prawidłową polaryzację – wyjście „+” wzmacniacza powinno być połączone z czerwonym zaciskiem głośnika.

Następnie należy zamontować moduł wzmacniacza, uprzednio przyklejając piankową taśmę uszczelniającą wokół krawędzi otworu w tylnej ścianie obudowy. Moduł należy przymocować za pomocą ośmiu śrub o długości 16 mm. Na **rysunku 22** i **fotografii 17** pokazano, jak powinien wyglądać po zamontowaniu.

Na końcu należy zamontować głośnik, oklejając wcześniej krawędź otworu piankową taśmą uszczelniającą. Do jego zamocowania należy użyć ośmiu śrub o długości 16 mm.

Do subwoofera zastosowano duże filcowe nóżki, chroniące podłogę przed zarysowaniem. Nie jest to lekki sprzęt.

Przed rozpoczęciem odsłuchu należy przeprowadzić prosty test nowego subwoofera i upewnić się, że wszystkie funkcje działają prawidłowo. Jeśli subwoofer współpracuje z aktywnymi głośnikami monitorującymi (EdW 5-6/2025), należy zapoznać się z instrukcją dotyczącą regulacji jego poziomu, aby prawidłowo dopasować go do zestawu. ■

Phil Prosser

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

Charakterograf diod półprzewodnikowych – przystawka do oscyloskopu

W treści artykułu opisano dwa proste układy elektroniczne umożliwiające wykreślanie na ekranie oscyloskopu charakterystyk napięciowo-prądowych diod półprzewodnikowych. Główną zaletą tych układów jest ich prosta konstrukcja i niski koszt wykonania. Samodzielny montaż tych urządzeń może stanowić ciekawą alternatywę dla zakupu fabrycznych charakterografów dostępnych na popularnych chińskich portalach internetowych w wielokrotnie wyższych cenach. Urządzenia te można szczególnie polecić pasjonatom elektroniki i elektroakustyki.

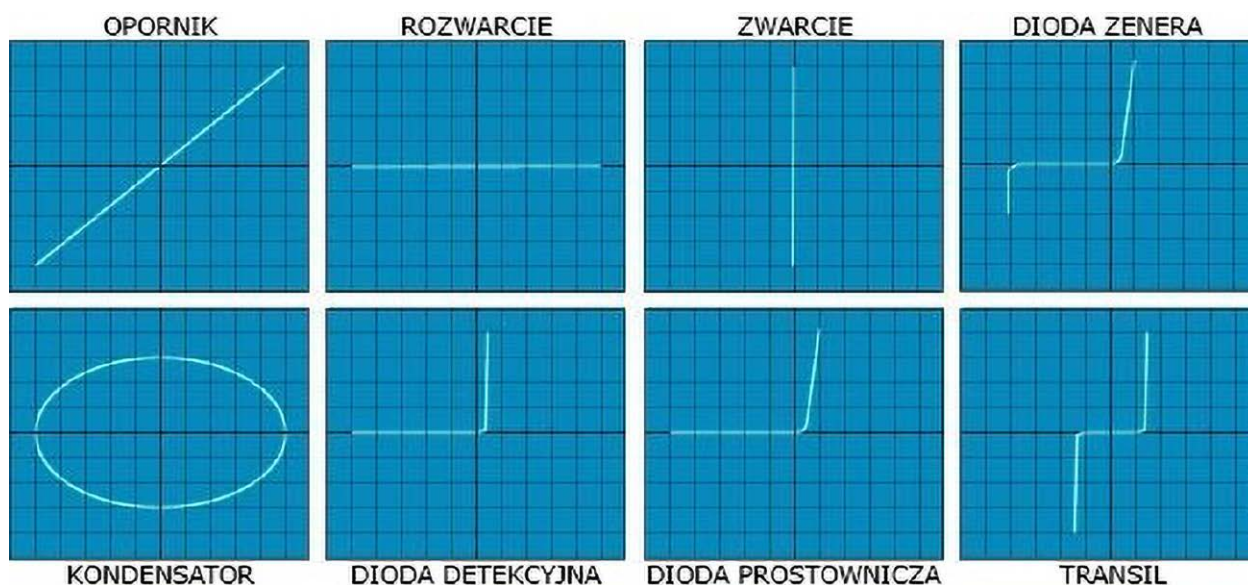
Wprowadzenie

Charakterografy diod niskonapięciowych i wysokonapięciowych są urządzeniami służącymi do wykreślania na ekranie oscyloskopu charakterystyk napięciowo-prądowych biernych i czynnych elementów elektronicznych celem szybkiego sprawdzenia ich parametrów w warunkach produkcyjnych. Mogą być one wykorzystywane w procesie kontroli jakości elementów elektronicznych na etapie przyjęcia towaru do magazynu lub wydania towaru z magazynu na linię produkcyjną. W niektórych przypadkach możliwe jest także sprawdzanie elementów elektronicznych po zamontowaniu ich na obwodach drukowanych, o ile w układzie nie są do nich dołączone równolegle inne elementy elektroniczne, które mogą mieć wpływ na wyniki pomiarów.

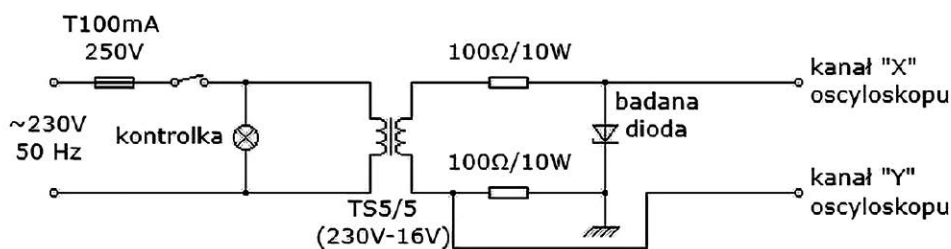
Opis techniczny urządzeń

Układy elektroniczne obydwu charakterografów są odseparowane galwanicznie od sieci zasilającej (w przypadku charakterografu wysokonapięciowego występuje nawet podwójna separacja galwaniczna) oraz zabezpieczone przed uszkodzeniem bezpiecznikami topikowymi zwłocznymi. Moc oporników wykorzystanych w układach została dobrana nadmiarowo aby zapobiec ich nadmiernemu nagrzewaniu podczas przeprowadzania długotrwałych pomiarów. Natężenie prądu elektrycznego płynącego przez badany element elektroniczny nie przekracza kilku (w przypadku charakterografu wysokonapięciowego) lub kilkudziesięciu (w przypadku charakterografu niskonapięciowego) miliamperów. Dokładna wartość natężenia prądu elektrycznego jest

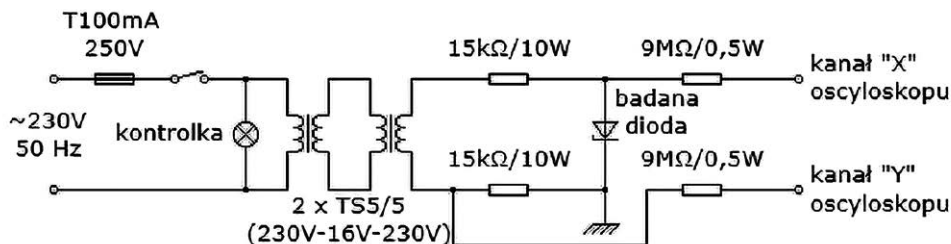
zależna od parametrów badanego elementu elektronicznego, tym niemniej rezystancje oporników występujące w układach elektronicznych obydwu charakterografów zostały dobrane w taki sposób aby maksymalna moc wydzielana na badanym elemencie elektronicznym nie przekraczała pół wata. W charakterografie wysokonapięciowym zastosowano dodatkowo dwa dziewięćmegaohmowe oporniki, które wraz z impedancjami wejściowymi oscyloskopu tworzą oporowe dzielniki napięcia tłumiące amplitudę sygnału pomiarowego dziesięciokrotnie. Obecność sygnału pomiarowego na gniazdach bananowych charakterografów sygnalizują neony, które są wbudowane we włączniki sieciowe. Rozstaw gniazd bananowych jest standardowy dla typowych urządzeń elektronicznych produkcji zachodniej, takich jak



Rysunek 1. Charakterystyki napięciowo-prądowe typowych elementów elektronicznych wyświetlane na ekranie oscyloskopu



Rysunek 2. Schemat ideowy charakterografu w wykonaniu niskonapięciowym



Rysunek 3. Schemat ideowy charakterografu w wykonaniu wysokonapięciowym

mierniki, zasilacze laboratoryjne czy generatory i wynosi dziewiętnaście milimetrów, dzięki czemu istnieje możliwość podłączania tam oprócz pętli pomiarowych także innych przejściówek zgodnych z tym standardem.

Obsługa urządzeń

1. Uruchomić oscyloskop (np. RIGOL DS1052E), wyłączyć podstawę czasu i ustawić oscyloskop w tryb pracy „XY”.
2. Czulość kanałów X i Y ustawić na 5 V/cm (w przypadku charakterografu niskonapięciowego) lub 100 V/cm (w przypadku charakterografu wysokonapięciowego).
3. Mnożnik kanałów X i Y ustawić na 1:1 (w przypadku charakterografu niskonapięciowego) lub 1:10 (w przypadku charakterografu wysokonapięciowego).
4. Polaryzację kanału X ustawić jako zgodną, natomiast polaryzację kanału Y ustawić jako przeciwną.

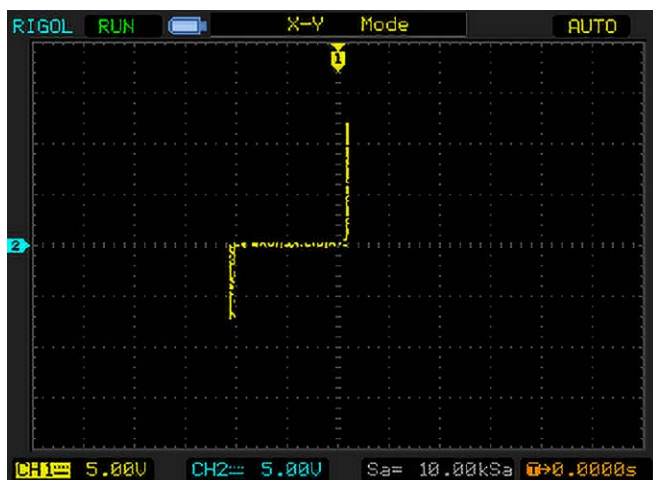
5. Podłączyć charakterograf do oscyloskopu przy pomocy ekranowanych przewodów wyposażonych w złącza BNC.
6. Podłączyć do gniazd bananowych charakterografu pęsetę pomiarową.
7. Podłączyć charakterograf do sieci zasilającej i uruchomić go przy pomocy włącznika sieciowego (po załączeniu zasilania powinna zaświecić się neonówka).
8. Przy pomocy pęsety pomiarowej dotknąć obydwie elektrody badanego elementu i obserwować przebieg charakterystyki na ekranie oscyloskopu.

Uwagi dotyczące BHP

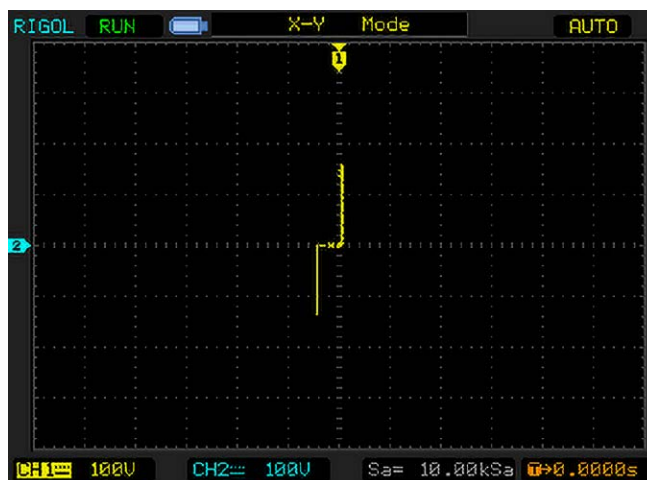
1. Charakterograf wysokonapięciowy może być używany wyłącznie przez pracowników upoważnionych do wykonywania pracy na danym stanowisku i posiadających świadectwo kwalifikacyjne SEP ze względu na to,

że na wyprowadzeniach pęsety pomiarowej występuje niebezpieczne napięcie o wartości przekraczającej 200 V RMS, co stwarza zagrożenie porażenia prądem elektrycznym.

2. Zabronione jest wykonywanie pomiarów diod LED przy pomocy obydwu charakterografów, ponieważ elementy te mogą ulec trwałemu uszkodzeniu na skutek przebicia (w przypadku charakterografu wysokonapięciowego) lub przeciążenia (w przypadku charakterografu niskonapięciowego).
3. Przed przystąpieniem do wykonywania pomiarów należy uważnie zapoznać się z kartą katalogową badanego elementu elektronicznego. Charakterograf wysokonapięciowy może współpracować wyłącznie z diodami, których dopuszczalne napięcie w kierunku zaporowym jest wyższe od 325 V. W przypadku



Rysunek 4. Charakterystyka napięciowo-prądowa dziesięciowoltowej diody Zenera zmierzona przy pomocy charakterografu niskonapięciowego



Rysunek 5. Charakterystyka napięciowo-prądowa czterdziestopięciowoltowej diody Zenera zmierzona przy pomocy charakterografu wysokonapięciowego



Rysunek 6. Pęseta pomiarowa firmy „AXIOMET” typu AX-TLP-TW-01

germanowych diod detekcyjnych wykorzystywanych w układach wielkiej częstotliwości i diod Schottky’ego wykorzystywanych w układach przełączających wskazane jest używanie wyłącznie charakterografu niskonapięciowego, ze względu na to, że diody te charakteryzują się wyjątkowo niskim dopuszczalnym napięciem w kierunku zaporowym.

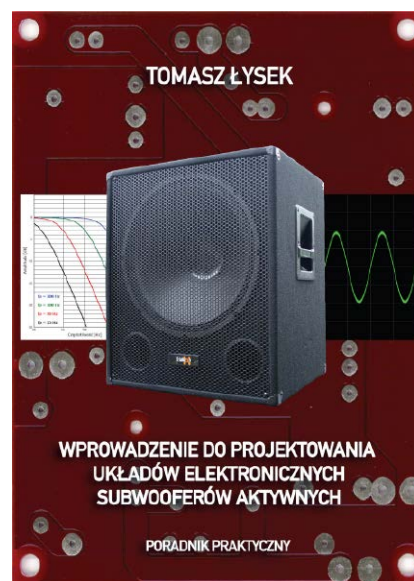
Źródła zakupu elementów

Poniżej zamieszczono wykaz najbardziej potrzebnych elementów wraz ze źródłami internetowymi. Pozostałe drobne elementy można zakupić w sklepach stacjonarnych:

- a. LT-207+503 H03VVH2-F 0.75/2C BLK 1M LIAN DUNG – kabel (https://www.tme.eu/pl/details/sn14-2_07_1bk/kable-zasil-komputerowe-i-universalne/lian-dung/lt-207-503-h03vvh2-f-0-75-2c-blk-1m/) – 2 sztuki)

- b. 42R511122 K+B – złącze zasilające AC (<https://www.tme.eu/pl/details/gz-01/zlaczka-iec-60320/k-b/42r511122/>) – 2 sztuki)
- c. PRW0AWJW101B00 ROYALOHM – Rezystor drutowy (<https://www.tme.eu/pl/details/ax10w-100r/rezystory-mocy/royalohm/prw0awjw101b00/>) – 2 sztuki)
- d. PRW0AWJP153B00 ROYALOHM – rezystor drutowy (<https://www.tme.eu/pl/details/ax10w-15k/rezystory-mocy/royalohm/prw0awjp153b00/>) – 2 sztuki)
- e. H3114 ZINC YIZN Jiangsu Tengyu Electronics co. – złącze BNC (<https://www.tme.eu/pl/details/bnc-006/zlaczka-bnc/yizn-jiangsu-tengyu-electronics-co/h3114-zinc/>) – 4 sztuki)
- f. D-3230 AXIOMET – złącze laboratoryjne bananowe 4 mm (<https://www.tme.eu/pl/details/d-3230/gniazda-bananowe-4mm/axiomet/>) – 2 sztuki)
- g. D-3231 AXIOMET – złącze laboratoryjne bananowe 4 mm (<https://www.tme.eu/pl/details/d-3231/gniazda-bananowe-4mm/axiomet/>) – 2 sztuki)
- h. R3-12E SCI – gniazdo (<https://www.tme.eu/pl/details/r3-12e/gniazda-bezpiecznikowe-na-panel/sci/>) – 2 sztuki)
- i. TS 5/5 INDEL – transformator sieciowy (https://www.tme.eu/pl/details/ts5_5/transformatory-z-mocowaniem/indel/ts-5-5/) – 3 sztuki)
- j. 50272 GOOBAY – kabel (<https://www.tme.eu/pl/details/cable-505-50-1/kable-polaczeniowe-koncentryczne/goobay/50272/>) – 4 sztuki)

REKLAMA



Rysunek 7. Okładka książki pt. „Wprowadzenie do projektowania układów elektronicznych subwooferów aktywnych. Poradnik praktyczny”

- k. AX-TLP-TW-01 AXIOMET – pęseta pomiarowa (<https://www.tme.eu/pl/details/ax-tlp-tw-01/przewody-pomiarowe-pojedyncze/axiomet/>) – 2 sztuki)
- l. CF1/2W-4M7 SR PASSIVES – rezystor węglowy (https://www.tme.eu/pl/details/cf1_2w-4m7/rezystory-tht/sr-passives/) – 4 sztuki)
- m. 522.707 ESKA – bezpiecznik topikowy (<https://www.tme.eu/pl/details/zct-0.1a/bezpieczniki-5x20mm-zwloczne/eska/522-707/>) – 2 sztuki)
- n. Z90 KRADEX – obudowa uniwersalna (<https://www.tme.eu/pl/details/s/z-90/obudowy-universalne/kradex/z90/>) – 2 sztuki)
- o. R13-112B-02-BR-2D-N-2 SCI – ROCKER (<https://www.tme.eu/pl/details/r13112b02br2n2/przelaczniki-typu-rocker/sci/r13-112b-02-br-2d-n-2/>) – 2 sztuki)

Książka o układach elektronicznych do subwooferów aktywnych

Zapraszam do zapoznania się z moją najnowszą książką pt. „Wprowadzenie do projektowania układów elektronicznych subwooferów aktywnych. Poradnik praktyczny”:

- <https://www.youtube.com/watch?v=j7HMMxHFCwg>
- <https://www.youtube.com/watch?v=KIo1eqxj4AE>
- <https://www.youtube.com/watch?v=gpQe89R5HEK> ■

mgr inż. Tomasz Łysek



Migające diody LED i śliniący się inżynierowie (31)

Jak to mi się często zdarza, nie mogę się w tej chwili opanować z podekscytowania. Chyba wspomniałem w zeszłym roku (2021; przypis redaktora), że zostałem zaproszony do wygłoszenia przemówienia na FPGA Forum 2022 w Norwegii (www.fpga-forum.no). W tym prestiżowym wydarzeniu biorą udział wszyscy liczący się w norweskiej branży układów programowalnych, między innymi projektanci, kierownicy projektów, kierownicy techniczni, naukowcy, studenci ostatniego roku oraz dystrybutorzy układów FPGA.

Norwegia, ho!

Tegoroczne ([w roku 2022; przypis redaktora](#)) Forum miało pierwotnie odbyć się w lutym, ale... covid... więc przełożono je na wrzesień. Jestem szczególnie wzruszony, ponieważ już kiedyś wygłaszałem przemówienie inauguracyjne podczas tego wyjątkowego wydarzenia – w odległej przeszłości, którą nazywaliśmy rokiem 2012. Z jednej strony to zaledwie 10 lat temu, co byłoby tylko drobnym skokiem w mojej machinie czasu, gdybym tylko... mógł uruchomić to ustrojstwo (tam, gdzie mieszkam, w Huntsville w stanie Alabama, zwyczajnie nie można dostać części). Z drugiej strony technika pędzi naprzód, jak to ma w zwyczaju, i wiele się zmieniło. Na przykład niektóre firmy zajmujące się układami FPGA, które były z nami w 2012 roku (np. TierLogic i Tabula), poddały się; poszły na dno i opuściły ziemski padół (zawsze lubiłem takie metafory). Z drugiej strony, do imprezy postanowiło dołączyć kilku energicznych nowicjuszy (np. Efinix i Renesas). No i dochodzi jeszcze fakt, że obecnie tkwimy po uszy w sztucznej inteligencji (AI), uczeniu maszynowym (ML) i uczeniu głębokim (DL). Ponadto jesteśmy zasypywani nowymi odmianami rzeczywistości (sam mam słabość do truskawkowej), w tym: rzeczywistością rozszerzoną (AR), rzeczywistością zredukowaną (DR), rzeczywistością wirtualną (VR) i wirtualnością rozszerzoną (AV) – by wymienić tylko kilka. Uff! Wystarczy powiedzieć, że nie wątpię, iż uda mi się jakoś zebrać myśli i znaleźć temat do rozmowy (tak jak w przypadku mojej drogiej, starej matki – prawdziwą sztuką jest sprawienie, byśmy przestali mówić).

Na Forum, oprócz delegacji z Norwegii, przybywają często uczestnicy z sąsiednich krajów, w tym ze Szwecji i z Finlandii. Mam również nadzieję spotkać przedstawicieli firmy Testonica (www.testonica.com), której siedziba znajduje się w Estonii – jednym z krajów bałtyckich. Ci ludzie

dysponują interesującą technologią o nazwie Quick Instruments. Jeśli projektujecie płytkę drukowaną z układem FPGA, to wystarczy, że podacie ogólny zarys systemu – w tym typy innych urządzeń na płytce oraz ich mapy pinów i rejestrów – a Quick Instruments automatycznie wygeneruje i skompiluje odpowiedni zestaw testów oprogramowania układowego, umożliwiając w ten sposób układowi FPGA przeprowadzenie zaawansowanego autotestu na poziomie płytki.

W przededniu nowej ery

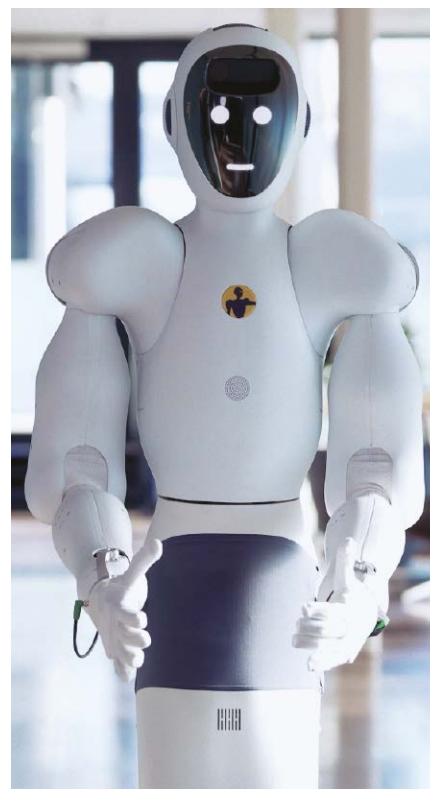
Forum FPGA tradycyjnie odbywa się w Trondheim – trzeciej pod względem liczby mieszkańców gminie w Norwegii. Trondheim jest między innymi siedzibą Norweskiego Uniwersytetu Nauki i Techniki (NTNU). Wspominam o tym z dwóch powodów. Po pierwsze, dzień przed moim wystąpieniem wygłoszę gościnny wykład dla grupy studentów NTNU studiów magisterskich z zakresu elektrotechniki i elektroniki (www.ntnu.edu). Po drugie, w Norwegii ma siedzibę firma Halodi Robotics (www.halodi.com), którą przedstawiłem w poprzednim artykule (*Practical Electronics*, sierpień 2022 r.; EdW 3/2026), a jej humanoidalne roboty EVE (**rysunek 1**) można spotkać na większości norweskich uniwersytetów. Mam nadzieję, że EVE zechce wziąć udział w moim wykładzie. Jeśli tak się stanie, możecie być pewni, że porobię zdjęcia i opowiem o tym w kolejnym artykule.

Robot EVE ma rozmiar zbliżony do ludzkiego i dysponuje 23 stopniami swobody. Dzięki dwuramiennej konstrukcji EVE może każdym ramieniem unieść ładunek o masie 8 kg. EVE potrafi również przykucnąć i podnieść przedmiot z podłogi lub wyjąć z szafki, a także sięgnąć w górę i umieścić przedmiot na półce.

Najbardziej zadziwiające jest dla mnie to, że roboty EVE są już wdrażane do produkcji

i działają na całym świecie w różnych zastosowaniach – od handlu detalicznego po ochronę. Na przykład w handlu detalicznym roboty EVE mogą wykonywać takie zadania, jak poruszanie się po alejkach supermarketu, wykrywanie produktów, które znalazły się w niewłaściwym miejscu, i odkładanie ich na odpowiednie półki, uzupełnianie zapasów, kompletowanie zamówień dla klientów zamawiających online z odbiorem osobistym, identyfikowanie i zgłaszanie towaru rozsypanego czy rozlanego. Lista możliwości jest znacznie dłuższa.

Podobnie jest w przypadku zastosowań związanych z bezpieczeństwem. Roboty EVE



Rysunek 1. Robot humanoidalny EVE. Zdjęcie: Halodi Robotics

mogą samodzielnie poruszać się po budynkach w poszukiwaniu intruzów o nieznanych twarzach, sprawdzać, czy drzwi są zamknięte, a światła wyłączone... i tak dalej. Na początku tego roku amerykańska firma ochroniarska ADT podpisała największe jak dotąd na świecie zamówienie na roboty humanoidalne, zamawiając od firmy Halodi 140 egzemplarzy EVE.

Wszystko zaczyna się od jednego serwomechanizmu

Robot EVE jest niesamowity. Ale wszystkie tego typu projekty zaczynają się od pojedynczego siłownika, takiego jak serwomechanizm – i właśnie o tym obecnie porozmawiamy. Jak wspomnieliśmy w poprzednim wydaniu, zaczęliśmy od hobbystycznego mikroserwomechanizmu o masie 9 g („9 g” oznacza przybliżoną masę serwomechanizmu, a nie wartość momentu obrotowego, jaki jest w stanie wytworzyć).

Pierwotnie planowałem rozebrać jedno z tych cudeniek, aby pokazać Wam, co się w nim kryje, więc zamówiłem na Amazonie zestaw pięciu mikroserwomechanizmów SG90 (dlaczego zmieniłem plany, wyjaśnię za chwilę). Istnieje mnóstwo dostawców tego typu produktów (<https://amzn.to/3O32sC4>), ale „dostajesz to, za co płacisz”, więc uważajcie, co zamawiacie.

Serwomechanizmy, na które się zdecydowałem, tanie i efektywne, umożliwiają obrót o 180°. Jak już wspominałem wcześniej, serwomechanizmy te działają w ten sposób, że nasz sterownik – w tym przypadku Arduino Uno – wysyła do serwomechanizmu cykliczne impulsy o dodatniej polaryzacji. Okres powtarzania impulsów nie jest szczególnie istotny – liczy się ich szerokość – ale powszechnie stosuje się okres „odświeżania” 20 ms, co oznacza, że wysyłamy impulsy do serwomechanizmu 50 razy na sekundę (50 Hz).

Jeśli chodzi o szerokość impulsów – wartość 1,5 ms powoduje, że serwomechanizm ustawia się w pozycji domyślnej (środkowej). W przypadku zamówionych przeze mnie serwomechanizmów dołączona specyfikacja techniczna podaje, że impuls o długości 1,0 ms odpowiada obrotowi o -90° (maksymalnie w lewo), natomiast impuls 2,0 ms – obrotowi o +90° (maksymalnie w prawo).

Impuls 1,5 ms zazwyczaj ustawia tego typu serwomechanizmy w pobliżu pozycji środkowej, natomiast dla poszczególnych egzemplarzy mogą obowiązywać różne szerokości impulsu dla położenia skrajnych. Na przykład zgodnie z informacjami podanymi przez Motors for Makers (<https://amzn.to/3IovsD0>), niektóre serwomechanizmy akceptują impulsy o długości od 0,7 ms (pełny obrót w lewo) do 2,3 ms (pełny obrót w prawo) – przez „pełny” rozumiemy maksymalny obrót obsługiwany przez dany serwowymotor.

Najważniejsze jest, aby nasze impulsy nie były ani zbyt wąskie, ani zbyt szerokie, tak aby nie przekraczały zakresu serwomechanizmu, ponieważ może to spowodować zerwanie zębatek, co bynajmniej nie jest pożądane.

Warto też pamiętać, że serwomechanizmy te są przeznaczone do zastosowań hobbystycznych, więc ich parametry techniczne i tolerancje mogą być znacznie mniej precyzyjne w porównaniu z urządzeniami klasy przemysłowej. Chodzi o to, że nawet w partii rzekomo „identycznych” serwomechanizmów mogą występować znaczne różnice, dlatego przed użyciem warto sprawdzić właściwości każdego z nich. A skoro o tym mowa...

Przekręć gałkę

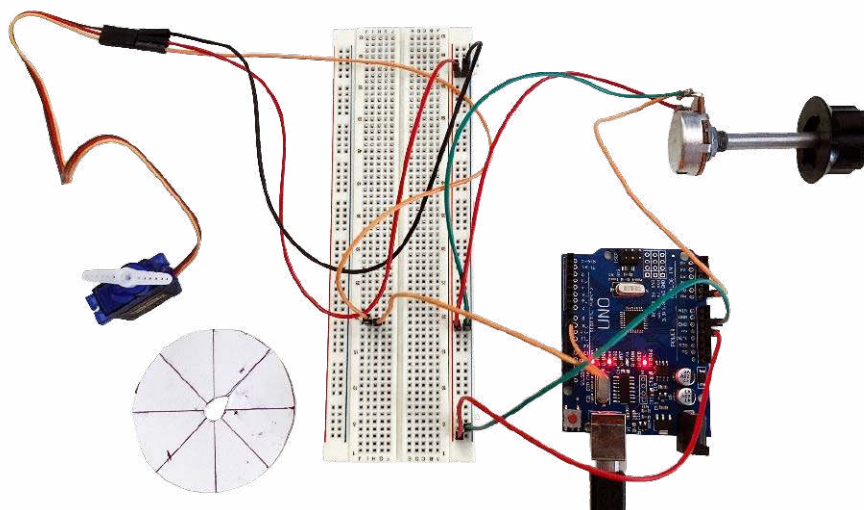
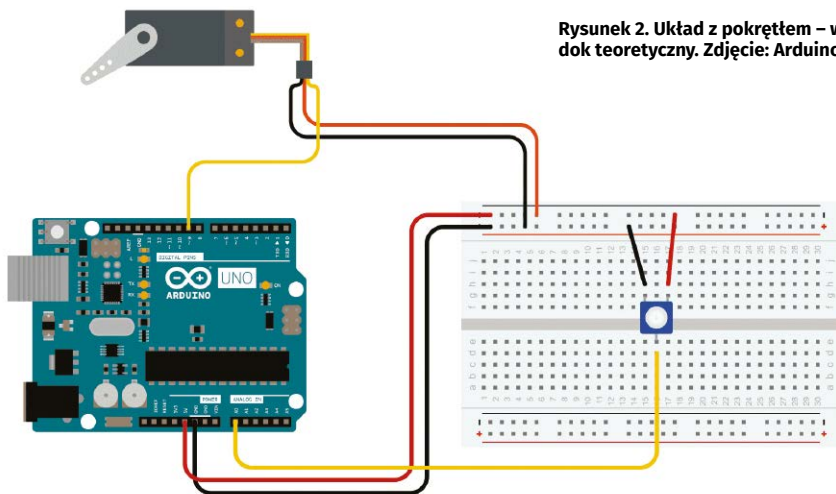
Starsi Czytelnicy mogą pamiętać brytyjski serial komediowy „The Goon Show”, w którym

występowali Spike Milligan, Harry Secombe, Peter Sellers i Michael Bentine. Ten zwarowany program, lubiany przez księcia Karola, nadawany był pierwotnie w latach 1951–1960, a potem powtarzano go od czasu do czasu. Pamiętam odcinek, w którym jedna z postaci powiedziała: „Przekręć gałkę ze swojej strony”. Inna postać odpowiedziała: „Nie ma żadnej gałki po tej stronie”. A pierwsza odparła: „W drzwiach, idioto!” Trzeba było tam być i słyszeć ich głosy. Wciąż się w duchu śmieję.

W poprzednim artykule wspomnieliśmy, że Arduino jest wyposażone w bibliotekę serwomechanizmów (<https://bit.ly/3O3rIZu>). Stworzyliśmy również układ i program „Zamiatanie”, które powodowały, że nasz serwomechanizm poruszał się tam i z powrotem. Teraz zamierzamy zrealizować to, co twórcy Arduino nazywają układem „Gałka”.

Cieszy fakt, jak nieskazitelnie wszystko wygląda w teorii (rysunek 2). Godne uwagi jest również to, jak bardzo „poszarpane” wydają się te elementy w rzeczywistości

Rysunek 2. Układ z pokrętkiem – widok teoretyczny. Zdjęcie: Arduino



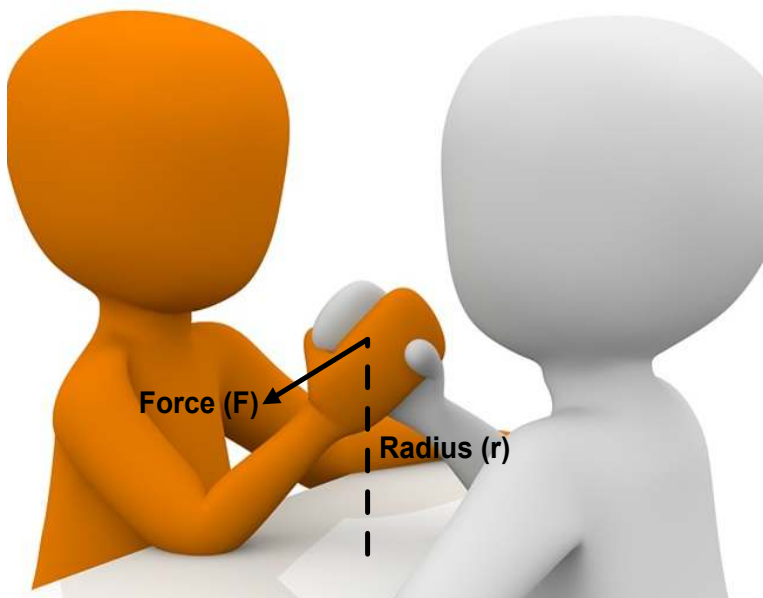
Rysunek 3. Układ z pokrętkiem – widok rzeczywisty. Po lewej stronie niebieski serwowymotor

```

1 #include <Servo.h>
2
3 Servo MyServo;
4
5 int PotPin = A0;
6 int PotVal;
7
8 void setup()
9 {
10  MyServo.attach(9, 1000, 2000);
11 }
12
13 void loop()
14 {
15  PotVal = analogRead(PotPin);
16  PotVal = map(PotVal, 0, 1023, 0, 180);
17  MyServo.write(PotVal);
18  delay(15);
19 }

```

Rysunek 4. Przykładowy program do obsługi układu z pokrętkiem



Rysunek 5. Wizualizacja zasady momentu obrotowego

(rysunek 3). W tym przypadku na płytce prototypowej używam nie potencjometru montażowego pokazanego w układzie teoretycznym, lecz dość masywnego potencjometru osiowego, widocznego w prawym górnym rogu zdjęcia.

Przyjrzyjmy się teraz programowi, którego użyjemy do zdejmowania charakterystyki naszego serwomechanizmu (rysunek 4). Pierwszym krokiem w wierszu 1. jest dołączenie biblioteki serwomechanizmu. W wierszu 3. tworzymy obiekt serwomechanizmu. Tutaj nazwaliśmy go MyServo, ale odpowiednia będzie każda inna poprawna nazwa (niebędąca słowem kluczowym). W wierszu 5. deklarujemy zmienną całkowitą (int) o nazwie PotPin, którą przypisujemy do analogowego wejścia A0 w Arduino, a następnie w wierszu 6. deklarujemy kolejną zmienną typu całkowitego o nazwie PotVal, której użyjemy do przechowywania wartości odczytanych z potencjometru.

W wierszu 10. używamy metody `attach()` do przypisania zmiennej serwomechanizmu do pinu Arduino. W tym przypadku podajemy trzy argumenty: 9, 1000 i 2000. Pierwszy argument jest obowiązkowy i określa pin, którego chcemy użyć do sterowania serwomechanizmem. Musi to być jeden z pinów obsługujących modulację szerokości impulsu (PWM).

Przypis redaktora: Biblioteka Servo w Arduino nie korzysta ze sprzętowego PWM. Sygnał sterujący jest generowany programowo z wykorzystaniem timerów mikrokontrolera, dlatego serwomechanizmy mogą być podłączane do różnych pinów cyfrowych, nie tylko do pinów oznaczonych jako PWM. Źródło:

[Arduino Servo Library Documentation \(docs.arduino.cc/libraries/servo/\)](https://www.arduino.cc/libraries/servo/)

W tym przykładzie używamy pinu 9. Drugi i trzeci argument, podane w mikrosekundach, są opcjonalne. Argument drugi określa szerokość impulsu odpowiadającą minimalnemu (0°) kątowi serwomechanizmu, natomiast argument trzeci zadaje szerokość impulsu odpowiadającą kątowi maksymalnemu (180°). Używamy wartości odpowiednio 1000 μ s i 2000 μ s (1 ms i 2 ms). Co ciekawe, domyślne wartości tych parametrów, obowiązujące jeśli sami ich nie zadamy, to odpowiednio 544 i 2400, co znacznie wykracza poza zakres, który uznaliśmy za prawidłowy.

W pętli głównej, w wierszu 15., zaczynamy od odczytania wartości z potencjometru i zapisania jej w zmiennej PotVal. Ponieważ Arduino Uno ma 10-bitowy przetwornik analogowo-cyfrowy (ADC), otrzymujemy wartości z przedziału 0...1023. Metoda `write()`, udostępniona przez bibliotekę serwomechanizmu Arduino, wymaga wartości z zakresu 0...180, odpowiadających odpowiednio kątom 0° ... 180° . Dlatego w wierszu 16. używamy funkcji `map()` Arduino, by przekonwertować wartości 0...1023 z potencjometru na ich odpowiedniki w zakresie 0...180. W wierszu 17. wpisujemy tę nową wartość do serwomechanizmu, a w wierszu 18. dajemy opóźnienie 15 ms, aby serwomechanizm miał czas na reakcję.

Zagadkowa zagadka

Na rysunku 3 w lewym dolnym rogu widać małe kółko z papieru, znakowane co 45° . To tylko coś, co szybko skleciłem na potrzeby tej dyskusji. Gdybym podchodził

do tego poważnie, stworzyłbym dokładniejsze narzędzie znakowane co 5° .

Kiedy umieściłem tę kartkę pod obracającym się ramieniem serwomechanizmu i uruchomiłem program przedstawiony wyżej, obracanie potencjometrem od jednej skrajnej pozycji do drugiej powodowało ruch ramienia serwomechanizmu jedynie o 90° . Zmieniłem wartości argumentów `min` i `max` w metodzie `attach()` na odpowiednio 700 i 2300. Zapewniło to ruch ramienia serwomechanizmu o około 135° . Na koniec wypróbowałem domyślne wartości biblioteki serwomechanizmów, czyli 544 i 2400, co dało ruch o kąt zbliżony do 180° , aczkolwiek w skrajnych pozycjach słyszałem już delikatne „zgrzytanie”.

To wszystko całkowicie mnie zaskoczyło i zdezorientowało. Niewielkie odchylenia jestem w stanie zaakceptować, ale dla czego w karcie katalogowej podano wartości minimalną i maksymalną wynoszące 1000 i 2000, skoro użycie wartości 544 i 2400 pozwala uzyskać znacznie większy zakres ruchu? Jeśli macie jakieś przemyślenia na ten temat, chętnie je poznam.

Przypis redaktora: Wartości 544...2400 μ s są domyślnym zakresem parametrów funkcji `attach()` w bibliotece Arduino Servo, a nie wymaganiem samego serwomechanizmu. Typowy zakres katalogowy dla serw modelarskich wynosi około 1,0 ms...2,0 ms (często rozszerzany do około 0,7 ms...2,3 ms). W tekście występuje ponadto niekonsekwencja (544 i 2500 zamiast 544 i 2400). Źródła: Arduino Servo Library Reference ([arduino.cc](https://www.arduino.cc/)), dokumentacje serw modelarskich (np. TowerPro SG90).

Nie jest to takie proste

Obawiam się, że właśnie w tym miejscu artykułu możemy spodziewać się zgrzytania zębami i rozdzierania szat, ponieważ znany wydawca *Practical Electronics*, Matt Pulzer, wciąż powtarza mi, że bym „trzymał się elektroniki”, a ja powtarzam sobie: „Ale te inne rzeczy są takie interesujące!”.

Na przykład... Termin „maszyna prosta” odnosi się do urządzenia mechanicznego, które służy do zmiany kierunku lub wielkości siły. Sześć klasycznych maszyn prostych, zgodnie z definicją renesansowych myślicieli, to: dźwignia, blok, koło pasowe, równia pochyła, klin oraz śruba.

Mógłbym bez końca rozprawiać o tych cudeńkach, ale mam inne sprawy na głowie. Szczególnie interesujący jest dla nas fakt, że za proste maszyny można uznać przekładnię zębatą, składającą się z dwóch lub więcej kół wyposażonych w zachodzące na siebie zęby, tak że gdy jedno koło zębate obraca się, pozostałe koła obracają się w przeciwnym kierunku.

Może warto zauważyć, że niektórzy – zwłaszcza w Wielkiej Brytanii – nieformalnie mówią na koła zębate „cogs”, choć pierwotnie termin ten odnosił się do zęba koła zębatego.

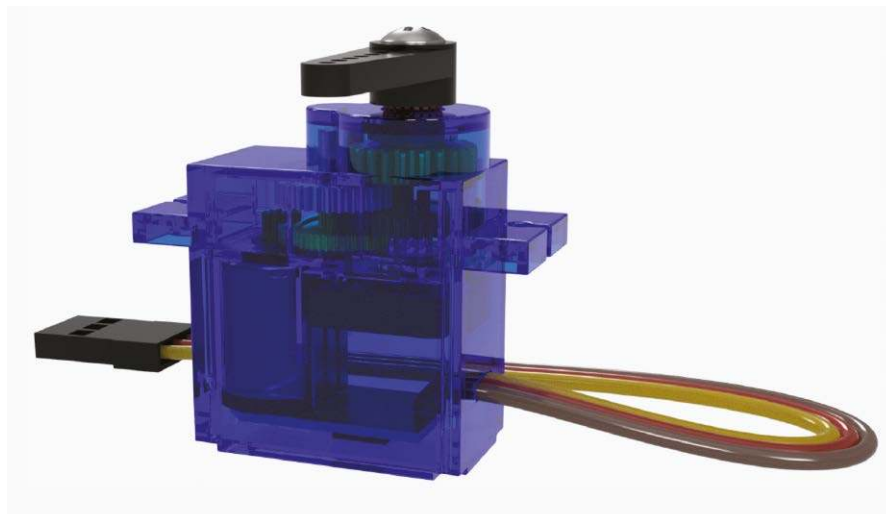
Układ kół zębatych może się na pierwszy rzut oka wydawać rzeczą prostą, ale w mechanizmie przekładni kryje się dużo więcej. Moglibyśmy spędzić całe tygodnie, zagłębiając się w te zawiłości, ale nie zrobimy tego, bo czuję, że wszechwidzące oko redaktora Matta zaczyna spoglądać w moją stronę.

Możesz sobie poprzekładać!

Obawiam się, że w tym momencie będziemy musieli wprowadzić kilka półtechnicznych zagadnień. Zaczniemy od **momentu obrotowego**, oznaczonego symbolem T , który jest obrotowym odpowiednikiem siły liniowej i który możemy traktować jako „siłę skręcającą”. Najłatwiej wyobrazić to sobie jako zawody w siłowaniu się na rękę (**rysunek 5**).

Jeśli założymy, że obaj zawodnicy mają równe szanse, to siły, które wywierają, znoszą się nawzajem – ale nie o to nam chodzi. Jeśli weźmiemy pod uwagę osobę znajdującą się bliżej, to moment obrotowy jest równy iloczynowi promienia i siły: $T=r \cdot F$.

Można na to spojrzeć z dwóch stron. Po pierwsze, założmy, że próbujecie otworzyć ciężkie drzwi na zawiasach, naciskając na nie. W tym przypadku maksymalna siła, jaką możecie wywrzeć (czyli jak mocno możecie naciskać), jest stała. Promień to odległość od zawiasów do punktu, w którym naciskacie. Jeśli jest to blisko zawiasów, to promień jest



Rysunek 6. Klasyczny serwomechanizm z ramieniem przymocowanym do wyjścia przekładni zębatej

mały, więc generowany moment obrotowy będzie niewielki, co oznacza, że otwarcie drzwi będzie trudne. Dla porównania – jeśli naciskacie na drzwi daleko od zawiasów, to promień jest większy, co oznacza, że – przy tej samej sile co poprzednio – możecie wytworzyć większy moment obrotowy, dzięki czemu (a) otwarcie drzwi będzie dziecinnie proste i (b)... przez drzwi będzie przepływał lekki prąd powietrza.

