

ELEKTRONIKA

dla wszystkich

nr 10/2023 (333) • październik • www.elportal.pl

DIY PLUS
tylko dla prenumeratorów

PROJEKTY dla elektroników

- ▶ Ozdoby Bożonarodzeniowe
– Hej, podziwiajcie.
Ośiem dekoracji świątecznych LED
- ▶ MiniHEART: miniaturowy symulator bicia serca
- ▶ Gdy chcesz, aby Ci nie przeszkadzać.
Jestem zajęty – dla wszystkich!

DIY dla wszystkich

- ▶ Pierwszy taki telefon, wielkości palca z ekranem dotykowym E-INK
- ▶ Zdalnie sterowany samochodzik zabawka
- ▶ System nadzoru temperatury baterii w pojazdach elektrycznych

TUTORIALE

- ▶ Praktyczny kurs op-ampów
- ▶ Audio OUT: Wzmacniacz audio do Theremina
- ▶ Chirurgia obwodowa: Problemy z symulacją cyfrowego układu dzielenia przez dwa
- ▶ Know-how: Fazowa regulacja mocy
- ▶ Porady laboratoryjne: Punkt pracy tranzystora bipolarnego
- ▶ KickStart: Cewki indukcyjne, czym są?
- ▶ Pokój Nauczycielski
- ▶ Ekscytacje Maxa: Migające diody LED i śliniacy się inżynierowie

Wyhoduj to sam



Ośiem dekoracji świątecznych LED

Zacznij budować już teraz, to zdążysz na święta

ISSN 1425-1698 Indeks 33362X
9 477142 5169238
16,90 zł (w tym 8% VAT)

EP.com.pl

Największy portal dla elektroników konstruktorów



Król automatyki
jest w Tobie

AutomatykaB2B.pl

FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

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

www.piekarz.pl





Najbardziej popularne kity AVT

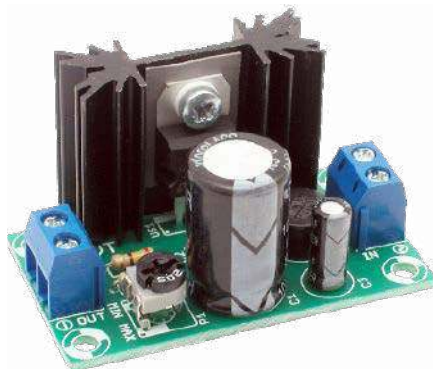
Poznaj listę **TOP 100** na www.elportal.pl/kityavt



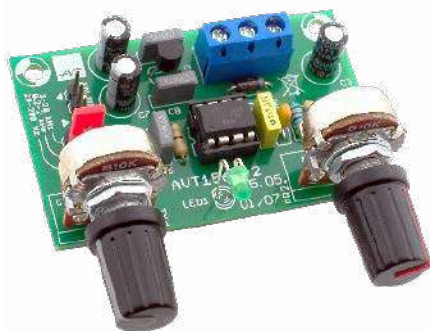
AVT1476 Automatyczny włącznik zmierzchowy
<https://sklep.avt.pl/avt1476.html>



AVT735 Regulator mocy PWM 10 A
<https://sklep.avt.pl/avt735.html>



AVT1066 Miniaturowy zasilacz uniwersalny z LM317
<https://sklep.avt.pl/avt1066.html>



AVT1569 Generator akustyczny 20 Hz...20 kHz
<https://sklep.avt.pl/avt1569.html>



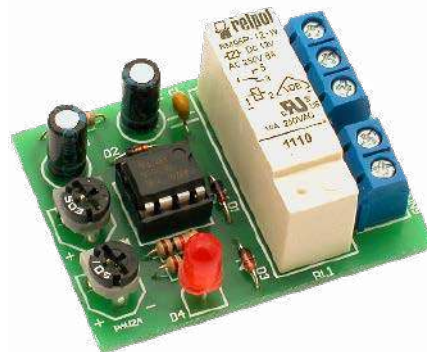
AVT1023 Przedwzmacniacz gramofonowy o charakterystyce RIAA
<https://sklep.avt.pl/avt1023.html>



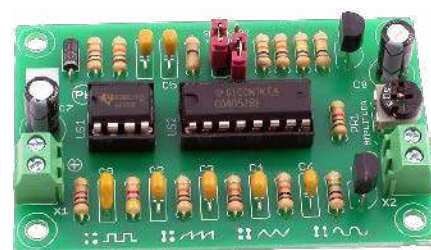
AVT5540 Radio FM z RDS
<https://sklep.avt.pl/avt5540.html>



AVT1594 Wzmacniacz mocy 2x45 W z STK4182
<https://sklep.avt.pl/avt1594.html>



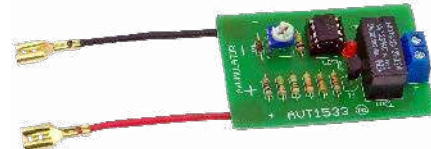
AVT1459 Uniwersalny układ czasowy
<https://sklep.avt.pl/avt1459.html>



AVT1327 Mini generator funkcyjny
<https://sklep.avt.pl/avt1327.html>



AVT1597/3 Wzmacniacz audio z układem TDA2050 35 W
<https://sklep.avt.pl/wzmacniacz-audio-z-ukladem-tda2050-zestaw-do-samodzielnego-montazu.html>



AVT1533 Zabezpieczenie akumulatora 12 V przed rozładowaniem
<https://sklep.avt.pl/avt1533.html>



AVT1661 Elektroniczna kostka do gry
<https://sklep.avt.pl/avt1661.html>



Pełna oferta na: sklep.avt.pl

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

PRENUMERATA

NA START
DO 6 WYDAŃ
GRATIS!

Cena drukowanej prenumeraty rocznej na start wynosi 185,90 zł
Przy zamówieniu prenumeraty dwuletniej za 304,20 zł
oszczędność wynosi równowartość sześciu wydań EdW

PO 5 LATACH
ZA PÓŁ CENY

Przedłuż prenumeratę drukowaną po zalogowaniu się do swojego panelu na www.UlubionyKiosk.pl/logowanie, gdzie znajdziesz atrakcyjną ofertę, która uwzględnia przysługujące Ci zniżki za lojalność. Po 5 latach nieprzerwanej prenumeraty **otrzymasz rabat 50% na drukowaną prenumeratę dwuletnią**

NOWOŚĆ! PRENUMERATA EdW+

Rozpocznij przygodę z elektroniką – poznaj jej podstawy, zamawiając roczną prenumeratę drukowaną EdW wraz z Praktycznym Kursem Elektroniki (PKE)

Do wysyłki prenumeraty dołączymy zestaw edukacyjny EDW A09 KPL, na który składają się:

1. projekt – układ elektroniczny samodzielnie uruchamiany przez kursanta. Wszystkie układy są montowane na dołączonej płytce stykowej, do której wkłada się „nóżki” elementów na wcisk,
2. pendrive z wykładami i materiałami multimedialnymi kursu PKE.
3. zasilacz płytek stykowych AVT3072 C
4. oraz zasilacz impulsowy 12 V, 1,4 A