Drugim sposobem spojrzenia na tę kwestię może być perspektywa serwomechanizmu wyposażonego na wyjściu w ramię (**rysunek 6**). W tym przypadku moment obrotowy, jaki może wytworzyć silnik, jest stały, więc siła, jaką może wyrzucić serwomechanizm, wynika ze wzoru $F=T/r$. Wynika z tego, że im większy jest promień ramienia (w punkcie, w którym podłączamy do niego układ przegubów), tym mniejszą siłą może ono wyrzucić.

Kolejnym pojęciem jest **prędkość kątowna**, oznaczana symbolem ω , która opisuje kąt, o jaki obiekt obraca się w danym czasie. W przypadku silników elektrycznych kąty mierzy się w stopniach ($^{\circ}$), a prędkość kątową – w obrotach na minutę (obr./min. lub RPM).

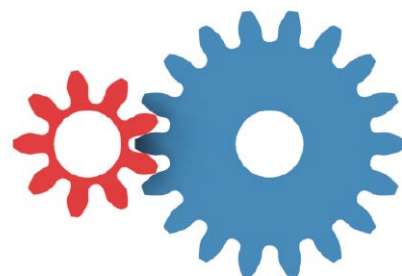
Przypis redaktora: W układzie SI jednostką prędkości kątowej jest rad/s, jednak w praktyce technicznej (zwłaszcza w dokumentacji silników i serwomechanizmów) powszechnie stosuje się prędkość obrotową, wyrażaną w obr./min. (RPM). Zgodnie z normą ISO 80000-3, są to dwie różne wielkości fizyczne, choć w inżynierii często utożsamiane lub przeliczane.

Gdy dwa koła zębate są zazębiane, koło z mniejszą liczbą zębów musi wykonać więcej obrotów niż koło z większą liczbą zębów. Rozważmy dwa koła zębate, z których jedno ma 9 zębów, a drugie 18 (**rysunek 7**).

W tym przypadku mniejsze koło wykona dwa obroty na każdy jeden obrót koła większego.

To, co dzieje się dalej, zależy od tego, które z kół zębatych jest kołem napędzającym (wejściowym), a które kołem napędzanym (wyjściowym). Jeśli kołem wejściowym jest koło większe, to – w przypadku naszego przykładu z kołami o 9 i 18 zębami – koło mniejsze (wyjściowe) będzie obracać się z dwukrotnie większą prędkością, ale z tylko połową momentu obrotowego (siły skręcającej). Dla porównania: jeśli kołem wejściowym jest koło mniejsze – co jest przypadkiem częstszym (i tak właśnie działa nasz serwomechanizm) – wówczas koło większe (wyjściowe) będzie obracać się z prędkością o połowę mniejszą, ale z dwukrotnie większym momentem obrotowym.

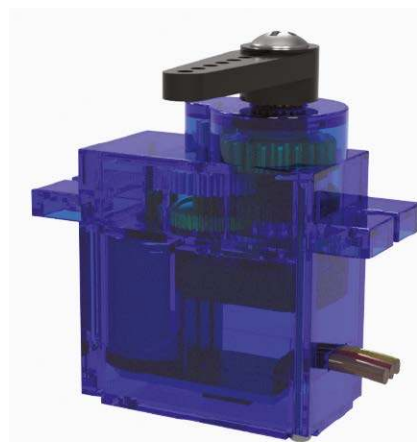
Termin „przełożenie” odnosi się do względnego momentu obrotowego między dwoma kołami zębatymi i zazwyczaj wyraża się go w postaci $X:1$, gdzie X oznacza proporcjonalny wzrost momentu obrotowego. W naszym przykładzie przełożenie wynosi 2:1. Gdyby większe koło zębate miało 27 zębów, przełożenie wynosiłoby 3:1 i tak dalej. Gdy połączone są dwa lub więcej kół zębatych, całość nazywamy „układem przekładniowym”.



Rysunek 7. Dwa koła zębate o 9 i 18 zębami. Rysunek: Steve Manley

Jeśli na przykład układ przekładniowy 2:1 jest połączony z układem 3:1, który z kolei jest połączony z układem 4:1, wówczas wynikowe przełożenie całego układu przekładniowego wyniesie $(2 \cdot 3 \cdot 4):1 = 24:1$.

Przekładnie są zazwyczaj dość wydajne, choć rzeczywiste straty mocy zależą od ich konstrukcji, materiałów i jakości wykonania. Aby uprościć sprawę, na potrzeby niniejszej dyskusji przyjmijmy, że sprawność przekładni wynosi 100%.



Rysunek 8. Serwomechanizm w obudowie (u góry) i bez obudowy (u dołu). Rysunek: Steve Manley

Aby dopełnić tę część naszych rozważań: w kontekście przekładni moc (P) jest równa iloczynowi momentu obrotowego i prędkości obrotowej: $P = T \cdot \omega$. W przypadku naszej przekładni z 9-zębowym wejściem i 18-zębowym wyjściem, $P_i = T_i \cdot \omega_i$ oraz $P_o = T_o \cdot \omega_o$ (odpowiednio: **moc wejściowa i wyjściowa; przypis redaktora**). Zakładając 100% sprawność, $P_o = P_i$, co oznacza, że $T_o \cdot \omega_o = T_i \cdot \omega_i$. To z kolei oznacza, że $T_o/T_i = \omega_i/\omega_o$.

Ponadto jeśli liczbę zębów na kole zębatym oznaczymy przez N , to zależność między prędkością kątową a ilością zębów można zapisać jako $\omega_i/\omega_o = N_o/N_i$. Podsumowując to wszystko możemy stwierdzić, że nasze przełożenie przekładni wynosi $2:1 = T_o/T_i = \omega_i/\omega_o = N_o/N_i$. Zwróćcie uwagę na podobieństwo do stosunków napięć, prądów i liczby zwojów transformatora elektrycznego.

Nacieszcie oczy

Jak już wspominałem, planowałem pierwotnie rozebrać jeden z moich tanich serwomechanizmów, zrobić zdjęcia i pokazać Wam, co znajduje się w środku. Mój przyjaciel Steve Manley stwierdził jednak, że może zrobić to znacznie lepiej, korzystając ze swojego oprogramowania do projektowania wspomagane komputerowo (CAD) Fusion 360.

Steve zaczął od zakupu kilku mikroserw AZ-Delivery MG90S na Amazonie (<https://amzn.to/3ywy4Kk>). Następnie rozmontował jedno z nich, zmierzył wszystkie elementy mikrometrem i przez mikroskop skrupulatnie policzył zęby na każdym z kół zębatych tworzących układ przekładniowy. Efektem była seria zdjęć, które wywołały u mnie łączy radości (**rysunek 8**).

Steve stwierdził, że wmawiano mu, iż wszystkie koła zębate w tym serwomechanizmie są metalowe. Okazało się jednak, że metalowe było tylko koło wyjściowe, natomiast pozostałe były z plastiku. Z komentarzy na Amazonie wynika, że z tego oszustwa niezadowolonych było również parę innych osób. Jak zawsze należy pamiętać o zasadzie caveat emptor – kupujący ponosi ryzyko.

Termin „przekładnia złożona” odnosi się do dwóch lub więcej kół zębatych, które są ze sobą połączone i w związku z tym obracają się z tą samą prędkością. Na rysunku 8 przekładniami złożonymi są koła zielone, niebieskie i czerwone, ponieważ każde z nich składa się z dwóch kół zębatych.

W przypadku tego serwomechanizmu mamy trzy wałki: wałek silnika, wałek pośredni i wałek wyjściowy. Należy zwrócić uwagę, że zielone i niebieskie koła

zębate złożone nie są fizycznie połączone z wałkiem pośrednim, ale po prostu obracają się na nim. Podobnie czerwone koło zębate nie jest fizycznie połączone z wałkiem wyjściowym. Jedynymi kołami zębatymi, które są fizycznie połączone z jakimiś wałkami, są: mosiężne koło zębate wejściowe (połączone z wałkiem silnika) oraz złote koło zębate wyjściowe (połączone z wałkiem wyjściowym).

Zwróćcie uwagę na czarne ramię umieszczone w górnej części serwomechanizmu. Element ten nazywany jest również „rogiem”. Rogi występują w różnych kształtach i rozmiarach. Zwróćcie też uwagę, że średnice wałków rosną od strony wejściowej do wyjściowej (wałek silnika = 1 mm, wałek pośredni = 1,16 mm, wałek wyjściowy = 1,36 mm). Podobnie rozmiar zębów i grubość kół zębatych zwiększają się w miarę przechodzenia przez układ przekładniowy. Wszystkie te zmiany są konieczne, aby dostosować się do rosnącego momentu obrotowego od wejścia do wyjścia.

Zielony element umieszczony poniżej czerwonego koła zębatego zawiera potencjometr przymocowany do wałka wyjściowego. Potencjometr ten służy do pomiaru aktualnego położenia kątownego wałka. W serwomechanizmach są stosowane również inne sposoby.

W dolnej części serwomechanizmu, pod silnikiem, jest zamontowana mała płytka drukowana zawierająca elektronikę sterującą. Oprócz stabilizatora napięcia i paru elementów dyskretnych płytka ta zawiera układ sterujący KC9702A, który interpretuje sygnał PWM pochodzący z zewnątrz, wykorzystuje wartość z potencjometru do określenia różnicy między pożądaną a rzeczywistą pozycją wałka wyjściowego i odpowiednio steruje silnikiem prądu stałego.

Co? Chyba żartujesz!

Zestawienie parametrów kół zębatych i przełożeń tworzących układ przekładniowy tego serwomechanizmu przedstawiono na **rysunku 9**. Jeśli pomnożymy wszystkie przełożenia $(5,2222 \cdot 3,8000 \cdot 4,0000 \cdot 3,2857)$, otrzymamy całkowite przełożenie układu napędowego od wejścia do wyjścia wynoszące około 260:1.

Na początku może to nie wywołać u Was dreszczyku emocji. Jednak po chwili namysłu zgodzicie się, że wyniki są bardziej imponujące.

Zacznijmy od tego, że moment obrotowy samego silnika jest znikomy. Steve zauważył na przykład, że może z łatwością zablokować silnik, ściskając jego obracający się wałek kciukiem i palcem wskazującym.

	Shaft	Gear	# Teeth	Ratio
	Input	Input	9	
	Intermediate	1A	47	5.2222
		1B	10	
	Output	2A	38	3.8000
		2B	8	
	Intermediate	3A	32	4.0000
		3B	7	
	Output	Output	23	3.2857

Rysunek 9. Elementy układu przekładni serwomechanizmu

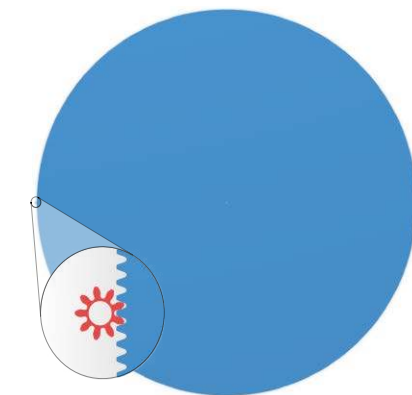
Steve pracuje obecnie nad filmikiem animowanym przedstawiającym działanie tego serwomechanizmu. Nie wiem, skąd wziął te dane, ale w swoim skrypcie zakłada, że prędkość obrotowa silnika wynosi 26 000 obr./min., a moment obrotowy – 0,0072 kG·cm. Ponadto prędkość obrotowa (ω) ramienia, zgodnie z danymi w karcie katalogowej, wynosi 60°/0,1 s. Zatem $\omega = 60^\circ/0,1 \text{ s} = 360^\circ/0,6 \text{ s} = 1 \text{ obr.}/0,6 \text{ s} = 100 \text{ obr.}/60 \text{ s} = 100 \text{ obr.}/\text{min.}$

Zakładając, że przełożenie przekładni wynosi 260:1, jak ustaliliśmy powyżej, wówczas redukcja prędkości od wejścia do wyjścia wyniesie $26\,000 \text{ obr.}/\text{min.} / 260 = 100 \text{ obr.}/\text{min.}$, co zgadza się z danymi technicznymi. Natomiast wzrost momentu obrotowego wyniesie $0,0072 \text{ kG}\cdot\text{cm} \cdot 260$, co daje na ramieniu serwomechanizmu moment obrotowy 1,87 kG·cm.

Przypis redaktora: momenty obrotowe podano tu w kG·cm. W układzie SI obowiązującą jednostką momentu obrotowego jest niutonometr (Nm). 1 Nm = 10,19 kG·cm.

Jeśli pominąć wszelkie „wypukłe” elementy wystające z góry i z boków, zastosowanie tego układu przekładniowego pozwala zmieścić serwomechanizm w małej obudowie o wymiarach zaledwie około 23 mm szerokości, 29 mm wysokości i 12 mm głębokości. Przypuśćmy, że powracamy do systemu z dwoma kołami zębatymi, od którego zaczęliśmy (rysunek 7). Ponadto założymy, że chcemy zachować nasz oryginalny silnik z kołem zębatym o 9 zębach, i że chcemy utrzymać przełożenie 260:1. Oznacza to, że nasze koło wyjściowe potrzebowałoby $9 \cdot 260 = 2340$ zębów, co – przy założeniu takiego samego modułu zęba jak w kole 9-zębowym – oznaczałoby średnicę około 468 mm (rysunek 10).

Myślę, że wszyscy zgodzimy się, iż w małych zdalnie sterowanych modelach koła zębate o średnicy zbliżonej do pół metra nie znalazłoby zastosowania. Nie wiem jak Was, ale mnie zadziwia fakt, że dzięki kilku niewielkim kołom zębatym możemy zmniejszyć



Rysunek 10. Bez układów przekładniowych modele hobbystyczne byłyby sporo większe niż ich rzeczywiste pierwowzory

półmetrowego potwora do rozmiaru około ośmiu centymetrów sześciennych.

W następnym odcinku

Obawiam się, że nadszedł czas, abym przekazał ten felieton do centrali i zmierzył się z „gniewem Matta” (właśnie wysłałem lokaja, by przyniósł mi specjalne spodnie do „drżenia ze strachu”). W następnym artykule, oprócz pokazania animacji Steve’a przedstawiającej z niesamowitymi szczegółami działanie wnętrza serwomechanizmu, przejdziemy do sterowania wieloma serwomechanizmami jednocześnie poprzez joysticki. Zanim jednak nastąpi ten szczęśliwy dzień, jak zawsze czekam na Wasze fascynujące komentarze, wnikliwe pytania i mądre sugestie. ■

Clive „Max” Maxfield

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, wrzesień 2022 (www.epemag3.com)

REKLAMA

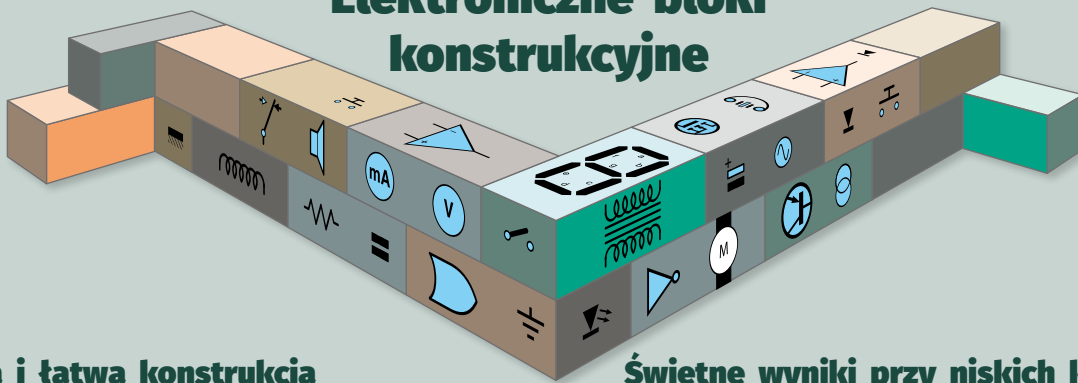
m.technik
Ciekawi świata są zawsze młodzi

Nie było zamknięcia.
Jest **nowe otwarcie!**

Ciekawszy,
na czasie,
na topie...
wiecznie młody!

przejrzyj i kupisz na stronie
www.ulubionykiosk.pl





Szybka i łatwa konstrukcja

Świetne wyniki przy niskich kosztach

Wybór i stosowanie siłowników, część 2

Przekształcanie sygnałów elektronicznych w ruch fizyczny – liniowy i obrotowy

W zeszłym miesiącu przyjrzelśmy się niektórym sposobom przekształcania sygnałów elektronicznych w ruch fizyczny. Omówiliśmy silniki prądu stałego i ich zastosowanie w siłownikach liniowych i serwomechanizmach RC. W tym miesiącu kontynuujemy temat, zajmując się silnikami z przekładnią i silnikami krokowymi.

Silniki z przekładnią

Motoreduktory to małe silniki prądu stałego, wyposażone w przekładnię połączoną z ich osią główną. Przekładnia zmniejsza prędkość i zwiększa moment obrotowy, dzięki czemu motoreduktory znajdują wiele zastosowań. Większość tych silników ma zeszlifowany wał wyjściowy, o przekroju w kształcie litery D, więc można na nich zamocować koła z odpowiednio ukształtowanym otworem wewnętrznym (rysunek 1).

Mikro silniki z przekładnią mają przekrój 10×12 mm i odsłonięte koła redukcyjne. Są one dostępne w bardzo szerokim zakresie przełożeń, od 5:1 do 1000:1. Motoreduktory „20D” wykorzystują zamknięte przekładnie o średnicy 20 mm o przełożeniach w przedziale od 29:1 do 154:1. Dostępnych jest wiele motoreduktorów o średnicach 25 mm i 37 mm z różnymi przełożeniami.

Jeśli zamierzasz eksploatować motoreduktor przez długi czas, upewnij się, że wybierasz model ze szczotkami węglowymi. Niektóre modele mają szczotki metalowe, które zużywają się szybciej. Dostępnych jest wiele takich silników, z różnymi przekładniami, długościami wału wyjściowego, na różne napięcia i maksymalne prądy zasilania. Wartości wyjściowego momentu obrotowego mogą przekraczać 50 kg-cm.

Niektóre silniki są wyposażone w enkodery, które przekazują informacje zwrotne dotyczące prędkości i kierunku obrotów silnika.

Jeśli zamierzasz używać motoreduktora w typowym zastosowaniu (np. do obracania kołami w modelu samochodu), najtańszym rozwiązaniem jest zakup zestawu zawierającego silnik, koło napędowe, sprzęgło i wspornik montażowy (rysunek 2).



Rysunek 1. Silnik z przekładnią – zwróć uwagę, że przekładnia jest przymocowana bezpośrednio do korpusu silnika, co zapewnia kompaktową konstrukcję całości. Widoczny jest typowy wał wyjściowy z przekrojem w kształcie litery D. Prezentowany motoreduktor jest przeznaczony do pracy przy napięciu 24 V i obraca się z prędkością 25 obr./min (150°/s) bez obciążenia. Kosztuje zaledwie 6 funtów



Rysunek 2. Silnik z przekładnią, wspornikiem montażowym, sprzęgłem i kołem, nadaje się do napędzania robotów, pojazdów i podobnych ruchomych urządzeń – prezentowany motoreduktor przy napięciu 6 V pobiera prąd o natężeniu 0,5 A i obraca się z prędkością 100 obr./min. Koszt wynosi 11 funtów (dzięki uprzejmości Banggood)



Rysunek 3. Zespół silnika i przekładni ze starej wiertarki elektrycznej zasilanej bateryjnie odznacza się dużym momentem obrotowym. Połączenie mechaniczne można wykonać poprzez zamocowanie wału w uchwycie

Jeśli szukasz dużego i mocnego motoreduktora, a dysponujesz ograniczonym budżetem, wykorzystaj starą wiertarkę elektryczną zasilaną akumulatorowo. Zużyte wiertarki wykorzystujące akumulatory niklowo-wodorkowe są bardzo często wyrzucane. Wewnątrz znajduje się silnik szczotkowy prądu stałego, połączony z przekładnią planetarną (**rysunek 3**). Ten mechanizm można łatwo wykorzystać w swoich konstrukcjach. Takie zespoły silników i przekładni zapewniają duży moment obrotowy, nawet przy niskim napięciu zasilania, co pozwala na regulację prędkości w szerokim zakresie, przy zachowaniu dużej mocy niezbędnej do pokonania rzeczywistych obciążeń. Sprzęgło montowane w wielu wiertarkach może służyć do ochrony silnika w przypadku zablokowania wału wyjściowego.

Podsumowując, w przypadku gdy wymagany jest ciągły napęd o wysokim momencie obrotowym, silnik z przekładnią jest doskonałym wyborem. Ogromny wybór i dostępność złączy, uchwytów montażowych i akcesoriów oznacza również, że można zrealizować wiele różnych projektów bez konieczności obróbki skrawaniem lub skomplikowanych prac montażowych. Należy jednak pamiętać, że tanie motoreduktory hobbystyczne mają ograniczoną trwałość – często stosuje się w nich tuleje ślizgowe zamiast łożysk kulkowych, a ich przekładnie nie zawsze umożliwiają łatwe smarowanie.

Silniki krokowe

W przeciwieństwie do ciągłego ruchu obrotowego, zapewnianego przez silniki szczotkowe prądu stałego, silniki krokowe – jak sama nazwa wskazuje – poruszają się dyskretnymi krokami, z których każdy stanowi niewielką stałą część pełnego obrotu. Silniki krokowe są stosowane w profesjonalnych i hobbystycznych drukarkach 3D oraz obrabiarkach CNC. Są one również wykorzystywane w drukarkach, aparatach fotograficznych i wielu innych produktach konsumenckich i przemysłowych (**rysunek 5**).

Silniki krokowe mają trzy główne zalety. Pierwszą z nich jest możliwość ustawienia wału wyjściowego w określonym położeniu. Oznacza to, że jeśli wał wyjściowy zostanie obrócony o – powiedzmy – 50 kroków, można dokładnie określić jego końcowe położenie. Należy jednak pamiętać, że silnik krokowy nie ma wbudowanego czujnika położenia i pracuje bez sprzężenia zwrotnego, na zasadzie zliczania kolejnych kroków. Funkcja do określania faktycznego stanu mechanizmu napędowego jest realizowana przez czujniki znajdujące się poza silnikiem.

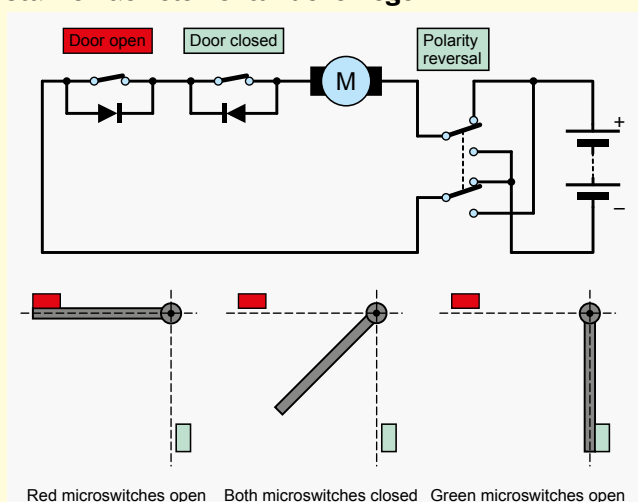
Drugą istotną zaletą silnika krokowego jest duży moment obrotowy przy niskich prędkościach ruchu. Dzięki temu może on skutecznie uruchamiać jakiś mechanizm z pozycji spoczynkowej, bez konieczności stosowania przekładni. Ponadto pracą silnika krokowego można precyzyjnie

Automatyczny wyłącznik rozwierający się w skrajnych ustawieniach elementu ruchomego

Informacje zwrotne dotyczące położenia siłownika mogą mieć złożony charakter. Zazwyczaj potrzebny jest mikrokontroler z oprogramowaniem PID. Ale co zrobić, jeśli chcesz po prostu, aby silnik prądu stałego wyłączył się po osiągnięciu określonej pozycji? W takim przypadku możesz skorzystać z podejścia stosowanego w siłownikach liniowych. Przyjrzyjmy się temu bliżej, jest to eleganckie i proste rozwiązanie.

Wyobraźmy sobie, że chcemy wyłączyć siłownik, gdy drzwi są całkowicie otwarte lub zamknięte. W obwodzie zasilania silnika zastosowano dwa normalnie zwarte mikroprzełączniki krańcowe, których dźwignie lub przyciski są ustawione w taki sposób, że jeden przełącznik otwiera się, gdy drzwi są całkowicie otwarte, a drugi przełącznik otwiera się, gdy drzwi są całkowicie zamknięte. Równoległe do styków przełączników dołączone są diody.

Aby otworzyć drzwi, należy zasilić obwód. Następnie, gdy drzwi zostaną całkowicie otwarte, mikroprzełącznik odcina zasilanie silnika. Aby zamknąć drzwi, wystarczy zmienić polaryzację zasilania silnika. Dioda dołączona równoległe do mikroprzełącznika odpowiadającego położeniu „drzwi otwarte” umożliwia przepływ prądu z pominięciem tego przełącznika, dzięki czemu siłownik może się cofnąć. Gdy drzwi całkowicie się zamkną, dzieje się to samo, ale tym razem otwiera się drugi przełącznik odcinający zasilanie. Należy pamiętać, że diody muszą być dostosowane do obciążeń prądowych silnika.



Rysunek 4. Prosty sposób na automatyczne wyłączenie zasilania silnika prądu stałego w pozycjach całkowicie otwartej i całkowicie zamkniętej



Rysunek 5. Silniki krokowe są stosowane w bardzo wielu urządzeniach konsumenckich i przemysłowych, co ułatwia ich pozyskiwanie. Są one również łatwo dostępne w handlu (dzięki uprzejmości Adafruit)

sterować. Zasilający go układ elektroniczny może regulować prędkość obrotową przez ustalenie odpowiedniej liczby kroków na sekundę. To sprawia, że działanie sterownika jest oparte na regulacji częstotliwości.

Silniki krokowe są najczęściej stosowane w mechanizmach, w których ważne jest precyzyjne określenie pozycji lub prędkości ruchu. Jednak w wielu zastosowaniach hobbystycznych silniki krokowe, ze względu na ich dużą wytrzymałość mechaniczną, służą do prostego napędu urządzeń. Większość średnich i dużych silników krokowych jest wyposażona w łożyska kulkowe. Istotną zaletą jest brak szczotek i komutatora, które ulegają zużyciu, dzięki temu silniki krokowe są niezawodne i mają długą żywotność. Jeśli potrzebujesz silnika o dobrych parametrach, użycie silnika krokowego może być lepszym rozwiązaniem niż zastosowanie konwencjonalnego silnika szczotkowego prądu stałego.

Jednak silniki krokowe mają jedną istotną wadę – aby osiągnąć wszystkie te zalety, konieczne jest zastosowanie elektronicznych układów sterujących. W przeciwieństwie do zwykłych serwomechanizmów, sterownik silnika krokowego musi zapewnić zasilanie uzwojeń nawet wtedy, gdy wał silnika spoczywa w bezruchu.

Silniki krokowe są dostępne w wielu rozmiarach i konfiguracjach – od małych jak moneta po duże jak kubek do kawy. Jednak wielu hobbystów używa silników krokowych, o znormalizowanych rozmiarach zgodnych z NEMA (National Electrical Manufacturers Association). Rozmiar NEMA odnosi się do wymiarów kołnierza montażowego silnika, z których wynika między innymi rozstaw otworów montażowych. Tak więc silnik krokowy NEMA 14 ma odległość 1,4 cala między otworami montażowymi kołnierza, NEMA 17 ma odległość 1,7 cala między otworami montażowymi i tak dalej (rysunek 6).

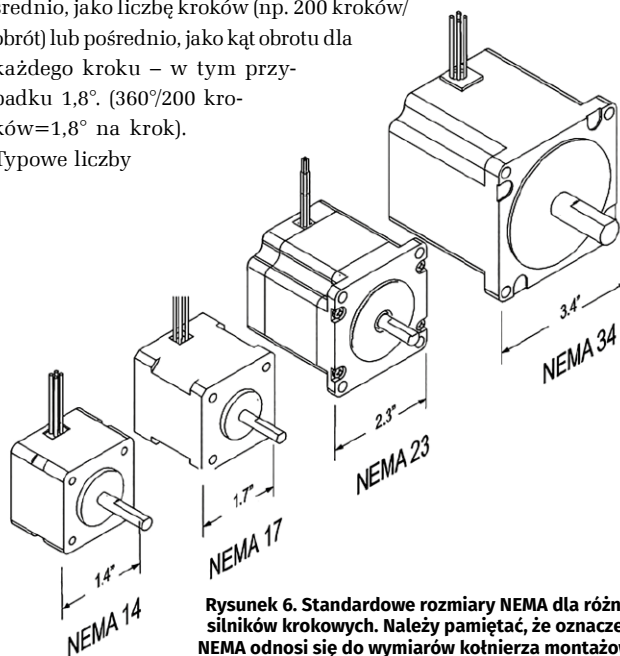
Należy pamiętać, że oznaczenie NEMA odnosi się wyłącznie do fizycznych rozmiarów płyty czołowej – nie odnosi się do właściwości mechanicznych ani elektrycznych silnika. Niemniej jednak, rozmiar NEMA 17 jest powszechnie stosowany w drukarkach 3D i mniejszych maszynach CNC. Większe rozmiary NEMA są bardziej powszechne

w maszynach CNC wykorzystywanych w zastosowaniach przemysłowych. Małe silniki krokowe są przydatne w zastosowaniach robotycznych i animatronicznych.

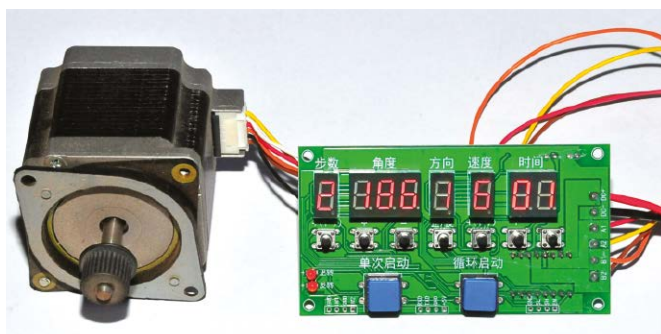
Podobnie jak motoreduktory, silniki krokowe są dostępne z zeszlifowanymi wałkami wyjściowymi, o przekroju w kształcie litery D. Dostępnych jest wiele adapterów, łączników, wałków, przekładni i dźwigni, które ułatwiają konstrukcję mechaniczną projektu.

Kolejną cechą silnika krokowego, którą należy wziąć pod uwagę, jest liczba kroków składających się na pełny obrót. Można ją wyrazić bezpośrednio, jako liczbę kroków (np. 200 kroków/obrót) lub pośrednio, jako kąt obrotu dla każdego kroku – w tym przypadku $1,8^\circ$. ($360^\circ/200$ kroków = $1,8^\circ$ na krok).

Typowe liczby



Rysunek 6. Standardowe rozmiary NEMA dla różnych silników krokowych. Należy pamiętać, że oznaczenie NEMA odnosi się do wymiarów kołnierza montażowego, a nie do parametrów elektrycznych silnika



Rysunek 7. Prezentowany samodzielny, programowalny sterownik pozwala na łatwe uruchomienie silnika krokowego. Można ustawić prędkość, kierunek i kąt obrotu, a moduł ma możliwość zapisania dziewięciu programów, które mogą być uruchamiane sekwencyjnie

kroków na obrót to 24, 48 i 200. Do innych cech silnika krokowego należą napięcie, prąd znamionowy na fazę, oraz moment obrotowy, utrzymujący wał silnika w nieruchomej pozycji. Na przykład silnik krokowy rozmiaru NEMA-17 może mieć następujące cechy: prąd 1,7 A na fazę przy napięciu 2,8 V i moment trzymania 3,7 kG·cm. Więcej informacji na temat momentu obrotowego można znaleźć w wydaniu z poprzedniego miesiąca. Należy pamiętać, że prąd uzwojeń jest utrzymywany także w bezruchu, choć w wielu sterownikach bywa on programowo ograniczany w celu zmniejszenia nagrzewania.

Istnieje wiele konfiguracji elektrycznych silników krokowych, w wyniku których powstają silniki 4-przewodowe, 6-przewodowe i 8-przewodowe. Można je zasilac na różne sposoby – np. szeregowo i równolegle, jednobiegunowo i dwubiegunowo, pełnym krokiem i mikrokrokiem. Szczegółowe omówienie różnych technik napędzania silników krokowych można znaleźć w artykule „Using Stepper Motors Parts 1–5” (*Korzystanie z silników krokowych*, części 1–5), PE, październik 2019 – luty 2020.

Jeśli używasz silnika krokowego w zastosowaniach wymagających precyzyjnego pozycjonowania lub utrzymywania określonej prędkości obrotowej, sugeruję zakup silnika krokowego odpowiednio dobranego do danego zastosowania, a także sterownika zalecanego przez dostawcę. Silniki krokowe mogą być wykorzystywane do zwyczajnego napędu jakichś mechanizmów, lub do ich pozycjonowania z wysoką precyzją i dużym momentem obrotowym. Do tego wymagane jest użycie odpowiednich sterowników, interfejsów i oprogramowania.

Jeśli jednak chcesz jedynie obracać silnikiem, to w handlu dostępne są niedrogie sterowniki w postaci płytek z układami elektronicznymi. Na przykład płytka Banggood/AliExpress „DC 8 V-27 V Programmable Stepper Motor Driver Controller Board Step/Angle/Direction/Speed/Time Adjustable 42/57 Phase” kosztuje od 10 do 20 funtów i współpracuje z większością silników krokowych podłączonych w konfiguracji czteroprzewodowej (**rysunek 7**). Obszernie omówiliśmy tę płytkę w publikacji *Electronic Building Block* z grudnia 2021 r.

Silniki krokowe dobrze sprawdzają się w przypadku konieczności precyzyjnego pozycjonowania napędzanych mechanizmów. Są one również przydatne tam, gdzie wymagana jest wysoka wytrzymałość. Gdybym miał za zadanie skonstruować działającą przez całą dobę ruchomą reklamę sklepową, użyłbym silnika krokowego.

Elektromagnesy

W sytuacjach, w których potrzebny jest siłownik o szybkim, krótkim ruchu liniowym, elektromagnes jest najtańszą i najwygodniejszą opcją.

Elektromagnes składa się z cewki nawiniętej na wydrążonym karkasie. Wewnątrz karkasu znajduje się ruchomy, żelazny tłok. Po podłączeniu zasilania do cewki tłok zostaje wciągnięty do wnętrza karkasu.

Podczas tego ruchu sprężyna zwrotna zostaje ściśnięta. Po odłączeniu zasilania od cewki sprężyna powraca do poprzedniego położenia i wysuwa tłok. Elektromagnesy działają bardzo szybko, więc mogą powodować natychmiastowe przemieszczenie jakichś elementów.

Do parametrów charakteryzujących elektromagnesy należą:

- Napięcie robocze
- Pobór prądu w stanie aktywnym (należy pamiętać, że po załączeniu zasilania prąd narasta wykładniczo, dążąc do wartości ustalonej wynikającej z rezystancji cewki)
- Maksymalna wartość siły (zwykle w niutonach).

Większość elektromagnesów jest przystosowana do pracy ciągłej – oznacza to, że mogą być stale zasilane i utrzymywane w pozycji wciągniętej. Należy pamiętać, że elektromagnesy są projektowane do pracy z prądem stałym lub przemiennym – nie są one bezpośrednio zamienne i wymagają odpowiedniej konstrukcji.

Należy upewnić się, że każdy zakupiony elektromagnes ma tłok pozwalający na łatwe wykonanie połączenia mechanicznego. Wiele elektromagnesów ma tłok ze spłaszczonym końcem, co ułatwia połączenie mechaniczne za pomocą prostego sworznia lub śruby i nakrętki. Elektromagnesy nie są drogie, można je także pozyskać ze starych magnetofonów lub odtwarzaczy wizyjnych.

Jeśli używasz tranzystora lub podobnego elementu do sterowania elektromagnesem, pamiętaj, że konieczne jest użycie diody zwrotnej połączonej równolegle z uzwojeniem, biorącej na siebie prąd indukowany w momencie wyłączenia zasilania.

Elektromagnesy nadają się do zastosowań takich jak zamki drzwiowe i szafkowe, ale mogą być również wykorzystywane do podnoszenia wskaźników wizualnych lub do automatyzacji pracy mechanicznych uchwyty i gałek.

Wnioski

Jeśli chcesz stworzyć system mechaniczny sterowany elektronicznie, masz do wyboru wiele możliwości. Obecnie dostępnych jest wiele niedrogich sterowników elektronicznych, a sektor detaliczny oferuje szeroki wybór części mechanicznych. Jednak sukces Twojego projektu będzie w dużej mierze zależał od wybranych siłowników i sposobu ich sterowania.

W przypadku zastosowań wymagających dużej prędkości ruchu i niskiego momentu obrotowego dobrze sprawdza się prosty silnik szczotkowy prądu stałego, który bezpośrednio napędza obciążenie (np. wentylator). Jeśli wymagany jest duży moment obrotowy (np. napędzanie wciągarki lub robota kołowego pokonującego wzniesienia), silnik z przekładnią jest tanim i skutecznym rozwiązaniem. W przypadku naprawdę dużych obciążeń warto rozważyć zastosowanie mechanizmu podobnego do napędu wycieraczek samochodowych.



Rysunek 8. Typowy elektromagnes ogólnego przeznaczenia ze sprężyną zwrotną. Ten model jest przeznaczony do pracy przy napięciu 12 V, pobiera prąd o natężeniu 2 A i ma maksymalną siłę przyciągania 20 N (około 2 kg). Jego koszt jest niski – tylko około 4 funtów

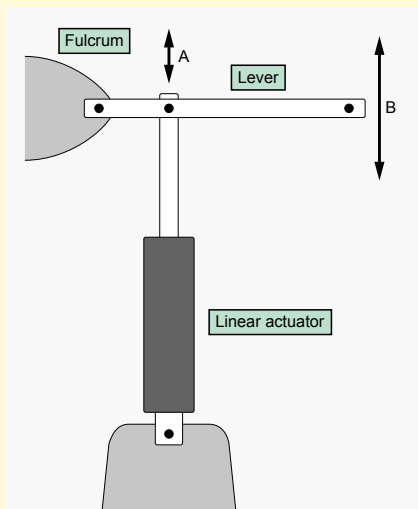
Zastosowanie dźwigni

W przypadku stosowania siłowników liniowych (patrz wydanie z poprzedniego miesiąca) i elektromagnesów szczególnie ważne są dźwignie. Dźwignie pozwalają zamieniać siłę na prędkość i skok ruchu – zwiększenie prędkości i zakresu ruchu odbywa się kosztem dostępnej siły, i odwrotnie – większa siła oznacza mniejszą prędkość i krótszy skok. Jeśli brzmi to znajomo, to dlatego, że działa to tak samo jak przekładnia. Można powiedzieć, że przekładnia składa się z ciągle obracających się dźwigni.

Powiedzmy, że chcemy stworzyć aerodynamiczny hamulec powietrzny do samochodu wyścigowego, taki, który podnosi się bardzo szybko, gdy kierowca naciska pedał hamulca (hamulec powietrzny wytwarza duży opór aerodynamiczny, dzięki czemu bardzo szybko spowalnia samochód przy dużych prędkościach – i to wszystko bez obaw o poślizg opon na drodze lub torze).

Przemieszczanie hamulca aerodynamicznego w górę lub w dół najlepiej byłoby realizować za pomocą siłownika liniowego, który po odłączeniu zasilania blokuje się w aktualnej pozycji. Jednak takie siłowniki są dość powolne, a hamulec aerodynamiczny powinien działać tak szybko, jak to możliwe. Czy ten proces można usprawnić stosując dźwignię? Jeśli umieścimy siłownik i punkt podparcia dźwigni (punkt obrotu) blisko siebie, zamienimy powolny ruch siłownika na szybszy (rysunek 9). Jednak wraz ze wzrostem prędkości wzrośnie również siła, jaką musi wytwarzać siłownik.

Załóżmy, że siłownik liniowy rozciąga się całkowicie w ciągu 5 sekund, a my chcemy ustawić dźwignię tak, aby hamulec aerodynamiczny wysuwał się całkowicie w ciągu



Rysunek 9. Dzięki zastosowaniu dźwigni możemy zwiększyć skok i prędkość działania siłownika liniowego – odległość B będzie większa niż odległość A, a obie zostaną osiągnięte w tym samym czasie. Nie ma jednak nic za darmo – siła, jaką musi wytworzyć siłownik liniowy, wzrasta wprost proporcjonalnie do długości dźwigni

zaledwie 0,2 sekundy. Jest to stosunek 25:1, więc siła działająca na siłownik liniowy wzrasta w tym samym stosunku. Zależy to od prędkości samochodu i powierzchni hamulca aerodynamicznego, ale założmy, że przy pełnym wysunięciu na hamulec aerodynamiczny działa siła 300 N (czyli siła, która próbuje go zamknąć). 300 N pomnożone przez 25=7500 N! W rzeczywistości wymagana siła nie będzie aż tak duża (ponieważ przy każdym kącie mniejszym niż pełne otwarcie hamulec pneumatyczny będzie wywierał mniejszą siłę), ale proste obliczenia pokazują,



Rysunek 10. Testowy, tylny, aerodynamiczny hamulec powietrzny w samochodzie wyścigowym. Podczas hamowania panel umieszczony między bocznymi żebrami jest podnoszony przez obrót o około 80 stopni. Wymagana szybkość podnoszenia i występujące siły spowodowały, że żaden dostępny na rynku siłownik elektryczny nie był w stanie tego dokonać, niezależnie od zastosowanej dźwigni. Ale był to w stanie zrobić cylinder pneumatyczny (wskazany strzałką). Jaki był wynik? Przy bardzo dużej prędkości hamulec działał doskonale, ale niestety przy prędkościach dopuszczalnych przez kodeks drogowy był raczej słaby

że żaden przeciętny siłownik liniowy nie spełni tego wymagania. W rzeczywistości, w tej sytuacji użyłem cylindra pneumatycznego na sprężone powietrze o dużym ciśnieniu. Co prawda nie spełniał on wymagań dotyczących siły, ale był bliski osiągnięcia tego celu.

Tak więc, chociaż dźwignie są bardzo skuteczne w działaniu, w opisywanych przez nas sytuacjach należy pamiętać, że nie ma nic za darmo – zawsze trzeba będzie pójść na kompromis między prędkością, skokiem a wymaganą siłą siłownika.

Jeśli wymagany jest ruch po linii prostej, siłownik liniowy zapewni trwałe i bardzo wydajne wyniki. Dzięki wbudowanym wyłącznikom krańcowym lub sterowaniu elektronicznemu opartemu na sprzężeniu zwrotnym można również łatwo ustawić maksymalne i minimalne wysunięcie ciężła. Każdy z tych siłowników opartych na silniku prądu stałego może być również sterowany metodą PWM, a dzięki zastosowaniu bardziej zaawansowanego sterownika można precyzyjnie regulować maksymalną prędkość i przyspieszenie w ruchu siłownika.

Jeśli wymagana jest precyzja pozycjonowania, a wymagany moment obrotowy i obciążenia nie są zbyt duże, zaś obrót ma odbywać się w zakresie mniejszym niż 180°, serwomechanizm RC może zapewnić wszystkie te funkcje – precyzyjnie i w niskiej cenie. Serwomechanizmy RC są dobrym przykładem złożonego produktu o dobrym stosunku jakości do ceny. W prostych przypadkach mogą być sterowane za pomocą mikrokontrolera, jednak gdy wymagana jest precyzyjna kontrola prędkości i przyspieszenia ruchu, wymagają użycia bardziej zaawansowanego kontrolera. Należy pamiętać, że serwomechanizmy RC nie są zazwyczaj przeznaczone do długotrwałej pracy. Przykładem może być rozsuwanie zasłon w oknach kilkaset razy w ciągu roku.

Silniki krokowe sprawdzają się tam, gdzie najważniejsza jest dokładność pozycjonowania – w maszynach CNC, drukarkach 3D i podobnych urządzeniach. Nie należy jednak zapominać o czymś, co wydaje się być nieco pomijane – jeśli chcesz, aby system sterowania

silnikiem krokowym działał z dużą dokładnością, silnik i sterownik muszą być dobrze dopasowane do wykonywanego zadania. Z drugiej strony, jeśli potrzebujesz po prostu bardzo wytrzymałego silnika, który ma pracować tysiące godzin w roku, silnik krokowy również może to zapewnić.

To stwierdzenie sprowadza nas z powrotem do motoreduktorów. Muszę przyznać, że gwałtowny wzrost podaży tych silników mnie zaskoczył. W handlu dostępne są konstrukcje o różnych parametrach elektrycznych, przełożeniach i konfiguracjach wału wyjściowego. Dzięki doborowi odpowiedniego przełożenia, możliwości pracy przy niskim napięciu zasilania i zastosowaniu czujników położenia w układzie ze sprzężeniem zwrotnym, motoreduktory mogą konkurować z silnikami krokowymi w wielu zadaniach, a jednocześnie są znacznie prostsze w użyciu. Z drugiej strony mogą obracać się wystarczająco szybko, aby spełniać wymagania dotyczące wysokich prędkości ruchu.

A elektromagnesy? Aby działały, wystarczy podłączyć lub odłączyć zasilanie prądem przemiennym lub stałym – nie ma nic prostszego! Naprawdę można dobrać siłownik odpowiedni do każdego projektu. ■

Julian Edgar

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, listopad 2022 (www.epemag3.com)



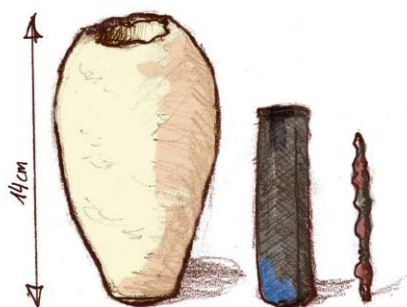
Fot. www.pexels.com/photo/2047905/

Wynalazcy i ich wynalazki w elektronice, część 3

W dwóch poprzednich odcinkach opisaliśmy licznych wynalazców, którzy wnieśli istotny wkład w rozwój elektroniki. Ich prace umożliwiły powstanie wielu nowoczesnych technologii. Sporo znaczących osiągnięć pojawiło się również na uniwersytetach, w firmach i innych zespołach ludzkich. Osiągnięcia te opisujemy w niniejszej – trzeciej i ostatniej – części.

Niniejszy artykuł, ostatni z tej serii, obejmuje znaczące wynalazki, których nie można przypisać konkretnej osobie, ponieważ albo nie znamy jej nazwiska, albo była ona częścią zespołu. W przeciwieństwie do poprzednich dwóch części, gdzie wynalazki były uporządkowane według dat urodzenia wynalazców, informacje uszeregowaliśmy według roku powstania wynalazku lub dokonania odkrycia.

Przypis redaktora: część opisanych wynalazków dotyczy ściśle Australii, co jest zrozumiałe, biorąc pod uwagę, że artykuł pochodzi z australijskiego czasopisma.



Rysunek 1. Wizerunek „baterii bagdadzkiej”. Źródło: <https://w.wiki/7FNe>

Sum elektryczny, ~2750 p.n.e.

Starożytny egipski mural w grobowcu architekta Ti w Sakkarze w Egipcie odnosi się do suma elektrycznego, o którym później Pliniusz i Plutarch mówili, że leczy bóle stawów i inne dolegliwości. Może to być jedno z najwcześniejszych odkryć związanych z elektrycznością.

Bateria bagdadzka, ~150 p.n.e. – 650 n.e.

Antyczna „bateria z Bagdadu” (rysunek 1) jest przez niektórych uważana za ogniwo baterii elektrycznej. Mogła mieć jednak inne przeznaczenie i brak dowodów na to, że była używana jako elektryczna bateria. Patrz artykuł na temat baterii w Silicon Chip, wydanie ze stycznia 2022 r. (www.siliconchip.au/Series/375).

Światłowodowy, 27 p.n.e.

Wiadomo, że już starożytni Rzymianie przeciągali szkło w długie elastyczne włókna, takie, jakie są współcześnie używane w światłowodach przeznaczonych do telekomunikacji i transmisji światła.

Latarnia morska, kabel transatlantyczny, 1858

Pierwszą latarnią morską z elektrycznym źródłem światła była South Foreland Lighthouse koło Dover w Wielkiej Brytanii.

Wykorzystywała ona łukową lampę węglową opracowaną przez Fredericka Hale’a Holmesa i do 1860 roku była poddawana testom. W 1872 roku otrzymała stałą instalację elektryczną.

Prąd był wytwarzany przez dwa silniki parowe opalane koksem, napędzające cztery prądnice współdzielone z sąsiednią latarnią morską. Badania nad oświetleniem elektrycznym dla latarni morskich prowadził wówczas Michael Faraday. Holmes zademonstrował mu instalację latarni.

Również w 1858 roku położono pierwszy transatlantyczny kabel telegraficzny. Działał on tylko przez trzy tygodnie. Potrzebował dwóch minut na przesłanie jednego znaku, czyli średnio 10 minut na słowo.

Podmorski kabel telegraficzny, 1859

Między australijskim stanem Wiktoria a wyspą Tasmanią został położony podmorski kabel telegraficzny. Był to wówczas najdłuższy kabel podmorski. Został on wyłączony w 1861 roku.

Telegraf transkontynentalny w USA, 1861

W USA oddano do użytku transkontynentalną linię telegraficzną.

Funkcjonalny kabel transatlantyczny, 1866

Transatlantycki kabel telegraficzny o większej funkcjonalności. Informacje mogły być przesyłane z prędkością ośmiu słów na minutę.

Międzynarodowy kabel telegraficzny, 1872

Uruchomiono międzynarodowy kabel telegraficzny między australijskim miastem Darwin a wyspą Jawą.

Połączenie telegraficzne Australia – Nowa Zelandia, 1876

Uruchomiono połączenie telegraficzne między Australią a Nową Zelandią.

Australijska transkontynentalna linia telegraficzna, 1877

W Australii między miastami Port Augusta a Albany została uruchomiona transkontynentalna linia telegraficzna. Miała długość 3196 km.

Pierwsze międzynarodowe połączenie telefoniczne, 1881

Pierwsza międzynarodowa rozmowa telefoniczna odbyła się między abonentami w Nowym Brunzwicku w Kanadzie i Maine w USA.

W Sydney w Australii otwarto centralę telefoniczną obsługującą 12 abonentów.

Elektrownia publiczna, 1882

W Londynie zbudowano pierwszą wielką elektrownię publiczną – Holborn Viaduct (znaną również jako Edison Electric Light Station). Wytwarzała ona moc 93 kW przy napięciu stałym 110 V, a jej generator napędzany był silnikiem parowym. Elektrownia zastąpiła używany poprzednio mały generator napędzany kołem wodnym w Godalming w hrabstwie Surrey, który generował jedynie 7,5 kW.

W Nowym Jorku została otwarta elektrownia Pearl Street. Miała sześć generatorów, każdy o mocy 100 kW, była zasilana parą, a jej ciepło odpadowe było wykorzystywane w miejscowej sieci grzewczej.

Telefonia i energetyka wodna, 1883

W australijskim mieście Adelajda została otwarta centrala telefoniczna z 48 abonentami, a w pobliskim Port Adelaide – centrala z 21 abonentami.

W kopalni cyny Mount Bischoff została uruchomiona pierwsza w Australii elektrownia wodna, zasilająca około 50 żarówek Swana.

Grafofon (fonograf), 1887

W Laboratorium Volta (założonym przez A.G. Bella), Chichester A. Bell i Sumner Tainter ulepszyli fonograf Edisona, używając wosku zamiast folii cynowej jako nośnika. Potwierdziło się tym samym przekonanie Alexandra Grahama Bella, że wosk jest lepszym nośnikiem zapisu.

Panowie założyli American Graphophone Company i zaczęli sprzedawać swój produkt, który odniósł sukces komercyjny.

Publiczne dostawy energii elektrycznej, 1888

Pierwszym miastem w Australii z publiczną siecią elektryczną, przeznaczoną dla oświetlenia łukowego i żarowego (240 V prądu stałego), było Tamworth w Nowej Południowej Walii.

Trójfazowy prąd przemienny, 1889

Miasto Young w Australii otrzymało zasilanie w systemie trójfazowym (jednym z wczesnych zastosowań tej technologii) do oświetlenia ulic, sklepów, biur i domów.

Elektrownia wodna prądu przemiennego, 1891

W Lauffen am Neckar rozpoczęła pracę pierwsza niemiecka elektrownia trójfazowego prądu przemiennego. Wytwarzano tam napięcie 15 kV i przesyłano je na odległość 175 km do Międzynarodowej Wystawy Elektrotechnicznej we Frankfurcie.

W Ames w amerykańskim stanie Kolorado została uruchomiona prawdopodobnie pierwsza na świecie komercyjna elektrownia wodna prądu przemiennego. Miała moc 3,75 MW i wytwarzała napięcie jednofazowe 3 kV 133 Hz. **Przypis redaktora:** Moc elektrowni Ames (1891) wynosiła około 100 KM (≈75 kW), a nie 3,75 MW. Wartość megawatowa odnosi się do późniejszych modernizacji tej instalacji. Źródła: IEEE, Ames Hydroelectric Generating Plant (1891) – Engineering and Technology History Wiki (ETHW); National Park Service – Ames Hydroelectric Plant.

Elektrownia ta wciąż działa, choć nie używa już oryginalnych generatorów. Natomiast inna elektrownia, pochodząca z 1905 roku, nadal wykorzystuje generator General Electric z 1904 roku wytwarzający napięcie 2,4 kV przy prądzie 1082 A (2,6 MW).

Normalizacja jednostek elektrycznych, 1893

Na Międzynarodowym Kongresie Elektrycznym w Chicago (USA) zostały opracowane normy i definicje jednostek elektrycznych: oma, ampera i wolta.

Publiczna elektrownia wodna, 1895

W Launceston w Australii uruchomiono pierwszą publiczną elektrownię wodną. Zasilala ona oświetlenie uliczne. W 1921 roku elektrownię przekształcono w trójfazową. Miała moc 2 MW i była używana do 1956 roku.

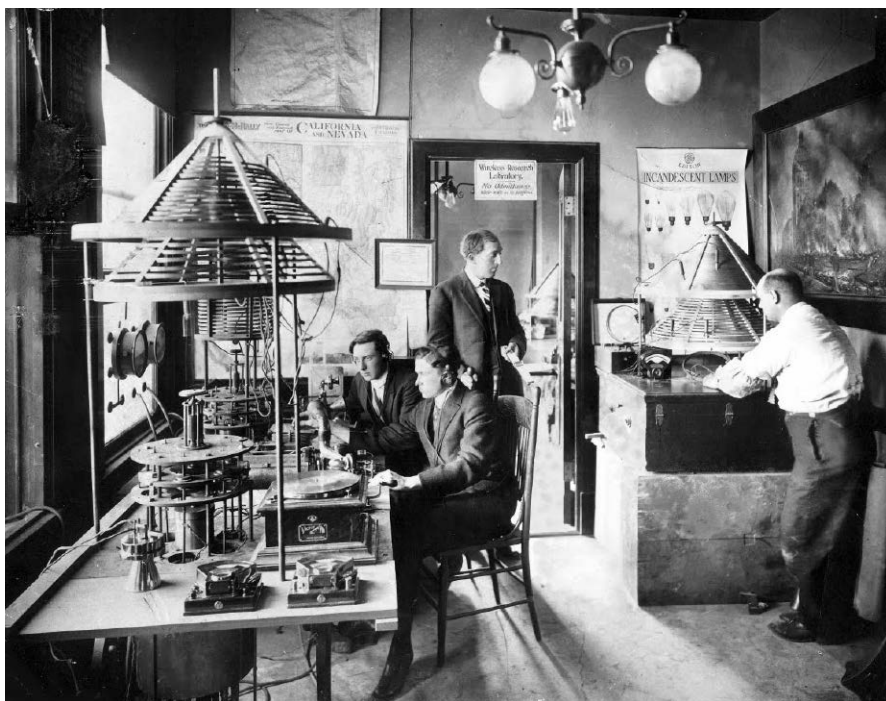
Telegraf międzynarodowy, 1902

Między Australią a Kanadą nawiązano połączenie telegraficzne, biegnące przez wyspy Fidżi i Norfolk.

Transmisja alfabetem Morse'a, 1906

Pierwszą w Australii oficjalną transmisję alfabetem Morse'a przeprowadziła firma Marconi Company na odcinku z Queenscliff do Devonport. Niektórzy twierdzą, że transmisje radiowe alfabetem Morse'a zostały przeprowadzone już w 1897 roku przez profesora Williama Henry'ego Bragga z Uniwersytetu w Adelajdzie, samodzielnie lub z G.W. Selbym z Melbourne.

Do 1906 roku Australia miała 46 elektrowni o łącznej mocy 36 MW.



Rysunek 2. Laboratorium radiowe Charlesa Herrolda w San Jose w stanie Kalifornia, około 1912 r. Herrold (stoi w drzwiach) nadawał z tego miejsca jako „radio KQW”. Źródło: <https://www.wiki/7EFw>

Produkcja żarówek z żarnikiem wolframowym, 1907

Firma Tokyo Electric Co. (poprzednik firmy Toshiba) rozpoczęła na małą skalę produkcję żarówek wolframowych. Pełnoskalową produkcję fabryka rozwinęła w 1910 roku.

Radiofonia, 1909

W stanie Kalifornia, za sprawą inżyniera Charlesa Davida Herrolda (1875–1948), rozpoczęła nadawanie stacja radiowa KQW (**rysunek 2**). Działała w celach eksperymentalnych, promocyjnych i szkoleniowych. Od około 1912 r. nadawano według programu wiadomości i audycje muzyczne. W tym czasie wiele innych stacji nadawało wyłącznie alfabetem Morse'a. Herrold otrzymał licencję komercyjną w 1921 roku. Stacja istnieje do dziś pod akronimem KCBS.

Amatorskie częstotliwości radiowe, 1912

W 1912 r. rząd Stanów Zjednoczonych uchwalił ustawę radiową przyznającą radioamatorom pasmo częstotliwości powyżej 1,5 MHz, ponieważ częstotliwości te uznawano za nieużyteczne. Doprowadziło to do odkrycia przez radioamatorów propagacji radiowej na falach krótkich poprzez odbicie od jonosfery. W 1921 r. przeprowadzono jednokierunkową transmisję przez Atlantyk, a w 1923 r. transmisję dwukierunkową (patrz www.siliconchip.au/link/abnv).

Transkontynentalne połączenie telefoniczne, 1915

W USA została przeprowadzona pierwsza transkontynentalna rozmowa telefoniczna na odległość 5794 km. Była możliwa dzięki nowo wynalezionemu wzmacniaczowi lampowemu.

Telefony z tarczą obrotową, 1919

Firma Bell System w USA wyprodukowała pierwsze telefony z tarczą obrotową.

Komercyjna radiofonia, 1920

W mieście Pittsburgh w USA rozpoczęła nadawanie pierwsza na świecie licencjonowana komercyjna stacja radiowa o akronimie KDKA.

Podwójnie zwinięte włókno wolframowe, 1921

Junichi Miura z Tokyo Electric Co. skonstruował pierwszą żarówkę z podwójnie zwiniętym żarnikiem wolframowym, wykorzystując technikę opracowaną przez Benbowa (1917). Żarówka weszła do produkcji na małą skalę w 1930 roku, a do produkcji masowej w 1936 roku.

„Telefon komórkowy”, 1922

Przeprowadzono wczesne eksperymenty z „telefonem komórkowym”. Było to przenośne dwukierunkowe radio (radio-telefon), wykorzystujące antenę parasolową, a jako uziemienie – hydrant przeciwpożarowy.

Do radia była transmitowana muzyka ze stacji bazowej. Patrz film na YouTube zatytułowany „Pierwszy na świecie telefon komórkowy (1922)” – <https://youtu.be/ILiLaRXHU0>.

Transatlantyckie połączenie telefoniczne, 1926

Przeprowadzono pierwsze transatlantyckie połączenie telefoniczne.

Radio samochodowe oraz system „Phonovision”, 1927

Powstało pierwsze masowo produkowane radio samochodowe – „Philco Transitone”. Dokładny rok jego powstania jest przedmiotem sporów. Przed powstaniem tego modelu radioodbiorniki musiały być specjalnie przystosowywane do użytku w samochodach.

John Logie Baird stworzył pierwszy odtwarzacz płyt wizyjnych o nazwie „Phonovision”. Sygnał wyjściowy z tarczy Nipkowa – mechanicznego skanera telewizyjnego – był nagrywany na płytę gramofonową. Rozdzielczość wynosiła tylko 30 linii, a częstotliwość – 5 klatek na sekundę. Niektóre z tych nagrań zostały odnalezione. W latach 1982–87 stworzono oprogramowanie do odzyskiwania obrazów z płyt „Phonovision”. Patrz strona internetowa www.siliconchip.au/link/abnw z filmem „30-line TV video recordings news feature” („reportaż z nagraniami telewizji 30-liniowej”), a także inne filmy: <https://youtu.be/J2mb4R9W9TI>, [www.siliconchip.au/link/abnx](https://youtu.be/G3CFkK5OORw), <https://youtu.be/G3CFkK5OORw>.

Start i lądowanie samolotu bez widoczności, 1929

Pierwszy start i lądowanie samolotu wyłącznie według przyrządów przeprowadził porucznik James Doolittle w dwupłatowcu Consolidated NY-2. Samolot był wyposażony w wysokościomierz Kollsmana, żyroskop kierunkowy Sperry'ego i sztuczny horyzont. Dysponował odbiornikiem radiolatarni znacznikowych National Bureau of Standards oraz specjalnym odbiornikiem radiowym ze wskaźnikiem wibracyjnym Radio Frequency Laboratories.

Synteza związków ferrytowych, 1930

Yogoro Kato i Takeshi Takei z Politechniki Tokijskiej po raz pierwszy zsyntetyzowali związki ferrytowe. Materiały te były i są wykorzystywane w cewkach indukcyjnych, transformatorach, elektromagnesach, dławikach przeciwzakłóceń, rdzeniach pamięci komputerowych, taśmach magnetycznych, materiałach pochłaniających promieniowanie radarowe, głośnikach i magnesach.

W 1937 roku firma TDK jako pierwsza zastosowała rdzenie ferrytowe w radioodbiornikach, dzięki czemu stały się one mniejsze i lżejsze. Do końca drugiej wojny

światowej TDK była jedyną firmą dostarczającą rdzenie ferrytowe.

Płyty długogrające, 1931

RCA wprowadziła pierwsze dostępne w handlu płyty długogrające (LP; „long play”). Miały one średnicę 30 cm, były odtwarzane z prędkością 33⅓ obrotów na minutę i zawierały do 11 minut dźwięku na stronę (tyle samo, co standardowa szpula filmowa o długości 305 metrów). Ze względu na koszty sprzętu odtwarzającego, a także nadejście Wielkiego Kryzysu, płyty LP stały się rynkową porażką. Dopiero wersja z 1948 roku odniosła sukces komercyjny.

Magnetofon „Magnetophon K1”, 1935

Niemiecka firma AEG zaprezentowała „Magnetophon K1” – pierwszy użyteczny magnetofon (**rysunek 3**). Wykorzystywał on niemetalową taśmę pokrytą materiałem magnetycznym (tlenkiem żelaza).

Rodzaj taśmy został pierwotnie oparty na pomysły Fritza Pfeumera (patrz wpis w poprzednim odcinku artykułu), a następnie rozwinięty przez Friedricha Matthiasa. Trwała głowica została zaprojektowana przez Eduarda Schüllera, który również zbudował egzemplarze prototypowe magnetofonu.

Jakość dźwięku była słaba, dopóki Walter Weber (1907–1944) nie odkrył (przez przypadek!) techniki prądu podkładu wielkiej częstotliwości, której zastosowanie znacznie poprawiło jakość dźwięku. Magnetofon posiadał wszystkie podstawowe funkcje, które stały się standardem w późniejszych magnetofonach analogowych. Film przedstawiający podobny model – „Magnetophon FT4”



Rysunek 3. Magnetofon AEG „Magnetophon K1” w trakcie transportu na berlińską wystawę radiową w 1935 r. Źródło: <https://museumofmagneticsoundrecording.org/ManufacturersAEGMagnetophon.html>

– można obejrzeć na stronie <https://youtu.be/cLjD0B6QoaM>.

Telewizja wysokiej rozdzielczości, 1939

Niemiecka firma Fernseh AG zademonstrowała 1029-liniową telewizję wysokiej rozdzielczości – jak na ówczesne możliwości techniczne – przeznaczoną do wyświetlania map wojskowych. System ten wymagał pasma o szerokości 15 MHz, dlatego też telewizja HDTV nie została wprowadzona powszechnie. Zmieniło się to w latach 90. wraz z pojawieniem się nadawania cyfrowego.

Komeracyjna radiofonia FM, standard NTSC, 1941

Komeracyjna radiofonia na falach ultrakrótkich (FM) formalnie zaczęła działalność w USA, choć wcześniej odbywały się transmisje eksperymentalne. Używano pasma 42...50 MHz, podzielonego na 40 kanałów. W 1945 roku radiofonia FM została przeniesiona do pasma 88...106 MHz z 80 kanałami, a następnie rozszerzona do 108 MHz i 100 kanałów.

Pojawił się monochromatyczny standard telewizyjny NTSC.

Komputer cyfrowy „Colossus”, 1943

Zbudowano pierwszy brytyjski programowalny komputer cyfrowy – „Colossus”.

Komputer cyfrowy „ENIAC”, 1945

Zbudowano amerykański komputer „ENIAC” – pierwszy na świecie programowalny komputer cyfrowy ogólnego przeznaczenia.

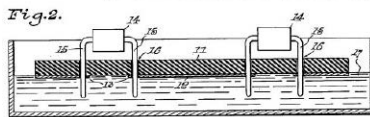
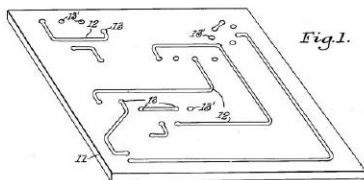
Wynaleziono również „Merrill” – elektroniczny system wyważania kół dla samochodów.

Australijska radiofonia na falach ultrakrótkich, 1947

W latach 1947–1961 w Australii miało miejsce eksperymentalne nadawanie programów na falach ultrakrótkich (FM). Ilość odbiorców była bardzo ograniczona (odbiorniki były drogie). Nadawanie wstrzymano, aby zwolnić pasmo dla telewizji, a później przywrócono je w paśmie, którego nikt inny na świecie nie używał. Na szczęście w 1975 roku radiofonię FM przeniesiono do powszechnie używanego pasma 88...108 MHz.

Płyty winylowe, 1948

Firma Columbia Records zaczęła używać polichloru winylu do produkcji płyt dźwiękowych. Płyty takie były trwalsze niż wcześniejsze płyty szelakowe. Na winylu można było umieścić znacznie drobniejsze rowki („mikrorowki”). Pozwoliło to na jednej stronie płyty o średnicy 30 cm uzyskać czas odtwarzania około 22 minut (istniały również płyty 25 cm). Format płyty opracował Peter Carl Goldmark (1906–1977).



Rysunek 4. Rysunek z patentu armii amerykańskiej nr 2,756,485 z 1956 r. dotyczącego produkcji płytek drukowanych

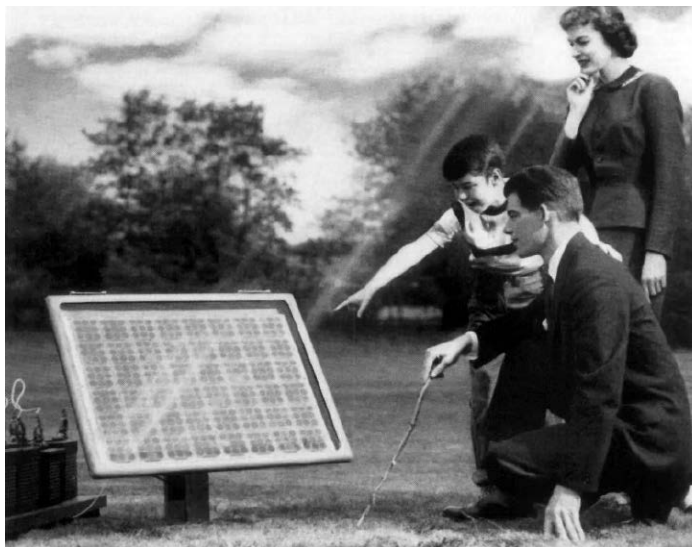
Płyty 45 obr./min., 1949

Firma RCA, konkurent Columbii, wprowadziła płytę 45 obr./min. o średnicy 18 cm, przeznaczoną jako format zastępczy dla płyt 78 obr./min., o zbliżonym czasie zapisu wynoszącym około pięciu minut na stronę. Ostatecznie „muzyka wysokiej jakości” była dystrybuowana na płytach 33 1/3 obr./min., a „muzyka popularna” na płytach 45 obr./min. Oba formaty istnieją do dziś.

Magnes trwały, płytki drukowane, 1950

W firmie Philips przypadkowo odkryto heksaferyt baru. Stał się on materiałem powszechnie stosowanym do wyrobu magnesów trwałych.

W 1956 roku Armia Stanów Zjednoczonych złożyła wniosek o przyznanie patentu USA nr 2,756,485 na „proces montażu układów elektrycznych” (rysunek 4). Konsekwencją tego wynalazku stała się masowa produkcja obwodów drukowanych. Przypis redaktora: Płytki drukowane były stosowane już wcześniej (od lat 30. XX wieku). Wspomniany patent dotyczył technologii ich montażu i produkcji, która przyczyniła się do ich upowszechnienia. Źródło: P. Eisler, Printed Circuits (1959).



Rysunek 5. Reklama pierwszego ognia słonecznego o praktycznym znaczeniu, produkcji firmy Bell, pochodząca z 25 kwietnia 1954 roku. Jego sprawność wynosiła 6%. Źródło: www.onthisday.com/photos/1st-solar-battery

Energia jądrowa, telewizja kolorowa i inne, 1951

Firma Sony wypuściła na rynek magnetofon „H-1”. Był to pierwszy model przewidziany do użytku domowego. Ważył 13 kg.

W miejscowości Arco (USA) został uruchomiony pierwszy reaktor jądrowy – „EBR-1”. Mógł on zasilać cztery żarówki 200-watowe.

Firma CBS w USA zademonstrowała system telewizji kolorowej oparty na elektromechanicznym systemie sekwencyjnym (FSC). W systemie istniało bardzo niewiele odbiorników. Standard ten został wycofany, a jego miejsce zajął standard NTSC.

System rozpoznawania mowy, gra wideo, 1952

Zademonstrowano pierwszy system rozpoznawania mowy, który był w stanie z 90% dokładnością rozpoznać cyfry od zera do dziewięciu wypowiedziane przez jednego mówcę. System nosił nazwę „Audrey” („Automatic Digit Recognition machine” – „aparat do automatycznego rozpoznawania cyfr”). Został skonstruowany przez H.K. Davisa w Bell Laboratories.

Alexander Shafro Douglas (1921–2010) stworzył na Uniwersytecie Cambridge w Anglii pierwszą grę wideo. Nazywała się „OXO” i była wersją gry w „kółko i krzyżyk”.

Maser, atomowa łódź podwodna, NTSC, 1953

Na Uniwersytecie Columbia powstał pierwszy „maser” (laser mikrofalowy – skrót MASER oznacza „Microwave Amplification by Stimulated Emission of Radiation” czyli „wzmacnianie mikrofalowe poprzez stymulowaną emisję promieniowania”). Został zbudowany przez Charlesa H. Townesa, Jamesa P. Gordona i Herberta J. Zeigera. Masery są wykorzystywane jako wysoce stabilne

wzorce częstotliwości i wzmacniacze dla sygnałów mikrofalowych o wyjątkowo niskim poziomie szumów. Mogą również wytwarzać fale elektromagnetyczne na częstotliwościach mikrofalowych i nie tylko.

Zwodowano pierwszy okręt podwodny o napędzie atomowym – USS „Nautilus”.

Opublikowano standard telewizji kolorowej NTSC.

Ogniwo fotowoltaiczne, komputer tranzystorowy, 1954

W Bell Laboratories zostało opracowane pierwsze użyteczne ogniwo fotowoltaiczne (rysunek 5). Jego twórcami byli Calvin Souther Fuller (1902–1994), Daryl Chapin (1906–1995) i Gerald Pearson (1905–1987).

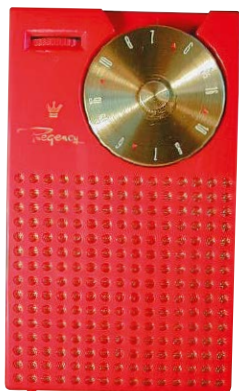
W USA miała miejsce pierwsza na świecie komercyjna transmisja telewizji kolorowej (NTSC). Przez długi czas większość programów była jednak czarno-biała ze względu na brak materiałów źródłowych w kolorze, a także z uwagi na wysokie ceny odbiorników kolorowych.

Firma Bell Labs zbudowała dla Sił Powietrznych Stanów Zjednoczonych pierwszy na świecie w pełni tranzystorowy komputer. „TRADIC” („TRAnsistor DIGital Computer” ew. „TRansistorized Airborne DIGital Computer” czyli „cyfrowy komputer na tranzystorach”) zawierał 684 tranzystory ostrzowe Bell 1734 i 10358 diod germanowych.

Do sprzedaży trafiło pierwsze przenośne radio tranzystorowe – Regency TR-1 (rysunek 6). Zawierało ono cztery tranzystory NPN firmy Texas Instruments i baterię 22,5 V. Kosztowało 49,95 dolarów, czyli równowartość dzisiejszych 850 dolarów (tyle mniej więcej płać dziś za ten odbiornik kolekcjonerzy!). Patrz artykuł w Silicon Chip z kwietnia 2013 (www.siliconchip.au/Article/3761).

Programowany syntezytor muzyczny i inne, 1955

Powstał syntezytor RCA „Mark I”, który był pierwszym programowanym syntezytorem



Rysunek 6. Regency TR-1 – pierwsze dostępne w handlu przenośne radio tranzystorowe

muzycznym. Na stronie www.siliconchip.au/link/abny znajduje się interesujący artykuł na temat jego działania.

Pojawił się pierwszy bezprzewodowy pilot do telewizora – „Zenith Flash-matic”. Wykorzystywał on światło widzialne i musiał być kierowany na jedną z czterech fotokomórek umieszczonych w rogach ekranu. Sterował różnymi funkcjami: włączaniem/wyłączaniem, wyciszaniem i zmianą kanału.

Napęd dyskowy IBM 350, magnetowid VRX-1000 i inne, 1956

Do sprzedaży trafił pierwszy komercyjny napęd dyskowy – IBM 350 (rysunek 7). Miał pojemność 3,75 MB i ważył około tony.

Firma Ampex wprowadziła do użytku studyjnego pierwszy komercyjny magnetowid monochromatyczny – VRX-1000 (rysunek 8). Wykorzystywał on taśmę dwucalową (5,08 cm) w formacie Quadraplex. Kosztował wówczas 50 000 dolarów, co dziś stanowi równowartość około 840 000 dolarów. Główną innowacją zastosowaną w tej maszynie było nagrywanie poprzeczne. Obraz wideo był zapisywany w poprzek taśmy, a nie wzdłuż, co pozwoliło użyć rozsądnej prędkości taśmy – 38 cm na sekundę. Przed wprowadzeniem magnetowidu jedynym nośnikiem umożliwiającym nagrywanie programów telewizyjnych była taśma filmowa. Patrz szczegółowy artykuł na temat nagrywania Quadraplex: Silicon Chip, marzec 2021; www.siliconchip.au/Article/14782.

Położono pierwszy transatlantyczny kabel telefoniczny – TAT-1 (Transatlantic No. 1). Mógł on przenosić 35 połączeń telefonicznych jednocześnie. Kanałem 36. transmitowano 22 połączenia telegraficzne.

Pocisk R-7, satelita Sputnik 1 i inne, 1957

Skonstruowano syntezytor RCA „Mark II” (rysunek 9), następcę modelu „Mark I”. Był znacznie łatwiejszy

do programowania. Zawierał dwa czytelniki taśmy dziurkowanej. Na taśmach były zakodowane kompozycje muzyczne. Przebiegi wynikowe działania maszyny były nagrywane na pokrytych lakierem płytach gramofonowych. Patrz https://youtu.be/rGN_VzEIZ1I oraz www.siliconchip.au/link/abnz.

W Shippingport w Pensylwanii (USA) rozpoczęła działalność na dużą skalę pierwsza na świecie cywilna elektrownia jądrowa. Jej głównym celem było wytwarzanie energii elektrycznej, ale stanowiła ona również potwierdzenie koncepcji reaktora powielającego, który w czasie, gdy wytwarza energię, przekształca również paliwo nierozszczepialne w rozszczepialne. Reaktor działał do 1982 roku.

Wprowadzono pierwszy międzykontynentalny pocisk balistyczny – radziecki R-7 „Semiorka”. Pociski takie były później przez wiele krajów wykorzystywane ponownie jako rakiety do wystrzeliwania satelitów i innych misji kosmicznych.

Ukazała się komercyjna wersja języka programowania FORTRAN. Jest on nadal używany przez matematyków, naukowców i inżynierów.

Związek Radziecki wystrzelił raketą R-7 „Semiorka” pierwszego sztucznego satelitę Ziemi „Sputnik 1”. Patrz szczegółowy artykuł na temat nadajników radiowych Sputnika w Silicon Chip (www.siliconchip.au/Series/407) oraz w EdW 7/2025 i 8/2025.

Magnetowid kolorowy, modemy, rozrusznik serca, 1958

Wystrzelono pierwszego amerykańskiego satelitę Ziemi „Explorer 1”.

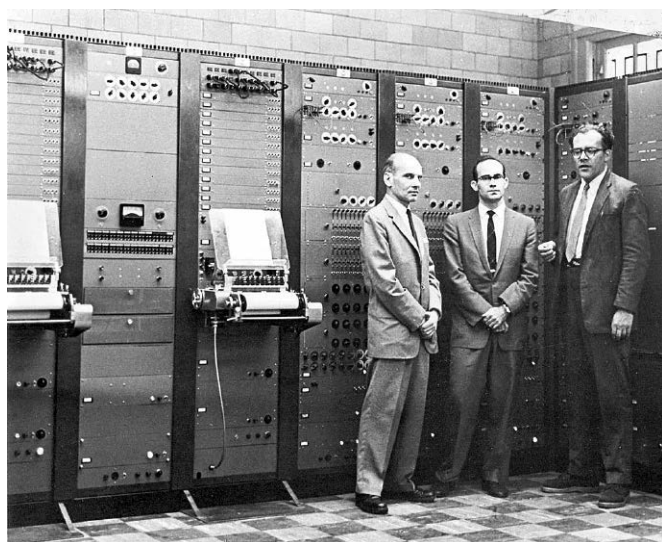
Na rynku pojawił się magnetowid kolorowy Ampex VR-1000B. Obsługiwał on wiele międzynarodowych standardów wideo. Broszurę produktu można zobaczyć na stronie www.siliconchip.au/link/abp0.



Rysunek 7. Dwa napędy dyskowe IBM 350 w budynku Red River Arsenal należącym do armii Stanów Zjednoczonych. Źródło: <https://www.wiki/7EFy>



Rysunek 8. Inżynier John Radis z firmy CBS obsługujący magnetowid Ampex VRX-1000, 30 listopada 1956 roku. Było to pierwsze użycie tego urządzenia podczas programu telewizyjnego. Źródło: www.quadvideotapegroup.com/2015/12/



Rysunek 9. Syntezytor muzyczny RCA „Mark II”. Uwagę zwracają czytelniki dziurkowanej taśmy papierowej. Źródło: https://electronicmusic.fandom.com/wiki/RCA_Synthesizer (CC-BY-SA)

W 1958 roku masowo produkowano dla armii amerykańskiej modemy (modulatory/demodulatory) telefoniczne. Typ „Bell 101” (rysunek 10) był używany w systemie obrony powietrznej SAGE. Technologia modemów (o prędkości 110 bit/s) została udostępniona publicznie w 1959 roku. Modemy były rozwinięciem konstrukcji multiplekserów do dalekopisów używanych m.in. przez serwisy informacyjne w latach 20. XX wieku.

Pojawił się pierwszy wszczepialny rozrusznik serca.

Uruchomiono pierwszy australijski reaktor jądrowy „HIFAR” („High Flux Australian Reactor”, „australijski reaktor wielkoprzepływowy”), przeznaczony do badań i produkcji radioizotopów. Działał on do 2007 roku.

William Higinbotham (1910–1994) w nowojorskim Brookhaven National Laboratory stworzył drugą w historii grę komputerową.



Rysunek 10. Modem Bell 101, wyprodukowany przez AT&T w 1958 r. Źródło: <https://history-computer.com/modem-complete-history-of-the-modem/>

Nazywała się „Tennis for Two” i była odmianą gry „Pong”.

Płytki uniwersalna, MOSFET, proces planarny i inne, 1959

Został uruchomiony pierwszy amerykański międzykontynentalny pocisk balistyczny – „SM-65 Atlas”. Był również używany w Programie Mercury do wynoszenia astronautów na orbitę okołozemską.

Wynaleziono wierconą płytkę uniwersalną, przeznaczoną do prototypowania układów elektronicznych.

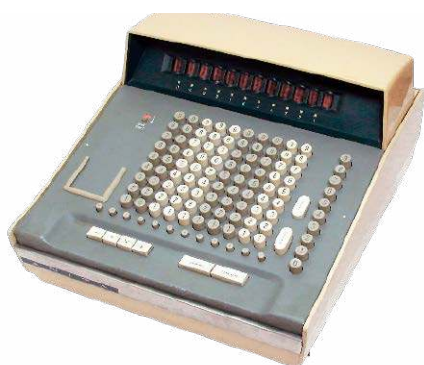
Wprowadzono pierwszą komercyjną kserokopiarę na zwykły papier – „Xerox 914”. Patrz film na stronie <https://youtu.be/9xZYcWsh8t0>.

Mohamed Atalla i Dawon Kahng z Bell Laboratories wynaleźli tranzystor MOSFET (tranzystor polowy z izolowaną bramką; metal-tlenek-półprzewodnik). Patrz artykuł o tranzystorach w Silicon Chip z maja 2022 r. (www.siliconchip.au/Article/15305).

Jean Amédée Hoerni (1924–1997) wynalazł proces planarny wytwarzania półprzewodnikowych układów scalonych.



Rysunek 11. Satelita ECHO 2 w trakcie testowania i kontroli, górujący nad otaczającymi go ludźmi. Pierwsza transmisja poprzez tego satelitę odbyła się z Kalifornii do New Jersey w 1960 roku. Źródło: NASA



Rysunek 12. Kalkulatory Anita MkVII i VIII (na zdjęciu) zostały wprowadzone na rynek jednocześnie w 1961 roku. Pierwszy elektroniczny model nosił numer VII, ponieważ wcześniejsze numery nadano kalkulatorom mechanicznym.
Źródło: <https://w.wiki/7EFz> (GNU FDL)

Więcej informacji na temat historii układów scalonych można znaleźć w artykułach w Silicon Chip z czerwca-sierpnia 2022 r. (www.siliconchip.au/Series/382).

Projekt NASA Echo, podzespoły SMD, 1960

NASA zainicjowała satelitarny projekt Echo. „Echo 1” i „Echo 2” były eksperymentalnymi biernymi satelitami komunikacyjnymi (**rysunek 11**) w formie nadmuchiwanych balonów o średnicy 30 m. Dostarczały one cennych danych na temat oporu aerodynamicznego oraz innych informacji. **Przypis redaktora:** Echo 1 został wysłany w 1960 roku, natomiast Echo 2 w 1964 roku. Były to bierne satelity (odbłyśniki fal radiowych), a nie aktywne przekaźniki sygnału. Źródło: NASA History Division – Project Echo.

IBM po raz pierwszy zademonstrował technologię montażu powierzchniowego (SMD), używając jej do budowy niewielkiego komputera. Technologia SMD została później wykorzystana w latach 60. w komputerach pojazdów kosmicznych w programach Saturn IB i Saturn V.

Kalkulatory elektroniczne ANITA, diody LED, 1961

Pierwszymi elektronicznymi kalkulatorami były ANITA Mark VII i Mark VIII (**rysunek 12**), skonstruowane w 1961 roku. Wykorzystywały one lampy elektronowe i lampy z zimną katodą. Pierwszym kalkulatorem elektronicznym na półprzewodnikach był Friden EC-130 z 1963 roku.

W laboratorium SERL w Baldock (Wielka Brytania) J. W. Allen i R. J. Cherry opracowali wczesne eksperymentalne diody LED emitujące światło widzialne. **Przypis redaktora:** Za pierwszą praktyczną diodę LED emitującą światło widzialne powszechnie uznaje się konstrukcję Nicka Holonyaka Jr. z 1962 roku. Wcześniejsze prace miały charakter eksperymentalny. Źródła: IEEE Spectrum,

„The Transistor Laser Turns 50” (2012); The Nobel Prize, „The Nobel Prize in Physics 2014 – Advanced Information”.

Złącze Josephsona, Telstar 1 i inne, 1962

Steve Russell (1937~) wynalazł trzecią grę komputerową. Była ona rozgrywana na komputerze mainframe Digital PDP-1. Nazywała się „Spacewar!” (**rysunek 13**). Patrz wideo zatytułowane „Spacewar! (1961) – pierwsza cyfrowa gra komputerowa”: <https://youtu.be/CwZAKJ8Y6YU>.

Zaobserwowany został „efekt Josephsona”, choć z początku nie rozpoznano go prawidłowo. Wykrycie tego efektu doprowadziło do powstania obwodu nadprzewodzącego zwanego „złączem Josephsona”, który znalazł zastosowanie m.in. w komputerach kwantowych, standardach napięcia i w cyfrowym przetwarzaniu sygnałów. Został on nazwany na cześć Briana Davida Josephsona (1940~).

Powstał system rozpoznawania mowy IBM „Shoebbox”. Mógł rozpoznać 16 wypowiedzianych słów (cyfry i operatory arytmetyczne). Był częścią obsługiwanego głosowo kalkulatora drukującego (**rysunek 14**).

Na orbitę eliptyczną (nie geostacjonarną) został wysłany satelita komunikacyjny Telstar 1. W 1963 roku został wysłany Telstar 2. Oba satelity były eksperymentalne. Żaden z nich nie jest już używany, choć oba wciąż znajdują się na orbicie. W tym samym 1962 roku Telstar 1 przeprowadził pierwszą transatlantycką satelitarną transmisję telewizyjną. Za jego pośrednictwem były również przesyłane dane między dwoma komputerami IBM 1401.

Kaseta kompaktowa Philips, kod ASCII i inne, 1963

Został ukończony podmorski kabel telefoniczny COMPAC („Commonwealth Pacific Cable System”; „kabel Wspólnoty Narodów pod Pacyfikiem”), łączący Australię, Nową Zelandię oraz Kanadę poprzez Hawaje



Rysunek 14. System rozpoznawania głosu/kalkulator IBM „Shoebbox”. Źródło: IBM

i wyspę Fidżi. Jego elementy działały od 1961 roku. Koncentryczny kabel mógł obsłużyć 80 połączeń telefonicznych lub 1760 dalekopisowych. Zastąpił on telefoniczne połączenia przez radio, które trzeba było rezerwować, a w przypadku złych warunków transmisji przesuwac na późniejszy termin.

Firma Philips wprowadziła na rynek pierwszą kasety magnetofonową – „Compact Cassette”. Patrz artykuł na ten temat w Silicon Chip, wydanie z lipca 2018 r. (www.siliconchip.au/Article/11136).

Została przeprowadzona pierwsza satelitarna transpacyficzna transmisja telewizyjna. Odbiła się ona między Japonią a USA za pośrednictwem eksperymentalnego satelity komunikacyjnego „Relay 1”, znajdującego się na orbicie eliptycznej.

Firma Nottingham Electric Valve Company w Wielkiej Brytanii wypuściła „Telcan” (**rysunek 15**) – magnetowid do użytku domowego. Wykorzystywał on ćwierćcalową dźwiękową taśmę magnetofonową, przesuwaną z prędkością 305 cm na sekundę, co jak na ten typ taśmy było prędkością bardzo dużą. Na każdej z dwóch ścieżek można było nagrać do 20 minut czarno-białego zapisu. Szerokość pasma nagrywania wynosiła 2,6 MHz, co zapewniało rozdzielczość 405 linii. Magnetowid był sprzedawany głównie jako zestaw do samodzielnego montażu za 60 funtów, co dziś stanowi równowartość około 2000 dolarów. Stanowił komercyjną porażkę. Więcej szczegółów można znaleźć na stronie www.siliconchip.au/link/abp1.

Opublikowano pierwszą wersję standardu kodowania znaków ASCII.

Kabel TPC-1, system faksowy Xerox, język BASIC i inne, 1964

Uruchomiono kabel komunikacyjny „Trans-Pacific TPC-1”, łączący Japonię, Guam, Hawaje i kontynentalne Stany Zjednoczone. Realizował on 128 jednoczesnych połączeń telefonicznych.



Rysunek 13. Gra „Spacewar!” uruchomiona na komputerze PDP-1. Źródło: <https://w.wiki/7EF5> (CC-BY-2.0)



Rysunek 15. Magnetowid domowy Telcan, wykorzystujący ćwierćcalową dźwiękową taśmę magnetofonową, sprzedawany głównie jako zestaw do montażu. Źródło: www.nottinghampost.com/news/history/20-best-things-nottingham-given-192680



Rysunek 16. Sony CV-2000, pierwszy masowo produkowany magnetowid przeznaczony na rynek krajowy. Wykorzystywał półcalową taśmę na szpuli. Źródło: www.smecc.org/sony_cv_series_video.htm

Firma Xerox wprowadziła pierwszy nowoczesny system faksowy, który nazwała LDX – „Long Distance Xerography” („kserografia na odległość”).

Opublikowano język programowania komputerów BASIC.

Robert Moog (1934–2005) zbudował pierwszy prototyp syntezatora muzycznego „Moog”. Modele dostępne handlowo zaczęły być produkowane w 1967 roku.