Cena prenumeraty EdW+ wynosi **280,90 zł**

TYLKO prenumeratorzy* otrzymują pełny dostęp do:

ARCHIWUM

cyfrowego archiwum EdW na elportal.pl/archiwum



projektów w zbiorze DIY+ na elportal.pl/diy

* Promocja z dostępem do archiwum EdW oraz projektów DIY+ dotyczy płatnej prenumeraty drukowanej lub płatnej e-prenumeraty EdW zamawianej na www.UlubionyKiosk.pl bądź przelewem na konto Wydawnictwa AVT. Po odnotowaniu płatności wysyłamy mailowo kod dostępu, za pomocą którego zalogujesz się na elportal.pl

Zamów prenumeratę lub e-prenumeratę na www.UlubionyKiosk.pl/prenumerata

Kontakt ws. prenumeraty: 22 257 84 22 (godz. 10.00–14.00), prenumerata@avt.pl
Kontakt merytoryczny ws. kursu PKE: kity@avt.pl • Konto bankowe: AVT-Korporacja sp. z o.o.
03-197 Warszawa, ul. Leszcynowa 11, ING Bank Śląski 18 1050 1012 1000 0024 3173 1013



8



17



24



32



88

Projekty dla elektroników:

Ozdoby Bożenarodzeniowe – Hej, podziwiajcie.

Ośmiu dekoracji świątecznych LED 8

Wyhoduj to sam. Sterowane cyfrowo,

jednopojemnikowe rozwiązanie do upraw w pomieszczeniach 17

MiniHEART: miniaturowy symulator bicia serca 24

Gdy chcesz, aby Ci nie przeszkadzać. Jestem zajęty – dla wszystkich! 32

Tutoriale:

Praktyczny kurs op-ampów 36

Audio OUT: Wzmacniacz audio do Theremina, część 1 40

Chirurgia obwodowa:

Problemy z symulacją cyfrowego układu dzielenia przez dwa 43

Know-how. Fazowa regulacja mocy 48

Porady laboratoryjne: Punkt pracy tranzystora bipolarnego 54

Edukacja w EdW dla szkół i uczelni: Wykład 11.

Ujemne sprzężenie zwrotne i historia wzmacniaczy operacyjnych 60

KickStart: Część 3: Cewki indukcyjne, czym są? 69

Pokój Nauczycielski 74

Ekscytacje Maxa:

• Migające diody LED i śliniacy się inżynierowie 77

• Sprytne porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania .. 81

DIY dla wszystkich:

Pierwszy taki telefon, wielkości palca z ekranem dotykowym E-INK 82

Zdalnie sterowany samochodzik zabawka 86

System nadzoru temperatury baterii w pojazdach elektrycznych 88

DIY PLUS

Inteligentny dwucewkowy sterownik przekaźnika

blokującego – moduł przekaźnika bistabilnego 91

Zabezpieczenie nadprądowe/nadrozładowcze

do akumulatorów kwasowo-ołowiowych 91

Rubryki stałe:

Prenumerata 3

Od wydawcy 5

Poczt 6

A za miesiąc w październikowym EdW



* Cyfrowy moduł FX do gitary

Z tym cyfrowym modulem efektów będziesz tworzyć dźwięk jak profesjonalny muzyk. Wygenerujesz unikalne dźwięki po podłączeniu do różnych instrumentów, takich jak elektryczna gitara solowa, gitara basowa, skrzypce lub wiolonczela. Równie ciekawe efekty uzyskasz podłączając moduł do wyjścia przedwzmacniacza mikrofonowego.

* Dwie świąteczne Gwiazdki LED

Każda z tych dwóch ozdób będzie wyglądać spektakularnie na choince lub gdziekolwiek indziej. Doskonale sprawdzają się także samodzielnie, wymagając do działania jedynie zasilania z gniazda USB.

* Instalowanie i używanie MPLAB X

Chociaż oprogramowanie i sprzęt Arduino ułatwiły tworzenie projektów z mikrokontrolerami, to wielu zaawansowanych użytkowników woli używać profesjonalnych narzędzi, np. MPLAB X, szczególnie w przypadku układów serii PIC. W artykule przeprowadzimy Cię przez proces tworzenia pierwszego projektu w MPLAB X.

* Tani system automatyki z użyciem uP PIC16F676

W tym projekcie proponujemy tanie i proste rozwiązanie, które można zastosować w automatyce domowej jak i w procesach przemysłowych. Możliwości systemu pozwalają na zdalne włączenie bądź wyłączenie dowolnego odbiornika energii elektrycznej.

* Plus zwykła porcja intrygujących projektów DIY.

* Plus wiele artykułów w Twoich ulubionych cyklach Tutoriali.

**W kioskach
od 1 października**

Krótki zarys historii sklepu AVT

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

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

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

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

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

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

Wiesław Marciniak

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

Kieszonkowy generator audio

Ponieważ tekst wspomina filtr eliptyczny (filtr „pi” z podwójnym „daszkim” L+C), chciałem sprawdzić, co to takiego.

I tu przeżyłem szok – mimo poszukiwań nie znalazłem ogólnego schematu ideowego (dotyczyło to również filtrów sygnowanych nazwiskami: Bessela, Czebyszewa itd.). Było wszystko – charakterystyki, wzory na obliczanie pasma, ale schematu – brak.

Załączam skrypt ćwiczeń „Filtry elektryczne” jako przykład, jak NIE powinno się uczyć studentów. Dla mnie osobiście jego treść jest SZOKUJĄCA pod kątem braku umiejętności przekazywania wiedzy. Miałem z tym trochę do czynienia jak autor dwóch rozdziałów w skrypcie PW i prowadzący ćwiczenia dla studentów Inżynierii Chemicznej PW, ale to było dawno temu (plik do pobrania: <https://tiny.pl/cr9zq>).

Andrzej Nowicki

Nienaprawialny sprzęt

Może nie nadążam za rozwojem cywilizacji technicznej, ale irytuje mnie polityka producentów, która ma swoją nazwę „planned obsolescence” (planowana przestarzałość) i utrudnianie lub wręcz uniemożliwianie napraw sprzętu elektronicznego. Rozumiem dążenie producentów do ciągłego kreowania popytu na nowe modele urządzeń elektronicznych, ale przecież takie praktyki są szkodliwe społecznie. Rośnie dzięki temu PKB, ale kosztem drenowania kieszeni klientów i cierpi na tym ekologia. Czy środowisko elektroników, również poprzez takie media jak EdW i EP, nie powinno jakoś nawoływać do opamiętania się i żądać od organów władzy, by tworzyły prawo przeciwstawiające się takim praktykom?

AC

Red. Elektrycy na całym świecie dostrzegają to i są szczególnie wrażliwi na ten problem. Na przykład redaktor naczelny magazynu „Silicon Chip”, Nicholas Vinen, poświęcił temu tematowi swój wstępniak w tegorocznym lutowym wydaniu SC, w którym przedstawił szereg pomysłów dla prawodawców w Australii i zachęca ich do działania.

Oto parę szczególnie rażących przykładów takich praktyk:

- firma Apple zaprojektowała swoje iPhone'y w taki sposób, że trudno jest dokonywać samodzielnych napraw lub wymieniać baterie. Ponadto Apple blokuje dostęp do części zamiennych – ma podpisaną umowę na wyłączność z producentem chipów, firmą Intersil;
- firma Lenovo była krytykowana za umieszczenie w swoich laptopach systemu BIOS, który blokował nieautoryzowane części zamienne, co utrudniało naprawy laptopów przez osoby trzecie;
- firma HP (Hewlett-Packard) jest krytykowana za stosowanie programów, które blokują użytkownikom możliwość korzystania z zamiennych tuszy i tonerów, a także utrudniające jest serwisowanie drukarek.

Do tej samej kategorii spraw można zaliczyć dążenie producentów do zarabiania na ładowarkach. Dlatego każdy producent robi wszystko, żeby do jego sprzętu pasowała tylko jego ładowarka. W Unii Europejskiej jest wielka determinacja, żeby doprowadzić do ujednolicenia standardów wtyczek zasilających, żeby jedną ładowarką udało się obsłużyć dowolny sprzęt.

Trzeba jednak dostrzegać, że współczesna technologia wielkoseryjnej produkcji sprzętu elektronicznego często nie pozostawia fizycznych możliwości ingerencji serwisanta. Poza tym niektóre masowo wytwarzane gadżety elektroniczne są tak tanie, że ich naprawa nie ma sensu.

Wiesław Marciniak

Nowy sklep AVT

Jestem częstym klientem (gościem) sklepu AVT przy Leszczyńskiej 11. Lubię go odwiedzać, jest niemal wszystko, czego dusza elektronika zapagnie. Mieszkam w Warszawie, więc odwiedzam chętnie sklep stacjonarny i rzadko korzystam ze sklepu internetowego AVT.

Wczoraj po raz pierwszy od paru tygodni skorzystałem i tu niespodzianka – sklep całkiem się zmienił. Nie cierpię zmian, więc się trochę zdenerwowałem. Byłem jednak zaskoczony jak gładko mi szło poszukiwanie potrzebnych mi elementów. Widać profesjonalny nadzór nad budową sklepu. Ciekaw jestem, która firma była w to zaangażowana.

Red. Dziękujemy wszystkim autorom listów na temat nowego sklepu AVT. Były też uwagi i propozycje usprawnienia obsługi zdalnej, między innymi poprzez użycie AI. Przekazaliśmy załodze sklepu.

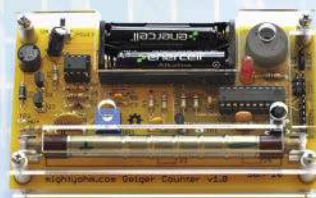
Patronat AVT

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

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

Elektor Bestsellers

SAVE UP TO
26% NOW!



www.elektor.com/sale/deals





Ozdoby Bożenarodzeniowe – Hej, podziwiajcie Osiem dekoracji świątecznych LED

Nasza mała dekoracyjna choinka LED z listopadowego numeru Silicon Chip-a w 2019 r. cieszyła się tak dużą popularnością, że postanowiliśmy zaprezentować nie jedną, nie dwie, ale siedem kolejnych świątecznych dekoracji, które możesz zmontować! Są one małe, tanie i łatwe do złożenia, więc dasz radę z łatwością zbudować w tym roku wszystkie osiem na Boże Narodzenie; albo chociaż kilka z nich.

Czytelnicy Silicon Chip-a zbudowali setki naszych Drzewek Bożenarodzeniowych (Tiny LED Christmas Tree) z numeru SC z listopada 2019 r. (siliconchip.com.au/Article/12086). Niektórzy kupili dziesięć lub więcej zestawów! Sami zrobiliśmy nawet kilka w domu dla naszych własnych choinek.

Nic dziwnego, że pozostają one tak popularnym projektem, ponieważ są łatwym sposobem na atrakcyjne udekorowanie choinki świetlnie wyglądającymi animowanymi ozdobami.

Zastanowił nas jednak list od Anthonyego i Annabel (luty 2020 r.). Dowiedzieliśmy się z niego, że dzieci w wieku zaledwie dziewięciu lat z powodzeniem zbudowały choinkę Tiny LED Xmas Tree.

Teraz nie ma wymówki, aby nie korzystać z dobrodziejstw konstrukcji SMD!

Nie tylko to. Anthony i Annabel podsunęli również doskonały pomysł, abyśmy dla dodatkowego urozmaicenia wykorzystali ten sam schemat elektryczny w różnych kształtach i kolorach.

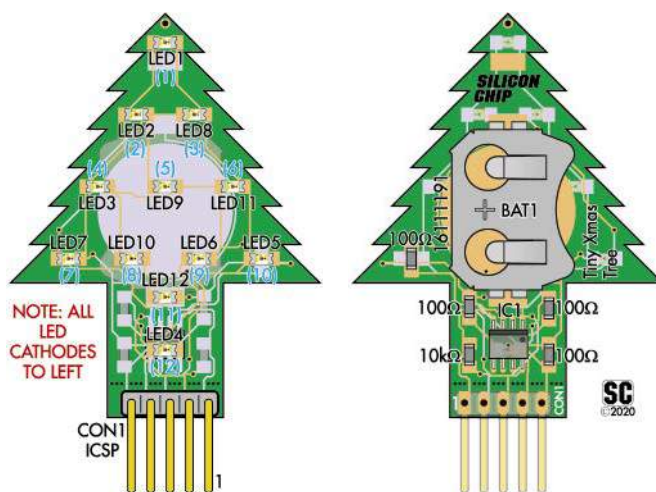
I właśnie to teraz zrobiliśmy. Te różnorodne ozdoby świąteczne są bardzo małe, ale idealne do dekoracji choinki. Wymyśliliśmy unikalne wzory ruchomych świateł pasujące do każdej ozdoby, a także dodaliśmy interaktywność, aby św. Mikołaj i jego renifery również mogły być częścią zabawy.

Obwód

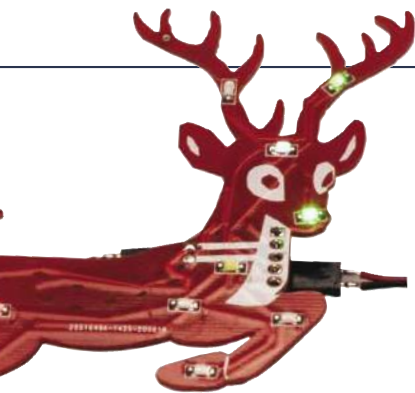
Schemat ideowy obwodu dla naszych nowych ozdób, pokazany na rysunku 1, jest zasadniczo taki sam jak w przypadku Świątecznego Drzewka opisanego w 2019 r.

Jeśli chcesz uzyskać więcej szczegółów na temat konkretnych założeń projektowych, zalecamy zapoznanie się z odnośnikiem.