Satelita geosynchroniczny Intelsat I inne, 1965

Luigi Dadda (1923–2012) wynalazł sprzętowy dwójkowy układ mnożący, pomyślany jako element komputerowych jednostek arytmetycznych. Był on mniejszy i szybszy niż poprzednia implementacja układu mnożącego.

Firma Sony wypuściła model „CV-2000” (CV = „consumer video”) – pierwszy masowo produkowany domowy magnetowid (rysunek 16). Nagrywał on obraz czarno-biały i używał taśmy szpulowej 13 mm. Znalazł szerokie zastosowanie w instytucjach finansowych i edukacyjnych. Brak regulacji śledzenia głowicy sprawiał, że nie można było wymieniać taśm między różnymi egzemplarzami magnetowidu. Zostało to poprawione w późniejszych wersjach urządzenia.

Wysłano pierwszego dostępnego w sprzedaży satelitę geosynchronicznego – „Intelsat I”. Przekazywał on 240 rozmów telefonicznych albo jeden program telewizyjny. Był używany przez ponad cztery lata, dopóki nie został wyłączony. Reaktywowano go tymczasowo na potrzeby misji Apollo 11 i kolejny raz w 1990 roku, z okazji 25. rocznicy. Nadal znajduje się na orbicie.

Telekopiarka Magnafax (faks), 1966

Firma Xerox wprowadziła na rynek pierwszy łatwy w obsłudze faks – „Magnafax Telecopier”, wykorzystujący standardową linię telefoniczną.

Standard PAL, bankomat ATM, satelita WRESAT i inne, 1967

We Francji został opracowany i zatwierdzony standard telewizji kolorowej SECAM. Przypis redaktora: był to pierwszy w Europie standard telewizji kolorowej, wkrótce przyjęty również w większości krajów „bloku wschodniego”, m.in. w Polsce.

W Wielkiej Brytanii rozpoczęto nadawanie telewizji kolorowej w standardzie PAL.

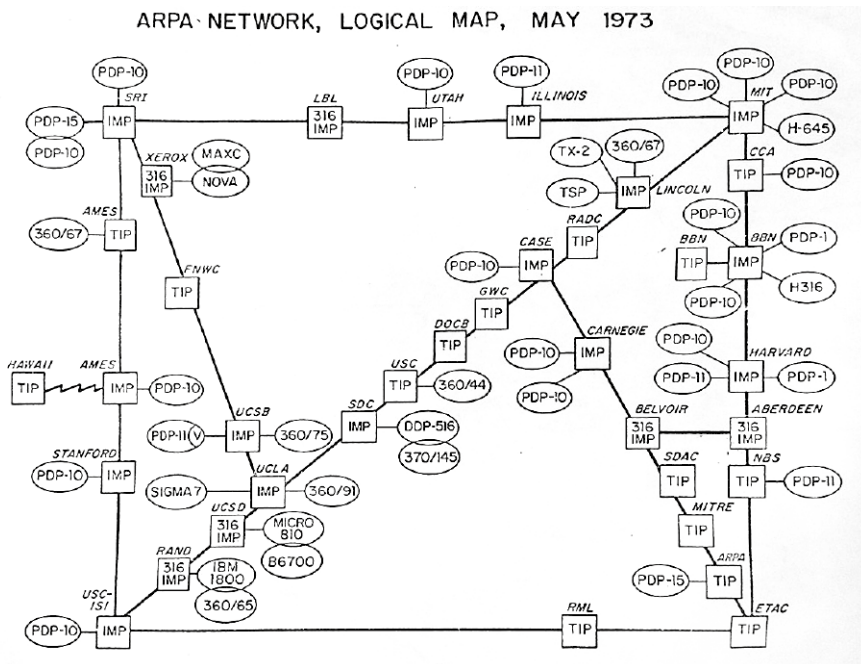
W banku Barclay’s w Enfield (Wielka Brytania) został zainstalowany pierwszy na świecie bankomat. Jego wykorzystanie polegało na wprowadzaniu czeków uprzednio wystawionych przez kasjera, oznaczonych radioaktywnym węglem ¹⁴C, co umożliwiało ich odczyt przez aparat i potwierdzenie ich autentyczności.

Został wysłany pierwszy lokalny satelita Australii – „WRESAT”. Patrz artykuł na jego temat w Silicon Chip z października 2017 r (www.siliconchip.au/Article/10822).

Wyświetlacze LCD, standard dla faksów, 1968

Zespół naukowców z RCA Laboratories zademonstrował wyświetlacz ciekłokrystaliczny (LCD) – matrycę 18x2 – wykorzystujący tryb dynamicznego rozpraszania (DSM) wynaleziony przez George’a Heilemeiera (zobacz wpis w zeszłym odcinku artykułu).

Międzynarodowy Związek Telekomunikacyjny opublikował standard faksowy Grupy 1. Zgodne z tym standardem urządzenia potrzebowały około sześciu minut na przesłanie jednej strony z prędkością 38 linii na centymetr.



Rysunek 17. Mapa sieci ARPANET – poprzednika Internetu – w stanie z 1973 roku Źródło: <https://w.wiki/7FPK>

Pamięć dynamiczna MOS, system Unix, ARPANET i inne, 1969

Firma Advanced Memory Systems Inc. rozpoczęła komercyjną produkcję pamięci dynamicznej (MOS DRAM – „metal oxide semiconductor dynamic random access memory” czyli „pamięć dynamiczna o dostępie swobodnym w technologii MOS”). Pamięci te, o pojemności 1024 bity, były sprzedawane wybranym odbiorcom. W tym samym roku Intel wyprodukował układ pamięci 1103, również o pojemności 1024 bity, i sprzedawał go na otwartym rynku. Układy te były używane w popularnych komputerach, takich jak HP 9800 i PDP-11. Więcej informacji na temat rozwoju pamięci DRAM można znaleźć w artykułach na temat pamięci komputerowych w wydaniach Silicon Chip ze stycznia i lutego 2023 r. (www.siliconchip.au/Series/393) oraz w EdW 9/2025.

Wprowadzono komputerowy system operacyjny Unix.

Ukazał się pierwszy komercyjny zegarek kwarcowy – Seiko „Quartz-Astron 35SQ”. Miał on dokładność ± 5 sekund na miesiąc. Żywotność baterii wynosiła około jednego roku.

Agencja Zaawansowanych Projektów Badawczych (ARPA) Departamentu Obrony Stanów Zjednoczonych utworzyła sieć komputerową pracującą z komutacją pakietów – ARPANET (rysunek 17). Sieć ta ostatecznie przekształciła się w dzisiejszy Internet.

Faks cyfrowy, kalkulator kieszonkowy, 1970

Firma Dacom wyprodukowała pierwszy cyfrowy faks – „DFC-10”, który stosował kompresję danych i mógł przesyłać jedną stronę w czasie krótszym niż minuta.

Opublikowano język programowania Pascal.

Pojawił się pierwszy dostępny w handlu kalkulator kieszonkowy – Canon „Pocketronic” (rysunek 18). Konstrukcja ta była zainspirowana prototypowym kalkulatorem



Rysunek 18. Canon Pocketronic – pierwszy dostępny handlowo ręczny kalkulator elektroniczny. Źródło: <https://w.wiki/7EG3> (CC-BY-SA-4.0)



Rysunek 19. Kenbak-1 – pierwszy komputer osobisty z 1971 roku. Źródło: <https://w.wiki/7FPV> (CC-BY-SA-4.0)

Texas Instruments „Cal Tech” z 1967 roku. „Pocketronic” zawierał trzy układy scalone MOS produkcji Texas Instruments. Nie miał wyświetlacza, a wyniki były drukowane na taśmie papierowej. Więcej informacji można znaleźć na stronie www.siliconchip.au/link/abp4.

Mikroprocesor Intel 4004, komputer Kenbak-1, pamięć EPROM i inne, 1971

Pojawił się pierwszy dostępny handlowo mikroprocesor – Intel 4004. Był przeznaczony głównie do kalkulatorów i zliczarek banknotów. **Przypis redaktora: użyto go w kalkulatorze stołowym Busicom 141-PF.**

Departament Obrony USA sfinansował pięcioletni program mający na celu stworzenie maszyny rozpoznającej mowę, która byłaby w stanie rozpoznawać w zdaniach słowa ze zbioru tysiąca słów. Zbudowano maszynę o nazwie Harpy, która rozpoznawała 1011 słów. Patrz plik PDF: www.siliconchip.au/link/abp5.

Firma Docutel wprowadziła maszynę „Total Teller” – bankomat, który mógł przyjmować depozyty, przelewać pieniądze z jednego konta na drugie i wypłacać gotówkę. Używał plastikowych kart, działał w trybie offline i miał mechaniczny wyświetlacz prezentujący komunikaty na specjalnym cylindrze. Do 1975 roku na całym świecie zainstalowano 3000 bankomatów, z czego 80% pochodziło od Docutela. W 1982 roku Docutel połączył się z firmą Olivetti.

Pojawił się pierwszy komputer osobisty (bez mikroprocesora) – „Kenbak-1” (rysunek 19). Sprzedano tylko 40 sztuk.

Intel wypuścił pierwszą pamięć EPROM („Erasable Programmable Read Only Memory”; „kasowalna programowalna pamięć stała”). Wynalazł ją Dov Frohman. Kostka Intel 1702 mogła przechować 256 bajtów danych.

Firma Sony wprowadziła na rynek kasety wideo w formacie „U-matic”. Był to pierwszy komercyjny format dla kaset

wideo. Wykorzystywał taśmę 3/4 cala (19 mm). Początkowo był przeznaczony na rynek konsumencki, okazał się jednak zbyt drogi. Upowszechnił się w instytucjach i przemyśle, a także w branży telewizyjnej. Patrz cykl na temat kaset wideo (Silicon Chip, marzec-czerwiec 2021; www.siliconchip.au/Series/359).

Mikroprocesor Intel 8008, język C, gra „Pong”, niebieskie diody LED i inne, 1972

Na rynek został wprowadzony magnetowid kasetowy Philips N1500. Ostatnie egzemplarze wyprodukowano w 1979 roku.

W sprzedaży pojawiły się ośmiocalowe dyskiety komputerowe.

Opublikowano standard faksowy Grupy 2 Międzynarodowego Związku Telekomunikacyjnego. Zgodnie z tym standardem urządzenia potrzebowały około trzech minut na przesłanie strony z rozdzielczością 38 linii na cm.

Wprowadzono Cartrivision – konsumencki format kaset wideo. Magnetowidy w tym standardzie były wbudowywane w drogie modele telewizorów. System okazał się komercyjną porażką. Zobacz www.angelfire.com/alt/cartrivision/.

Kod systemu operacyjnego Unix został przepisany na język programowania C. Rok 1972 można zatem uznać za datę, w której C stał się „językiem głównego nurtu”. Język C został opracowany głównie w latach 1969–1973. Jest nadal szeroko stosowany – w swojej oryginalnej formie i w odmianach, takich jak C++ i C#.

Pojawił się pierwszy mikroprocesor do komputerów osobistych – 8-bitowy Intel 8008.

Skonstruowano pierwszy na świecie kieszonkowy kalkulator naukowy – HP-35.

Ukazała się gra „Pong” – pierwsza gra komputerowa, która odniosła komercyjny sukces. W wydaniu Silicon Chip z czerwca 2021 r. opublikowaliśmy projekt odtworzenia tej gry w oryginalnej postaci, a w sierpniu 2021 r. jej

zmodernizowaną, zminiaturyzowaną wersję. Więcej informacji można znaleźć na stronie www.pong-story.com.

Herbert Paul Maruska (1944~) z firmy RCA wynalazł pierwszą niebieską diodę LED. Firma była jednak wówczas pogrążona w chaosie i cały projekt został anulowany. Ponadto przyrząd był zbyt słaby, aby znaleźć praktyczne zastosowanie. Patrz www.siliconchip.au/link/abp6. Ostatecznie w 2014 roku Nagrodę Nobla za wynalezienie w 1993 roku niebieskich diod LED o wysokiej jasności otrzymali Isamu Akasaki (1929-2021), Hiroshi Amano (1960~) i Shūji Nakamura (1954~) z Uniwersytetu Nagoja w Japonii. Współczesne białe diody LED są zazwyczaj realizowane jako niebieskie diody LED pokryte luminoforem.

Program SPICE, Ethernet, graficzny interfejs użytkownika, 1973

Firma Micral wypuściła pierwszy komputer osobisty. Opierał się na mikroprocesorze Intel 8008.

Wprowadzono program komputerowy do symulacji analogowych obwodów elektronicznych: SPICE („Simulation Program with Integrated Circuit Emphasis” czyli „program symulacyjny z naciskiem na symulację układów scalonych”). Program ten i jego pochodne (takie jak LTspice) jest nadal szeroko stosowany.

Robert Melancton Metcalfe (1946~) wraz ze swoim zespołem w Xerox Palo Alto Research Center (Xerox PARC) w Kalifornii wynalazł sieć Ethernet. Jest ona jednym z kluczowych elementów, na których opiera się funkcjonowanie Internetu.

Motorola zaprezentowała telefon komórkowy. Upowszechnienie tego wynalazku zajęło jednak trochę czasu.

W Bell Labs zademonstrowano pierwszy strojony laser (o długości fali dającej się zmniejszać w pewnych granicach).

Na rynku pojawił się komputer Xerox Alto (**rysunek 20**) – pierwszy komputer z graficznym interfejsem użytkownika i myszą (**rysunek 21**) – na dziesięć lat przed ukazaniem się modelu Apple Lisa. Kosztował 32 000 dolarów, czyli równoważność dzisiejszych 330 000 dolarów. Jak widać na zdjęciach, miał on monitor zorientowany pionowo.

Komputer EDUC-8, system operacyjny CP/M i inne, 1974

Firma Electronics Australia opublikowała urządzenie, które w tamtym czasie uważano za pierwszy na świecie komputerowy zestaw do składania – EDUC-8 (**rysunek 22**). Później okazało się jednak, że wyprzedził go o miesiąc inny model – „Mark-8”. Projekt firmy EA został jednak uznany za lepszy.

Powstał pierwszy program typu WYSIWYG, („What You See Is What You Get” – „dostajesz to, co widzisz”; program pozwalający uzyskać wydruki wyglądające niemal identycznie jak obraz na ekranie monitora; przypis redaktora). Służył do sporządzania dokumentów, był więc wczesnym edytorem tekstu. Działał na komputerze Xerox Alto.

Wprowadzono komputerowy system operacyjny CP/M, wyparty później przez MS-DOS.

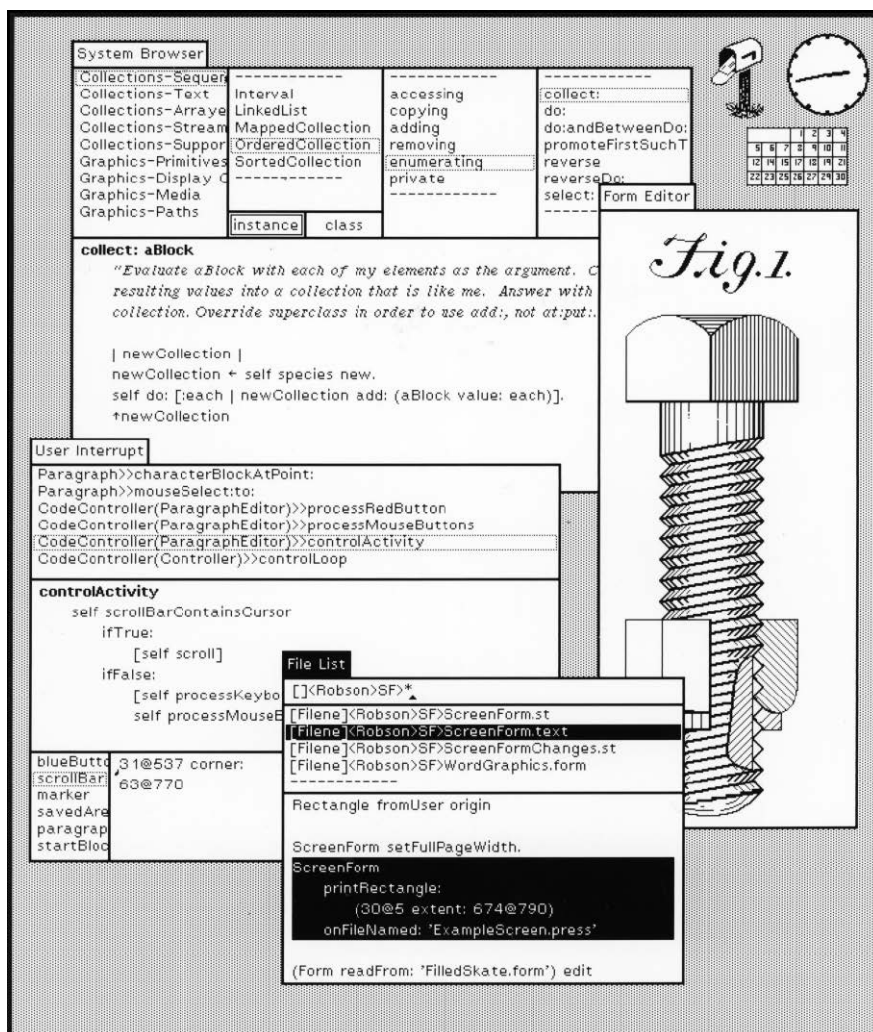
Cyfrowy aparat fotograficzny Kodak, system Betamax i inne, 1975

Został wynaleziony pierwszy samodzielny cyfrowy aparat fotograficzny. Jego twórcą był Steven Sasson (1950~) z firmy Kodak. Aparat miał rozdzielczość 100×100 pikseli, a obrazy były zapisywane cyfrowo na taśmie magnetofonowej. Zapis trwał 23 sekundy.

Pojawił się komputer osobisty Altair 8800 w formie zestawu do składania. Wydarzenie to jest przez wielu uważane za początek rewolucji mikrokomputerowej.



Rysunek 20. Komputer Xerox Alto z 1973 r. Źródło: <https://www.wiki/7EG4>



Rysunek 21. Graficzny interfejs użytkownika komputera Xerox Alto. Źródło: <https://interface-experience.org/objects/xerox-alto/>

Stworzono system magnetowidowy Betamax, przeznaczony do użytku domowego (szczególnie zawiera seria o nagrywaniu wideo opublikowana w Silicon Chip).

Garrett Brown wynalazł system „Steadicam”, używany do stabilizacji obrazu z kamery wideo poprzez odizolowanie ruchu operatora. System wyprodukowała firma Cinema Products Corporation. Patrz artykuł na ten temat w Silicon Chip w wydaniach z listopada i grudnia 2011 r.; www.siliconchip.au/Series/33.

System VHS, dyskietka 5,25 cala i inne, 1976

Opublikowano pierwszy edytor tekstu dla komputerów domowych, o nazwie „Electric Pencil” („ołówek elektryczny”), do użytku na komputerach takich jak Altair 8800, Sol-20, a później TRS-80 i IBM PC.

Powstał system magnetowidowy VHS, przeznaczony do użytku domowego.

Pojawiły się dyskietki komputerowe 5,25 cala.

Komputery Apple II, Commodore PET, TRS-80 i inne, 1977

W Turynie we Włoszech zostało zainstalowane pierwsze użytkowe łącze światłowodowe.

Pojawiły się komputery domowe: Apple II, Commodore PET i TRS-80. Były to modele bardzo znaczące w historii komputerów domowych.

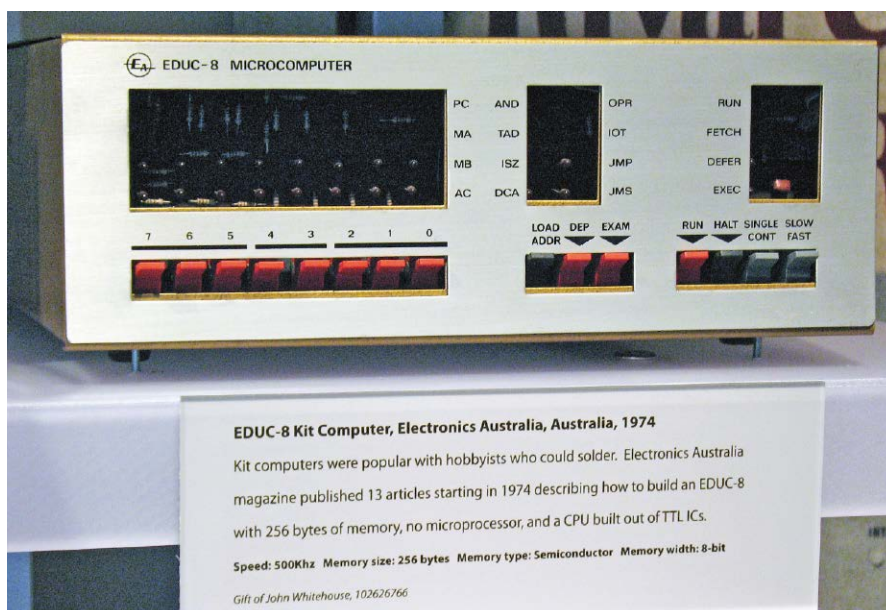
Synteza mowy, „LaserDisc” i inne, 1978

Firma Texas Instruments wypuściła pierwszy układ scalony syntezy mowy – TMS5100. Wykorzystywał on „predykcyjne kodowanie liniowe pobudzone wysokością dźwięku”, co znacznie zmniejszyło ilość informacji wymaganych do generowania mowy. Układ został wykorzystany w zabawce edukacyjnej „Speak & Spell”.

Na rynku pojawił się system płyt wizyjnych „LaserDisc”. W systemie tym urządzenia domowe odtwarzały filmy, ale nie mogły ich nagrywać. Technologia „LaserDisc” została później wykorzystana w standardach płyt kompaktowych CD, DVD i Blu-ray. System „LaserDisc” nigdy nie był zbyt popularny, jednak w tamtym okresie oferował dobrą jakość odtwarzania wideo, znacznie przewyższając pod tym względem standard VHS.

Fairlight CMI, sieci telefoniczne 1G i inne, 1979

W Australii opracowano cyfrowy syntezytor/sampler Fairlight CMI (Computer Musical Instrument). Opierał się na projekcie Tony’ego Furse’a, licencjonowanym przez Kima Ryrie i Petera Vogela związanych z australijskim czasopiśmie „Electronics Today International”. Była to jedna z pierwszych muzycznych „stacji roboczych” z samplingiem.



Rysunek 22. Komputer Electronics Australia EDUC-8. Źródło: <https://w.wiki/7EG6>

System uważano wówczas za rewolucyjny. Patrz film „Jak Fairlight CMI zmienił bieg muzyki” na stronie <https://youtu.be/jkiYy0i8FtA>.

Wydano bardzo popularny edytor tekstu „WordPerfect”.

Japońska firma Nippon Telegraph and Telephone (NTT) wdrożyła pierwszą sieć telefonii komórkowej 1G.

Philips i Grundig wypuścili konsumencki format kaset wideo „Video 2000”. Został on wycofany w 1988 roku.

Ukazał się program arkusza kalkulacyjnego „VisiCalc”. Był przeznaczony dla komputera Apple II. Okazał się aplikacją „zwalającą z nóg” i zaimplementowano go również na wielu innych komputerach. Był poprzednikiem programów takich jak Excel. Więcej informacji można znaleźć na stronie: <http://danbricklin.com/visicalc.htm>.

Komputer Commodore VIC-20 i inne, 1980

International Telecommunications Union wydała cyfrowy standard faksowy Grupa 3. Czas przesyłania strony został skrócony do 6-15 sekund (nie licząc nawiązywania połączenia). Standard obsługiwał zmienną rozdzielczość skanowania, aż do 157 linii na centymetr.

Miała miejsce premiera komputera domowego Commodore VIC-20.

System operacyjny MS-DOS 1.0, 16-bitowy przetwornik C/A i inne, 1981

Wraz z komputerem IBM PC ukazał się system operacyjny MS-DOS w wersji 1.0.

Pojawił się „Osborne 1”, uważany za pierwszy dostępny w handlu rzeczywiście przenośny komputer. Nie jest jednak jasne, czy może być uznany za pierwszy „laptop”. Do tego tytułu pretenduje wiele modeli.



Rysunek 23. Przenośny telefon Motorola DynaTAC 8000X. Źródło: <https://w.wiki/7FPn> (CC-BY-SA-3.0)

Na rynku pojawił się 16-bitowy jedno-układowy przetwornik cyfrowo-analogowy (DAC) PCM53/DAC700. Zaprojektowany przez Jimmy'ego Naylora i zespół projektowy Texas Instruments/Burr-Brown, stał się podstawowym elementem prawie wszystkich odtwarzaczy cyfrowych płyt kompaktowych (CD).

RCA wypuściła „Capacitance Electronic Disc” (CED) – analogowy system odtwarzania płyt wideo. Płyty były odczytywane za pomocą mechanicznego rysika, miały średnicę 30 cm i mogły na każdej ze stron przechować 60 minut zapisu w systemie NTSC. Produkt nie cieszył się popularnością i został wycofany, m.in. ze względu na konkurencję ze strony odtwarzaczy „LaserDisc”.

Odtwarzacz CD, Commodore 64, 1982

W Japonii ukazał się pierwszy odtwarzacz płyt Compact Disc (CD), konstrukcja opracowana wspólnie przez firmy Philips i Sony.

Na rynek wprowadzono komputer domowy Commodore 64.

Dyskietka 3,5 cala, język C++ i inne, 1983

W oparciu o specyfikację organizacji Microfloppy Industry Committee ukazały się pierwsze 3,5-calowe dyskietki komputerowe.

Pojawił się język programowania C++, obiektowa wersja języka C, dziś powszechnie używana.

Na rynek wypuszczono pierwszy komputer osobisty z wbudowanym dyskiem stałym: IBM PC XT. Dysk miał standardowo pojemność 10 MB.

Motorola wypuściła pierwszy przenośny telefon – „DynaTAC 8000X” (**rysunek 23**). Ważył prawie 1 kilogram. Czas ładowania akumulatora wynosił 10 godzin. Model był sprzedawany za 3995 dolarów (około 18750 dolarów na dzisiejsze pieniądze).

Dr Mitsuaki Oshima z firmy Panasonic wynalazł elektroniczną stabilizację obrazu. Pierwszą kamerę wideo wyposażoną w elektroniczną stabilizację obrazu Panasonic wypuścił sporo później – w 1988 roku.

Komputer Apple Macintosh, pamięć CD-ROM, 1984

Ukazał się komputer osobisty Apple Macintosh. W tym samym roku na rynek wprowadzono również komputer domowy Commodore Amiga.

Zapowiedziano ukazanie się płyty CD-ROM, służącej do przechowywania danych, a opartej na standardzie płyty Compact Disc (CD).

System rozpoznawania mowy IBM Tangora, 1985

Wprowadzono eksperymentalny system rozpoznawania mowy IBM „Tangora”. Działał na komputerze IBM PC AT. Rozpoznawał 20 000 słów i zamieniał je na tekst.

Format nagrywania wideo „Sony D-1”, 1986

Wprowadzono profesjonalny studyjny format cyfrowego zapisu wizji „Sony D-1”.

Nadprzewodniki wysokotemperaturowe, 1987

Odkryto nadprzewodniki działające w „wysokich” temperaturach. Obecnie istniejące tego typu nadprzewodniki działają w temperaturze około -135°C (przy normalnym ciśnieniu atmosferycznym).

Transatlantycki kabel światłowodowy „TAT-8”, 1988

Uruchomiono pierwszy transatlantycki kabel światłowodowy – „TAT-8”. Jego maksymalna przepustowość wynosiła 280 Mbit/s, co odpowiadało 40 000 jednoczesnym rozmowom telefonicznym. W 2002 roku został on wycofany z użycia. Jego maksymalna przepustowość została osiągnięta już po krótkim okresie eksploatacji, choć początkowo sądzono, że nie zostanie osiągnięta nigdy. Był naruszany przez rekiny, prawdopodobnie dlatego, że mogły one wyczuwać jego promieniowanie elektromagnetyczne. Kabel ten odegrał kluczową rolę w rozwoju Internetu, zapewniając szybkie połączenie między ośrodkiem CERN w Europie a uczelnią Cornell University w USA.

Odbiornik GPS, sieć WWW i inne, 1989

Pojawił się pierwszy ręczny odbiornik GPS o przeznaczeniu komercyjnym – „Magellan NAV 1000”.

Zademonstrowano metodę dostępu CDMA („Code Division Multiple Access”; „wielokrotny dostęp z dzieleniem kodów”), przeznaczoną dla systemów telefonii komórkowej.

Tim Berners-Lee wynalazł sieć World Wide Web („Sieć Ogólnosiwiatowa”). Została udostępniona publicznie w 1991 roku.

System rozpoznawania mowy „DragonDictate”, 1990

Udostępniono handlowo pierwsze oprogramowanie do rozpoznawania mowy – „DragonDictate”. Obecnie nosi ono nazwę „Dragon NaturallySpeaking” i jest własnością firmy Microsoft.

Sieci 2G, Linux, Python, 1991

Wprowadzono sieci telefonii komórkowej standardu 2G.

Ukazał się komputerowy system operacyjny Linux, będący darmową wersją „open-source” (z dostępnym kodem źródłowym) systemu Unix.

Zaprezentowano język programowania Python.

Kabel TASMAN2, system Windows 3.1 i inne, 1992

Uruchomiono pierwszy podmorski światłowód australijski – TASMAN2. Połączył on Australię z Nową Zelandią. Działał z prędkością około 1 Gb/s.

Ukazał się system operacyjny Windows 3.1 dla komputerów IBM PC. Oznaczał on odejście od tekstowego interfejsu użytkownika z wierszem poleceń, typowego dla systemu DOS, w stronę interfejsu graficznego.

Pojawił się „Newton MessagePad” firmy Apple – wczesny „cyfrowy asystent osobisty” z funkcją rozpoznawania pisma odręcznego. Pomógł on utorować drogę dla późniejszych „inteligentnych” urządzeń.

System Windows NT, program HAARP, 1993

Ukazał się system operacyjny Windows NT. Jego rdzeń stanowi dziś podstawę nowoczesnych wersji systemu Windows, takich jak 10 i 11. Graficzny interfejs użytkownika nadal przypominał jednak Windows 3.1.

Rozpoczęto realizację programu HAARP („High-frequency Active Auroral Research Program” czyli „program badań nad aktywnością zorzy polarnej w wysokich częstotliwościach”), mającego na celu prowadzenie badań nad górną atmosferą i jonosferą. Patrz artykuł na temat HAARP w październikowym wydaniu Silicon Chip z 2012 roku: www.siliconchip.au/Article/492.

Karty pamięci „CompactFlash” i inne, 1994

Firma SanDisk wyprodukowała pierwsze karty pamięciowe „CompactFlash”. Początkowo miały pojemność 2 MB. Był to pierwszy szeroko rozpowszechniony standard kart pamięci FLASH.

Firma Apple wypuściła na rynek komputery domowe i biurowe wykorzystujące 32-bitowe procesory PowerPC firmy IBM, co oznaczało odejście od wcześniej używanych procesorów Motorola.

IBM skierował na rynek „Simon Personal Computer” (SPC), pierwszy „smartfon” – choć termin ten jeszcze wtedy nie istniał. SPC miał ekran dotykowy LCD i mógł być używany do wykonywania i odbierania połączeń telefonicznych, wysyłania i odbierania faksów, e-maili i „wiadomości błyskawicznych”. Sprzedano 50 000 sztuk w cenie po 1099 dolarów (około 3500 dolarów dzisiaj).

System Windows 95, radio DAB, 1995

Ukazał się system operacyjny Windows 95, z graficznym interfejsem użytkownika przypominającym współczesne wersje Windows.

W Europie rozpoczęto nadawanie radia cyfrowego (DAB czyli „Digital Audio Broadcasting”; „cyfrowe radio dźwiękowe”).

Odtwarzacz DVD, „smartfon” PalmPilot, 1996

W Japonii wyprodukowano pierwszy cyfrowy odtwarzacz płyt wideo (DVD).

Opublikowano standard telewizji cyfrowej ATSC.

Pojawił się „PalmPilot”, wczesny poprzednik nowoczesnego smartfona.

Przenośny odtwarzacz MP3 „MPMan F10”, 1997

Wprowadzono standard telewizji cyfrowej DVB-T. Pierwsza transmisja odbyła się w Szwecji.

Pojawił się pierwszy przenośny odtwarzacz MP3 o nazwie „MPMan F10” firmy Saehan Information Systems.

Standard ADSL, 1998

Opracowano standard techniczny ADSL („cyfrowa asymetryczna linia abonencka”) ANSI T1.413 Issue 2. ADSL umożliwia szybkie przesyłanie danych przez standardowe linie telefoniczne z miedzianymi kablami. W Australii standard został wprowadzony w 2000 roku.

Urządzenia Bluetooth, 1999

Na rynek zostało wprowadzone pierwsze urządzenie Bluetooth.

Karty pamięci SD, system Windows 2000, 2000

Pojawiły się pierwsze karty pamięci SD („Secure Digital”), w wersjach o pojemności 32 MB i 64 MB.

Ukazał się system operacyjny Windows 2000, łączący rdzeń Windows NT z interfejsem graficznym Windows 95. Jest on podstawą wszystkich nowszych wersji systemu operacyjnego Windows.

Sieć komórkowa 3G, system Mac OS X, „iPod”, 2001

Wprowadzono komórkowe sieci telefoniczne standardu 3G. Oferowały one szybką mobilną transmisję danych o prędkości do 7,2 Mb/s.

Firma Apple wypuściła Mac OS X – zaprojektowany od nowa graficzny system operacyjny oparty na systemie FreeBSD. Mac OS X jest w produktach Apple do dziś głównym systemem operacyjnym.

Firma Apple wprowadziła na rynek „iPoda” – odtwarzacz MP3.

Transmisje telewizji cyfrowej ISDB-T, 2003
Japonia rozpoczęła nadawanie telewizji cyfrowej w standardzie ISDB-T.

Urządzenie „LongPen” do zdalnego składania podpisów, 2004

Margaret Atwood (1939~) wynalazła „LongPen” – urządzenie do zdalnego składania podpisów, przeznaczone głównie dla autorów w celu podpisywania egzemplarzy książek. Urządzenie zostało wprowadzone na rynek w 2006 roku. Nawiązuje ono do „teleautografu” Elishy Graya z 1888 roku.

Standard DMB, 2005

Korea Południowa przyjęła standard DMB („Digital Multimedia Broadcasting” czyli „cyfrowa transmisja multimedialna”),



Rysunek 24. Relief z egipskiej świątyni w Denderze – czyżby żarówka z I wieku p.n.e.? Zdjęcie wykonał w 2024 roku redaktor artykułu.

przeznaczony dla urządzeń mobilnych w celu strumieniowego przesyłania danych radiowych i telewizyjnych. System był rozwinięciem standardu nadawania radiowego DAB.

Standard DTMB, reaktor „OPAL”, 2006

W Chinach został przyjęty standard telewizyjny DTMB („Digital Terrestrial Multimedia Broadcast”; „naziemna cyfrowa transmisja multimedialna”).

Rozpoczęto realizację projektu „OPAL” – australijskiego lekko-wodnego reaktora basenowego. Miał zastąpić reaktor „HIFAR” przy badaniach i produkcji radioizotopów, np. do celów medycznych i zastosowań przemysłowych.

Apple „iPhone”, 2007

Firma Apple zaprezentowała „iPhone’a” – pierwszy prawdziwie nowoczesny smartfon. **Przypis redaktora:** Określenie to ma charakter umowny. Za pierwszy smartfon uznaje się często IBM Simon (1994), natomiast iPhone spopularyzował współczesny model urządzenia z ekranem dotykowym. Źródła: Encyclopaedia Britannica, hasło „smartphone”; Computer History Museum, „The First Smartphone”; Apple, informacja o premierze iPhone’a z 2007 roku. W 2000 roku pojawił się telefon BlackBerry, ale urządzenia tego typu, z wbudowaną klawiaturą, ostatecznie wyszły z użycia.

W Australii rozpoczęto nadawanie audycji w standardzie DAB+, zapewniającym jakość dźwięku zbliżoną jak w standardzie DAB, ale zajmującym węższe pasmo.

Sieci komórkowe 4G, 2009

Wprowadzono standard telefonicznych sieci komórkowych 4G (LTE).

Apple „iPad”, 2010

Firma Apple wypuściła „iPada” – wczesny tablet z ekranem dotykowym.

Sieci komórkowe 5G, 2019

Wprowadzono standard komórkowych sieci telefonicznych 5G.

Dodatkowe informacje

- „Obrazy fonogramów na papierze, 1250–1950”
– <https://youtu.be/TESkh3hX5oM>
- „Eksperymenty i obserwacje z elektrycznością”, wykonane w Filadelfii w Ameryce przez Benjamina Franklina, 1751 – www.siliconchip.au/link/abpf
- Film przedstawiający prosty projekt konstrukcyjny: „Voltaic Pile, the First Battery” („stos voltaiczny – pierwsza bateria”)
– <https://youtu.be/pW4UUOgJX6k>
- „Electric Incandescent Lighting” („elektryczne oświetlenie żarowe”) autorstwa Edwina Jamesa Houstona i Arthura Edwina Kennelly’ego, 1896
– www.siliconchip.au/link/abpe
- „The Progress of Invention in the Nineteenth Century” („rozwój wynalazków w XIX wieku”) autorstwa Edwarda W. Byrna, 1900 – www.siliconchip.au/link/abpc
- „A History of Wireless Telegraphy” („historia telegrafu bez drutu”) autorstwa J.J. Fahie, wydanie trzecie, 1902
– www.siliconchip.au/link/abpd
- „Jak brzmi nadajnik iskrownika?”
– <https://youtu.be/VMdYte66D2Y>
- Pierwsze woltomierze cyfrowe i rodziny automatycznego testowania
– www.hp9825.com/html/dvms.html
- Najstarsze zachowane nagranie wideo: „The Edsel Show”, telewizja CBS, 13. października 1957 roku
– <https://youtu.be/Ze0Az9tdkHg>
- Najstarsze zachowane kolorowe nagranie wideo: „Ofiarność”, telewizja WRC, 22. maja 1958 roku
– <https://youtu.be/4vBEMGTdDYc>

Procesor Apple Silicon (ARM), 2020

Apple wprowadziło na rynek komputery z własnymi procesorami Apple Silicon z serii M1, wykorzystującymi zuniifikowaną architekturę pamięci, co pozwoliło uzyskać wysoką wydajność przy niskim zużyciu energii. ■

dr David Maddison

Przypis redaktora:

Żarówka z Dendery, I wiek p.n.e.

W krypcie świątyni w egipskim mieście Dendera znajduje się relief przedstawiający zdaniem niektórych... antyczną żarówkę z kablem zasilającym, spoczywającą na wsporniku przypominającym współczesny porcelanowy izolator wysokiego napięcia (rysunek 24). Archeolodzy rzecz jasna tłumaczą znaczenie reliefu na inny sposób. Kto ma rację?

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

TAWOIA Glass (szkło kwarcowe)

<https://sklep.avt.pl/pl/menu/tawoia-glass-4505.html>



BESTSELLERY sklepu AVT – sklep.avt.pl

3 unikalne
serie
gniazdek
i
włączników

Rabat dla Czytelników EdW
przy zakupie podaj kod **EdW2505GW**

Kod ważny do 30.09.2025

-5%

Rabat dla Prenumeratorów EdW
przy zakupie podaj numer prenumeraty

-10%

Ceramic Loft (ceramika)

<https://sklep.avt.pl/pl/menu/seria-ceramic-loft-4190.html>

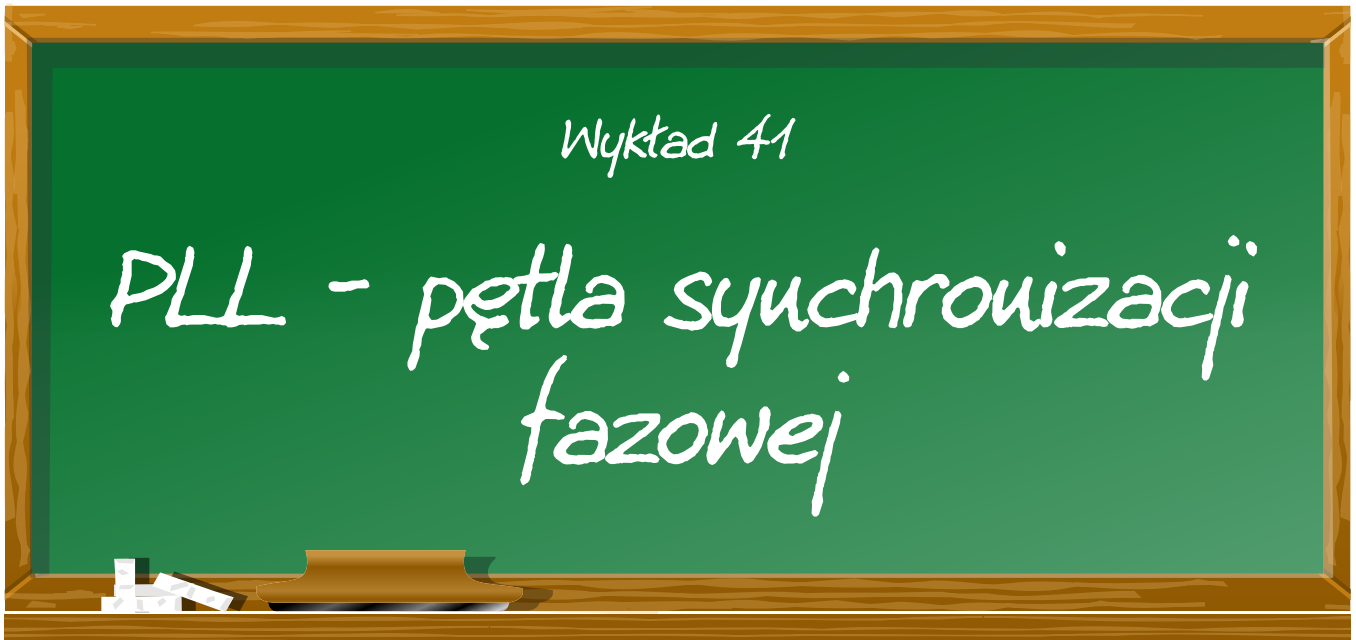


Retro PRL (bakelit)

<https://sklep.avt.pl/pl/series/retro-prl-3237.html>



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 pętli synchronizacji fazowej (PLL) można zsynchronizować fazę lub częstotliwość jednego sygnału z odpowiadającą wielkością drugiego sygnału. W artykule przedstawiamy szczegółowe omówienie tej interesującej techniki analogowej.

Czego można się nauczyć z tego artykułu

Przyznamy – to bardzo obszerny artykuł. Jednak jego zakres wykracza daleko poza samą pętlę PLL. Poznasz kilka bardzo ważnych zagadnień z elektroniki analogowej, takich jak detektory fazy, oscylatory sterowane napięciem (VCO) oraz zasady działania sprzężeń zwrotnych pomiędzy poszczególnymi blokami. Omawiane elementy składowe pętli synchronizacji fazowej spotyka się nie tylko w tego typu układach, ale również jako samodzielne, praktyczne rozwiązania.

Podstawowa zasada działania PLL

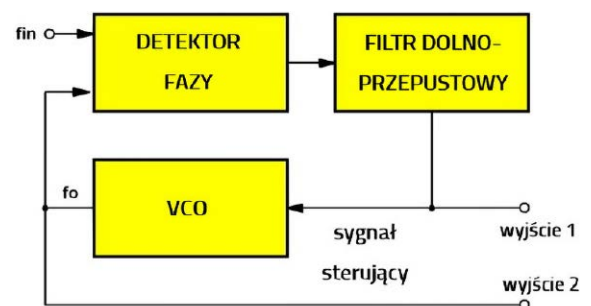
Dosłowne tłumaczenie skrótu PLL (Phase-Locked Loop) to pętla synchronizacji fazowej. Taka pętla umożliwia zsynchronizowanie częstotliwości lub fazy dwóch sygnałów zmiennych. Oznacza to, że układ automatycznie dostosowuje częstotliwość lub fazę drugiego sygnału do odpowiednich parametrów sygnału odniesienia.

Schemat blokowy PLL

Schemat blokowy pętli PLL przedstawiono na poniższym rysunku. Sercem układu jest VCO, czyli oscylator sterowany napięciem. Generuje on sygnał wyjściowy f_o , którego częstotliwość jest określana przez napięcie sterujące doprowadzone do wejścia VCO. Sygnał wyjściowy z tego oscylatora jest porównywany w tzw. detektorze fazy z fazą i częstotliwością sygnału wejściowego PLL. Detektor ten wytwarza napięcie stałe, którego wartość i biegunowość są proporcjonalne do różnicy częstotliwości lub fazy obu sygnałów wejściowych. W większości przypadków sygnał wyjściowy detektora fazy ma jednak postać przebiegu impulsowego i nie nadaje się bezpośrednio do sterowania VCO, który wymaga napięcia stałego. Dlatego pomiędzy tymi blokami stosuje się filtr dolnoprzepustowy, który przekształca sygnał impulsowy z wyjścia detektora fazy w wygładzone napięcie stałe. To właśnie to napięcie steruje częstotliwością oscylatora VCO.