W szczególności zapoznaj się z panelem na temat układu LED typu Charlieplexing (na stronie 48 wydania SC z listopada 2019 r., nazwa pochodzi od imienia projektanta obwodu, Charlesa Allena),



Na wypadek, gdybyś go uprzednio przegapił, oto projekt Choinki Tiny Christmas Tree z listopada 2019 roku, który zainspirował te nowe projekty. Wciąż jest to sam w sobie całkowicie realny i aktualny projekt. Może być używany samodzielnie lub w połączeniu z dowolną z nowych ozdób (siliconchip.com.au/Article/12086)



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/cfw5h>

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

aby dowiedzieć się, w jaki sposób sterujemy tak wieloma diodami LED z 8-końcówkowego mikroprocesora. W największym skrócie, polega to na ułożeniu diod w odpowiedniej matrycy i wykorzystaniu faktu, że wyjście mikroprocesora, oprócz stanu wysokiego i niskiego może znajdować się w stanie o wysokiej impedancji.

Obwód wykorzystuje IC1, 8-bitowy mikroprocesor PIC12F1572, zasilany bezpośrednio z 3 V litowej baterii pastylkowej CR2032. Ogniwo jest po prostu podłączone do styków zasilania IC1, styku 1 (VDD lub dodatni biegun zasilania) i styku 8 (VSS lub ujemny biegun zasilania).

Używamy PIC12F1572 z powodów wyjaśnionych w artykule „Nowy PIC” na stronie 83 SC, ale stworzyliśmy wsady

oprogramowania układowego, które pasują również do PIC12F675, więc możesz użyć dowolnego układu scalonego z wymienionych w tym projekcie.

Układ IC1 jest dostarczany w małej 8-stykowej obudowie SOIC SMD (small outline integrated circuit). Jest kompaktowy, ale wystarczająco łatwy w zastosowaniach. Cztery ze styków GPIO (wejście/wyjście ogólnego przeznaczenia) układu IC1 (styki 2, 3, 5 i 6) są podłączone do rezystorów 100 Ω , a następnie do matrycy 12 diod LED.

Pomiędzy każdą parą końcówek IC1 znajdują się dwie diody LED, jedna skierowana w jednym kierunku, a druga w przeciwnym. Sześć kombinacji par styków pomnożone przez dwie diody LED na kombinację daje w sumie 12 diod LED.

Możemy zaprogramować mikroprocesor tak, aby podłączyć styki GPIO do dodatniego („wysokiego”) lub ujemnego („niskiego”) zacisku baterii lub do żadnego z nich („stan wysokiej impedancji”). Dzięki różnym kombinacjom możemy zapalać po kolei każdą z diod LED.

Należy pamiętać, że pokazane tutaj numery diod LED niekoniecznie odpowiadają kolejności, w jakiej są one sterowane. Numery w nawiasach wskazują sposób, w jaki są one przyporządkowane w oprogramowaniu.

Zrobiliśmy to w ten sposób, ponieważ spodziewamy się, że osoby zmieniające treść oprogramowania uznają numery z oprogramowania (cyjan) za bardziej logiczne niż oznaczenia używane do montażu płytek drukowanych.

Te «programowe» numery odpowiadają również kolejności, w jakiej diody LED zostały ułożone w oryginalnej matrycy.

Staraliśmy się zachować tę kolejność w przypadku innych ozdób. W rezultacie można użyć tego samego oprogramowania, aby uzyskać różne wzory iluminacji dla każdego projektu.

Oprogramowanie

Aby wygenerować wzory za pomocą diod LED, programujemy mikroprocesor, aby ustawiał swoje styki GPIO w określonym stanie, czyli zapalał określone diody LED, a następnie przechodził w stan «uśpienia» na krótką chwilę (około 16 ms). Następnie «budził się», wyłączał diody LED, a potem ponownie „zasypiał” na około 64 ms.

Cykl ten powtarza się, a program decyduje, które diody LED są zapalane, aby wyświetlić interesujący wzór.

Utrzymywanie mikroprocesora w trybie uśpienia przez większość czasu minimalizuje zużycie energii. Ponieważ mikroprocesor „śpi” praktycznie przez cały czas, energia jest wykorzystywana głównie do zasilania diod LED.

A ponieważ diody LED są włączone średnio tylko przez około 20% czasu, bateria wystarcza na długi okres.

W zeszłym roku stwierdziliśmy na podstawie obliczeń, że typowa bateria powinna wystarczyć na około trzy miesiące. Nasz prototyp (wykorzystujący rezystory ograniczające prąd diod LED o wartości 1 k Ω) działał przez pięć miesięcy, zanim zaczął przygasać i gasnąć.



Oto wybór ozdób świątecznych, które wykonaliśmy na długo przed wielkim dniem. Oprócz samego Świętego Mikołaja (który oczywiście musi być czerwony!), niektóre pozostałe, dzięki uprzejmości kilku sprytnych producentów PCB, są dostępne w różnych kolorach (patrz lista części)

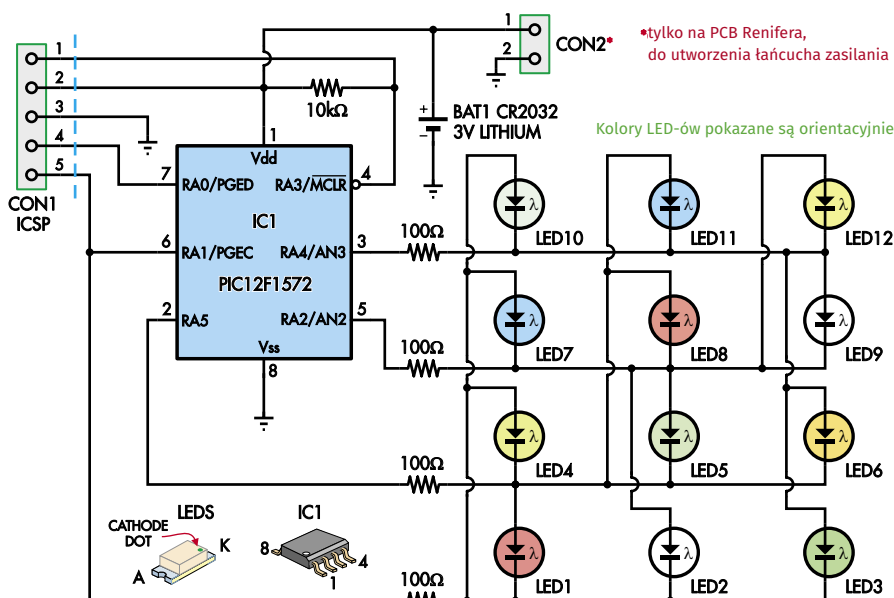
W związku z tym zalecamy konstruktorom stosowanie rezystorów 100 Ω, co zapewni około dwóch do trzech miesięcy żywotności; to więcej niż wystarczająco, aby przetrwać Święta Bożego Narodzenia i Nowy Rok.

Stworzyliśmy różne wzory sekwencji LED, aby jak najlepiej pasowały do każdej ozdoby, a także wzór pół-losowy, który może być używany dla każdej z ozdób. Ponieważ obwód jest w rzeczywistości taki sam, możesz wypróbować różne programy na różnych ozdobach, aby sprawdzić, czy wyświetlają one to, co lubisz.