Rola sprzężenia zwrotnego w pętli PLL

Pętla PLL jest układem ze sprzężeniem zwrotnym, w którym sprzężenie powoduje, że każda zmiana fazy lub częstotliwości sygnału wejściowego jest natychmiast przekazywana do generatora VCO. Generator ten możliwie szybko dostosowuje własną częstotliwość lub fazę, przywracając zgodność obu wielkości.



Zasadniczy schemat blokowy pętli fazowej PLL (© 2021 Jos Verstraten)

Podstawowy schemat układu PLL

Jak już wspomniano, najprostsza pętla PLL składa się z trzech bloków:

- detektora fazy na wejściu,
- filtru dolnoprzepustowego,
- generatora sterowanego napięciem (VCO).

Bloki te są połączone w sposób przedstawiony na rysunku powyżej. Taki układ stanowi podstawę każdej pętli PLL i występuje niezależnie od stopnia złożoności konkretnego rozwiązania.

W kolejnych rozdziałach najpierw omówione zostanie działanie poszczególnych bloków, a następnie ich współpraca w układzie ze sprzężeniem zwrotnym. To właśnie sprzężenie zwrotne nadaje pętli PLL jej unikalne właściwości.

Detektor fazy Trzy techniki

Do realizacji detektora fazy stosuje się różne rozwiązania. Najczęściej spotykane są trzy:

- bramka EXOR (XOR),
- detektor fazy sterowany zboczami (edge-controlled),
- mnożnik analogowy.

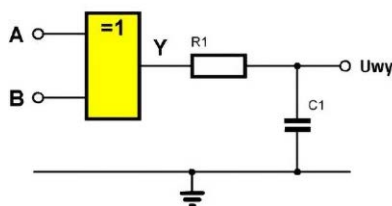
Bramka EXOR jako detektor fazy Zasada działania

W najprostszym rozwiązaniu detektor fazy w pętli PLL może być zrealizowany jako zwykła bramka EXOR. Symbole oraz tabelę prawdy takiego układu cyfrowego przedstawiono na rysunku poniżej. Podstawową właściwością bramki EXOR jest to, że na jej wyjściu pojawia się stan niski („L”), gdy oba sygnały wejściowe są jednakowe.

Jeżeli więc na oba wejścia podane są przebiegi prostokątne o identycznej częstotliwości i zgodne w fazie, to w każdej chwili mają one tę samą wartość („L” lub „H”). W rezultacie na wyjściu bramki stale występuje stan niski.

Nawet niewielka różnica fazy lub częstotliwości między sygnałami wejściowymi powoduje jednak, że pojawiają się chwile, w których oba sygnały są przeciwne. W takich momentach na wyjściu generowane są wąskie dodatnie impulsy. Ich obecność świadczy więc o różnicy fazy lub częstotliwości na wejściu.

Impulsy te muszą zostać przekształcone przez filtr dolnoprzepustowy w napięcie sterujące generatora VCO. Na rysunku pokazano najprostszy przykład takiego filtra, dołączonego do wyjścia Y bramki EXOR.



WEJŚCIE A	WEJŚCIE B	WYJŚCIE Y
L	L	L
H	L	H
L	H	H
H	H	L

Detektor fazy zrealizowany na bramce EXOR (© 2021 Jos Verstraten)

Działanie układu z bramką EXOR

Zasadę działania układu złożonego z bramki EXOR i filtru dolnoprzepustowego najłatwiej wyjaśnić graficznie. Na wykresach przedstawionych na dwóch poniższych rysunkach porównano sześć przypadków.

Podczas analizy tych przebiegów przyjmuje się, że stan logiczny „H” odpowiada napięciu dodatniemu, a stan „L” napięciu ujemnemu.

- W lewym wykresie dla SYTUACJI A nie występuje różnica częstotliwości ani fazy między sygnałami wejściowymi A i B bramki. Na wyjściu Y (czerwony) stale utrzymuje się stan „L”. Kondensator ładuje się do tego ujemnego napięcia. W rezultacie filtr daje na wyjściu duże ujemne napięcie stale Uwy (zielony).
- Na wykresach dla SYTUACJI B występuje niewielkie przesunięcie fazowe między sygnałami A i B. Na wyjściu bramki pojawiają się wąskie dodatnie impulsy Y, co powoduje, że kondensator podczas ich trwania ładuje się nieco mniej ujemnie. Średnie napięcie na kondensatorze Uwy(śr) staje się więc bardziej dodatnie, czyli mniej ujemne niż w przypadku sygnałów zgodnych w fazie.
- Na wykresach dla SYTUACJI C przesunięcie fazowe zwiększa się. W rezultacie dodatnie impulsy na wyjściu bramki stają się szersze, a kondensator w przerwach między okresami rozładowania ładuje się w większym stopniu. Średnie napięcie stale na nim staje się jeszcze mniej ujemne.

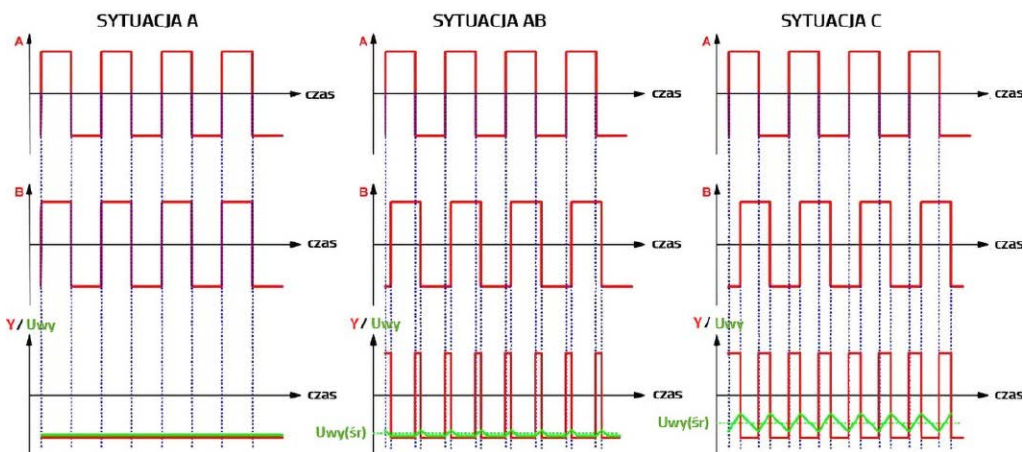
Przejdźmy teraz do wykresów przedstawionych na rysunku.

- Na wykresach dla SYTUACJI D pokazano, co się dzieje, gdy między sygnałami wejściowymi występuje przesunięcie fazowe dokładnie 90°. W rezultacie impulsy na wyjściu bramki stają się symetryczne. Kondensator ładuje się i rozładowuje przez jednakowy czas, co prowadzi do tego, że średnie napięcie na nim wynosi 0 V.
- Na wykresach dla SYTUACJI E przesunięcie fazowe jeszcze się zwiększa. Dodatnie impulsy na wyjściu bramki stają się teraz szersze niż ujemne, w wyniku czego średnie napięcie na kondensatorze przyjmuje wartość dodatnią.
- Stan ten jest jeszcze wyraźniejszy na wykresach dla SYTUACJI F. W tym przypadku przesunięcie fazowe wynosi dokładnie 180°. Wyjście bramki pozostaje wtedy stale w stanie „H”, co powoduje, że kondensator ładuje się do wysokiego dodatniego napięcia.

Podsumowanie

Działanie układu można podsumować za pomocą wykresu przedstawiającego zależność między średnim napięciem na kondensatorze Uwy(śr) a przesunięciem fazowym ϕ między oboma sygnałami wejściowymi. Pokazano to na rysunku poniżej.

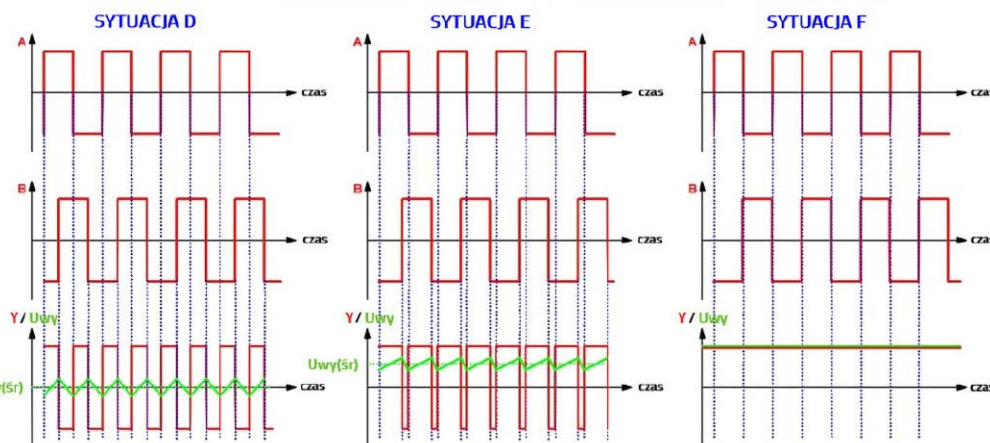
Napięcie na kondensatorze zmienia się liniowo od wartości maksymalnie ujemnej do maksymalnie dodatniej, gdy przesunięcie fazowe między sygnałami wejściowymi rośnie od 0° do 180° . Dla przesunięcia fazowego równego 90° napięcie na kondensatorze wynosi 0 V .



Działanie bramki EXOR z filtrem, część 1 (© 2021 Jos Verstraten)

Teraz zmieniają się częstotliwości obu sygnałów

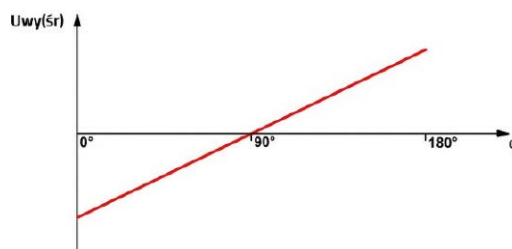
To, co się dzieje, gdy zmienia się nie faza, lecz częstotliwość obu sygnałów, przedstawiono graficznie na rysunku poniżej. Jeden z sygnałów ma stałą częstotliwość f_0 . Rozważania rozpoczynają się od przypadku, w którym sygnały wejściowe



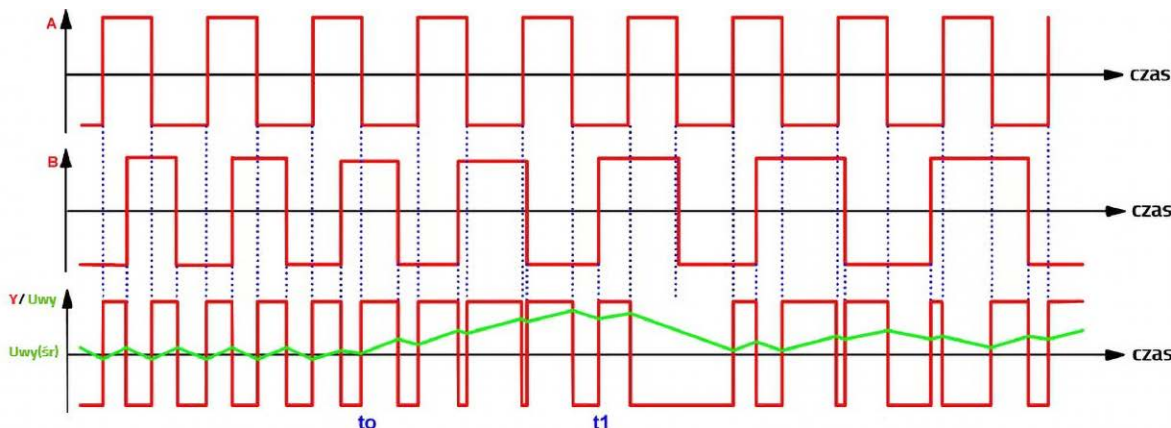
Działanie bramki EXOR z filtrem, część 2 (© 2021 Jos Verstraten)

mają tę samą częstotliwość, lecz są przesunięte w fazie o 90° . Napięcie na kondensatorze wynosi wtedy 0 V .

Jeżeli częstotliwość drugiego sygnału zacznie powoli maleć (chwila t_0), można zauważyć, że napięcie na kondensatorze staje się dodatnie. Jest to naturalne, ponieważ zmiana częstotliwości natychmiast powoduje powstanie przesunięcia fazowego. Jednak, gdy różnica częstotliwości między sygnałami dalej rośnie, okaże się, że w pewnym momencie (t_1 na wykresie) napięcie na kondensatorze przestaje rosnąć i zaczyna maleć – i to dość szybko. Również to jest logiczne: gdy częstotliwości obu sygnałów wejściowych zbyt różnią, ich okresy stają się na tyle niesynchroniczne, że raz



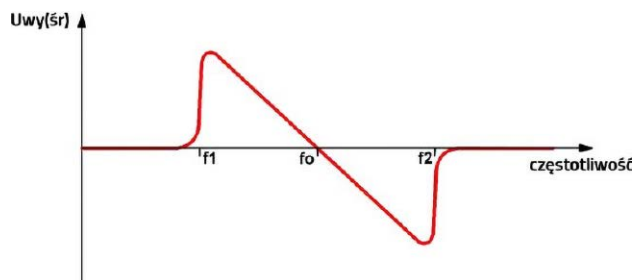
Zależność między napięciem na kondensatorze a przesunięciem fazowym między oboma sygnałami wejściowymi (© 2021 Jos Verstraten)



Przebieg napięcia na kondensatorze, gdy sygnały wejściowe mają różne częstotliwości (© 2021 Jos Verstraten)

występuje dodatnie przesunięcie fazowe, a innym razem ujemne. Kondensator jest więc naprzemiennie ładowany i rozładowywany, co powoduje, że średnie napięcie ponownie dąży do zera.

Analogiczny wykres można sporządzić dla przypadku, gdy częstotliwość sygnału wejściowego jest większa niż częstotliwość sygnału odniesienia. Wówczas na kondensatorze również pojawia się napięcie stałe, ale o wartości ujemnej. Zjawisko to utrzymuje się do momentu, gdy różnica częstotliwości stanie się zbyt duża – wtedy napięcie na kondensatorze ponownie zbliża się do zera.



Zależność między napięciem na kondensatorze a różnicą częstotliwości między oboma sygnałami wejściowymi (© 2021 Jos Verstraten)

Podsumowanie

Zależność między napięciem $U_{wy}(\bar{s})$ na kondensatorze a odchyłką częstotliwości drugiego sygnału można przedstawić graficznie, jak pokazano na rysunku poniżej. Zakres od f_1 do f_2 nazywa się zakresem synchronizacji (ang. capture range, zakres przechwytywania) układu i stanowi jedną z najważniejszych właściwości pętli PLL.

Przypis redaktora: W literaturze rozróżnia się dwa pojęcia: zakres przechwytywania (capture range), czyli zakres częstotliwości, w którym pętla PLL może wejść w stan synchronizacji, oraz zakres podtrzymania synchronizacji (lock range), czyli zakres, w którym pętla utrzymuje synchronizm po jego osiągnięciu. Zakres przechwytywania jest zazwyczaj węższy od zakresu podtrzymania.

Wady metody z bramką EXOR

Układ złożony z bramki EXOR i filtru dolnoprzepustowego działa bardzo dobrze i istnieje wiele pętli PLL wykorzystujących to rozwiązanie. Wadą tej metody jest jednak konieczność sterowania układu idealnymi przebiegami prostokątnymi. Ponadto rozwiązanie z bramką EXOR wymaga, aby sygnały były symetryczne w czasie, czyli miały dokładnie 50% wypełnienia.

Detektor fazy sterowany zboczami

Zasada działania układu

Detektor fazy sterowany zboczami jest układem, który wprawdzie wymaga sterowania sygnałami cyfrowymi, ale nie jest konieczne, aby były one symetryczne w czasie. Można więc podać na jego wejścia dwa wąskie impulsy, a układ i tak poprawnie porówna ich fazę i częstotliwość.

Wynika to bezpośrednio z faktu, że układ reaguje wyłącznie na zbocza narastające sygnałów. W praktyce mierzy on różnicę czasu między zboczem narastającym pierwszego sygnału a zboczem narastającym drugiego.

Na podstawie tej różnicy oraz kolejności pojawiania się zboczy generowane są dwa sygnały sterujące, którymi sterowane są dwa przełączniki elektroniczne. Przełączniki te ładują kondensator filtru lub powodują jego rozładowanie.

Schemat blokowy zasady działania

Schemat blokowy układu przedstawiono na rysunku poniżej. Detektor fazy jest złożonym układem, zawierającym więcej niż cztery przerzutniki typu flip-flop. Jego działanie jest skomplikowane, dlatego skupimy się tu wyłącznie na sposobie sterowania dwoma przełącznikami elektronicznymi S1 i S2.

Jest oczywiste, że gdy zamknięty zostanie przełącznik S1, kondensator C1 filtru RC ładuje się do dodatniego napięcia zasilania. Natomiast gdy zamknięty zostanie przełącznik S2, kondensator rozładowuje się do ujemnego napięcia zasilania.

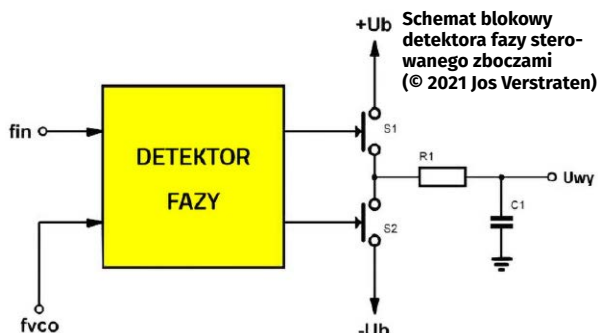
Jeżeli częstotliwość sygnału wejściowego jest znacznie mniejsza niż częstotliwość generatora VCO w pętli PLL, wówczas przełącznik S2 pozostaje niemal cały czas zamknięty. Kondensator filtru rozładowuje się do napięcia ujemnego, które w zamkniętej pętli PLL służy do zmniejszania częstotliwości generatora VCO.

Jeżeli natomiast częstotliwość sygnału wejściowego jest znacznie większa niż częstotliwość własna VCO, przełącznik S1 pozostaje prawie stale zamknięty. Kondensator filtru ładuje się wtedy do napięcia dodatniego, które powoduje zwiększenie częstotliwości generatora w pętli PLL.

Gdy częstotliwość sygnału wejściowego zbliża się do częstotliwości generatora VCO, detektor fazy sterowany zboczami zaczyna porównywać wzajemne położenie zboczy narastających obu sygnałów.

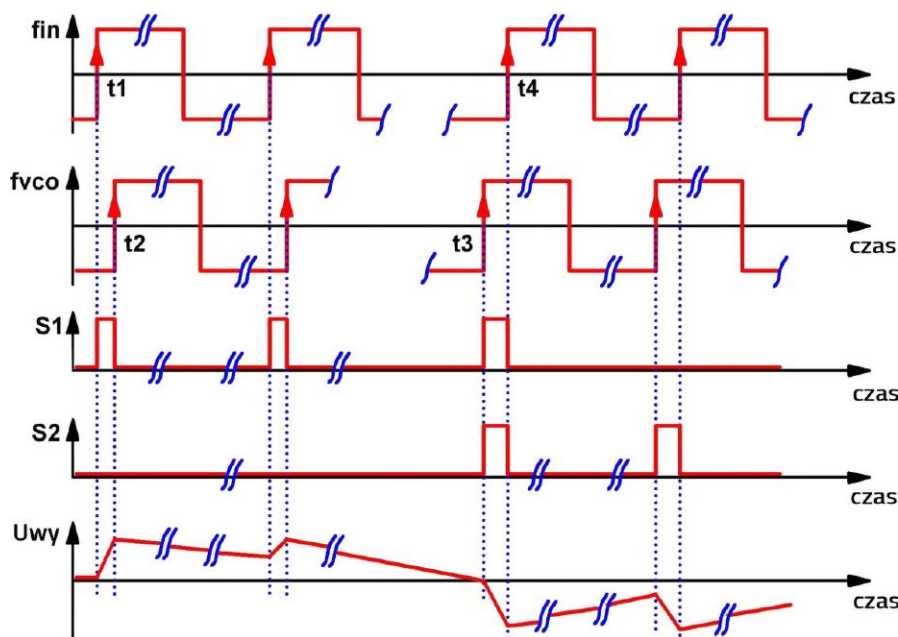
Graficzne wyjaśnienie działania układu

Jeżeli zbocze narastające sygnału wejściowego f_{in} pojawia się jako pierwsze, jak pokazano na lewych wykresach rysunku poniżej, zamykany jest przełącznik S1. Przełącznik ten zamyka się w chwili pojawienia się zbocza narastającego sygnału wejściowego i otwiera się przy pojawieniu się zbocza narastającego sygnału z VCO. Kondensator filtru ładuje się więc do napięcia dodatniego, którego wartość jest wprost proporcjonalna do różnicy czasu między tymi zboczami. Napięcie to powoduje niewielkie zwiększenie częstotliwości generatora VCO, dzięki czemu różnica faz między sygnałami maleje lub całkowicie zanika.



Jeżeli jako pierwsze pojawia się zbocze narastające sygnału z VCO, jak pokazano na prawych wykresach, zamykany jest przełącznik S2. Również w tym przypadku przełącznik otwiera się przy nadejściu zbocza narastającego drugiego sygnału. Kondensator zostaje wtedy rozładowany do napięcia ujemnego, a powstałe napięcie sterujące powoduje niewielkie zmniejszenie częstotliwości VCO, tak aby także w tym przypadku zredukować różnicę faz.

Jeżeli między sygnałami nie występuje ani różnica częstotliwości, ani fazy, oba zbocza narastające pojawiają się dokładnie w tym samym momencie i żaden z przełączników S1 ani S2 nie jest załączany. Na kondensatorze filtru nie powstaje napięcie, a generator VCO pracuje z własną częstotliwością.

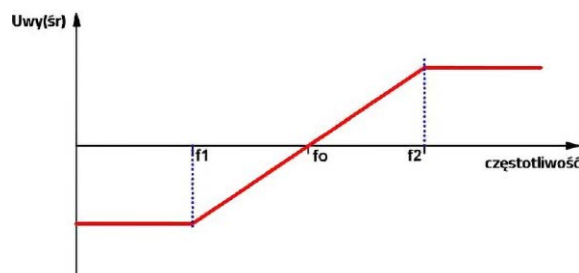


Graficzne wyjaśnienie działania detektora fazy sterowanego zboczami (© 2021 Jos Verstraten)

Mnożnik analogowy jako detektor fazy Czym jest mnożnik analogowy?

Właściwości oraz ogólny symbol mnożnika analogowego przedstawiono na rysunku poniżej. Mnożnik analogowy to układ zdolny do realizacji operacji matematycznej $C = A \cdot B$. W tym przypadku A i B są chwilowymi wartościami dwóch analogowych napięć wejściowych, a C jest chwilową wartością ich iloczynu.

Układ uwzględnia algebraiczne reguły znaków: plus razy plus daje plus, ale także minus razy minus daje plus.



Zależność między napięciem wyjściowym a różnicą częstotliwości i fazy obu sygnałów wejściowych (© 2021 Jos Verstraten)

Działanie układu

Dzięki tej właściwości mnożnik analogowy może być wykorzystany jako detektor fazy. Dużą zaletą tego rozwiązania jest możliwość pracy z przebiegami sinusoidalnymi. Zasada działania została wyjaśniona na podstawie wykresów przedstawionych na rysunku poniżej.

Na lewych wykresach pokazano przypadek, gdy oba sygnały wejściowe są zgodne w fazie. W chwilach t_0 , t_1 , t_2 , t_3 i t_4 sygnały wejściowe przechodzą przez zero, dlatego napięcie wyjściowe mnożnika również przyjmuje wtedy wartość zerową. W przedziale od t_0 do t_1 oba sygnały są dodatnie, więc napięcie wyjściowe również jest dodatnie. W przedziale od t_1 do t_2 oba sygnały są ujemne, a więc napięcie wyjściowe również pozostaje dodatnie. Gdy sygnały są zgodne w fazie, na wyjściu mnożnika otrzymujemy przebieg o zawsze dodatniej wartości.

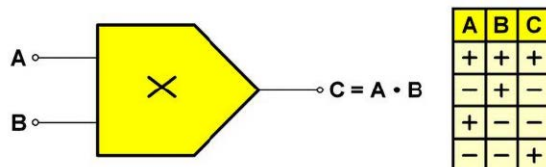
Podobnie jak wcześniej, wyjście mnożnika analogowego jest połączone z filtrem RC. Filtr ten wygładza dodatnie połówki sinusoidy, w wyniku czego na kondensatorze pojawia się dodatnie napięcie stałe (zielony wykres).

Na środkowych wykresach przedstawiono przypadek, gdy sygnały wejściowe są w przeciwfazie. Ponieważ chwilowe wartości napięć mają wówczas przeciwne znaki, napięcie wyjściowe mnożnika między przejściami przez zero jest ujemne. Kondensator filtru ładuje się więc do maksymalnego napięcia ujemnego.

Na prawych wykresach pokazano sytuację, gdy między sygnałami wejściowymi występuje przesunięcie fazowe 90° . Napięcie wyjściowe mnożnika przechodzi przez zero dwukrotnie częściej, ponieważ zera obu sygnałów nie pokrywają się. Pomiedzy tymi punktami napięcie wyjściowe jest naprzemiennie dodatnie i ujemne, gdyż znaki sygnałów wejściowych zmieniają się względem siebie. Na wyjściu powstaje więc przebieg przypominający sinusoidę o podwójnej częstotliwości, którego średnia wartość napięcia stałego wynosi 0.

Podsumowanie

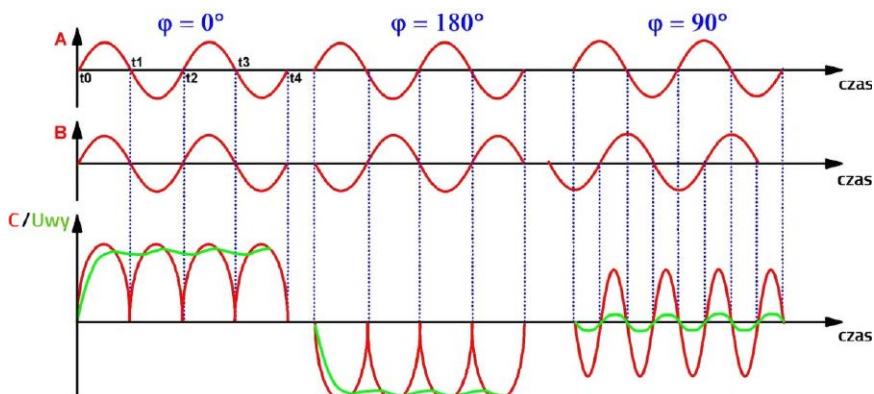
Zestawienie wyników trzech omówionych przypadków prowadzi do wykresu przedstawionego na rysunku poniżej. Otrzymana charakterystyka jest bardzo zbliżona do charakterystyki detektora fazy z bramką EXOR. Można więc uznać, że mnożnik analogowy jest bardzo dobrym detektorem fazy, którego istotną zaletą jest możliwość pracy z sygnałami sinusoidalnymi.



Mnożnik analogowy (© 2021 Jos Verstraten)

Rozwiązanie to ma jednak pewną wadę. Wartość napięcia wyjściowego mnożnika analogowego zależy od amplitudy sygnałów wejściowych. Jeżeli amplitudy te nie są stałe, również napięcie wyjściowe będzie się zmieniać.

Ponieważ napięcie wyjściowe filtra zależy bezpośrednio od wartości tego sygnału, oznacza to, że napięcie sterujące generatora VCO zależy nie tylko od relacji częstotliwości i fazy obu sygnałów wejściowych, ale także od ich amplitudy.

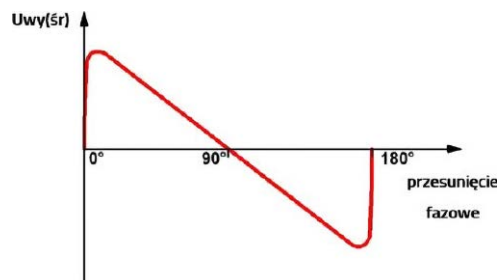


Graficzne wyjaśnienie działania mnożnika analogowego (© 2021 Jos Verstraten)

Generator sterowany napięciem (VCO) Różne rozwiązania

Opracowano różne metody sterowania częstotliwością generatora za pomocą napięcia stałego:

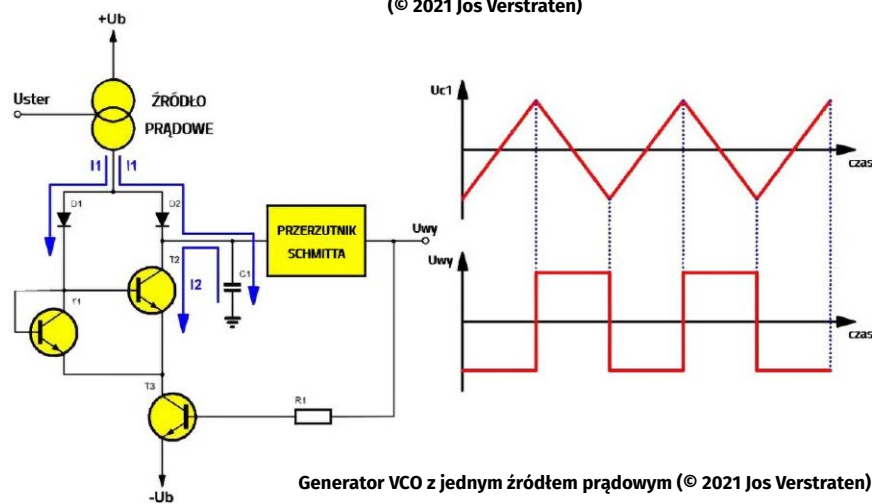
- VCO z jednym źródłem prądowym,
- VCO z integratorem i komparatorem,
- VCO symetryczny,
- VCO z wyjściem kwadraturowym.



Podsumowanie działania mnożnika analogowego (© 2021 Jos Verstraten)

VCO z jednym źródłem prądowym Schemat VCO z jednym źródłem prądowym

Jedna z najprostszych realizacji generatora VCO, wykorzystująca jedno stałe źródło prądowe, została przedstawiona na rysunku poniżej. Układ składa się ze sterowanego źródła prądowego, które może odprowadzać swój prąd wyjściowy przez dwie diody D1 i D2. Istotnym elementem jest kondensator C1. Na początku jest on ładowany prądem I1 ze źródła prądowego, a napięcie na nim rośnie liniowo. Do kondensatora dołączony jest jednak przerzutnik Schmitta z dwoma symetrycznymi progami przełączania.



Generator VCO z jednym źródłem prądowym (© 2021 Jos Verstraten)

Gdy napięcie na kondensatorze osiągnie górny próg, wyjście przerzutnika Schmitta przechodzi w stan „H”.

W efekcie tranzystor T3 zostaje wysterowany w stan przewodzenia. Prąd źródła prądowego jest teraz odprowadzany przez przewodzący tranzystor T1 do ujemnego napięcia zasilania. Złącza baza-emiter tranzystorów T1 i T2 są połączone równolegle, co powoduje, że oba tranzystory przewodzą w identyczny sposób. W rezultacie przez T2 płynie prąd dokładnie równy prądowi płynącemu przez T1.

Prąd ten może pochodzić jedynie z kondensatora, który zaczyna się rozładowywać prądem I2 równym prądowi płynącemu przez T1. Ponieważ prąd ten jest całkowicie określony przez sterowane źródło prądowe, kondensator rozładowuje się z takim samym prądem, z jakim był wcześniej ładowany. W efekcie napięcie na kondensatorze maleje liniowo z taką samą szybkością, z jaką wcześniej rosło.

Proces ten trwa aż do chwili, gdy napięcie na kondensatorze osiągnie dolny próg przełączania przerzutnika Schmitta. Wyjście tego bloku przechodzi ponownie w stan niski, tranzystor T3 zostaje zatkany, odcina dwa pozostałe tranzystory, a kondensator może ponownie rozpocząć ładowanie.

Wniosek

W rezultacie na kondensatorze powstaje w pełni symetryczne napięcie o przebiegu trójkątnym, a na wyjściu przerzutnika Schmitta – przebieg prostokątny, który jest oczywiście całkowicie zsynchronizowany z przebiegiem trójkątnym. Ten sygnał prostokątny jest wykorzystywany, gdy generator VCO pracuje w pętli PLL.

Napięcie to steruje jednym z wejść detektora fazy, najczęściej tym, które odpowiada za sygnał o stałej częstotliwości f_0 .

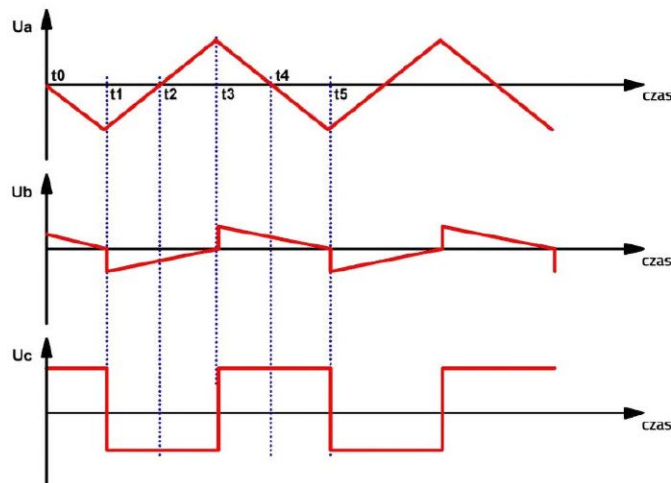
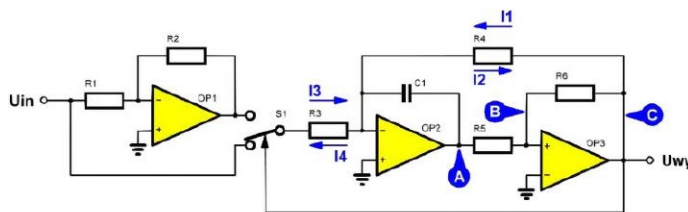
Układ integrator-komparator Zasada działania układu

Jedną z często stosowanych realizacji generatora VCO jest układ integrator-komparator. To ten sam układ, który stanowi podstawę większości analogowych generatorów funkcyjnych. Zasadniczy schemat blokowy przedstawiono na rysunku poniżej.

Układ zbudowany jest z trzech wzmacniaczy operacyjnych. Pierwszy (OP1) pracuje jako wzmacniacz odwracający o wzmocnieniu dokładnie -1 . Drugi (OP2) jest skonfigurowany jako integrator. Jego wejście jest sterowane poprzez przełącznik elektroniczny S1 – podawany jest na nie sygnał wejściowy (pozycja dolna) lub jego odwrócona wersja (pozycja górna).

Trzeci wzmacniacz operacyjny (OP3) pracuje jako komparator. Porównuje on napięcie wyjściowe integratora (punkt A) z masą.

Napięcie wejściowe U_{in} stanowi napięcie sterujące układu, które wyznacza częstotliwość generatora VCO i pochodzi z filtru detektora fazy.



Układ integrator-komparator (© 2021 Jos Verstraten)

Działanie układu

Założmy, że napięcie wejściowe wynosi 0 V. Oba styki przełącznika S1 znajdują się wtedy na poziomie 0 V. Przez rezystor R3 nie płynie prąd, ponieważ wejście odwracające wzmacniacza OP2 również znajduje się na potencjale masy. Wejście nieodwracające tego wzmacniacza jest bowiem bezpośrednio połączone z masą, a każdy wzmacniacz operacyjny z ujemnym sprzężeniem zwrotnym dąży do tego, aby różnica napięć między jego wejściami była równa zero.

Analiza działania układu opiera się na założeniu, że w chwili t_0 napięcie wyjściowe U_c komparatora jest dodatnie. Napięcie to wymusza przepływ prądu I_1 przez rezystor R4. Prąd ten zasila integrator OP2, powodując liniowy spadek napięcia wyjściowego tego układu.

Można to wyjaśnić w następujący sposób: prąd I_1 może płynąć dalej jedynie przez kondensator, ponieważ wzmacniacz operacyjny charakteryzuje się bardzo dużą impedancją wejściową. Prąd przepływający przez kondensator powoduje jego liniowe ładowanie. Ponieważ jednak lewa okładka kondensatora jest połączona z masą (wejściem odwracającym OP2), napięcie na jego prawej okładce musi maleć.

Miedzy punktami A i C znajduje się dzielnik napięcia złożony z rezystorów R5 i R6. Napięcie w punkcie B jest wyznaczone przez napięcia w punktach A i C, czyli na wyjściach wzmacniaczy OP2 i OP3. W chwili t_0 napięcie w punkcie B jest dodatnie. W miarę jak napięcie wyjściowe integratora maleje, napięcie w punkcie B również maleje.

W chwili t_1 napięcie w tym punkcie spada do 0 V. Komparator OP3 przełącza się, ponieważ jego wejście odwracające jest połączone z masą. W rezultacie napięcie wyjściowe U_c wzmacniacza OP3 przyjmuje wartość ujemną. Ma to dwa skutki. Po pierwsze, napięcie w punkcie B również przyjmuje wartość ujemną. Punkty A i C znajdują się teraz na potencjale ujemnym, więc węzeł dzielnika napięcia między nimi także musi mieć napięcie ujemne. Po drugie, przez rezystor R4 zaczyna płynąć prąd I_2 , równy co do wartości prądowi I_1 , lecz o przeciwnym kierunku.

Integrator zbudowany na wzmacniaczu OP2 jest teraz sterowany prądem o przeciwnym kierunku. W rezultacie kondensator C1 zaczyna się rozładowywać, a napięcie wyjściowe integratora OP2 rośnie. W miarę jak napięcie wyjściowe integratora rośnie, napięcie w punkcie B ponownie zbliża się do zera.

W chwili t_3 napięcie w punkcie B ponownie osiąga 0 V, a komparator przełącza się z powrotem, dając na wyjściu napięcie dodatnie.

Wpływ napięcia wejściowego U_{in}

Dotychczas zakładaliśmy napięcie wejściowe równe 0 V. Rozważmy teraz przypadek, gdy do wejścia układu zostanie przyłożone napięcie dodatnie. Układ OP1 wraz z przełącznikiem S1 powoduje wtedy zwiększenie prądu wejściowego integratora. Gdy komparator OP3 doprowadza prąd I_1 przez rezystor R4 do wejścia integratora OP2, przez rezystor R3 zaczyna dodatkowo płynąć drugi prąd I_3 . Oba prądy mają ten sam kierunek, co powoduje wzrost całkowitego prądu integratora.

Gdy napięcie wyjściowe komparatora zmienia stan, przełącza się również przełącznik S1, w wyniku czego przez R3 zaczyna płynąć prąd I_4 wypływający z integratora. W obu przypadkach prąd ładowania lub rozładowania kondensatora C1 zwiększa się w takim samym stopniu. W efekcie kondensator ładuje się i rozładowuje szybciej, a częstotliwość pracy układu rośnie.

Widać więc, że wzrost częstotliwości zależy wyłącznie od wartości rezystora R3 oraz wielkości napięcia wejściowego. Między wzrostem częstotliwości a dodatnim napięciem wejściowym istnieje zależność liniowa.

W analogiczny sposób można wykazać, że przy podaniu na wejście napięcia ujemnego częstotliwość układu maleje. Również w tym przypadku istnieje liniowa zależność między wartością napięcia a spadkiem częstotliwości.

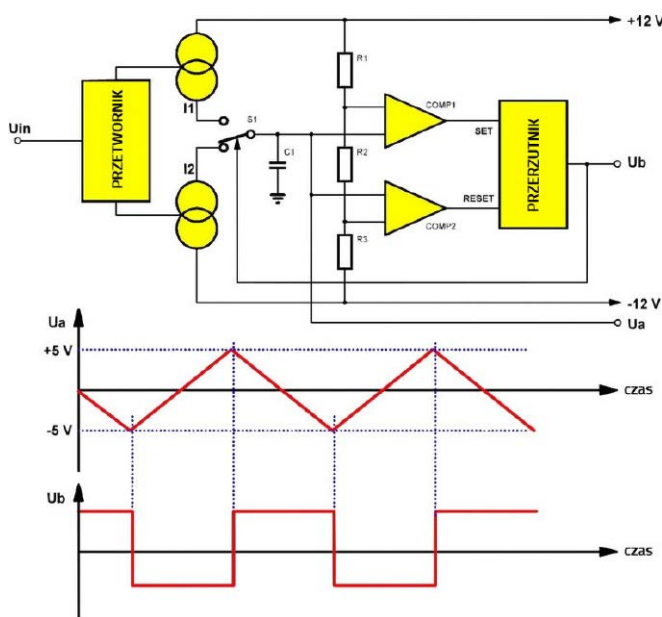
Podsumowanie

Układ VCO generuje na wyjściu wzmacniacza OP2 napięcie o przebiegu trójkątnym, a na wyjściu wzmacniacza OP3 napięcie prostokątne. Częstotliwość tych sygnałów jest określona przez wartości rezystorów R3 i R4, pojemność kondensatora C1 oraz wielkość napięcia wejściowego U_{in} .

Symetryczny generator VCO Schemat blokowy

W symetrycznym generatorze VCO, którego schemat blokowy przedstawiono na rysunku poniżej, wykorzystuje się dwa źródła prądowe, które naprzemiennie ładują i rozładują kondensator. Oba źródła są sterowane z sygnału wejściowego za pomocą przetwornika (konwertera). Układ ten zapewnia, że przez oba źródła płyną prądy o jednakowej wartości, lecz przeciwnych kierunkach. Wartość tych prądów zależy od wielkości napięcia wejściowego.

Napięcie na kondensatorze jest porównywane w dwóch komparatorach z symetrycznymi progami przełączania, w przedstawionym przykładzie $+5\text{ V}$ i -5 V . Wyjścia tych komparatorów sterują wejściami set i reset przerzutnika typu D. Wyjście Q tego przerzutnika steruje przełącznikiem elektronicznym S1, który przełącza źródła prądowe do kondensatora.



Symetryczny generator VCO (© 2021 Jos Verstraten)

Działanie układu

Załóżmy, że kondensator C1 jest całkowicie rozładowany, a przełącznik S1 znajduje się w górnym położeniu. Górne źródło prądowe dostarcza wtedy prąd I_1 , który powoduje liniowe ładowanie kondensatora. Po pewnym czasie napięcie na kondensatorze osiąga wartość $+5\text{ V}$. Górny komparator przełącza się, a jego wyjście ustawia (set) przerzutnik. Wyjście Q (U_b) przyjmuje wtedy stan dodatni. Napięcie to steruje przełącznikiem elektronicznym S1, powodując jego przełączenie.

Kondensator zaczyna się teraz rozładowywać przez dolne źródło prądowe prądem I_2 . Po pewnym czasie napięcie na kondensatorze osiąga wartość -5 V . Dolny komparator przełącza się, a jego wyjście zeruje (reset) przerzutnik. Wyjście Q przyjmuje stan ujemny, przełącznik ponownie zmienia położenie i znów włącza prąd ładowania I_1 .

Podsumowanie

Na kondensatorze C1 powstaje w pełni symetryczne napięcie o przebiegu trójkątnym. Na wyjściu przerzutnika otrzymujemy natomiast symetryczny przebieg prostokątny.

VCO z wyjściem kwadraturowym Czym jest wyjście kwadraturowe?

Trzy omówione wcześniej układy VCO generują napięcie trójkątne oraz prostokątne o określonej zależności fazowej. Gdy napięcie trójkątne rośnie, napięcie prostokątne ma stan niski. Gdy napięcie trójkątne maleje, przebieg prostokątny przełącza się w stan wysoki. Można więc powiedzieć, że między tymi dwoma sygnałami występuje przesunięcie fazowe rzędu 90° .

W niektórych zastosowaniach PLL pożądane jest jednak, aby dwa sygnały wyjściowe były zgodne w fazie (0°). Dlatego niektóre układy scalone PLL mają dodatkowe, trzecie wyjście, nazywane wyjściem kwadraturowym (quadrature output). Na wyjściu tym dostępny jest drugi przebieg prostokątny, który jest zgodny w fazie z przebiegiem trójkątnym.

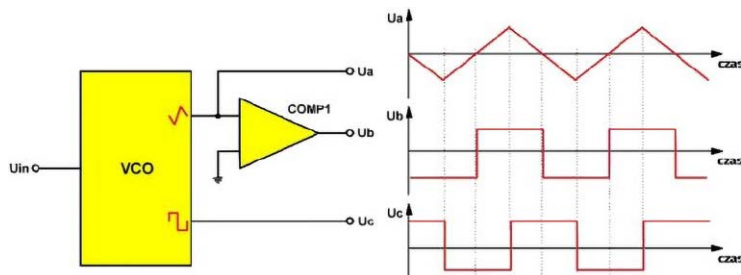
Zasada działania układu

Prosta metoda uzyskania tego trzeciego wyjścia została przedstawiona na rysunku powyżej. Wyjście trójkątne generatora VCO jest podłączone do dodatkowego komparatora, w którym przebieg trójkątny porównywany jest z masą.

Gdy napięcie trójkątne jest dodatnie, komparator generuje na wyjściu dodatnie napięcie U_b . Gdy natomiast przebieg spada poniżej zera, komparator przełącza się i na wyjściu pojawia się napięcie ujemne. W ten sposób otrzymujemy dodatkowy przebieg prostokątny, który jest zgodny w fazie z przebiegiem trójkątnym.