Ozdoby

Dostępnych jest siedem nowych ozdób świątecznych. Pięć z nich jest przeznaczonych do użytku indywidualnego, a dwie można zawiesić osobno lub połączyć, aby stworzyć centralny element (pełnowymiarowej) choinki. Oczywiście nadal można budować małe Choinki wg opisu opublikowanego w roku 2019; w sumie jest to osiem różnych ozdób.

Pięć nowych pojedynczych ozdób to Skarpeta, Czapka Świąteczna, Cukierkowa Laska, Gwiazda i Bombka. Ponieważ uważamy, że każda choinka wygląda świetnie pokryta bombkami, udostępniamy ten ostatni projekt w czterech różnych kolorach masek lutowniczych. Możesz zaopatrzyć się w asortyment kolorowych bombek i udekorować swoją choinkę w spektakularny sposób! Podobnie jak pierwotna Choinka, płytka Skarpety jest również dostępna w różnych kolorach.



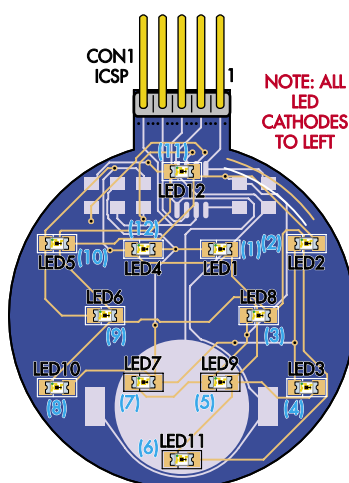
Małe ozdoby Bożenarodzeniowe z diodami LED

Rysunek 1. Schemat ideowy obwodu dla naszych małych ozdób świątecznych jest zasadniczo taki sam jak obwód dla małej choinki LED opublikowanej w listopadzie 2019 roku, aczkolwiek z nowszym mikroprocesorem PIC. Jest bardzo prosty i pozwala na zapalenie jednej diody LED na raz. Oprogramowanie pozwala zapalać te diody LED w dowolnej kolejności, a także zostały stworzone różne wersje sposobu zapalania diod w celu utworzenia fizycznej matrycy LED dla każdej ozdoby

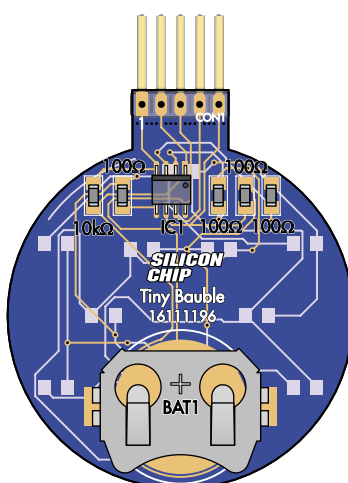
Dwie specjalne ozdoby to Renifer i Sanie Świętego Mikołaja. Można je wieszać pojedynczo, ale dodaliśmy również dodatkowe ścieżki do tych płytek drukowanych, dzięki czemu można je połączyć ze sobą, a przewody połączeniowe imitują uprzęż (naprzód, Rudolfie!).

Za pomocą tych przewodów można nawet podłączyć większy zestaw baterii, dzięki czemu można zaprzęgnąć cały tuzin reniferów, tak jak robi to Święty Mikołaj, a przy okazji nie martwić się, że zabraknie energii w Wigilię.

Podobnie jak w przypadku pierwotnego Choinki, wybór diod LED również zależy

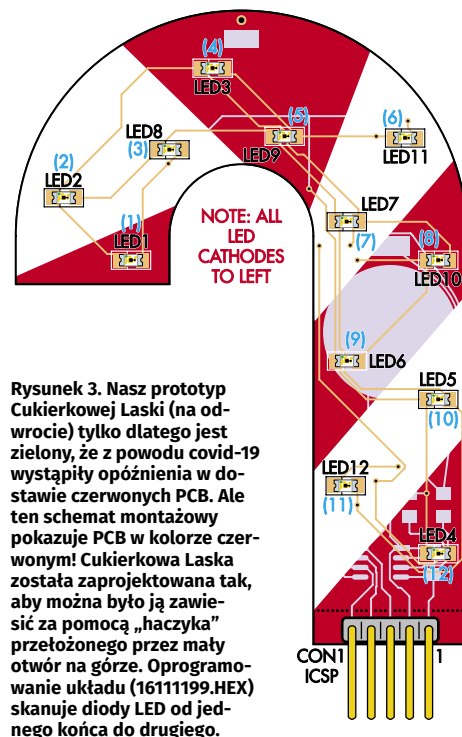


SC
©2020



Rysunek 2. Ponieważ przypuszczamy, że Bombka będzie jedną z najpopularniejszych ozdób, oferujemy ją w kolorze czerwonym, żółtym, zielonym i niebieskim. W ten sposób można tworzyć zestawy, a także zmieniać kolory diod LED. Oprogramowanie (1611196.HEX) cyklicznie zapala diody LED wokół Bombki lub można zamiennie użyć oprogramowania 1611190.HEX, aby uzyskać losowy, migoczący wzór.

UWAGA – wszystkie katody diod LED skierowane w lewo



Rysunek 3. Nasz prototyp Cukierkowej Laski (na odwrocie) tylko dlatego jest zielony, że z powodu covid-19 wystąpiły opóźnienia w dostawie czerwonych PCB. Ale ten schemat montażowy pokazuje PCB w kolorze czerwonym! Cukierkowa Laska została zaprojektowana tak, aby można było ją zawiesić za pomocą „haczyka” przetożonego przez mały otwór na górze. Oprogramowanie układu (1611199.HEX) skanuje diody LED od jednego końca do drugiego.

wyłącznie od Ciebie. Zbudowaliśmy nasze prototypy z losową mieszanką czerwonych, zielonych i białych diod LED, ale można również dodać do nich żółte, bursztynowe, różowe, cyjanowe lub niebieskie. Nasze zestawy są dostarczane ze „standardowymi” kolorami, ale można również zamówić dodatkowe zestawy diod LED w tych innych kolorach za pośrednictwem naszego sklepu internetowego (szczegółowe informacje można znaleźć na liście części).

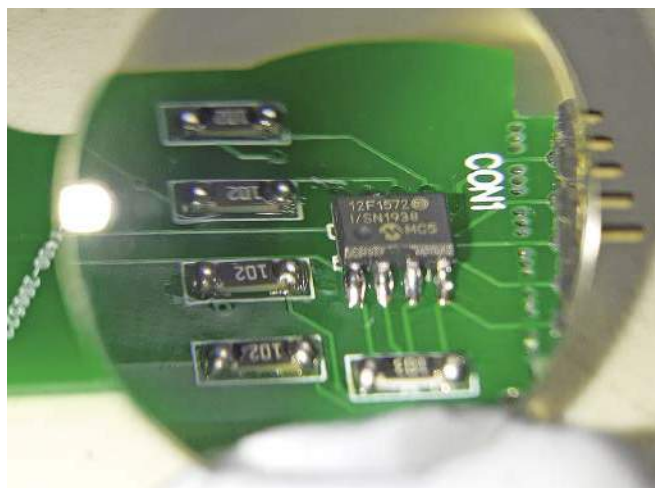
Na przykład, możesz zbudować niebieską bombkę i ozdobić ją niebieskimi diodami LED. Ale dopóki Rudolf ma czerwony nos, nie ma to znaczenia! <https://www.youtube.com/watch?v=TCraG9KREdg>.