Działanie pętli PLL Wprowadzenie

Na podstawie dotychczas zgromadzonej wiedzy można wyjaśnić działanie układu ze sprzężeniem zwrotnym,



Generator VCO z wyjściem kwadraturowym (© 2021 Jos Verstraten)

złożonego z detektora fazy, filtru i generatora VCO. Najłatwiej zrobić to w sposób graficzny. W dalszym ciągu zakładamy, że detektor fazy został zrealizowany za pomocą bramki EXOR.

PLL z sygnałem wejściowym o częstotliwości f_o

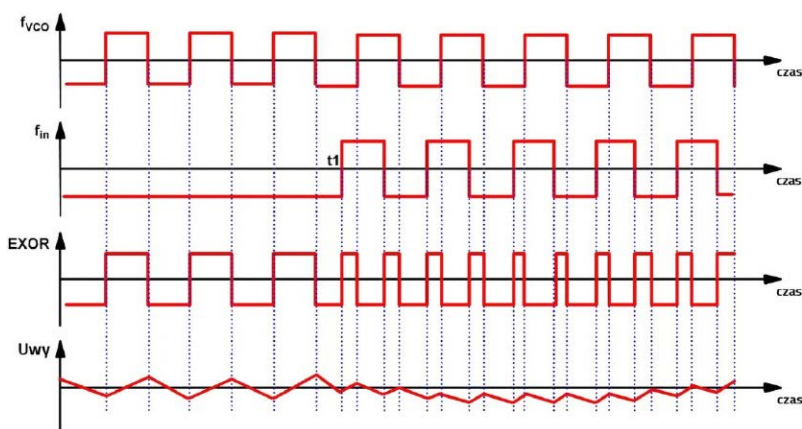
To, co dzieje się, gdy na wejście pętli PLL podany zostaje sygnał o częstotliwości równej własnej częstotliwości generatora VCO (f_o), przedstawiono na rysunku poniżej.

Przed chwilą t_1 sygnał wejściowy nie jest jeszcze obecny. Detektor fazy jest wówczas sterowany wyłącznie sygnałem z wyjścia VCO. Na wyjściu bramki EXOR pojawia się więc ten sam sygnał – symetryczny przebieg prostokątny.

Po przejściu przez filtr sygnał ten zostaje przekształcony w napięcie stałe o średniej wartości 0 V. Napięcie to steruje VCO, jednak ponieważ wynosi 0 V, nie wpływa na jego częstotliwość. Generator pracuje więc z własną częstotliwością f_o .

W chwili t_1 na drugim wejściu detektora fazy pojawia się sygnał wejściowy o tej samej częstotliwości f_o . W tym momencie między sygnałami występuje dowolna relacja fazowa. W rezultacie na wyjściu bramki EXOR pojawiają się wąskie impulsy dodatnie lub ujemne. Napięcie na kondensatorze zaczyna rosnąć lub maleć, w zależności od tej relacji fazowej. Napięcie to steruje VCO, powodując wzrost lub spadek jego częstotliwości.

W wyniku tej zmiany częstotliwości zmienia się również relacja fazowa między sygnałami.



Reakcja pętli PLL na nagłe pojawienie się sygnału wejściowego o częstotliwości f_o . (© 2021 Jos Verstraten)

Sprężenie zwrotne zaczyna działać

Układ ze sprzężeniem zwrotnym powoduje, że relacja fazowa ustawia się tak, aby napięcie sterujące VCO ponownie wynosiło 0 V. Jak wspomniano przy omawianiu detektora EXOR, ma to miejsce wtedy, gdy między sygnałami ustali się stałe przesunięcie fazowe równe 90° . Pętla PLL reguluje więc częstotliwość wyjściową VCO tak, aby spełniony był ten warunek.

Podsumowanie

Jeżeli na wejście pętli PLL podany zostanie sygnał o częstotliwości równej częstotliwości własnej generatora VCO, układ samoczynnie ustawi się tak, że między sygnałem wejściowym a sygnałem wyjściowym VCO będzie występować stałe przesunięcie fazowe wynoszące 90° .

PLL z sygnałem wejściowym o zmienionej częstotliwości

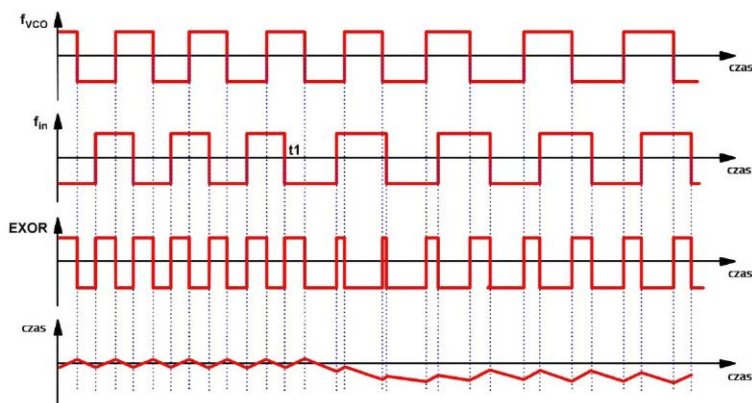
To, co dzieje się, gdy na wejście pętli PLL podany zostaje sygnał o częstotliwości różniącej się od własnej częstotliwości generatora VCO (f_o), przedstawiono na rysunku poniżej.

Do chwili t_1 częstotliwość sygnału wejściowego jest równa f_o , czyli częstotliwości własnej VCO. Napięcie wyjściowe filtru wynosi 0 V, a układ znajduje się w stanie równowagi.

W chwili t_1 częstotliwość sygnału wejściowego ulega zmianie. W rezultacie powstaje przesunięcie fazowe różne od 90° , a detektor fazy zaczyna generować wąskie impulsy dodatnie lub ujemne. Impulsy te wpływają na napięcie na kondensatorze w znany sposób – napięcie stałe przestaje być równe 0 V.

Napięcie to steruje generatorem VCO, powodując zwiększenie lub zmniejszenie jego częstotliwości. Z wykresów wynika, że pętla PLL dąży teraz do nowego stanu równowagi. Układ dostosowuje swoją częstotliwość do częstotliwości sygnału wejściowego.

W efekcie generator VCO musi być stale sterowany napięciem różnym od zera, a pętla PLL wytwarza napięcie stałe będące miarą różnicy częstotliwości między f_o a f_1 , czyli nową częstotliwością sygnału wejściowego.



Zachowanie pętli PLL po podaniu na wejście sygnału o częstotliwości różnej od f_o . (© 2021 Jos Verstraten)

Zastosowania pętli PLL Wprowadzenie

Pętla PLL znajduje zastosowanie w wielu układach analogowych. W tym rozdziale przedstawiono krótki przegląd wybranych zastosowań.

Demodulacja sygnałów FM

Generator VCO jest dostrajany do częstotliwości nośnej demodulowanego sygnału. Każde odchylenie częstotliwości sygnału wejściowego, wynikające z modulacji FM, powoduje zmianę napięcia sterującego VCO. Napięcie to jest proporcjonalne do odchyłki częstotliwości, a więc w istocie stanowi odwzorowanie sygnału modulującego, który został nałożony na falę nośną.

W ten prosty sposób można odzyskać sygnał FM bez konieczności stosowania dużych i kosztownych cewek.

Demodulacja sygnałów PSK

PSK (Phase Shift Keying) to metoda, w której sygnały cyfrowe są przekształcane w przebiegi zmienne. Przejście ze stanu „L” do „H” i odwrotnie jest realizowane przez zmianę fazy sygnału o 180°. Ponieważ pętla PLL potrafi wykrywać przesunięcia fazowe, może być wykorzystana do odzyskania informacji cyfrowej z sygnału modulowanego fazowo.

Demodulacja sygnałów FSK

Inną metodą przesyłania sygnałów cyfrowych, na przykład przez linię analogową, jest kodowanie stanów „L” i „H” za pomocą dwóch różnych częstotliwości. Metoda ta nazywana jest FSK (Frequency Shift Keying). Pętla PLL jest naturalnym rozwiązaniem do wykrywania tych różnic częstotliwości i przekształcania ich z powrotem w użyteczny sygnał cyfrowy.

Powielanie częstotliwości

Rozszerzając podstawowy układ o dzielniki częstotliwości włączone między wyjście VCO a wejście detektora fazy, można uzyskać system umożliwiający mnożenie częstotliwości sygnału wejściowego przez określony współczynnik. Dzięki temu z dokładnej częstotliwości wzorcowej, np. rezonatora kwarcowego, można uzyskać wiele innych częstotliwości o tej samej dokładności.

Dobrym przykładem takiego zastosowania jest zestaw: generator HF Marconi TF2015 wraz z synchronizatorem cyfrowym Marconi TF2171 Digital Synchroniser. Generator TF2015 wytwarza modulowane i niemodulowane przebiegi sinusoidalne w zakresie od 10 MHz do 520 MHz, o amplitudzie do 200 mV. Moduł TF2171 zawiera wewnętrzny, bardzo stabilny wzorzec częstotliwości 5 MHz oraz zestaw cyfrowych dzielników częstotliwości, dzięki którym można – za pomocą siedmiu przełączników obrotowych – ustawić częstotliwość generatora TF2015 z dokładnością do 100 Hz.

Oba urządzenia pracują w układzie ze sprzężeniem zwrotnym wykorzystującym pętlę PLL, która steruje częstotliwością generowaną przez TF2015. W ten sposób uzyskuje się bardzo wysoką stabilność i dokładność częstotliwości, rzędu $\pm 1/10\,000\,000$.

Filtry strojone

Za pomocą rozbudowanych układów PLL można realizować filtry o bardzo stromych charakterystykach przejścia (zarówno w paśmie przepustowym, jak i zaporowym). W razie potrzeby można w ten sposób projektować filtry pasmowoprzepustowe o szerokości pasma rzędu zaledwie kilkudziesięciu Hz – co jest praktycznie nieosiągalne przy użyciu klasycznych technik analogowych.

Obróbka sygnału (signal conditioning)

Bardzo ważnym zastosowaniem pętli PLL jest możliwość odzyskania sygnału zanurzonego w szumie, którego poziom jest znacznie wyższy niż poziom samego sygnału.

Demodulacja stereofoniczna

Prawdopodobnie analogowa radiofonia FM zniknie z eteru w 2023 roku.

Przypis redaktora: Prognoza ta nie sprawdziła się. W roku 2026 analogowa radiofonia FM nadal funkcjonuje w wielu krajach, a instytucje nadzorujące rynek radiowy (np. Ofcom w Wielkiej Brytanii) zakładają jej utrzymanie co najmniej do około 2030 roku. Proces odchodzenia od FM przebiega wolniej i jest zróżnicowany regionalnie.

W nowoczesnych tunerach FM do dekodowania sygnałów stereo wykorzystuje się techniki PLL.

Podczas nadawania analogowego sygnału stereofonicznego oba kanały L i R są przesyłane w postaci sygnałów sumy i różnicy, czyli L+R oraz L-R. Sygnał L-R jest modulowany na tłumionej fali nośnej o częstotliwości 38 kHz. Aby umożliwić jej odtworzenie w odbiorniku, razem z sygnałem nadawany jest ton pilotujący o częstotliwości 19 kHz.

Za pomocą pętli PLL można z tego sygnału bardzo precyzyjnie odtworzyć nośną 38 kHz, bez konieczności stosowania strojonych filtrów z cewkami i kondensatorami. ■



Zestaw Marconi TF2015/TF2171 stanowi dobry przykład zastosowania pętli PLL (© Marconi Instruments)

Jos Verstraten

Od redakcji: Choć czasy pandemii koronawirusa wydają się dziś już nieco odległe, nie wszystkie wnioski z tamtego okresu warto odłożyć do lamusa. Jednym z nich jest po prostu dokładne mycie rąk – czynność banalna, a jednak zaskakująco często wykonywana „na skróty”. W przeciwieństwie do licznych, chwilowych „wynałazków pandemicznych”, które szybko popadły w zapomnienie, prezentowany projekt stanowi praktyczne przypomnienie, że nawet najprostsze zalecenia mają sens, o ile są konsekwentnie realizowane. A w tym przypadku elektronika może w tym skutecznie pomóc.

Zdalnie sterowany sygnalizator na czas pandemii koronawirusa

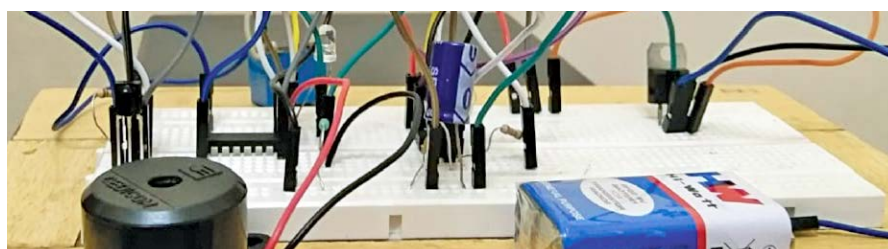
Pokazany tu projekt wpisuje się w zalecenia Światowej Organizacji Zdrowia (WHO) obowiązujące w czasie pandemii koronawirusa. Jednym z podstawowych zaleceń – obok zachowania dystansu społecznego – jest staranne mycie rąk (w kwestii „dystansu” pojawiały się zresztą także dość kuriozalne projekty DIY, np. z czujnikami odległości noszonymi na głowie i skierowanymi w cztery strony świata – przyp. red.).

Zgodnie z wytycznymi, ręce należy przez co najmniej 20 sekund dokładnie namydlać, a dopiero potem spłukać pod bieżącą wodą. Dwadzieścia sekund może wydawać się krótkim czasem, jednak w praktyce łatwo go nie dotrzymać bez dodatkowego przypomnienia. Można co prawda posłużyć się zegarkiem lub stoperem w smartfonie, lecz podczas mycia rąk nie jest to wygodne.

Zaproponowany układ sygnalizuje upływ zadanego czasu. Indykacja realizowana jest dwójako: za pomocą diody LED oraz – dla osób preferujących sygnał akustyczny – brzęczyka. Prototyp wykonany przez autora przedstawiono na **rysunku 1**, natomiast schemat ideowy układu pokazano na **rysunku 2**.

W układzie wykorzystano baterię 9 V, stabilizator 5 V typu LM7805 (IC1), timer NE555 (IC2), driver ULN2003 (IC3), przełącznik SPDT z cewką 5 V (RL1) oraz kilka prostych elementów pasywnych. Czas trwania impulsu wyznacza stała czasowa R2 i C2 i wynosi około 20 sekund. W timerze skonfigurowanym jako przerzutnik monostabilny obowiązuje zależność $T=1,1RC$. W analizowanym przypadku $R=R_2=18\text{ k}\Omega$ oraz $C=C_2=1000\text{ }\mu\text{F}$, zatem $T=1,1 \times 18\text{ k}\Omega \times 1000\text{ }\mu\text{F} \approx 20\text{ s}$.

W układzie wykorzystano również popularny odbiornik podczuwani TSOP1738. Układ ten w chwili odbioru sygnału generuje na wyjściu stan niski. Po naciśnięciu dowolnego przycisku pilota sygnał ten trafia na wejście pierwszego kanału drivera ULN2003. Za pośrednictwem przełącznika następuje wyzwolenie przerzutnika monostabilnego zrealizowanego na układzie NE555. W rezultacie wyjście układu IC2



Rysunek 1. Prototyp autora zmontowany na płycie uniwersalnej

przyjmuje stan wysoki na około 20 s, co powoduje zaświecenie diody LED1 oraz włączenie sygnału dźwiękowego brzęczyka.

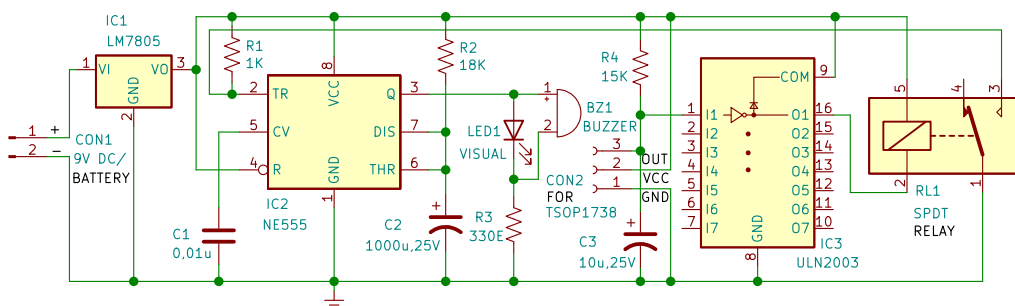
Seria odbiorników TSOP17XX pracuje na różnych częstotliwościach nośnych, co jest zakodowane w oznaczeniu „XX” układu. Elementy te współpracują z pilotami urządzeń powszechnego użytku, takich jak telewizory czy odtwarzacze DVD. TSOP1738, pracujący z częstotliwością 38 kHz, należy do najczęściej stosowanych. Układ zawiera fotodiodę PIN, przedwzmacniacz oraz filtr, co umożliwia poprawną detekcję sygnału zmodulowanego. Odbiornik TSOP17XX jest zamknięty w obudowie z trzema wyprowadzeniami, których kolejność pokazano na **rysunku 3**.

Trzecim układem zastosowanym w projekcie jest driver ULN2003 (IC3). Rozkład

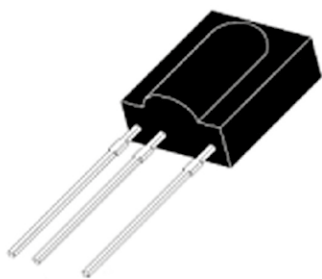
jego wyprowadzeń przedstawiono na **rysunku 4**. W prezentowanym układzie wykorzystano jeden kanał, który bezpośrednio steruje cewką przełącznika RL1 o napięciu 5 V.

Działanie układu jest proste. Pilot zdalnego sterowania należy umieścić w pobliżu umywalki. Aby rozpocząć odmierzenie czasu, wystarczy nacisnąć dowolny przycisk pilota. Świecenie diody LED oraz sygnał dźwiękowy brzęczyka informują, jak długo należy myć ręce przed ich spłukaniem. Na **rysunku 5** przedstawiono projekt jednostronnej płytki PCB, natomiast na **rysunku 6** – rozmieszczenie elementów na tej płytce.

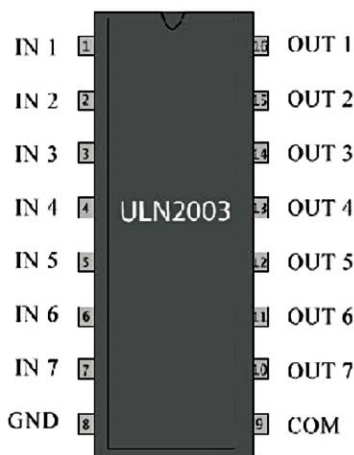
Po zmontowaniu układu należy umieścić go w odpowiednio przygotowanej obudowie. Na jej przedniej ścianie powinny znaleźć się odbiornik TSOP1738, dioda LED oraz



Rysunek 2. Schemat ideowy



Rysunek 3. Kolejność wyprowadzeń odbiornika TSOP1738



Rysunek 4. Wyprowadzenia układu ULN2003

złącze CON1. Wewnątrz, oprócz płytki PCB, należy zamontować baterię 9 V, natomiast brzęczyk można umieścić wewnątrz obudowy lub na jej tylnej ścianie. Przed odbiornikiem TSOP1738 należy wykonać okienko, tak aby promieniowanie podczerwone z nadajnika pilota mogło bez przeszkód do niego docierać. ■

Rakesh Jain

Ten artykuł był wcześniej opublikowany na łamach „EFY”, styczeń 2023 (efymag.com)

Od Redakcji EdW: Trudno oprzeć się wrażeniu, że jest to kolejny mało przekonujący projekt – i to nie tylko dlatego, że obostrzenia związane

Wykaz elementów:

Półprzewodniki:

IC1: LM7805 – stabilizator 5 V
IC2: NE555 – układ czasowy
IC3: ULN2003 – driver przełącznika

Rezystory:

(wszystkie 0,25 W, ±5%)
R1: 1 kΩ
R2: 18 kΩ
R3: 330 Ω
R4: 15 kΩ

Kondensatory:

C1: 100 nF ceramiczny
C2: 1000 µF/25 V elektrolityczny
C3: 10 µF/25 V elektrolityczny

Pozostałe:

CON1: złącze 2-pinowe
BZ1: brzęczyk
RL1: przełącznik SPDT 5 V
TSOP1738 odbiornik pilota podczerwieni
Bateria 9 V lub zasilacz 9 V DC

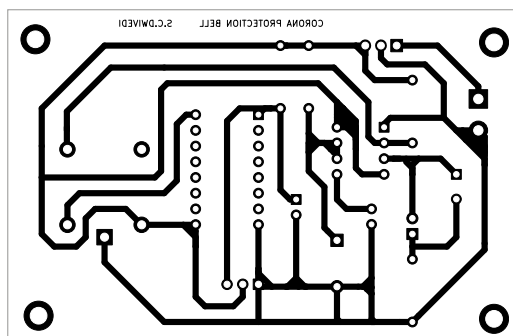
z koronawirusem zostały już zniesione. Pomijając jednak sens układu, można zadać pytanie: na ile sposobów można zrealizować 20-sekundowy przerzutnik monostabilny? Zapewne na setki, jeśli nie tysiące. A jednak trudno byłoby wymyślić rozwiązanie bardziej rozbudowane i zawiłe.

W roli przerzutnika monostabilnego pracuje tu „samotny” układ 555, który można wyzwać znacznie prościej. Autor zdecydował się jednak na wyzwalanie z pilota, co w tym zastosowaniu nie wydaje się szczególnie wygodne. Skoro tak, użycie odbiornika TSOP jest zrozumiałe, ale sygnał z jego wyjścia można byłoby doprowadzić do wejścia Trigger timera znacznie prostszą drogą – choćby przez zwykłą diodę (a nawet i bez niej układ prawdopodobnie działałby poprawnie). Po co więc dodatkowo komplikować tor sygnałowy układem ULN2003 i przełącznikiem?

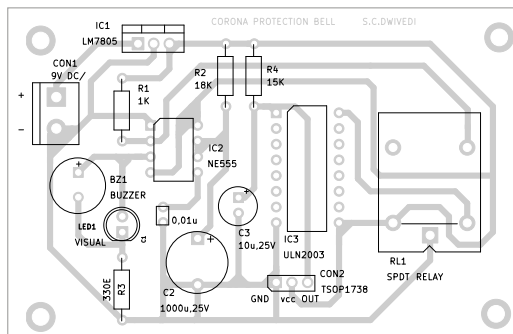
Jeżeli już zastosowano przełącznik, naturalne byłoby wykorzystanie go do sterowania obciążeniem. Tymczasem zarówno brzęczyk, jak i dioda LED są zasilane bezpośrednio z wyjścia NE555, które – szczególnie przy napięciu 5 V – ma ograniczoną obciążalność. Samo użycie timera 555 w konfiguracji monostabilnej jest oczywiście rozwiązaniem klasycznym. Odmierzany czas wynika ze stałej czasowej R2 i C2 i spełnia zależność $T \approx 1,1RC$ (ściślej: $\ln 3 - RC$), co autor poprawnie zauważa.

Pewnym utrudnieniem jest fakt, że wejście Trigger reaguje na stan niski. ULN2003 zawiera siedem tranzystorów Darlingtona (co widać na rysunku 4 – układ nie ma osobnego zasilania), a wykorzystano z niego tylko jeden tor, który dodatkowo odwraca sygnał. Ponieważ odbiornik TSOP1738 również generuje na wyjściu stan niski w chwili odbioru sygnału, konieczne stało się ponowne odwrócenie sygnału. Funkcję tę pełni przełącznik, wykorzystujący styki NC. W praktyce oznacza to jednak, że cewka przełącznika jest stale zasilana, a jedynie na krótką chwilę podczas wyzwalania zostaje odłączona. Przy zasilaniu bateryjnym trudno uznać to za rozwiązanie korzystne.

Obciążalność wyjść ULN2003 wynosi do 0,5 A, więc przy odpowiednio dobranym przełączniku układ będzie działał poprawnie. Nie zmienia to jednak faktu, że jest to rozwiązanie – mówiąc ogólnie – mało ekonomiczne energetycznie. W tym kontekście zastosowanie stabilizatora 7805 również nie wydaje się konieczne, ponieważ zarówno NE555, jak i ULN2003 mogą pracować bezpośrednio przy napięciu 9 V, a odmierzany czas jest w praktyce mało wrażliwy



Rysunek 5. Płytki PCB



Rysunek 6. Rozmieszczenie elementów na płytce PCB

na jego zmiany. W razie potrzeby prąd cewki przełącznika można byłoby ograniczyć dodatkowym rezystorem.

Powyższe uwagi nie są błędami w ścisłym znaczeniu, ale raczej pewnymi usterkami projektowymi. Można by również zastanawiać się nad brakiem diody zabezpieczającej równolegle do cewki przełącznika, jednak w tym przypadku jest ona już wbudowana w strukturę ULN2003. Samo użycie tego układu w tak ograniczonym zakresie może natomiast sprawiać wrażenie dość przypadkowego.

Pewne wątpliwości budzi także równoległe zasilanie brzęczyka i diody LED przez wspólny rezystor ograniczający prąd – skuteczność takiego rozwiązania zależy od parametrów zastosowanych elementów. Dodatkowo, interfejs pomiędzy wyjściem odbiornika TSOP a wejściem ULN2003 wymaga odpowiedniego ukształtowania sygnału, ponieważ TSOP generuje krótkie impulsy odpowiadające kodowi pilota. Obwód RC utworzony przez R4 i C3 pełni rolę układu wygładzającego, zapewniając długą stałą czasową narastania (rzędu 150 ms) i krótką opadania, co w praktyce pozwala na poprawne sterowanie przełącznikiem.

Pozostaje jednak zasadnicze pytanie: po co aż tak komplikować układ, który można zrealizować znacznie prościej? Odpowiedź na nie pozostanie zapewne tajemnicą autora. Można więc powiedzieć, że rozwiązanie jest niewątpliwie oryginalne, a przez to interesujące – choć jego praktyczny sens pozostaje dyskusyjny.

Automatyczny kran umywalkowy z oświetleniem LED

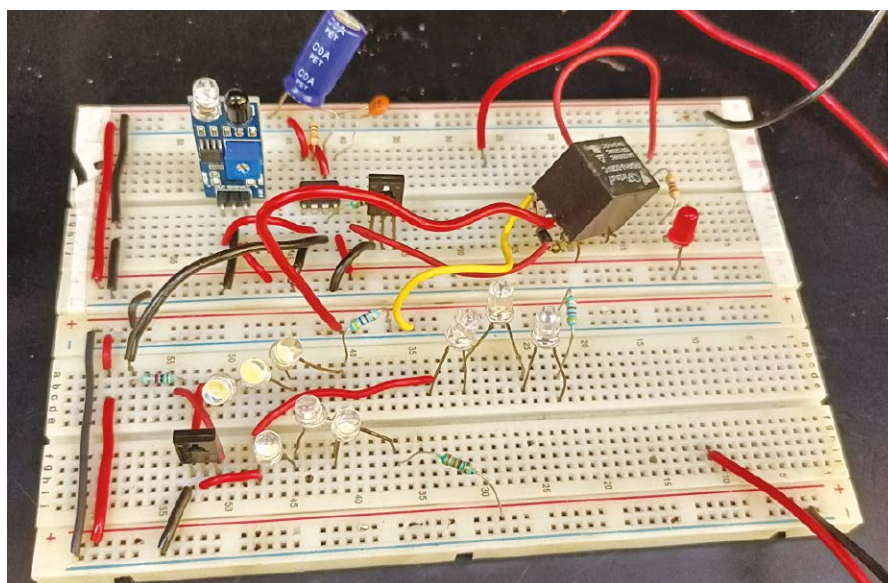
Chcąc w nocy skorzystać z umywalki, możesz mieć problem ze znalezieniem po omacku przełącznika w celu włączenia światła w łazience lub pomocniczego światła w okolicy umywalki. W takiej sytuacji, pomocnym może być zamontowanie automatycznego kranu z oświetleniem LED-owym. Prosty układ będący tematem bieżącego projektu jest sterownikiem kranu z elektrozaporem i podświetleniem z wykorzystaniem kilku białych diod LED.

Istotnym elementem decydującym o automatycznej pracy kranu jest wyposażenie sterownika w czujnik podczerwieni, który rozpoznaje, że osoba stoi przed umywalką lub zareaguje na ruch dłoni w okolicy kranu umywalkowego. Drugim istotnym elementem jest timer, który po odmierzeniu zadanego odcinka czasu spowoduje zamknięcie elektrozaporu i wyłączenie światła. Timer zrealizowano z wykorzystaniem popularnego układu zaprojektowanego właśnie do takich zastosowań – „wiecznie młodego” NE555. Czujka podczerwieni jest natomiast gotowym podzespołem dostępnym jako niewielki moduł. Moduł ten wyposażony jest w diodę nadawczą i odbiorczą, które pracują (podobnie jak popularne piloty zdalnego sterowania) w zakresie podczerwieni. Integralną częścią takiego modułu jest też komparator z możliwością regulacji czułości zadziałania. Tak skonstruowany układ można nazwać w pełni automatycznym sterownikiem kranu z zaworem elektromagnetycznym.

Opis i działanie układu

Na **rysunku 1** pokazano prototyp wykonanego przez autora projektu. Dla celów testowych układ zmontowano na uniwersalnej płytce stykowej. Natomiast na **rysunku 2** widzimy pełny schemat ideowy sterownika.

Układ wykonano z użyciem popularnych i tanich elementów. Najistotniejszymi są:

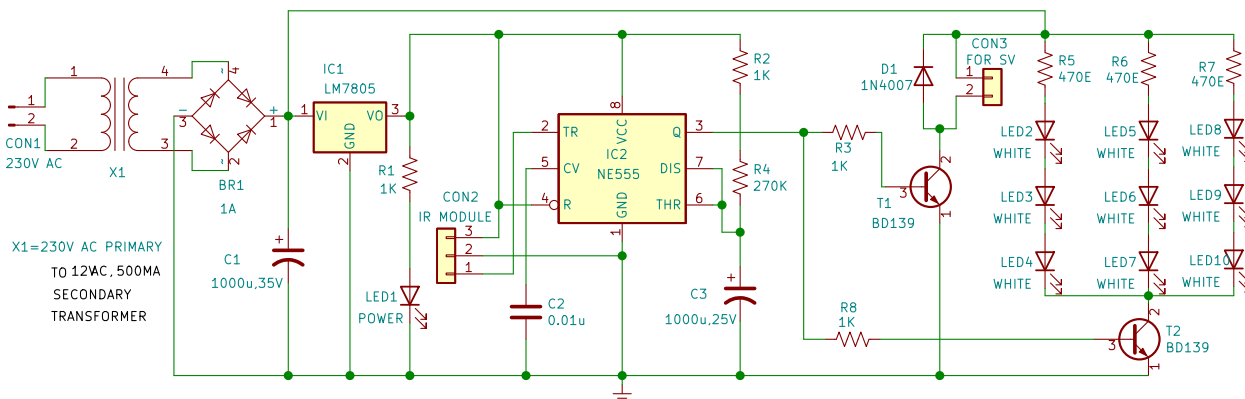


Rysunek 1. Prototyp wykonany przez autora

transformator sieciowy obniżający napięcie 230 V AC do około 12 V AC, o prądzie znamionowym do 0,5 A (X1), mostek prostowniczy (BR1), stabilizator liniowy 5 V LM7805 (IC1) i układ czasowy NE555 (IC2). Kluczowym podzespołem jest jednak moduł IR, który należy podłączyć do złącza oznaczonego CON2. Układ ten wymaga jedynie trzech przewodów, z których dwa to zasilanie +5 V i masa. Sygnał wyjściowy (na pinie 1) jest aktywny w stanie

niskim i jest podawany bezpośrednio na układ NE555. Układ ten skonfigurowano do pracy jako przerzutnik monostabilny, wyzwalany przez wejście TRIGGER. Wyzwolenie tego przerzutnika wymaga obniżenia potencjału tego wejścia poniżej 1/3 napięcia zasilania.

Jedynie część układu jest zasilana napięciem stabilizowanym +5 V. Są to układ NE555 oraz moduł IR (podczerwieni). Pozostałe



Rysunek 2. Schemat ideowy układu

Wykaz elementów:

Półprzewodniki:

IC1: LM7805 – stabilizator 5 V
IC2: NE555 – układ czasowy
T1, T2: BD139 – tranzystor NPN
BR1: mostek prostowniczy 1 A
D1: 1N4007 – dioda prostownicza
LED1: dioda LED 5 mm (czerwona lub zielona)
LED2...LED10: białe diody LED 5 mm

Rezystory:

(wszystkie 0,25 W, ±5%, węglowe)
R1...R3, R8: 1 kΩ
R4: 270 kΩ
R5...R7: 470 Ω

Kondensatory:

C1: 1000 µF/35 V elektrolityczny
C2: 100 nF ceramiczny
C3: 1000 µF/25 V elektrolityczny

Pozostałe:

CON1, CON3: złącze 2-pinowe
CON2: złącze 3-pinowe
SV: elektrozawór z cewką 12 V typ NC (normalnie zamknięty)
X1: transformator sieciowy 230 V AC, uzwojenie wtórne 12 V AC/500 mA

obwody o większym poborze prądu mogą być zasilane napięciem niestabilizowanym, z tętnieniami, przed stabilizatorem LM7805. Napięcie to jest filtrowane jedynie kondensatorem o dużej pojemności, podłączonym na wyjściu dwupołkowego mostka prostowniczego. W tym punkcie podłączono również diodę LED1, informującą o obecności napięcia zasilania.

Czas odmierzany przez przerzutnik monostabilny jest określony wartościami rezystorów R4 (oraz R2) i pojemnością kondensatora C3. Ich iloczyn stanowi stałą czasową wynoszącą około 4,5 minuty. Z uwagi na próg przełączania wejścia THRESHOLD (2/3 napięcia zasilania) stałą czasową RC należy przemnożyć przez ln3, czyli około 1,1. Daje to czas około 300 s, czyli 5 minut. W praktycznym zastosowaniu sterownika

automatycznego kranu czas ten może się okazać zbyt długi. Korekta do indywidualnych potrzeb jest możliwa przez zmianę wartości R4 i/lub C3. Wyjście układu NE555 (pin 3) w stanie aktywnym ma poziom wysoki i steruje dwoma kluczami tranzystorowymi NPN. Zastosowano tranzystory BD139. T1 steruje bezpośrednio elektrozaworem, a T2 włącza białe diody LED, skonfigurowane jako trzy gałęzie po trzy diody z rezystorami ograniczającymi 470 Ω.

Od redakcji EdW: Na schemacie (rysunek 2) zaznaczono, że elektrozawór należy podłączyć do złącza CON3, co może sugerować zastosowanie zaworu 5 V. Z uwagi jednak na zasilanie zaworu napięciem niestabilizowanym (co jest w pełni uzasadnione), jego cewka powinna być przewidziana na napięcie około 12 V. W każdej sytuacji konieczna jest także dioda antywrótnoległa D1, pokazana na rysunku 2. Jest to standardowe zabezpieczenie klucza (tranzystora T1) przy obciążeniu o charakterze indukcyjnym. W tym projekcie elektrozawór będzie prawdopodobnie elementem stanowiącym największe obciążenie dla zasilacza. Prąd ten nie jest pobierany przez stabilizator LM7805, dlatego nie ma potrzeby stosowania radiatora. Jednak w niektórych sytuacjach transformator o prądzie znamionowym uzwojenia wtórnego na poziomie 500 mA może się okazać „wąskim gardłem” tego projektu. W takim przypadku należy dobrać transformator po zmierzeniu sumarycznego prądu obciążenia, tj. prądu pobieranego przez część zasilaną stabilizatorem 7805, prądu elektrozaworu oraz prądu diod podświetlających. Jako

niedopatrzenie należy również uznać brak kondensatora elektrolitycznego na wyjściu stabilizatora 7805 (na schemacie).

Montaż i testowanie układu

Układ jest na tyle prosty, że można go zmontować na uniwersalnej płytce PCB. Docelowo należy także przewidzieć odpowiednią, szczelną obudowę. Kluczowym i najtrudniejszym elementem może okazać się montaż modułu czujnika podczerwieni (IR). Powinien on znajdować się pod kranem, a montaż należy wykonać tak, aby nie groziło zalanie czujnika wodą. Czułość powinna być wystarczająca, aby układ reagował na obecność dłoni w odległości 10...20 cm. Zbyt duża czułość jest również niepożądana. Należy upewnić się, że układ nie włącza się samoczynnie. Prace mechaniczne zapewne będą trudniejsze niż montaż obwodu elektronicznego. Należy zwrócić uwagę na montaż samego elektrozaworu oraz na ciśnienie wody w instalacji wodociągowej. Zawory tego typu są jednokierunkowe, to znaczy mają określony kierunek przepływu wody. Jeśli instalacja wyposażona jest w zbiornik wody, wystarczające ciśnienie zapewnia wysokość około 3 m słupa wody. W przypadku niespełnienia tych warunków może być konieczny także montaż pompy wody. ■

Suresh Chandra Dwivedi

Materiał filmowy do artykułu:
https://youtu.be/xGi3CU_Ndn8

Ten artykuł był wcześniej opublikowany na łamach „EFY”, lipiec 2024 (efymag.com)

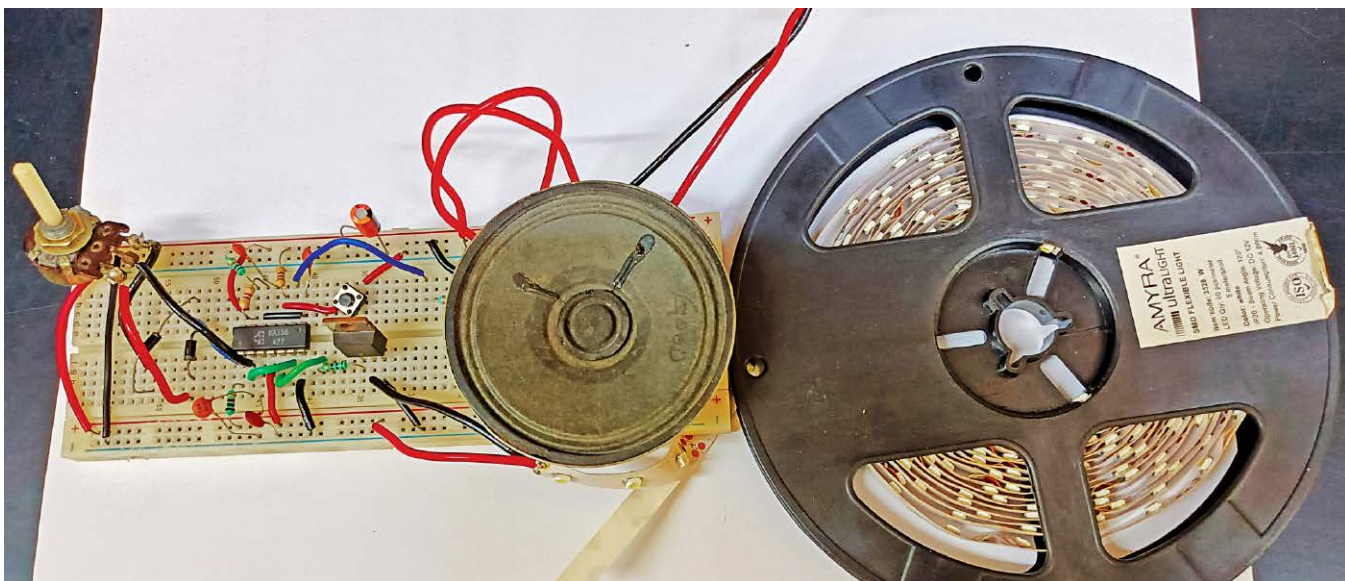
REKLAMA

Publikujemy dla projektantów i programistów elektroniki

ELPORTAL.pl

Ściemniacz taśmy LED i dzwonek

Projekt przedstawiony w bieżącym odcinku „Zrób to sam” zawiera dwa niezależne obwody: ściemniacz oświetlenia LED oraz dzwonek.



Rysunek 1. Prototyp zmontowany na uniwersalnej płytce stykowej

Regulacja jasności diod LED jest nieco trudniejsza niż w przypadku tradycyjnych żarówek (dziś należałoby powiedzieć – przestarzałych). Sytuacja pozostaje jednak stosunkowo prosta, jeśli wykorzystujemy taśmy LED zasilane napięciem stałym 12 V. Mimo to regulacja ciągła (polegająca na zmianie prądu diod) nie jest zalecana z kilku powodów. Powoduje ona zmianę barwy światła, a sam proces regulacji charakteryzuje się niską sprawnością energetyczną. Z tego względu zastosowano regulację impulsową, pozbawioną tych wad.

Zastosowanie opisanego układu pozwala zmniejszyć zużycie energii, a dodatkowo zapewnia atrakcyjny efekt wizualny. Druga część projektu, czyli dzwonek, ma natomiast klasyczne i szerokie zastosowanie.

Działanie układu oparto na układzie scalonym LM556, który zawiera w jednej obudowie dwa timery typu 555. Regulacja jasności odbywa się za pomocą potencjometru, natomiast dzwonek jest uruchamiany przyciskiem chwilowym (monostabilnym).

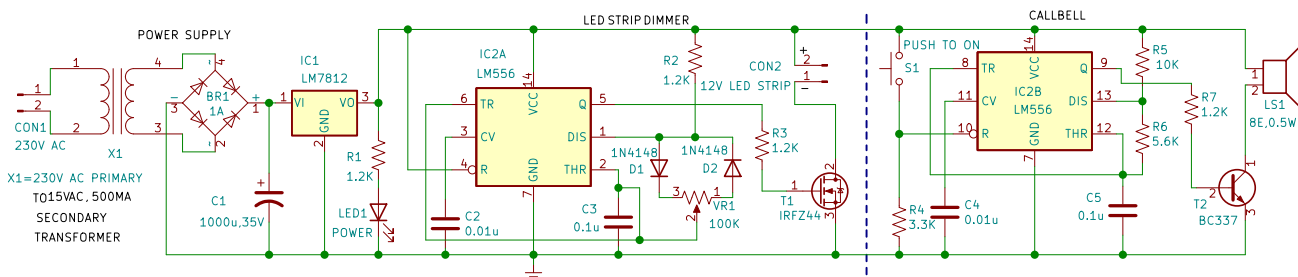
Na **rysunku 1** przedstawiono prototyp, który pozytywnie przeszedł testy w laboratorium redakcji „Electronics For You”.

Opis i działanie układu

Schemat ideowy ściemniacza i dzwonka przedstawiono na **rysunku 2**. Zawiera on wspólny obwód zasilania, wykonany w klasyczny sposób z wykorzystaniem transformatora sieciowego oraz liniowego stabilizatora napięcia.

Dzięki zastosowaniu scalonego stabilizatora LM7812 zasilacz zawiera niewiele elementów. Uzwojenie wtórne transformatora X1 obniża napięcie sieciowe 230 V AC do wartości 15 V AC, przy obciążalności do 500 mA. Napięcie stałe jest prostowane w mostku prostowniczym BR1 i filtrowane kondensatorem C1 o dużej pojemności. W tym punkcie obwodu tętnienia napięcia (o częstotliwości 100 Hz) są jednak znaczne. Dopiero za stabilizatorem IC1 uzyskuje się stabilizowane napięcie 12 V DC.

W części ściemniacza, oprócz układu scalonego LM556, pracuje tranzystor polowy IRFZ44 pełniący funkcję klucza. Część dzwonka wykonano w bardzo podobny sposób, jednak w tym przypadku elementem



Rysunek 2. Schemat ideowy ściemniacza i dzwonka

Wykaz elementów:

Półprzewodniki:

IC1: LM7812 – scalony stabilizator napięcia 12 V
IC2: LM556 – podwójny timer
T1: IRFZ44 – tranzystor MOSFET z kanałem typu N
T2: BC337 – tranzystor NPN
D1, D2: 1N4148 – dioda sygnałowa
LED1: dioda LED 5 mm

BR1: mostek prostowniczy 1 A

Rezystory:

(wszystkie 0,25 W, ±5% węglowe)
R1, R2, R3, R7: 1,2 kΩ
R4: 3,3 kΩ
R5: 10 kΩ
R6: 5,6 kΩ
VR1: 100 kΩ – potencjometr
Kondensatory:
C1: 1000 µF/35 V elektrolityczny
C2, C4: 0,01 µF ceramiczny
C3, C5: 0,1 µF ceramiczny

Pozostałe:

CON1, CON2: złącze 2-pinowe
S1: przycisk dzwonek (monostabilny)
LS1: głośnik 8 Ω/0,5 W
X1: transformator sieciowy, uzwojenie pierwotne 230 V, uzwojenie wtórne 15 V/500 mA
Taśma LED, 12 V

kluczującym jest tranzystor NPN typu BC337. Steruje on głośnikiem o impedancji 8 Ω i mocy 0,5 W, który generuje dźwięk dzwonka.