Budowa

Wiemy, że jesteś podekscytowany, więc przejdziemy od razu do budowy. W większości wszystkich siedem nowych ozdób jest bardzo podobnych. Zapoznaj się ze schematami montażowymi PCB, rysunkami 2...8, które pokazują, gdzie znajdują się komponenty po obu stronach każdej ozdoby.

Poniższe instrukcje dotyczą wszystkich ozdób, ale jeśli budujesz Renifera lub Sanie Świętego Mikołaja, podamy dodatkowe informacje o tym, jak można je połączyć.

Ponieważ każda ozdoba ma unikalne oprogramowanie, w przypadku budowania wielu typów należy unikać pomieszania wstępnie zaprogramowanych mikroukładów.



Dzięki rozstawowi końcówek układu PIC12F1572 równemu 1,27 mm, poszczególne styki obudowy SOIC-8 łatwo jest przylutować. Jeśli między stykami powstaną mostki lutowniczy, do jego usunięcia można użyć pasty topnikowej i plecionki lutowniczej. Pokazane tutaj krople lutu są znacznie większe niż potrzeba (za to jakie są eleganckie!), ale to działa; trochę za dużo lutu jest lepsze niż za mało!

Jeśli masz wstępnie zaprogramowane mikroprocesory PIC, nie będziesz musiał montować CON1, złącza programowania. W takim przypadku należy usunąć małą część płytki, przeznaczoną na złącze CON1, ponieważ łatwiej będzie to zrobić teraz niż później.

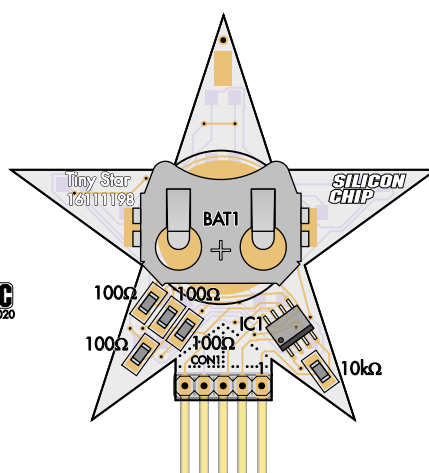
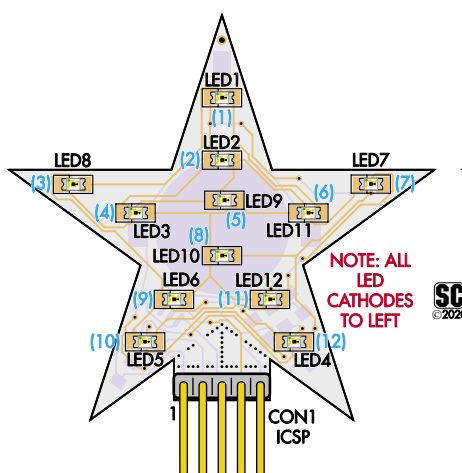
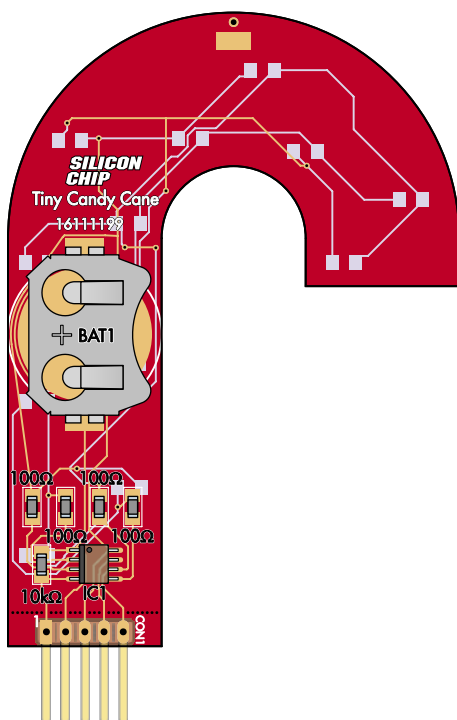
Wyjątkami są Bombka, Renifer i Sanie Świętego Mikołaja. Bombka ma odłamywaną zakładkę, ale jest to również najlepszy sposób na jej powieszenie, więc należy ją pozostawić. Renifer i Sanie Świętego Mikołaja nie mają odłamywanych złączy, ponieważ są one używane do okablowania w „uprzęży”.

W zależności od tego, jakie masz plany, możesz nie potrzebować uchwytu na ogniwo

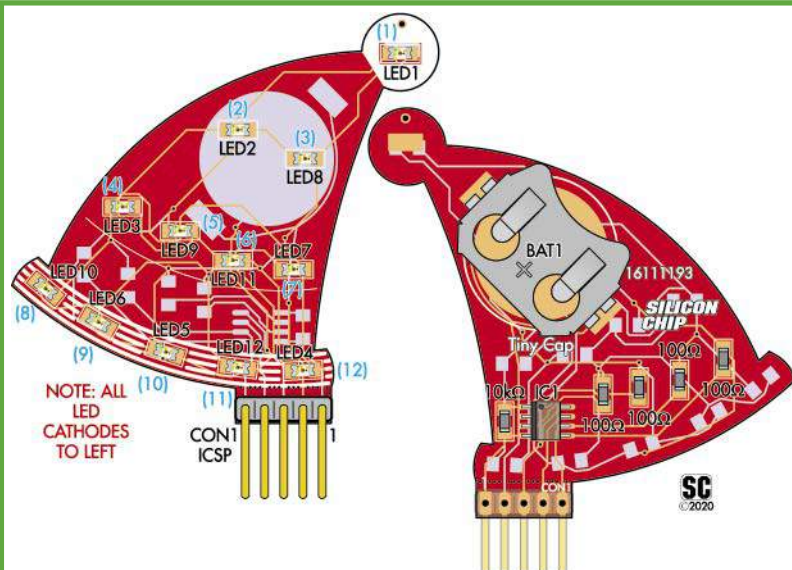
do Renifera lub Sań Świętego Mikołaja, ponieważ do zasilania tych ozdób może być użyta „uprzęży”.

W przypadku innych ozdób (Gwiazda, Skarpeta, Czapka i Cukierkowa Laska), jeśli nie musisz programować mikroukładów w obwodzie, ostrożnie ponacinaj PCB nożem modelarskim wzdłuż linii małych otworków. Zapewni to, że miedziane ścieżki nie oderwą się od płytki drukowanej. Następnie ostrożnie wygnij i odłam złączkę; odpowiednie do tego są płaskie szczypce.