Budowa ściemniacza taśmy LED

Kluczowym elementem jest timer IC2A, który skonfigurowano do pracy astabilnej. Częstotliwość pracy generatora pozostaje w przybliżeniu stała, natomiast w szerokim zakresie można regulować współczynnik wypełnienia. Generator działa na zasadzie cyklicznego ładowania i rozładowywania kondensatora C3, przy czym obwody ładowania i rozładowania są rozdzielone diodami D1 i D2. Zastosowanie potencjometru o dużej wartości umożliwia regulację współczynnika wypełnienia PWM w szerokim zakresie.

Taśmę LED wykorzystaną w prototypie pokazano na **rysunku 3**. Jej maksymalna moc wynosi około 4,8 W/m. Dioda LED1 sygnalizuje obecność zasilania zarówno w części ściemniacza, jak i dzwonka.

Od redakcji EdW: Być może pewnym niedopatrzeniem jest brak wyłącznika umożliwiającego całkowite wyłączenie diod LED. Nietrudno zauważyć, że regulacja PWM ma w praktyce bardzo szeroki zakres – od około 1,2% do niemal 100%. Przy tak małym współczynniku wypełnienia diody LED praktycznie przestają świecić. Nie jest to jednak rozwiązanie energooszczędne, dlatego zastosowanie wyłącznika byłoby wskazane.

Częstotliwość pracy generatora nie ulega istotnym zmianom, ponieważ wydłużeniu czasu włączenia (ON) towarzyszy skrócenie czasu wyłączenia (OFF) i odwrotnie. Istotne jest, aby częstotliwość pozostawała poza zakresem percepcji wzroku, co zapobiega dostrzegalnemu migotaniu. W tym

układzie częstotliwość kluczkowania wynosi około 145 Hz, co można uznać za wartość wystarczającą.

Pewne zastrzeżenia może budzić obwód zasilania. Sprawność zastosowanego zasilacza można oszacować na poziomie około 50%. O ile w wielu prostych konstrukcjach można to zaakceptować, to w przypadku układów oświetleniowych stanowi to istotną wadę. Można zatem rozważyć zastąpienie go impulsowym zasilaczem 12 V o wydajności grądowej dobranej do mocy diod LED.

W zasilaczu zastosowanym w projekcie ograniczeniem są transformator oraz liniowy stabilizator LM7812. Stabilizator ten będzie wymagał zastosowania radiatora. Z kolei tranzystor wykonawczy T1 pracuje w trybie kluczkowania i nie powinien się znacząco nagrzewać, nawet bez radiatora.

Autor podaje dopuszczalną moc taśmy LED na poziomie 4,8 W/m. Transformator o wydajności prądowej 0,5 A po stronie wtórnej pozwala zasilić niewiele ponad 1 m takiej taśmy. Dodatkowym niedopatrzeniem jest brak kondensatora elektrolitycznego na wyjściu stabilizatora IC1.

Budowa i działanie dzwonka

W dzwonku wykorzystano drugi timer zawarty w układzie scalonym LM556. Został on skonfigurowany do pracy w trybie astabilnym. W tym przypadku zarówno częstotliwość, jak i współczynnik wypełnienia są stałe.

W stanie spoczynku wejście RESET układu IC2B jest utrzymywane w stanie niskim przez rezystor R4. W układach LM555 (oraz LM556) wejście RESET jest aktywne w stanie niskim, co oznacza, że generator zostaje odblokowany dopiero po naciśnięciu przycisku S1.

Generator astabilny zrealizowany na IC2B działa analogicznie jak w części ściemniacza i opiera się na cyklicznym ładowaniu oraz rozładowywaniu kondensatora C5. Napięcie na kondensatorze zmienia się pomiędzy poziomami 1/3 i 2/3 napięcia zasilania, co wynika z progów przełączania wewnętrznych komparatorów. Znając te poziomy oraz wartości elementów R5, R6 i C5, można wyznaczyć częstotliwość pracy generatora, która wynosi około 680 Hz.

Stan wysoki na wyjściu 9 układu IC2B powoduje cykliczne włączanie tranzystora T2, w którego obwodzie kolektora znajduje się głośnik LS1. Dźwięk jest generowany tylko w czasie naciśnięcia przycisku S1.

Od redakcji EdW: W dzwonku można zaakceptować sposób sterowania głośnikiem przedstawiony na schemacie z rysunku 2. Przy krótkotrwałej pracy



Rysunek 3. Taśma LED wykorzystana w prototypie

układu najprawdopodobniej nie spowoduje to jego uszkodzenia. Warto jednak zauważyć, że głośnik pracuje w tym przypadku ze składową stałą napięcia i prądu, czego należy unikać.

Dodatkowo można oszacować, że współczynnik wypełnienia sygnału sterującego tranzystorem T2 jest stosunkowo duży i wynosi około 73%. Czas ładowania kondensatora C5 jest dłuższy niż czas jego rozładowania, a stan wysoki wyjścia Q występuje właśnie w tej fazie.

Głośnik zawiera cewkę drgającą, dlatego z punktu widzenia wzmacniacza małe obciążenie rezystancyjne. W przedstawionym układzie sterowania wskazane byłoby jednak zastosowanie diody podłączonej równolegle do głośnika, ze względu na jego składową indukcyjną.

Montaż i testowanie układu

Układ jest na tyle prosty, że można go zmontować na uniwersalnej płytce drukowanej. Docelowo warto przygotować odpowiednią obudowę. Oprócz montażu elementów na płycie należy w odpowiednim miejscu zamontować przycisk dzwonek oraz przykleić taśmę LED.

Elementy zewnętrzne, takie jak głośnik i potencjometr ściemniacza, można podłączyć za pomocą przewodów.

Przetestowanie układu polega na sprawdzeniu, czy działa on zgodnie z założeniami projektowymi. Po zakończeniu montażu należy podłączyć napięcie sieciowe do złącza CON1. Świecenie diody LED1 oznacza obecność napięcia zasilającego +12 V. ■

Suresh Chandra Dwivedi

Materiał filmowy do artykułu:
https://youtu.be/-eSaUjSw_U

Ten artykuł był wcześniej opublikowany na łamach „EFY”, wrzesień 2024 (efymag.com)



Krystian – Młodzi Entuzjaści Elektroniki, Szkoła Podstawowa nr 86, Wrocław

Dźwięk wypełnia pomieszczenie, ktoś coś mówi, ktoś inny stuka w stół, w tle gra muzyka. Z pozoru zwykłe odgłosy, które na co dzień umykają naszej uwadze. A co, jeśli spróbujemy je zobaczyć? Czy da się sprawić, by dźwięk przestał być tylko tym, co słyszymy, a stał się czymś widocznym? Podczas ostatnich zajęć przekonaaliśmy się, że tak – zbudowaliśmy układ, który zamienia dźwięk w światło.

Tym razem spotkaliśmy się w zupełnie innych okolicznościach niż miesiąc temu. Wtedy obserwowaliśmy ruch – robaczki napędzane drganiami silniczka potrafiły „ożyć” i poruszać się w nieprzewidywalny sposób. Teraz zamiast ruchu pojawiło się światło, a zamiast drgań mechanicznych – sygnały akustyczne. Zamiast zastanawiać się, w którą stronę pobiegnie robaczka, zaczęliśmy zadawać inne pytania: co się stanie, gdy do mikrofonu coś powiemy, klaśniemy albo włączymy muzykę? I czy układ będzie reagował za każdym razem tak samo?

Co zbudujemy tym razem?

Wśród projektów pojawił się kolorofon LED (AVTEDU644) – układ, który natychmiast przyciąga uwagę. Po podłączeniu zasilania i lekkim podkręceniu czułości diody zaczynają reagować na każdy dźwięk z otoczenia. Czasem migają gwałtownie, czasem płynnie się rozjaśniają, a czasem reagują tylko na wyraźniejsze impulsy.



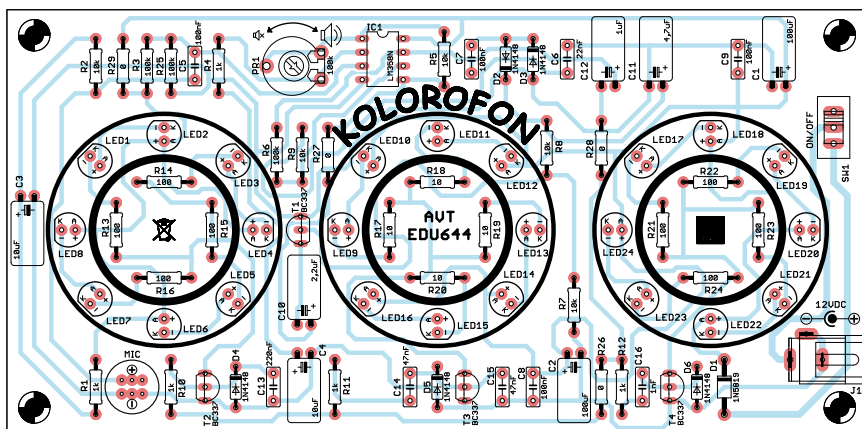
Rysunek 1. Kolorofon LED (kod AVTEDU644). Zmontowany układ

Najciekawsze zaczyna się jednak wtedy, gdy przestajemy traktować go tylko jako „efekt świetlny”. Wystarczy chwila, by zauważyć, że różne dźwięki wywołują różne reakcje. Jedne powodują szybkie błyski, inne spokojniejsze zmiany. Pojawia się naturalne pytanie: dlaczego tak się dzieje?

I tu zaczyna się prawdziwa zabawa. Kolorofon przestaje być tylko urządzeniem, a staje się narzędziem do eksperymentów. Wystarczy zmienić źródło dźwięku, jego głośność albo charakter, by zobaczyć, jak układ „interpretuje” sygnał. Można sięgnąć po coś, co większość z nas ma zawsze pod ręką.

Fotografia 1. Od lewej: Julek, Adrian, dwaj Kornele oraz Krystian. Młodzi Entuzjaści Elektroniki, Szkoła Podstawowa nr 86, Wrocław





Rysunek 2. Schemat montażowy układu

Smartfon, który zwykle służy do codziennych zastosowań, w tym przypadku staje się kieszonkowym laboratorium. Może generować dźwięki o określonej częstotliwości, zmieniać ich kształt i poziom, a tym samym pozwala świadomie sprawdzać działanie układu. To już nie tylko obserwacja – to eksperyment.

Na **rysunku 1** pokazano zmontowany układ, a na krótkim filmie pod adresem <https://youtu.be/7nDjTh7R9ZM> możesz zobaczyć działanie kolorofonu podczas reakcji na dźwięk.

Schemat montażowy

Schemat montażowy to rysunek, który pokazuje, gdzie dokładnie na płytce drukowanej należy umieścić każdy z elementów zestawu. Dzięki niemu łatwo odnaleźć właściwe miejsce dla rezystorów, kondensatorów, diod czy układów scalonych i innych podzespołów,

ponieważ wszystkie komponenty są oznaczone takimi samymi desygntorami, zarówno na schemacie montażowym, liście elementów jak i na schemacie ideowym. Ułatwia to bezbłędne i szybkie składanie układu, nawet osobom początkującym. Schemat montażowy pomaga również uniknąć pomyłek, takich jak wlutowanie elementu w niewłaściwe miejsce lub ustawienie go w złej orientacji. Schemat montażowy, który dodatkowo pokazuje układ ścieżek i padów, bardzo pomaga w kontroli poprawności wykonanych połączeń lutowanych. Dzięki temu łatwo ustalić, czy połączenia pomiędzy sąsiednimi polami są przewidziane w projekcie, czy też powstały przez pomyłkę, na przykład na skutek przypadkowego zwarcia ich cyną podczas nieostrożnego lutowania. Taki podgląd znacząco ułatwia wykrywanie błędów i zwiększa pewność, że układ został zmontowany prawidłowo. Schemat montażowy pokazano na **rysunku 2**.

Montaż rezystorów

Zgodnie z informacjami z listy elementów przylutuj rezystory o określonych wartościach rezystancji na odpowiednich pozycjach na płytce drukowanej. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Rezystor jest elementem, który ogranicza przepływ prądu w obwodzie elektrycznym. Dzieje się tak dlatego, że ma on określoną rezystancję – właściwość materiału wywołującą spadek napięcia podczas przepływu prądu. Energia elektryczna zamienia się w nim na ciepło, co jest naturalnym skutkiem przepływu prądu przez rezystancję.
- Rezystor nie ma biegunowości – działa tak samo niezależnie od kierunku przepływu prądu. Dlatego jego montaż na płytce nie wymaga uwzględnienia orientacji. Ważne jest jedynie, aby

Wykaz elementów:

R1, R4, R10...R12: 1 kΩ (brązowy-czarny-czerwony-żółty)
R2, R5, R7...R9: 10 kΩ (brązowy-czarny-pomarańczowy-żółty)
R3, R6, R25: 100 kΩ (brązowy-czarny-żółty-żółty)
R13...R16, R21...R24: 100 Ω (brązowy-czarny-brązowy-żółty)
R17...R20: 10 Ω (brązowy-czarny-czarny-żółty)
R26...R29: 0 Ω (czarny)
D2...D6: 1N4148
D1: 1N5819
Podstawka IC1: Podstawka pod IC1
LED1...LED8: dioda LED CZERWONA
LED9...LED16: dioda LED ŻÓŁTA
LED17...LED24: dioda LED NIEBIESKA
C6: 22 nF (może być oznaczony 223)
C5, C7...C9: 100 nF (może być oznaczony 104)
C13: 220 nF (może być oznaczony 224)
C14, C15: 47 nF (może być oznaczony 473)
C16: 1 nF (może być oznaczony 102)
C12: 1 μF
C10: 2,2 μF
C11: 4,7 μF
C3, C4: 10 μF
C1, C2: 100 μF
T1...T4: BC337 (lub podobny)
MIC: mikrofon elektretowy
SW1: włącznik
J1: gniazdo zasilania
PR1: potencjometr 100 kΩ + wałek regulacyjny
IC1: układ LM358

Zanim przystąpisz do montażu, zapoznaj się z instrukcjami dostępnymi online:



Lutowanie komponentów przewlekanych do płytki drukowanej

<https://elportal.pl/files/2026/02/15/3678-lutowanie-komponentow-przewlekanych-do-plytki-drukowanej.pdf>



Montaż przyjazny naprawom + Przycinanie nadmiaru wyprowadzeń

<https://elportal.pl/files/2026/02/15/3679-montaz-przyjazny-naprawom-przycinanie-nadmiaru-wyprowadzen.pdf>



Zagadnienia BHP związane z lutowaniem

<https://elportal.pl/files/2026/02/15/3680-zagadnienia-bhp-zwiazane-z-lutowaniem.pdf>



Zagadnienia BHP związane z przycinaniem nadmiaru wyprowadzeń

<https://elportal.pl/files/2026/02/15/3681-zagadnienia-bhp-zwiazane-z-przycinaniem-nadmiaru-wyprowadzen.pdf>



Zagadnienia BHP związane z uruchamianiem zmontowanego układu

<https://elportal.pl/files/2026/02/15/3682-zagadnienia-bhp-zwiazane-z-uruchamianiem-zmontowanego-ukladu.pdf>

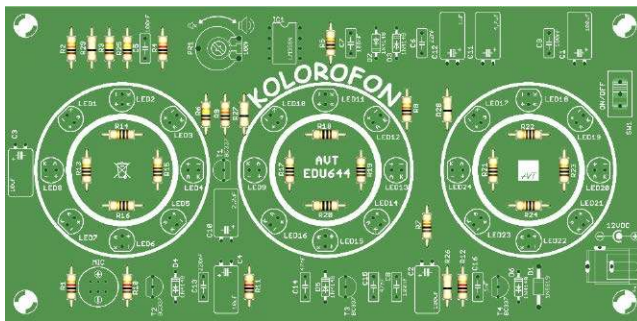
w danym miejscu umieścić właściwy rezystor o odpowiedniej wartości. Sam kierunek montażu pozostaje dowolny.

- Rezystory są zazwyczaj jednymi z najniższych elementów montowanych na płytce drukowanej. Ich przyłutowanie nie utrudnia późniejszego montażu wyższych komponentów, dlatego lutuje się je zazwyczaj w pierwszej kolejności.
- Spośród wszystkich komponentów znajdujących się w zestawie wyodrębni rezystory i odłóż je na osobną stertę.
- Pozostałe elementy odłóż na bok, a podczas montażu sięgaj wyłącznie po kolejne rezystory z wcześniej przygotowanej sterty.
- Za każdym razem, gdy weźmiesz do ręki kolejny rezystor, zmierz jego wartość za pomocą multimetru ustawionego w tryb omomierza. Odczytaną rezystancję zapamiętaj.
- Jeśli pomiar rezystorów sprawia Ci trudność, poproś o pomoc kolegę lub osobę prowadzącą zajęcia. Gdy masz dostęp do internetu, możesz też skorzystać z instrukcji dostępnej na stronie <https://elportal.pl/do-pobrania> – znajdziesz tam dokument „Pomiar wartości rezystorów za pomocą multimetru”, przygotowany jako materiał uzupełniający do EdW 11/2024.
- Na dołączonej do zestawu liście elementów odszukaj zmierzoną wcześniej wartość rezystancji (w Ω , $k\Omega$ lub $M\Omega$), a następnie odczytaj desygnator lub desygnatory przypisane do tej wartości.
- Zegnij wyprowadzenia rezystora i umieść go w płytce (**rysunek 3**) w miejscu oznaczonym właściwym desygnatorem.
- Zadbaj o to, aby każdy rezystor był włożony do płytki do końca i dobrze do niej przylegał. Estetyczny montaż nie tylko poprawia wygląd gotowego urządzenia, lecz także stabilizuje element w płytce, chroniąc go przed uszkodzeniami mechanicznymi, oraz ułatwia późniejszą diagnostykę i ewentualne naprawy.
- Przyłutuj element do płytki. Jeśli nie wiesz, jak to zrobić lub chcesz upewnić się, że wykonujesz tę czynność prawidłowo, przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.
- Usuń nadmiar wyprowadzeń za pomocą obcinaczek. Jeśli nie wiesz, jak to zrobić lub chcesz upewnić się, że wykonujesz tę czynność prawidłowo, przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.

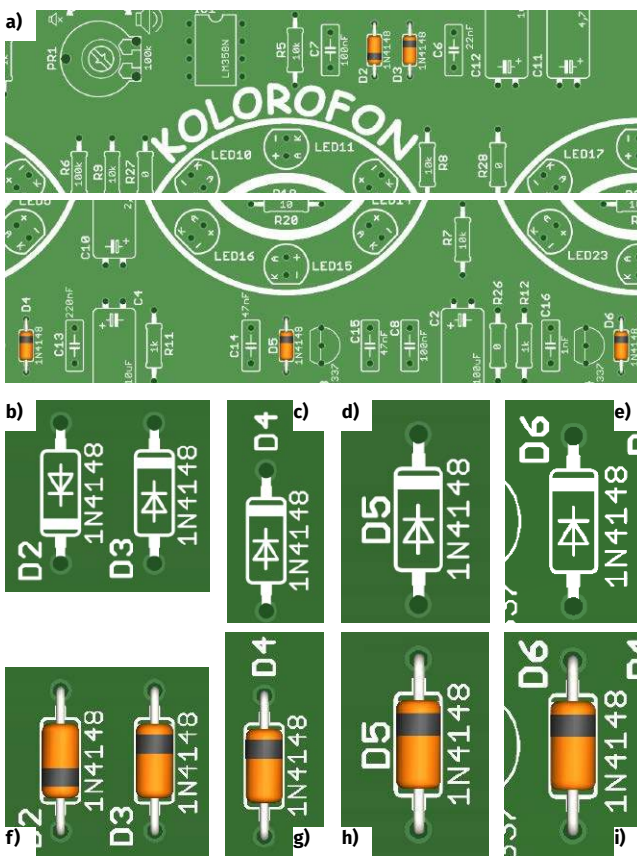
Montaż diod sygnałowych

Zgodnie z informacjami z listy elementów przyłutuj diody sygnałowe D2...D6 we wskazanych lokalizacjach. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Dioda sygnałowa jest elementem, który przewodzi prąd głównie w jednym kierunku, a blokuje jego przepływ w kierunku przeciwnym. Wynika to z właściwości złącza półprzewodnikowego, które przy odpowiedniej polaryzacji przewodzi, a przy odwrotnej – pozostaje zamknięte. Dzięki temu diody sygnałowe umożliwiają sterowanie kierunkiem przepływu prądu, ochronę układów oraz przetwarzanie sygnałów elektrycznych.
- Dioda ma biegunowość – przewodzi prąd tylko w jednym kierunku. Dlatego jej montaż na płytce wymaga zachowania właściwej orientacji. Na obudowie diody znajduje się pasek oznaczający katodę, który musi być ustawiony zgodnie ze znakiem na płytce drukowanej. Umieszczenie diody w złym kierunku spowoduje, że nie będzie działała prawidłowo w układzie.
- Diody sygnałowe są elementami o niewielkiej wysokości, dlatego również montuje się je we wczesnym etapie lutowania



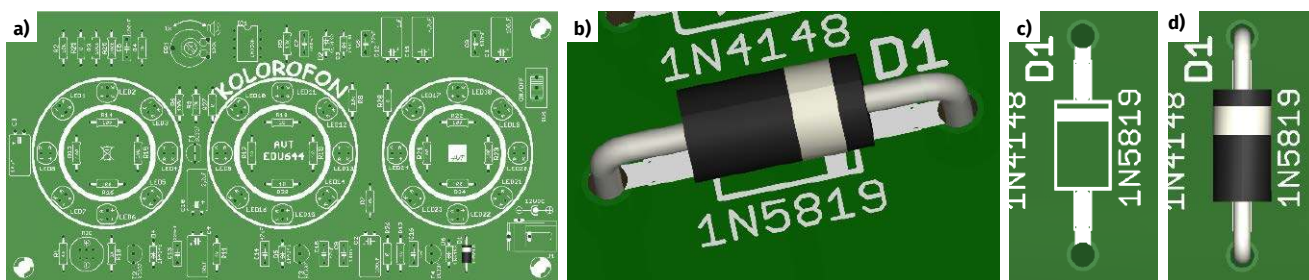
Rysunek 3. Rezystory o odpowiednich wartościach, oznaczone za pomocą kodu kolorowych pasków, zamontowane na właściwych pozycjach, zgodnie z listą elementów



Rysunek 4. a) lokalizacja elementów na płytce; b), c), d), e) obrys komponentów na płytce PCB uwzględniający pasek polaryzacji; f), g), h), i) poprawny montaż elementu z uwzględnieniem kierunku

– zwykle tuż po rezystorach lub podstawkach. Dzięki temu płytka pozostaje stabilna, a diody nie utrudniają późniejszego montażu wyższych komponentów.

- Za każdym razem, gdy weźmiesz do ręki kolejną diodę, odczytaj jej oznaczenie na obudowie i porównaj je z informacją w liście elementów. Diody różnią się między sobą typami i parametrami, dlatego ważne jest, aby upewnić się, że właściwa dioda zostanie zamontowana we właściwym miejscu na płytce.
- Warto zadbać o to, aby każda dioda była włożona do płytki do końca i dobrze do niej przylegała. Estetyczny montaż nie tylko poprawia wygląd gotowego urządzenia, lecz także stabilizuje element w płytce, chroniąc go przed uszkodzeniami mechanicznymi, oraz ułatwia późniejszą diagnostykę i ewentualne naprawy.



Rysunek 5. a) lokalizacja diody prostowniczej na płytce; b), c), d) poprawny montaż tego elementu z uwzględnieniem kierunku

- Jeśli rozpoznawanie typów diod sprawia Ci trudność, poproś o pomoc kolegę lub osobę prowadzącą zajęcia. Możesz też spróbować poszukać informacji na ten temat w Internecie, w książkach lub u agentów AI – pod warunkiem, że masz taką możliwość i potrafisz zweryfikować wiarygodność informacji z tych źródeł.
- Zegnij wyprowadzenia diody i umieść ją w płytce (**rysunek 4**) w miejscu oznaczonym właściwym desygnatorem, pamiętając o zachowaniu jej orientacji (patrz wyżej).
- Przylutuj element do płytki. Jeśli nie wiesz, jak to zrobić, albo chcesz upewnić się, że wykonujesz to prawidłowo, przeczytaj sekcję *Lutowanie komponentów przewlekanych do płytki drukowanej*. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.
- Usuń nadmiar wyprowadzeń za pomocą obcinaczek. Jeśli nie wiesz, jak to zrobić, albo chcesz upewnić się, że wykonujesz to prawidłowo, przeczytaj sekcję *Przycinanie nadmiaru wyprowadzeń*. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.

Montaż diody Schottky'ego

Kolejnym elementem do zamontowania jest dioda Schottky'ego D1 typu 1N5819. Jej montaż przebiega analogicznie jak w przypadku diod sygnałowych, dlatego możesz skorzystać ze wskazówek opisanych w sekcji powyżej. Warto jednak pamiętać, że dioda Schottky'ego różni się od typowej diody sygnałowej (np. 1N4148) znacznie mniejszym spadkiem napięcia w kierunku przewodzenia oraz szybszym działaniem. Dzięki temu lepiej sprawdza się w układach zasilania, gdzie ważne są małe straty. Z tego powodu diody te nie są zamienne – zastosowanie zwykłej diody krzemowej mogłoby pogorszyć działanie układu.

Montaż podstawki pod układ scalony

Przylutuj do płytki podstawkę pod układ IC1. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Podstawka pod układ scalony przewodzi prąd pomiędzy wyprowadzeniami układu a ścieżkami na płytce, ale sama w sobie

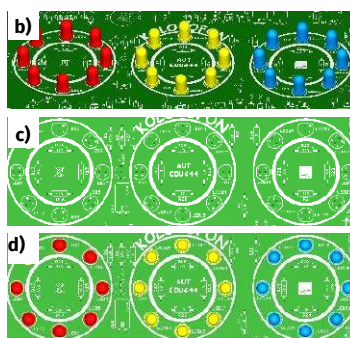
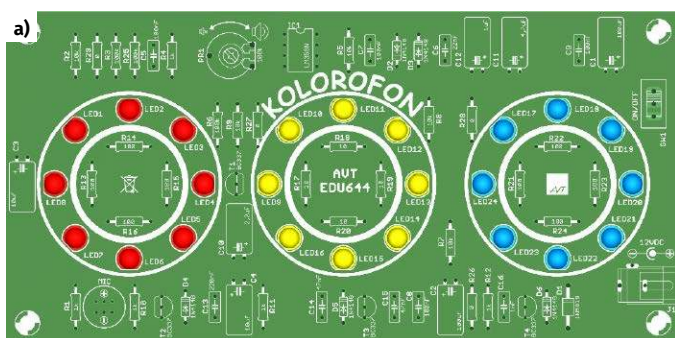
- nie pełni żadnej aktywnej funkcji elektrycznej – nie zmienia sygnałów, nie wzmacnia ich ani nie ogranicza. Jej głównym zadaniem jest zapewnienie wygodnego, bezpiecznego i wielokrotnego montażu układu scalonego bez ryzyka uszkodzenia jego wyprowadzeń lub pól lutowniczych na płytce. Podstawka umożliwia między innymi łatwą wymianę układu scalonego w razie jego uszkodzenia lub pomyłki podczas montażu.
- Podstawka pod układ scalony nie ma biegunowości w sensie elektrycznym, ale ma określoną orientację montażową. W jej obudowie znajduje się znacznik (najczęściej wycięcie lub kropka), który musi być ustawiony zgodnie ze znakiem na płytce drukowanej (**rysunek 6**). Prawidłowa orientacja podstawki jest konieczna, aby później poprawnie włożyć do niej układ scalony.
- Podstawki pod układy scalone należą do elementów o średniej wysokości, dlatego zazwyczaj montuje się je po przylutowaniu najniższych komponentów, takich jak rezystory. Dzięki temu płytka pozostaje stabilna, a podstawka nie zasłania dostępu do miejsc montażowych przeznaczonych dla pozostałych elementów.
- Przylutuj element do płytki. Jeśli nie wiesz, jak to zrobić lub chcesz upewnić się, że wykonujesz tę czynność prawidłowo, przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.
- Podczas lutowania pinów do płytki staraj się, by pomiędzy lutowanymi wyprowadzeniami nie powstały niechciane połączenia, czyli zwarcia, zwane również mostkami lutowniczymi. Jeśli podczas lutowania pojawiają się zwarcia, najłatwiej będzie, trzymając płytkę jedną ręką, ustawić ją pod kątem prostym względem blatu. Następnie należy ponownie podgrzać połączone pola lutownicze oraz przy pomocy grotu lutownicy i siły grawitacji pozwolić nadmiarowi cyny spłynąć na blat. Dzięki temu uwolnisz pady podstawki od zwarć.
- W przypadku podstawek nie ma potrzeby przycinania wyprowadzeń. Po przylutowaniu pozostaw je w oryginalnej długości.

Montaż diod LED

Zgodnie z informacjami z listy elementów na odpowiednich pozycjach przylutuj diody LED o właściwych kolorach. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

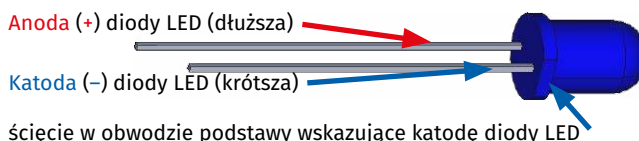


Rysunek 6. a) lokalizacja podstawki na płytce; b), c), d) poprawny montaż tego elementu z uwzględnieniem kierunku



Rysunek 7. a) lokalizacja diod LED na płytce; b), c), d) poprawny montaż elementów z uwzględnieniem kolorów i obrotów (zwróć uwagę na fragment prostej linii w obrysach diod LED wskazujący katodę)

- Dioda LED to element elektroniczny, który świeci, gdy płynie przez niego prąd w odpowiednim kierunku. Łączy w sobie działanie zwykłej diody – przewodzi prąd tylko w jedną stronę – oraz funkcję źródła światła. Dzięki temu LED-y mogą sygnalizować działanie układu, informować o stanie pracy urządzenia lub pracować w układach generujących efekty świetlne.
- Tak jak każda dioda, LED ma biegunowość. Oznacza to, że musi być podłączona we właściwym kierunku, inaczej nie zaświeci, a w szczególnym przypadku ulegnie uszkodzeniu. Jej katodę najczęściej oznacza ścięcie na obudowie oraz krótsza nóżka (rysunek 8). Przed montażem sprawdź, gdzie na PCB znajduje się oznaczenie katody, i ustaw diodę zgodnie z nim.
- Dioda LED jest elementem, który od razu przyciąga wzrok obserwatora, dlatego estetyka jej montażu ma duży wpływ na końcowy wygląd budowanego urządzenia. Warto zadbać o to, aby LED była ustawiona prostopadłe do płytki i równo do niej przylegała – nawet drobne odchylenia mogą być widoczne po uruchomieniu układu, szczególnie gdy diod jest więcej.
- Z uwagi na powyższe LED-y najlepiej montować na stosunkowo wczesnym etapie lutowania. Rezystory, podstawki pod układy scalone oraz diody prostownicze i sygnałowe są zazwyczaj nieco niższe, ale zaraz po nich warto umieścić na płytce diody LED. W tym momencie pole lutownicze jest wciąż dobrze dostępne, i nic nie zasłania miejsca montażu, co ułatwia przylutowanie LED-ów równo i estetycznie.
- Zanim włożysz diodę LED do płytki, sprawdź w liście elementów, jaki kolor LED-a powinien zostać zamontowany w danej lokalizacji. Same diody – zwłaszcza w bezbarwnych obudowach – mogą wyglądać bardzo podobnie lub wręcz identycznie, dlatego warto upewnić się, jaki kolor świecenia ma LED, który trzymasz w ręku.
- Jeśli w projekcie występuje kilka kolorów diod LED w bezbarwnych obudowach, zasadne jest ich wcześniejsze posegregowanie. Najprościej zrobić to za pomocą multimetru ustawionego w tryb badania diod lub ciągłości obwodu. Przyłożenie sond – czerwonej do anody diody LED i czarnej do jej katody (fotografia 2) – spowoduje lekkie świecenie diody LED, co pozwoli od razu ustalić jej kolor. Dzięki temu można przyporządkować poszczególne LED-y do właściwych grup i ułożyć je na osobnych stertach. Takie przygotowanie znacząco zmniejsza ryzyko pomyłek podczas montażu i gwarantuje prawidłowy efekt wizualny w gotowym urządzeniu.
- Jeśli upewniłeś się co do odpowiedniej polaryzacji i kolorów, przylutuj element do płytki. Jeśli nie wiesz, jak to zrobić lub chcesz upewnić się, że wykonujesz tę czynność prawidłowo, przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.



Rysunek 8. Opis wyprowadzeń diody LED („plusowe” wyprowadzenie dłuższe, „minusowe” krótsze)



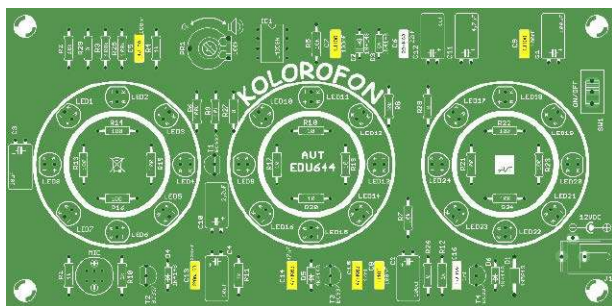
Fotografia 2. Sprawdzenie diody LED za pomocą multimetru ustawionego na funkcję testowania diod. Po przyłożeniu sondy czerwonej do anody, a czarnej do katody, sprawna dioda LED powinna się zaświecić. Jeśli dioda ma odpowiednio długie (jeszcze nie przycięte) wyprowadzenia można się wspomóc krokodylkami

- Usuń nadmiar wyprowadzeń za pomocą obcinaczek. Jeśli nie wiesz, jak to zrobić lub chcesz upewnić się, że wykonujesz tę czynność prawidłowo, przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.

Montaż kondensatorów stałych

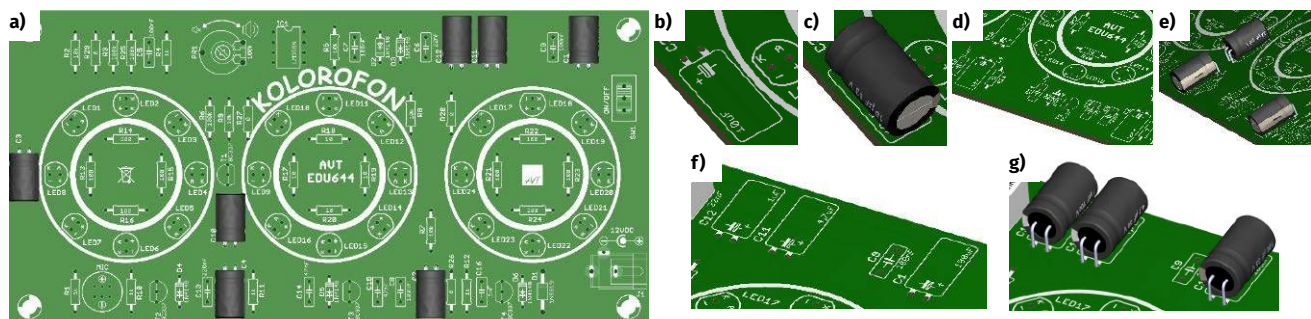
Zgodnie z informacjami z listy elementów przylutuj dostępne w zestawie kondensatory stałe we wskazanych lokalizacjach. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Kondensator stały (na przykład foliowy lub ceramiczny) to element elektroniczny służący do magazynowania ładunku elektrycznego. W układach może pełnić wiele funkcji: filtrować zakłócenia, stabilizować pracę obwodów, blokować składową stałą lub współpracować z rezystorami w układach czasowych. Kondensatory stałe mogą filtrować zakłócenia, stabilizować pracę obwodów, blokować składową stałą lub współpracować z rezystorami w układach czasowych.



Rysunek 9. Lokalizacja elementów na płytce oraz poprawny ich montaż. Elementy dostarczone w zestawie mogą różnić się typem, obudową oraz kolorem od tych pokazanych na rysunku

- Kondensatory stałe w większości przypadków nie mają biegunowości, co oznacza, że można je montować w dowolną stronę. Wyjątkiem są rzadko spotykane konstrukcje oznaczone dodatkowym paskiem lub symbolem elektrody zewnętrznej – jeśli takie oznaczenie występuje, należy postępować zgodnie z opisem w dokumentacji. W typowych projektach edukacyjnych kondensatory foliowe można traktować jako elementy niepolaryzowane.
- Przed montażem sprawdź zgodność pojemności kondensatora z miejscem, w którym ma zostać zamontowany. Odczytaj oznaczenie z jego obudowy i porównaj je z informacją w liście elementów. Zapoznaj się także z komentarzem dotyczącym sposobów oznaczania pojemności kondensatorów, obecnym na liście elementów – ułatwi to prawidłową identyfikację każdego elementu.
- Kondensatory foliowe mają zwykle dość sztywną obudowę i proste wyprowadzenia, dlatego łatwo je ustawić w odpowiedniej pozycji. Warto dopilnować, aby były ustawione prostopadle do płytki i dobrze do niej przylegały. Estetyczny montaż nie tylko poprawia wygląd gotowego urządzenia, ale również ułatwia późniejszą diagnostykę i ewentualne naprawy.
- Kondensatory stałe zalicza się do elementów średniej wysokości, dlatego montuje się je po elementach najniższych (rezystorach, diodach) i tranzystorach, a przed wysokimi kondensatorami elektrolitycznymi, przełącznikami czy złączami. Taka kolejność zapewnia stabilne ułożenie płytki oraz dobry dostęp podczas wkładania komponentów do płytki i ich lutowania.
- Po włożeniu kondensatora do odpowiednich otworów delikatnie odchyl jego wyprowadzenia, aby nie wypadł przy odwracaniu płytki. Jeśli wyprowadzenia są odpowiednio sztywne, wystarczy niewielkie zagięcie pod kątem kilku stopni.



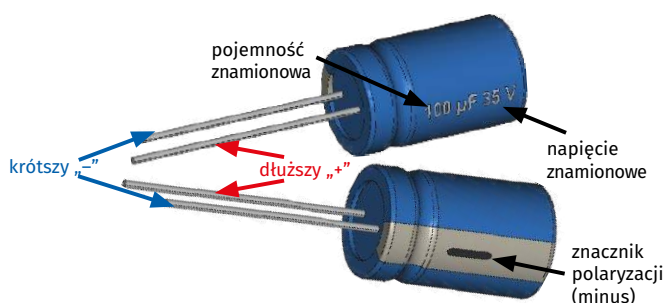
Rysunek 10. a) lokalizacje kondensatorów elektrolitycznych na płytce; b), c), d), e), f), g) poprawny montaż tych elementów z uwzględnieniem poprawnej polaryzacji

- Przylutuj element do płytki. Jeśli nie wiesz, jak to zrobić, albo chcesz upewnić się, że wykonujesz to prawidłowo, przeczytaj sekcję *Lutowanie komponentów przewlekanych do płytki drukowanej*. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.
- Usuń nadmiar wyprowadzeń za pomocą obcinaczek. Jeśli nie wiesz, jak to zrobić, albo chcesz upewnić się, że wykonujesz to prawidłowo, przeczytaj sekcję *Przycinanie nadmiaru wyprowadzeń*. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.

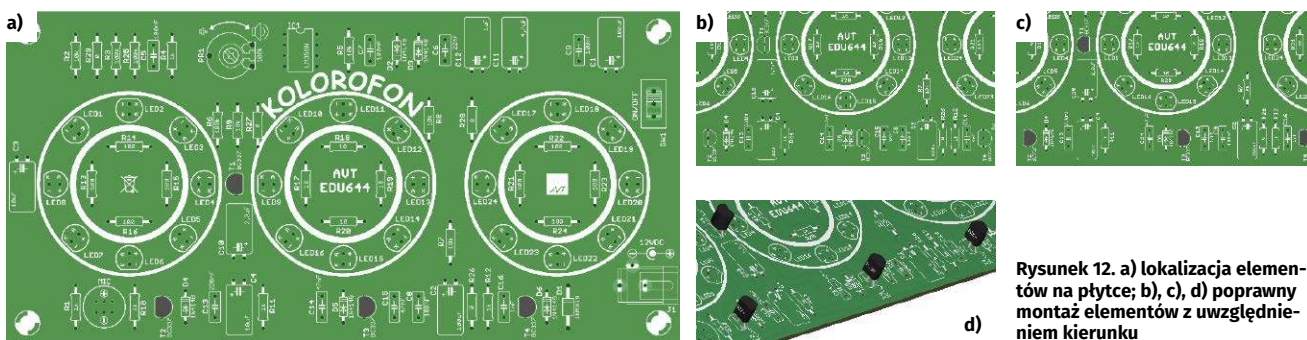
Montaż kondensatorów elektrolitycznych

Zgodnie z informacjami z listy elementów przylutuj kondensatory elektrolityczne o odpowiednich pojemnościach we wskazanych lokalizacjach. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Kondensator elektrolityczny to element spolaryzowany, który może magazynować stosunkowo duży ładunek elektryczny i pełnić w układzie różne funkcje: stabilizować napięcie, wygładzać tętnienia, filtrować zakłócenia lub dostarczać krótkotrwałych impulsów prądowych. Dzięki dużej pojemności w niewielkiej obudowie jest często stosowany w zasilaczach i układach energoelektronicznych.
- Kondensatory elektrolityczne mają zawsze określoną biegunowość. Na obudowie znajduje się wyraźne oznaczenie minusa (zwykle biały pasek), a dłuższa noga oznacza plus zasilania. Montaż odwrotny grozi uszkodzeniem kondensatora, a w skrajnych przypadkach nawet jego rozrwyaniem (rozszczeniem i dezintegracją). Zawsze upewnij się, że plus i minus znajdują się we właściwych otworach na płytce.



Rysunek 11. Na korpusie kondensatora elektrolitycznego odnajdziesz – między innymi – informacje o nominalnej pojemności oraz dopuszczalnym napięciu pracy a także znacznik polaryzacji



Rysunek 12. a) lokalizacja elementów na płytce; b), c), d) poprawny montaż elementów z uwzględnieniem kierunku

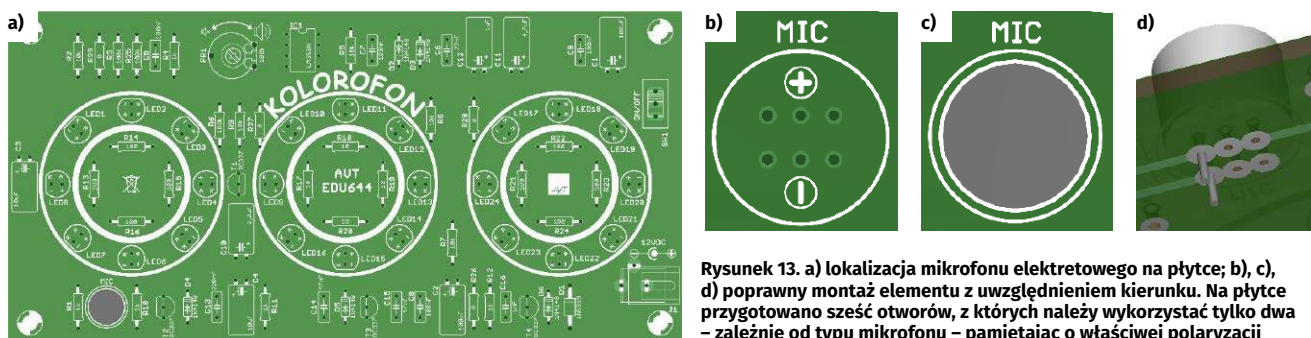
- Przed montażem sprawdź zgodność pojemności i napięcia kondensatora z miejscem, w którym ma zostać umieszczony. Odczytaj nadruk z obudowy (na przykład „100 μ F 35 V”) i porównaj go z informacją w liście elementów.
- Kondensatory elektrolityczne są wysokimi elementami, dlatego montuje się je dopiero po wlutowaniu wszystkich niższych komponentów, takich jak rezystory, diody, tranzystory czy kondensatory foliowe. Taka kolejność ułatwia pracę oraz zapewnia stabilne oparcie płytki podczas lutowania.
- Po umieszczeniu kondensatora w otworach sprawdź jeszcze raz jego orientację. W przypadku elementów polaryzowanych warto wyrobić sobie nawyk podwójnego sprawdzania przed przylutowaniem – to pozwala uniknąć potencjalnie niebezpiecznych błędów.
- Aby kondensator nie wypadł podczas obracania płytki, delikatnie odchyl jego wyprowadzenia na zewnątrz. Wystarczy niewielkie odgięcie pod kątem kilku stopni. Zbyt silne zaginięcie może utrudnić przycinanie nadmiaru wyprowadzeń, a także przyszyły ewentualny demontaż kondensatora.
- Przylutuj element do płytki. Jeśli nie wiesz, jak to zrobić lub chcesz upewnić się, że wykonujesz tę czynność prawidłowo, przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.
- Usuń nadmiar wyprowadzeń za pomocą obcinaczek. Jeśli nie wiesz, jak to zrobić lub chcesz upewnić się, że wykonujesz tę czynność prawidłowo, przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.
- **Pamiętaj o bezpieczeństwie. Kondensator elektrolityczny zamontowany odwrotnie lub podłączony do wyższego niż znamionowe napięcie może ulec uszkodzeniu, a nawet gwałtownie wybuchnąć. Dlatego przed pierwszym podłączeniem zasilania zawsze upewnij się, że został zamontowany poprawnie i ma właściwe parametry.**
- **Przed pierwszym podłączeniem napięcia do układu zawierającego kondensatory elektrolityczne obowiązkowo zakładaj okulary ochronne.**

Montaż tranzystorów NPN

Zgodnie z informacjami z listy elementów przylutuj tranzystory NPN, T1...T4 we wskazanych lokalizacjach. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Tranzystor bipolarny (BJT) to element elektroniczny, który umożliwia sterowanie przepływem prądu. Może wzmacniać sygnały, włączać lub wyłączać obwody, a także pełnić rolę klucza sterującego innymi elementami, takimi jak diody LED, przekaźniki czy brzęczyki.

- Istnieją dwa podstawowe typy tranzystorów: NPN i PNP. Choć na pierwszy rzut oka mogą wyglądać identycznie, nie są zamienne. Każdy z nich pracuje w inny sposób i wymaga innego sposobu podłączenia. Montaż tranzystora PNP w miejscu przewidzianym dla NPN (lub odwrotnie) spowoduje, że układ nie zadziała, a w najgorszym wypadku element ulegnie uszkodzeniu. Zawsze sprawdź w liście elementów, jaki typ tranzystora ma znaleźć się w danej lokalizacji.
- Tranzystor ma trzy wyprowadzenia: bazę, kolektor i emiter. Ich rozmieszczenie zależy od typu i producenta, dlatego przed montażem koniecznie porównaj układ nóżek z oznaczeniem na płytce oraz z dokumentacją elementu. Włożenie tranzystora w niewłaściwej orientacji uniemożliwi jego poprawne działanie. Jeśli chcesz zastosować zamiennik zachowaj szczególną ostrożność.
- Ponieważ tranzystory mają określony kształt obudowy (najczęściej półokrągły z jednej strony), ważne jest ich prawidłowe ustawienie. Należy włożyć je do płytki tak, aby płaska i zaokrąglona część obudowy dokładnie odpowiadała oznaczeniu na PCB. Niesymetryczny kształt obudowy wraz z odpowiednim obrysem komponentu na płytce pozwala uniknąć pomyłek oraz ułatwia montaż i późniejsze serwisowanie układu.
- Tranzystory należy montować po elementach najniższych, takich jak rezystory i diody, ale przed komponentami wysokimi (np. elektrolitami czy złączami). Dzięki temu ich nóżki będą łatwo dostępne do lutowania, a obudowa będzie stabilna podczas ustawiania w odpowiedniej pozycji.
- Zadbaj o estetyczny montaż tranzystora. Jego obudowa powinna być pionowa, a wyprowadzenia równomiernie dociśnięte do płytki. Przechylony tranzystor nie tylko wygląda nieestetycznie, ale też jest bardziej narażony na uszkodzenia mechaniczne.
- Po umieszczeniu tranzystora w płytce delikatnie odgnij jego dwie skrajne nóżki, aby nie wypadł podczas lutowania. Następnie przylutuj środkową nóżkę, ustaw do pionu dwie skrajne i je również przylutuj. Lutując, pamiętaj, aby nie przegrzewać wyprowadzeń – nadmierna temperatura może skrócić żywotność elementu.
- Jeśli chcesz upewnić się, że czynność lutowania wykonujesz prawidłowo, przeczytaj sekcję *Lutowanie komponentów przewlekanych do płytki drukowanej*. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.
- Usuń nadmiar wyprowadzeń za pomocą obcinaczek. Jeśli nie wiesz, jak to zrobić, albo chcesz upewnić się, że wykonujesz to prawidłowo, przeczytaj sekcję *Przycinanie nadmiaru wyprowadzeń*. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.



Montaż mikrofonu elektretowego

Zgodnie z informacjami z listy elementów przylutuj mikrofon elektretowy we wskazanej lokalizacji. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Mikrofon elektretowy jest przetwornikiem elektroakustycznym, który zamienia fale dźwiękowe na sygnał elektryczny o niewielkiej amplitudzie. Do poprawnej pracy wymaga zasilania (polaryzacji), zwykle poprzez rezystor podłączony do dodatniego napięcia zasilającego.
- Mikrofon elektretowy ma biegunowość, dlatego podczas montażu należy zwrócić szczególną uwagę na oznaczenia wyprowadzeń. W typowych konstrukcjach minus jest połączony z metalową obudową mikrofonu (fotografia 3). Nieprawidłowe podłączenie może uniemożliwić działanie układu, a w niektórych przypadkach doprowadzić do uszkodzenia elementu.
- Mikrofony elektretowe są elementami o niewielkiej wysokości, dlatego montuje się je zwykle po elementach najniższych (rezystory, diody), a przed wyższymi komponentami.
- Przed montażem upewnij się, że masz do czynienia z mikrofonem elektretowym, sprawdzając jego wygląd (metalowa, cylindryczna obudowa z otworem akustycznym) oraz listę elementów zestawu.
- Na dołączonej do zestawu liście elementów odszukaj mikrofon (MIC), a następnie sprawdź jego desygnator przypisany do konkretnego miejsca na płytce drukowanej.
- Umieść mikrofon w płytce, zwracając uwagę na zgodność biegunowości z oznaczeniami na laminacie. Wyprowadzenia powinny przechodzić przez otwory swobodnie, bez użycia siły.



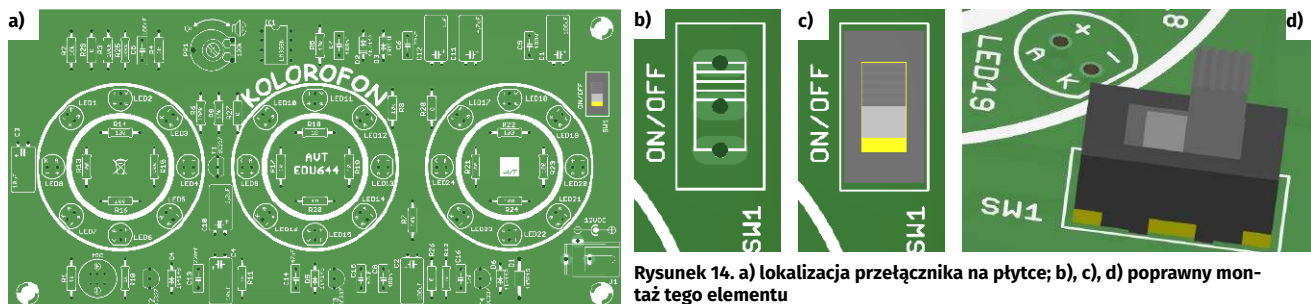
Fotografia 3. Wyprowadzenia mikrofonu elektretowego; wyprowadzenie połączone z metalowym korpusem stanowi biegun ujemny („minus“)

- Dociśnij mikrofon do płytki, tak aby przylegał stabilnie do laminatu.
- Przylutuj element do płytki. Jeśli nie wiesz, jak to zrobić lub chcesz upewnić się, że wykonujesz tę czynność prawidłowo, przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.
- Usuń nadmiar wyprowadzeń za pomocą obcinaczek. Jeśli nie wiesz, jak to zrobić lub chcesz upewnić się, że wykonujesz tę czynność prawidłowo, przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.

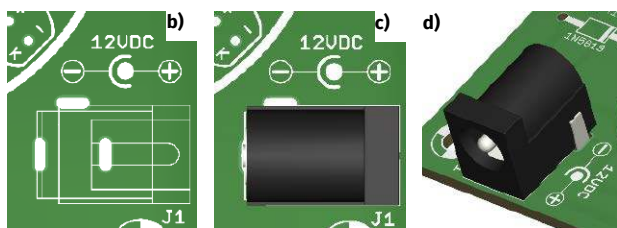
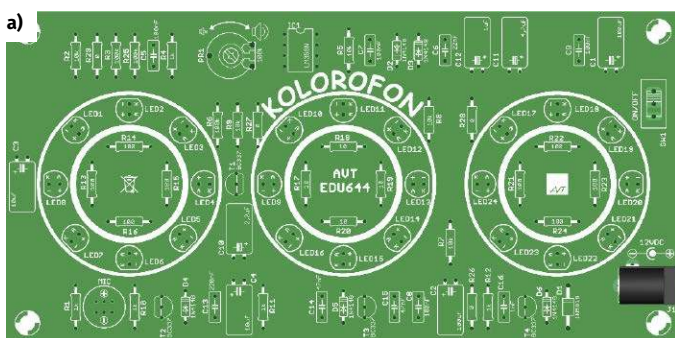
Montaż przełącznika zasilania

Zgodnie z informacjami z listy elementów zamontuj przełącznik SW1 na wskazanej lokalizacji. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Przełącznik zasilania SW1 to element, który przełącza połączenie pomiędzy pinem środkowym a jednym ze skrajnych, zależnie od położenia hebelka. W tego typu konstrukcji kierunek montażu nie ma żadnego znaczenia – niezależnie od tego, jak zostanie obrócony, będzie działał prawidłowo.
- Obrys na płytce PCB może zawierać dodatkowe linie symbolizujące położenie hebelka, ale służą one wyłącznie temu, by łatwo rozpoznać miejsce montażu. Nie są to oznaczenia biegunowości ani wymaganej orientacji.
- Włóż przełącznik do płytki tak, aby wszystkie trzy jego wyprowadzenia swobodnie przeszły na drugą stronę PCB. Piny przełącznika są sztywne i nie nadają się do wyginania, dlatego nie należy ich odchyłać w celu stabilizacji elementu.
- Ponieważ wyprowadzenia są sztywne, przełącznik trzeba przytrzymać podczas lutowania – można to zrobić ręką albo poprosić o pomoc kolegę lub opiekuna.
- Najpierw przylutuj środkowy pin przełącznika. Ten pojedynczy punkt lutowniczy pozwala ustabilizować komponent i kontrolować jego położenie względem płytki. Jeśli nie wiesz, jak to zrobić lub chcesz upewnić się, że wykonujesz tę czynność



Rysunek 14. a) lokalizacja przełącznika na płytce; b), c), d) poprawny montaż elementu



Rysunek 15. a) lokalizacja elementu na płytce; b), c) i d), poprawny montaż tego elementu

prawidłowo, przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.

- Po upewnieniu się, że przełącznik dobrze przylega do PCB, przylutuj pozostałe dwa wyprowadzenia. Dzięki temu unikniesz sytuacji, w której element zostanie przylutowany pod kątem lub z przerwą pomiędzy obudową a powierzchnią płytki.
- Prawidłowe luty powinny być gładkie, błyszczące i mieć kształt niewielkiego stożka. W przypadku wątpliwości przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.
- Po zakończeniu montażu sprawdź mechaniczne działanie hebelka. Przełącznik powinien poruszać się lekko i wyraźnie wskakiwać w dwie pozycje pracy.
- W przypadku tego typu przełącznika nie ma potrzeby przycinania jego wyprowadzeń. Po przylutowaniu pozostaw je w oryginalnej długości.

Montaż złącza zasilania

Zgodnie z informacjami z listy elementów zamontuj gniazdo zasilania J1 na wskazanej lokalizacji. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Złącze zasilania typu baryłkowego (ang. barrel jack) 5,5/2,1 mm to element służący do doprowadzenia napięcia z zewnętrznego zasilacza. Składa się z trzech wyprowadzeń: pinu środkowego (zwykle dodatniego), styku zewnętrznego (masa) oraz styku przełączanego, który może rozłączać obwód po włożeniu wtyku.
- Złącze nie jest symetryczne, dlatego jego orientacja na płytce ma znaczenie. Należy je zamontować zgodnie z obrysem na PCB, tak aby otwór na wtyk był skierowany na zewnątrz płytki (zgodnie z projektem obudowy lub krawędzią laminatu).
- Na płytce drukowanej, obudowie urządzenia lub tabliczce znamionowej polaryzacja złącza powinna być jednoznacznie oznaczona (rysunek 15).
- Włóż złącze do płytki tak, aby wszystkie wyprowadzenia swobodnie przeszły przez otwory. Element powinien przylegać

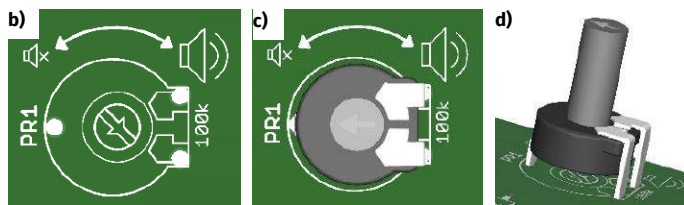
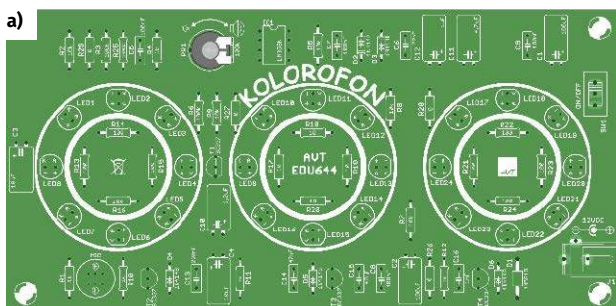
również do laminatu – zapewni to odpowiednią wytrzymałość mechaniczną podczas wkładania i wyjmowania wtyku.

- Ze względu na dużą masę i sztywność wyprowadzeń, złącze należy przytrzymać podczas lutowania. W przeciwnym razie może się odchylić i zostać przylutowane pod kątem.
- Najpierw przylutuj jedno z wyprowadzeń. Pozwoli to ustabilizować element i ewentualnie skorygować jego położenie względem płytki.
- Po sprawdzeniu, że złącze jest ustawione prosto i przylega do PCB, przylutuj pozostałe wyprowadzenia. Pady mogą wymagać nieco większej ilości cyny i dłuższego nagrzewania ze względu na większą powierzchnię metalową.
- Prawidłowe luty powinny być gładkie, błyszczące i dokładnie obejmować zarówno wyprowadzenie, jak i pole lutownicze. W przypadku wątpliwości przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.
- Po zakończeniu montażu sprawdź mechaniczne działanie złącza, wkładając wtyk zasilacza. Powinien wchodzić z wyraźnym oporem i stabilnie się trzymać, bez luzów.
- W przypadku tego typu złącza nie ma potrzeby przycinania wyprowadzeń – są one fabrycznie krótkie i przystosowane do montażu w płytce drukowanej.

Montaż potencjometru

Zgodnie z informacjami z listy elementów przylutuj potencjometr we wskazanej lokalizacji. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Potencjometr montażowy typu PT-10 jest elementem regulacyjnym, który umożliwia płynną zmianę rezystancji w obwodzie. Posiada trzy wyprowadzenia: dwa skrajne odpowiadają za całkowitą rezystancję, a środkowe (suwak) umożliwia uzyskanie regulowanej rezystancji zależnej od położenia osi.
- Orientacja potencjometru ma znaczenie mechaniczne. Należy zamontować go zgodnie z obrysem na PCB, tak aby oś regulacyjna była dostępna od właściwej strony płytki.



Rysunek 16. a) lokalizacja elementu na płytce; b), c) i d) poprawny montaż tego elementu z uwzględnieniem kierunku

- Obrys na płytce może wskazywać położenie korpusu oraz kierunek regulacji. Nie są to oznaczenia biegunowości, ale pomagają poprawnie ustawić element względem pozostałych komponentów.
- Włóż potencjometr do płytki tak, aby wszystkie trzy wyprowadzenia oraz ewentualne piny pozycjonujące swobodnie przeszły przez otwory. Nie należy wyginać wyprowadzeń na siłę.
- Ze względu na sztywność wyprowadzeń potencjometr należy przytrzymać podczas lutowania, aby nie został przylutowany pod kątem.
- Najpierw przylutuj jedno z wyprowadzeń (najlepiej środkowe), co pozwoli ustabilizować element i w razie potrzeby skorygować jego ustawienie.
- Po upewnieniu się, że potencjometr przylega równo do laminatu, przylutuj pozostałe wyprowadzenia. Dzięki temu uzyskasz solidne i trwałe połączenie mechaniczne.
- Prawidłowe luty powinny być gładkie, błyszczące i mieć kształt niewielkiego stożka. W przypadku wątpliwości przeczytaj sekcję *Zanim przystąpisz do montażu...* oraz zapoznaj się z dostępnymi poradnikami, w szczególności z instrukcjami BHP.
- Po zakończeniu montażu sprawdź działanie potencjometru, obracając jego oś. Ruch powinien być płynny i bez zacięć.
- Jeśli potencjometr jest wyposażony w osobny wałek regulacyjny, należy go zamontować po zakończeniu lutowania. Wałek wciska się w gniazdo osi potencjometru aż do wyczuwalnego oporu.
- Podczas montażu wałka nie należy używać nadmiernej siły ani narzędzi – może to uszkodzić mechanizm potencjometru.
- Po zamontowaniu wałka sprawdź jego osadzenie oraz działanie – powinien obracać się razem z osią potencjometru, bez luzów i bez ryzyka wypadnięcia.

Pierwsze podłączenie zasilania do płytki

Zanim włożymy układ scalony do podstawki, warto na chwilę podłączyć zasilanie do jeszcze „nieukończonych” płytki i sprawdzić, czy napięcie na złączu zasilania nie spada do zera – taki spadek oznaczałby błąd montażowy lub zwarcie w obwodzie zasilania.

Przed przystąpieniem do opisanych poniżej czynności koniecznie zapoznaj się z dokumentem Zagadnienia BHP związane z uruchamianiem zmontowanego układu



Rysunek 17. Pomiar napięcia na odpowiednich wyprowadzeniach elementów D1 i D6

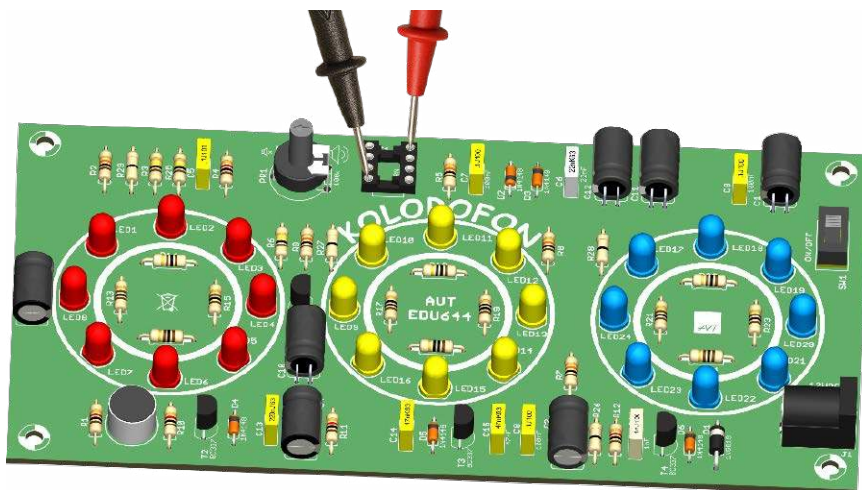
<https://elportal.pl/files/2026/02/15/3682-zagadnienia-bhp-zwiazane-zuruchamianiem-zmontowanego-ukladu.pdf>

Na początku ustawiamy włącznik w pozycji OFF i mierzymy napięcie na wejściu, czyli bezpośrednio na stykach gniazda J1. Jeśli pojawia się tam spodziewana wartość (napięcie znamionowe zasilacza), oznacza to, że źródło zasilania działa poprawnie. Następnie przełączamy włącznik zasilania oznaczony na płytce desygnatorem SW1 w pozycję ON i sprawdzamy, czy napięcie nie zanika. Spoglądając na układ ścieżek na płytce – mając ją w ręku bądź korzystając ze schematu montażowego (rysunek 2) – można zauważyć, że biegun ujemny zasilania połączony jest w pierwszej kolejności z anodą diody D6 a biegun dodatni, za pośrednictwem włącznika SW1 do anody diody D1. Tam też można zmierzyć obecność tego napięcia (rysunek 17).

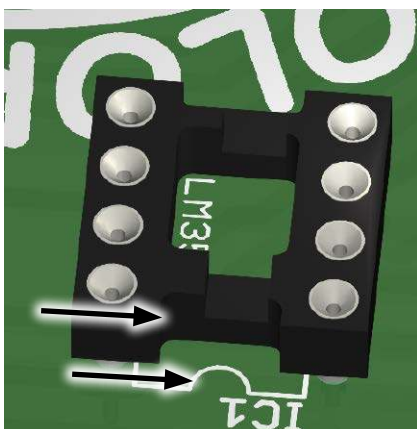
Zanik napięcia wskazywałby na usterkę montażową zbudowanego układu, na przykład przypadkowe zwarcie powstałe przez połączenie cyną dwóch sąsiednich pól

lutowniczych. W namierzeniu takich pomyłek bardzo pomaga wspomniany przed chwilą schemat montażowy (rysunek 2), na którym dokładnie widać przebieg ścieżek, rozmieszczenie padów oraz to, które połączenia powinny istnieć, a których być nie powinno. Po odnalezieniu usterki należy ją oczywiście wyeliminować. Gdy napięcie wejściowe jest prawidłowe i stabilne, upewniamy się, że po ustawieniu włącznika w pozycję ON pojawia się ono również na odpowiednich pinach podstawki pod układ scalony.

Numery pinów zasilających zazwyczaj da się odczytać ze schematu ideowego danego układu. Na rysunku 20 widać, że zasilanie podawane jest za pomocą rezystora R29 o wartości 0 Ω (zwórka) na wyprowadzenie numer 8 układu scalonego IC1, a minus zasilania podłączony jest do wyprowadzenia 4 tego układu. W celu sprawdzenia obecności napięcia zasilającego na podstawce przeznaczonej dla układu scalonego IC1 sondę czerwoną multimetru nastawiono na funkcję woltomierza napięć stałych



Rysunek 18. Pomiar napięcia zasilającego układ scalony na odpowiednich wyprowadzeniach podstawki po ustawieniu przełącznika SW1 w pozycję ON. Użyj schematu ideowego, montażowego lub noty katalogowej układu, by ustalić właściwe wyprowadzenia lub skorzystaj z rysunku powyżej



Rysunek 19. Przed zamontowaniem układu scalonego w podstawce należy upewnić się, że znaczniki kierunku montażu – na płytce, w podstawce i obudowie układu – znajdują się w tej samej pozycji

(VDC) w zakresie „do 20 V” przykładamy do otworu 8, a sondę czarną do otworu numer 4 tej podstawki (**rysunek 18**). Powinniśmy odczytać dodatnią wartość liczbową zbliżoną do napięcia nominalnego zasilacza.

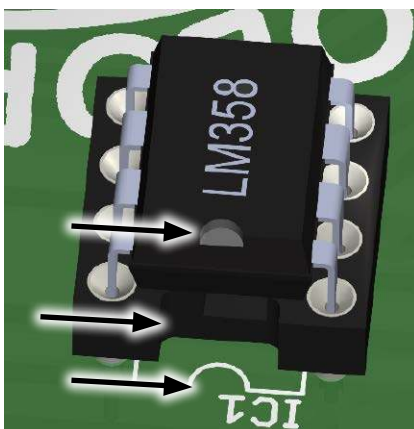
Takie testowe podłączenie zasilania oraz pomiary wykonane na źródle zasilania i na odpowiednich pinach podstawki pozwalają wstępnie zweryfikować poprawność montażu oraz uchronić najdroższe i najbardziej wrażliwe elementy przed uszkodzeniem.

Montaż układu scalonego w podstawce

Jeśli woltomierz podczas pomiaru napięcia na źródle zasilania oraz na odpowiednich pinach podstawki wskazał właściwe wartości (tu około 12 V) możesz odłączyć zasilacz, wyłączyć przełącznik SW1 i zamontować układ w podstawce. W przeciwnym wypadku, musisz odłączyć zasilanie, odnaleźć i naprawić usterkę montażową.

Włożenie układu scalonego do podstawki wydaje się łatwe, ale trzeba przy tym zachować uwagę i ostrożność. Ważne jest, aby układ był skierowany we właściwą stronę (co wytłumaczę za moment) oraz by wszystkie jego nóżki trafiły dokładnie w otwory podstawki. Nóżki nie mogą się wygiąć ani tym bardziej złamać. A gdyby jednak któraś się urwała, czasem udaje się ją zastąpić kawałkiem odciętego pinu z innego elementu już przylutowanego do płytki.

Drugą, obok ostrożności podczas montażu sprawą, o jaką należy zadbać, to właściwy kierunek montażu układu scalonego w podstawce. W tym celu należy przypilnować, by kropka lub wycięcie na układzie scalonym, wskazujące kierunek montażu, pokrywało się z pozostałymi znacznikami w podstawce oraz na warstwie opisowej PCB (**rysunek 19**). Gdy lokalizacja znacznika



na układzie scalonym zgadza się z pozostałymi, można przystąpić do wciśnięcia układu w podstawkę.

Podsumowanie montażu

Po ukończeniu montażu sprawdź, proszę, czy wszystkie połączenia lutowane są błyszczące i nie ma zimnych lutów oraz czy żadne sąsiednie pola lutownicze nie są ze sobą błędnie połączone.

W przypadku wątpliwości dotyczących poprawności wykonanych połączeń lutowanych skorzystaj z poradnika *Lutowanie komponentów przewlekanych do płytki drukowanej* (<https://elportal.pl/files/2026/02/15/3678-lutowanie-komponentow-przewlekanych-do-plytki-drukowanej.pdf>).

Ale właściwie... dlaczego to działa? Albo inaczej...

Tym razem, zamiast szczegółowo omawiać funkcję każdego elementu, spojrzymy na układ całościowo i wyróżnimy jego najważniejsze bloki funkcjonalne. To podejście odbiega od klasycznego „szkolnego” analizowania schematów, ale dzięki temu może być znacznie ciekawsze. Zamiast powielać teorię, potraktujmy kolorofon jako wdzięczny obiekt do eksperymentów i samodzielnych obserwacji – być może to dobry kierunek także na przyszłość. Schemat kolorofonu LED (**rysunek 20**) jest nieco bardziej rozbudowany niż zwykle, dlatego takie ujęcie wydaje się szczególnie uzasadnione.

W górnej części schematu znajduje się tor zasilania z gniazdem J1, wyłącznikiem SW1 oraz diodą D1 zabezpieczającą układ przed odwrotną polaryzacją. Kondensatory filtrujące zapewniają stabilne napięcie pracy. Po lewej stronie umieszczono mikrofon elektretowy wraz z elementami wejściowymi,

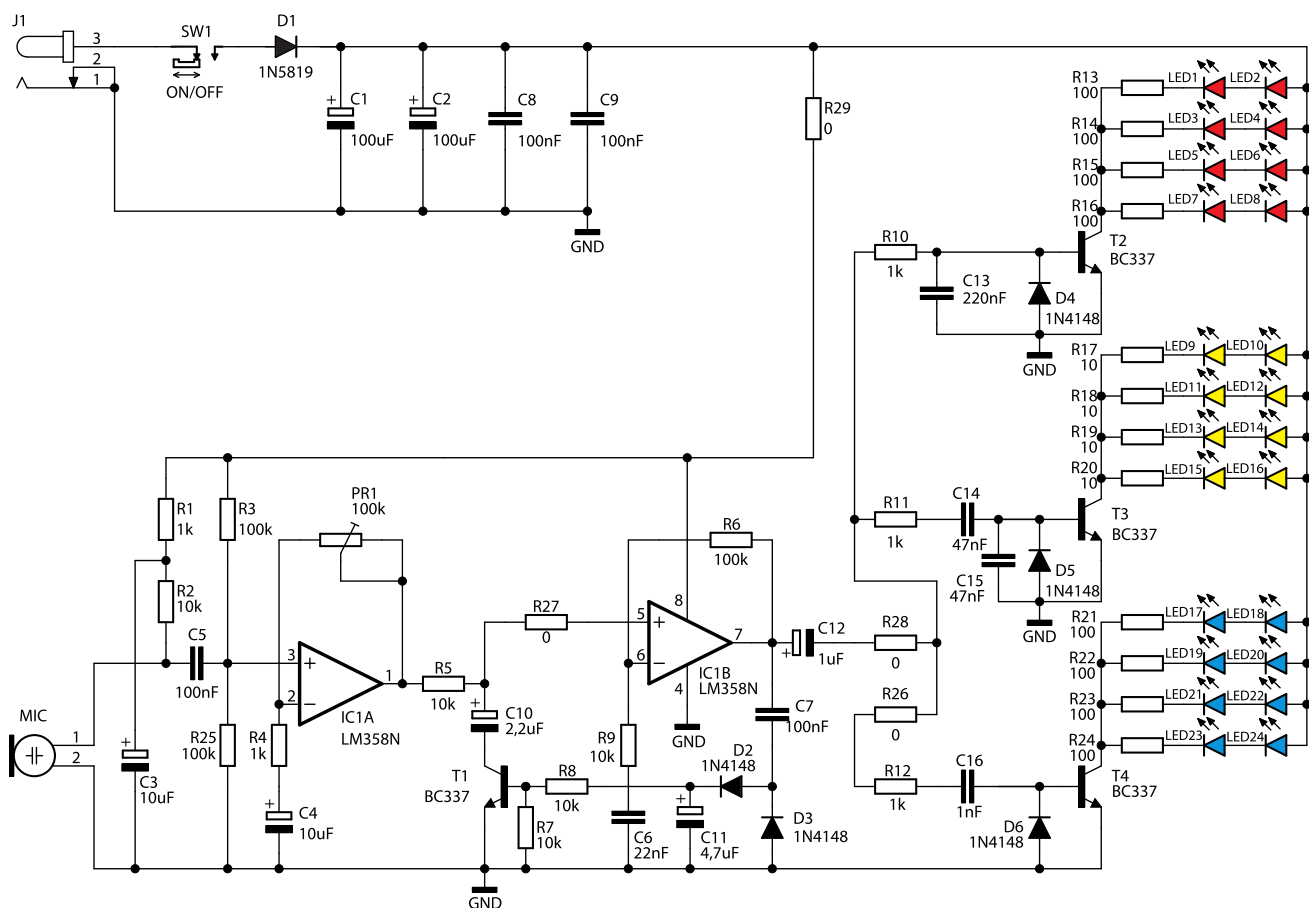
które doprowadzają sygnał akustyczny do pierwszego wzmacniacza operacyjnego. W tej części znajduje się także potencjometr PR1, umożliwiający regulację czułości urządzenia.

Sygnał akustyczny, po wzmocnieniu i wstępnym ukształtowaniu, trafia równolegle do trzech torów sterujących tranzystorami T2...T4. Tory te różnią się stałymi czasowymi obwodów RC (rezystor-kondensator), dlatego każda grupa diod LED reaguje inaczej na zmiany sygnału – jedne bardziej na szybkie impulsy, inne na wolniejsze zmiany amplitudy. Nie jest to jednak klasyczny podział na trzy ściśle wydzielone pasma częstotliwości, jak ma to miejsce w typowych kolorofonach z filtrami pasmowymi, lecz uproszczone rozwiązanie wykorzystujące właściwości obwodów RC do kształtowania odpowiedzi czasowej, dające atrakcyjny efekt świetlny.

Całość jest klasycznym układem analogowym, w którym wszystkie operacje – wzmocnienie, kształtowanie przebiegu i sterowanie – realizowane są za pomocą wzmacniaczy operacyjnych, tranzystorów oraz elementów biernych. Takie rozwiązanie ma dużą wartość dydaktyczną, ponieważ w przejrzysty sposób pokazuje, jak z prostych bloków funkcjonalnych można zbudować efektowny układ reagujący na dźwięk.

Układ może być także dobrą bazą do samodzielnych eksperymentów. Warto na przykład porównać działanie trzech torów sterujących, zmieniając wartości kondensatorów i obserwując wpływ na sposób świecenia diod LED. Proste obliczenie stałych czasowych $\tau = R \cdot C$ pozwala oszacować, które gałęzie będą reagowały szybciej, a które wolniej. Ciekawym doświadczeniem jest również podanie na wejście sygnału z generatora (zamiast mikrofonu) i obserwacja reakcji układu dla różnych częstotliwości oraz amplitud.

Do takich prób można z powodzeniem wykorzystać smartfon z prostą aplikacją generującą sygnał audio, np. Frequency Sound Generator lub podobny generator funkcji. Aplikację tego typu umożliwiają generowanie przebiegów sinusoidalnych, prostokątnych czy trójkątnych w szerokim zakresie częstotliwości. Podając sygnał z telefonu (np. przez zbliżenie głośnika do mikrofonu), można wykonać kilka prostych doświadczeń: zmieniając częstotliwość sygnału, obserwować, która grupa diod reaguje najsilniej; zmieniając amplitudę (głośność), sprawdzić próg zadziałania poszczególnych torów; porównać reakcję układu dla



Rysunek 20. Schemat ideowy układu

różnych przebiegów (sinus, prostokąt); a także wykonać przemianę częstotliwości (sweep) i obserwować płynne zmiany efektu świetlnego.

Dodatkowo można zmierzyć napięcia w wybranych punktach układu (np. na wyjściu wzmacniacza operacyjnego i bazach tranzystorów), aby zobaczyć, jak zmienia się kształt sygnału w kolejnych blokach. Połączenie takich obserwacji z prostymi obliczeniami stałych czasowych pozwala w praktyce zaobserwować wpływ obwodów RC na przetwarzanie sygnału oraz lepiej zrozumieć, dlaczego układ reaguje różnie na szybkie i wolne zmiany dźwięku.

Podsumowanie

Zmontowany kolorofon LED to układ prosty, a jednocześnie bardzo efektowny

– dwa wzmacniacze operacyjne, kilka tranzystorów i diod LED wystarczy, by zamienić dźwięk w światło reagujące na otoczenie. Każdy sygnał akustyczny – muzyka, mowa czy pojedyncze stuknięcie – natychmiast znajduje swoje odzwierciedlenie w pracy diod.

Warto jednak zwrócić uwagę na coś jeszcze. Do eksperymentów z takim układem nie potrzeba specjalistycznego sprzętu laboratoryjnego. Wystarczy smartfon. To niewielkie urządzenie może pełnić rolę generatora sygnałów, źródła dźwięku, a nawet prostego analizatora. Pozwala w łatwy sposób sprawdzić, jak układ reaguje na różne częstotliwości, amplitudy czy przebiegi. Innymi słowy – mamy w kieszeni narzędzie, które w prosty sposób zamienia zabawę w światło eksperymentowanie.

Budując kolorofon, nie tylko ćwiczysz lutowanie i poznajesz działanie układów analogowych. Otrzymujesz także punkt wyjścia do własnych prób – takich, które można przeprowadzić od razu, bez dodatkowych urządzeń. To właśnie w takich doświadczeniach najłatwiej zobaczyć, jak teoria przekłada się na praktykę.

Mamy nadzieję, że ten projekt dostarczył Ci satysfakcji i zachęcił do dalszego eksperymentowania. Bo elektronika zaczyna się tam, gdzie pojawia się ciekawość – a najlepsze pomysły często rodzą się z prostych prób.

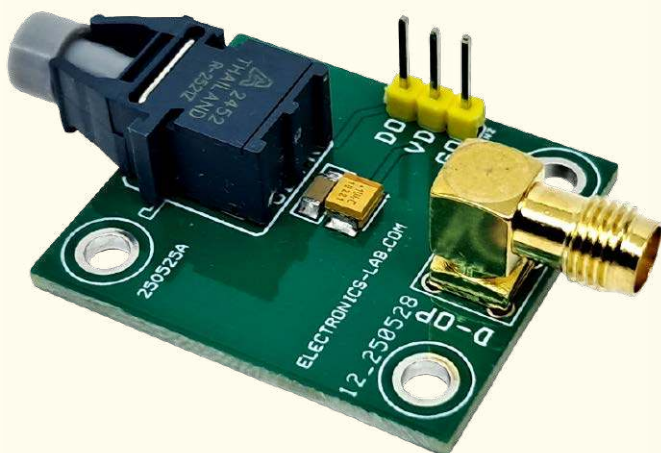
Do zobaczenia w następnym numerze! ■

Mariusz Ciszewski

REKLAMA

Mnóstwo doskonałych projektów, tylko na:

EPcom.pl



Wszechstronny odbiornik światłowodowy Versatile Link

Odbiornik HBBR-2521Z umożliwia odbiór danych TTL z prędkością do 5 MBaud na dystansie do 53 metrów. Wykorzystuje zasilanie 5 V i charakteryzuje się wysoką odpornością na zakłócenia elektromagnetyczne (EMI) oraz izolacją napięciową. Moduł ten pozwala użytkownikom na stworzenie uniwersalnego łącza światłowodowego dla aplikacji wymagających niskokosztowych rozwiązań. Odbiornik odbiera dane przez łącze światłowodowe i wyprowadza sygnały w standardzie TTL. Obsługuje dystanse transmisji do 19 metrów przy użyciu standardowego światłowodu oraz do 53 metrów przy zastosowaniu kabla o podwyższonych parametrach.

Wszechstronny nadajnik światłowodowy Versatile Link

Nadajnik HBBR-1521Z umożliwia przesyłanie danych TTL z prędkością do 5 MBaud na dystansie do 53 metrów. Wykorzystuje zasilanie 5 V oraz diodę LED 660 nm, oferując wysoką odporność na zakłócenia elektromagnetyczne (EMI) i izolację napięciową. Moduł ten pozwala użytkownikom na ustanowienie uniwersalnego łącza światłowodowego dla aplikacji wymagających tanich rozwiązań o wysokiej odporności. Przesyła dane źródłowe TTL przez połączenie światłowodowe, obsługując dystanse do 19 metrów ze standardowym kablem oraz do 53 metrów z kablem o podwyższonych parametrach.

Niektóre projekty aktualnie dostępne tylko dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl:

1. Bezprzewodowy kontroler samochodu-roboty Bluetooth
2. Kontroler ramienia robotycznego wykorzystujący bezprzewodowego pada z konsoli PS3
3. 7-segmentowy mini zegar wykorzystujący PIC16F628A i DS1307 RTC
4. Światło LED oparte na czujniku zbliżeniowym
5. OLEDUINO – wyświetlacz OLED kompatybilny z Arduino
6. Inteligentny regulator lutownicy – precyzyjny regulator grzałki
7. Wskaźnik poziomu paliwa z wyświetlaczem OLED
8. Bezprzewodowy odbiornik wilgotności i temperatury
9. Przerwania zewnętrzne (sprzętowe) i przerwania zegara w MicroPython
10. Laserowy czujnik odległości z wyświetlaczem OLED i RP2040
11. Inklinometr z 17-segmentowym wyświetlaczem słupkowym
12. Izolowany repeater USB – USB 2.0
13. Knight Rider Light – 16 diod LED dużej mocy (kompatybilny z Arduino)
14. Dźwięk do kolorowych efektów świetlnych (kompatybilny z Arduino)
15. Nowy i ulepszony licznik Geigera – teraz z Wi-Fi!
16. Detektor zalania
17. Lampa nastrojowa LED o dużej mocy
18. Kontroler dzwonów kościelnych
19. Arduino Nano – włączanie/wyłączanie urządzeń za pomocą pilota na podczerwień (dwa kanały)
20. Lampa sufitowa LED z czujnikiem ruchu PIR – kompatybilna z Arduino
21. Inteligentny ściemniacz LED z Bluetooth – 4-kanałowy włącznik/wyłącznik Bluetooth
22. Czterokanałowy izolator cyfrowy, wzmacniony, szybki, o niskim poborze mocy
23. Sterowanie prędkością, kierunkiem i zatrzymaniem silnika DC z modułem RF NRF24L01
24. Nadajnik zdalnego sterowania z pojedynczym joystickiem wykorzystujący NRF24L01
25. 8-kanałowy zdalny nadajnik RF z protokołami: Holtek i szeregowym
26. 8-kanałowy zdalny odbiornik RF z protokołami: Holtek i szeregowym
27. Pojemnościowy czujnik wilgotności do konwertera wyjścia analogowego
28. Mostek H dla wysokiej mocy szczotkowego silnika prądu stałego z czujnikiem prądu
29. Przetwornica DC-DC buck 12...75 V na 10 V na wyjściu
30. Czujnik prądu low-side 10 µA...10 mA
31. Kontroler ramienia robota z bezprzewodowym pilotem PS3
32. Termiczny czujnik masowego przepływu powietrza – anemometr statotemperaturowy
33. Precyzyjny wzmacniacz transimpedancjowy z przełączanym integratorem
34. Wysokowydajny monofoniczny wzmacniacz audio klasy D o mocy 20 W
35. Kontroler pełnego mostka z przesunięciem fazowym i prostowaniem synchronicznym wykorzystujący UCC28950

Miesięcznik „Elektronika dla Wszystkich” (12 numerów w roku) jest wydawany we współpracy z kilkoma redakcjami zagranicznymi



Wydawnictwo:
AVTKorporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Redaktor naczelny:
Mariusz Ciszewski
mariusz.ciszewski@elportal.pl

Adres redakcji:
03-197 Warszawa, ul. Leszczyńska 11
e-mail: edw@elportal.pl, www.elportal.pl

Dział reklamy:
Katarzyna Gugala
katarzyna.gugala@elportal.pl, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobański
jakub.sobanski@elportal.pl

Sekretarz redakcji:
Dariusz Welik
dariusz.welik@elportal.pl

Copyright AVTKorporacja 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, 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)
www.ulubionykiosk.pl

AT-AD269S

Mikroskop cyfrowy
z ekranem 10 cali,
powiększenie do 5000×,
5 obiektywów i endoskop
ANDONSTAR AD269S-M



AT-AD409PRO

Mikroskop do lutowania
z profesjonalnym
metalowym stojakiem,
ekran 10,1 cala,
powiększenie do 300×, HDMI
ANDONSTAR AD409Pro



BESTSELLERY sklepu AVT – sklep.avt.pl

Mikroskopy cyfrowe dla elektroników

Rabat dla Czytelników EdW
przy zakupie podaj kod **EdW2505MC**

Kod ważny do 30.09.2025

-3%

Rabat dla Prenumeratorów EdW
przy zakupie podaj numer prenumeraty

-6%

AT-AD246S-M

Mikroskop cyfrowy 7 cali
z powiększeniem:
60...240×, 18...720×,
1560...2040×
ANDONSTAR AD246S-M



AT-AD407

Mikroskop cyfrowy 7 cali,
powiększenie do 270×
ANDONSTAR AD407



AT-AD249S-M

Mikroskop cyfrowy 10 cali
z powiększeniem:
60...240×, 18...720×, 1560...2040×
ANDONSTAR AD249S-M



AT-AD210

Mikroskop cyfrowy 5...260×
z wyświetlaczem 10,1 cala
ANDONSTAR AD210



Subscribe to Elektor's newsletter and get the chance to

WIN

a Raspberry Pi Pico W board



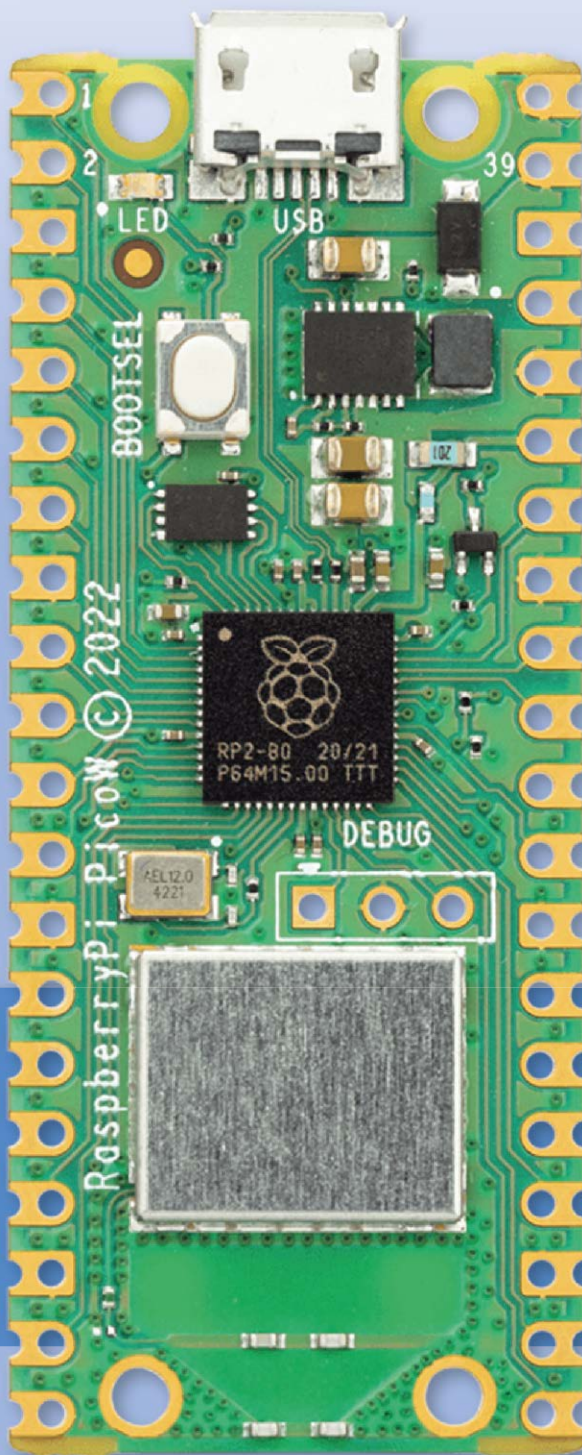
www.elektor.com/eda



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