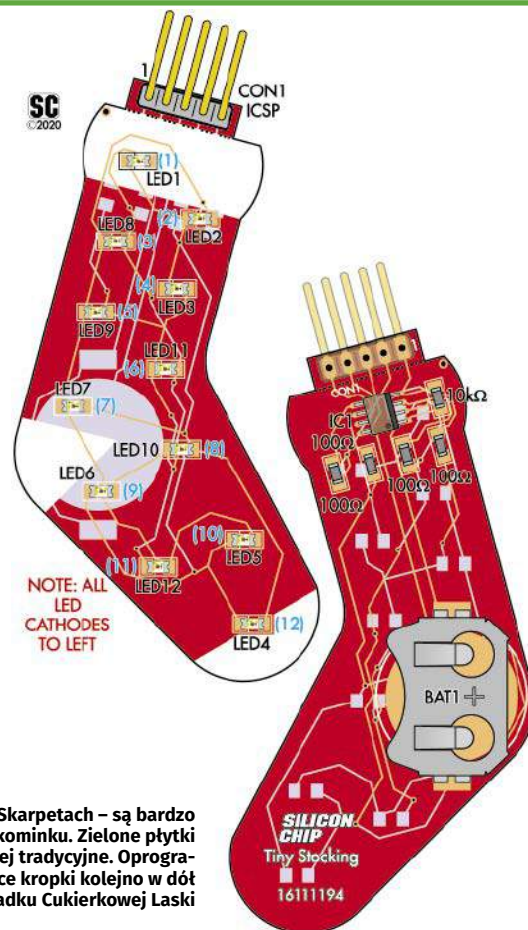
Złączka powinna się odłamać dość czysto, ale w razie czego można usunąć zadziory za pomocą pilnika. Jeśli to możliwe, wykonaj



Rysunek 4. Gwiazda z białym sitodrukiem na PCB jest jednym z bardziej efektownych wariantów i będzie wyglądać świetnie na tle zielonej choinki. Jest to również jedna z bardziej kompaktowych płytek PCB. Oznacza to, że niektóre ścieżki znajdują się blisko miejsca, w którym odłamuje się sekcja CON1. Oprogramowanie Gwiazdy (16111198.HEX) zapala diody LED promieniście od środka PCB do każdej końcówki po kolei.



Rysunek 5. Chociaż niektóre diody LED na Czapce ustawione są pod nieco innym kątem, katody nadal znajdują się po lewej stronie. Zakładka CON1 znajduje się bardzo blisko niektórych diod LED w prawym dolnym rogu, więc w razie potrzeby należy ją bardzo ostrożnie usunąć. Oprogramowanie 1611193.HEX zapala cyklicznie każdą diodę LED w dolnym rzędzie po kolei, podobnie jak w oryginalnej ozdobie Choinka



Rysunek 6. Nie oczekuj, że dostaniesz duże prezenty w tych Skarpetach – są bardzo małe! Jeśli nie masz miejsca na choince, możesz powiesić je na kominku. Zielone płytki PCB będą wyglądać efektownie, podczas gdy czerwone są bardziej tradycyjne. Oprogramowanie (1611194.HEX) zapala diody LED przesuwając świeące kropki kolejno w dół po każdej stronie, podobnie jak w przypadku Cukierkowej Łaski

wszystkie te czynności na zewnątrz, nosząc maskę, ponieważ pył z PCB może być drażniący.

Lutowanie

Jest to prawdopodobnie najbardziej krytyczna część montażu. Do lutowania małych części lutowanych powierzchniowo zalecamy posiadanie lutownicy z cienkim grotem, pęsety, pasty topnikowej, plecionki lutowniczej („knota”) i lupy. Do przytrzymania płytki PCB podczas lutowania można użyć kulki

kleju, takiego jak Blu-tack (rodzaj kitu samoprzylepnego wielokrotnego użytku, dostępnego na znanym portalu).

Topnik lutowniczy wytwarza dym i przykry zapach po podgrzaniu, więc przydaje się również wyciąg oparów lutowniczych lub alternatywnie, praca przy otwartym oknie.

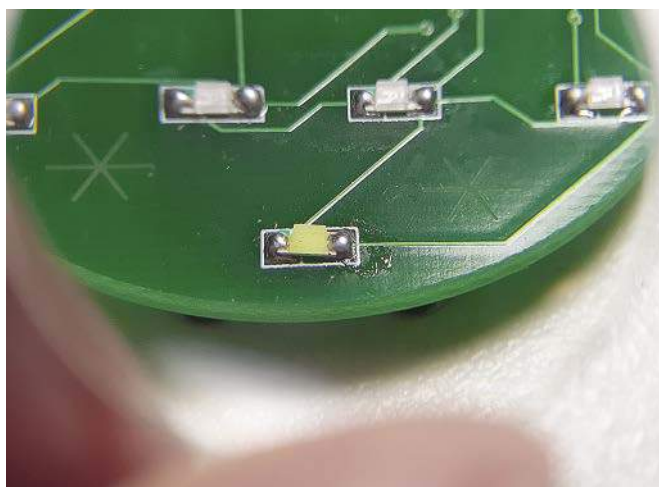
Najlepiej jest mieć czysty obszar roboczy z dużą ilością miejsca i światła. Małe części SMD są znane z tego, że wyskakują z uchwytu pęsety. Jeśli pole robocze nie

jest uporządkowane, nie ma nadziei na znalezienie upuszczonej części!

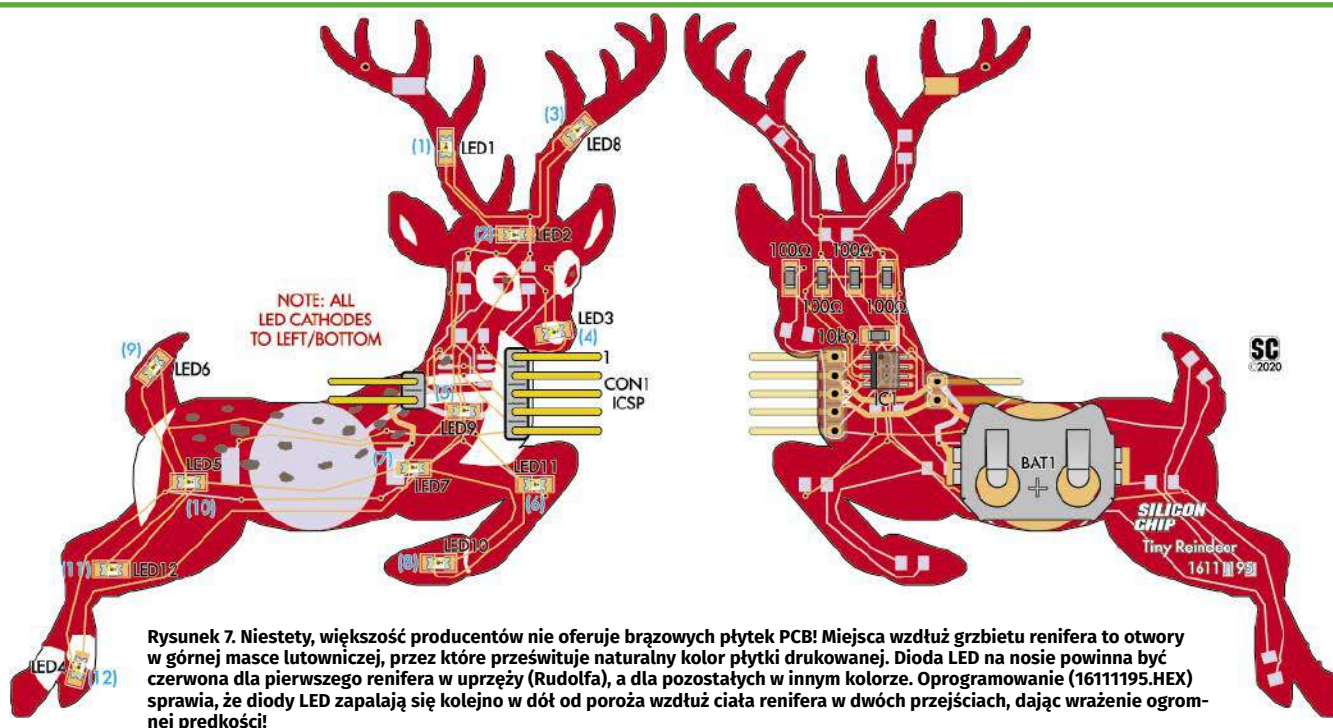
Dobłą techniką pracy z elementami SMD jest lutowanie jednej końcówki w celu zgrubnego umieszczenia (przymocowania) części. W razie potrzeby ponownie przelutuj to połączenie i wyreguluj pęsetą położenie części, aż element będzie płasko przylegał do płytki drukowanej, a wszystkie styki będą prostopadłe do swoich pól lutowniczych. Kolejno ostrożnie nałóż lut na pozostałe styki, a następnie wróć i odśwież pierwszy styk, nakładając nieco więcej świeżego lutu.

Dobrym pomysłem jest również nałożenie pasty topnikowej na pola lutownicze i styki przed ich lutowaniem. Pomaga to przenieść lut z grotu lutownicy do pól lutowniczych i styków. Użyj lupy, aby sprawdzić lutowane połączenia. Pomiędzy polem lutowniczym a stykiem podzespółu SMD powinno być dobre, gładkie złącze, ale nie nakładaj za dużo lutu, aby uniknąć mostkowania z sąsiednimi stykami. Obejrzyj nasze przykładowe zdjęcie, aby zobaczyć zbliżenie dobrego połączenia lutowanego. Lut powinien być gładki i błyszczący.

Należy również zwracać szczególną uwagę na schematy montażowe, aby sprawdzać postępy podczas lutowania każdej ozdoby.



Diody LED z przodu PCB mają rozmiar 1206 i przy długości 3,2 mm i szerokości 1,6 mm są wystarczająco łatwe w lutowaniu za pomocą większości standardowych grotów lutownicy. Zwróć uwagę na mały zielony trójkąt w lewym górnym rogu każdej diody LED, wyrównany z małym białym znakiem katody widocznym pod częścią



Rysunek 7. Niestety, większość producentów nie oferuje brązowych płytek PCB! Miejsca wzdłuż grzbietu renifera to otwory w górnej masce lutowniczej, przez które prześwituje naturalny kolor płytki drukowanej. Dioda LED na nosie powinna być czerwona dla pierwszego renifera w uprzęży (Rudolfa), a dla pozostałych w innym kolorze. Oprogramowanie (1611195.HEX) sprawia, że diody LED zapalają się kolejno w dół od poroża wzdłuż ciała renifera w dwóch przejściach, dając wrażenie ogromnej prędkości!

UWAGA – należy uważać na ogniwa pastylkowe, gdy w pobliżu znajdują się małe dzieci!

Podobnie jak w przypadku każdego projektu wykorzystującego ogniwa pastylkowe, należy zachować ostrożność, aby upewnić się, że nie ma szans, aby dostały się one w ręce małego dziecka. Mogłyby natychmiast włożyć je do ust, a połknięcie ich może poważnie zaszkodzić. Jeśli masz małe dzieci (poniżej pięciu lat), przykryj ozdoby przezroczystą rurką termokurczliwą lub przyklej baterię na miejscu (np. za pomocą neutralnie utwardzanego przezroczystego silikonu), aby nie można jej było łatwo usunąć.

Jeśli budujesz wiele ornamentów (a dla czego miałbyś tego nie robić?), możesz robić je pojedynczo lub równolegle – to zależy od Ciebie. Upewnij się tylko, że jeśli robisz je równolegle, nie pomylisz części do różnych ozdób, np. kolorów diod LED.

Zanim przejdiesz dalej, zastanów się, które kolorowe diody LED chcesz umieścić na każdej ozdobie. Ponieważ po wyjęciu z opakowania diody LED mogą wyglądać identycznie, najlepiej jest albo zamontować wszystkie diody LED każdego koloru za jednym razem, albo wyjąć tylko tyle, ile potrzebujesz w danym momencie.

Jeśli zgubisz jakąś kolorową diodę LED, większość multimetrów cyfrowych ustawionych na tryb testowania diod zaświeci diodę LED SMD, jeśli dotkniesz sondą do obu jej końców. Uważaj tylko, aby nie naciskać zbyt mocno, bo możesz oderwać delikatne styki! Jeśli dioda się nie zaświeci, spróbuj zamienić miejscami sondy.

Zazwyczaj czarna sonda spowoduje zaświecenie diody po dotknięciu styku w miejscu oznaczonym zieloną kropką (katoda).

Dalsza część instrukcji opisuje sposób montażu pojedynczej ozdoby.

Zacznij od zamontowania układu IC1 z tyłu płytki drukowanej. Sprawdź orientację układu

scalonego, szukając małej kropki w jednym rogu i skosu wzdłuż jednej krawędzi. Te dwie cechy muszą pokrywać się z linią zaznaczoną sitodrukiem na PCB i pokazaną na powiązanym schemacie montażowym płytki. Kropka powinna również znajdować się najbliżej wycięcia pokazanego na obrysie układu scalonego. Ponadto pole lutownicze PCB dla styku 1 jest prostokątne, podczas gdy pozostałe są zaokrąglone.

Użyj techniki lutowania wspomnianej powyżej, aby unieruchomić układ scalony na miejscu za pomocą pojedynczej przylutowanej końcówki. Nie przejmuj się, jeśli wykonasz mostek lutowniczy; skup się na upewnieniu się, że układ scalony jest prawidłowo umieszczony, a wszystkie osiem wyprowadzeń jest wyrównanych na odpowiednich polach. Ponownie rozpuść lut i dopasuj położenie, jeśli jest to konieczne, a następnie przylutuj pozostałe styki.

Jeśli masz mostek lutowniczy między dwoma lub więcej wyprowadzeniami, nałóż pastę topnikową i przykryj mostek kawałkiem plecionki lutowniczej. Przyciśnij lutownicę do plecionki, a gdy lut się rozpuści, ostrożnie odciągnij ją. Jeśli jest bardzo dużo lutu, może być konieczne powtórzenie tego procesu.

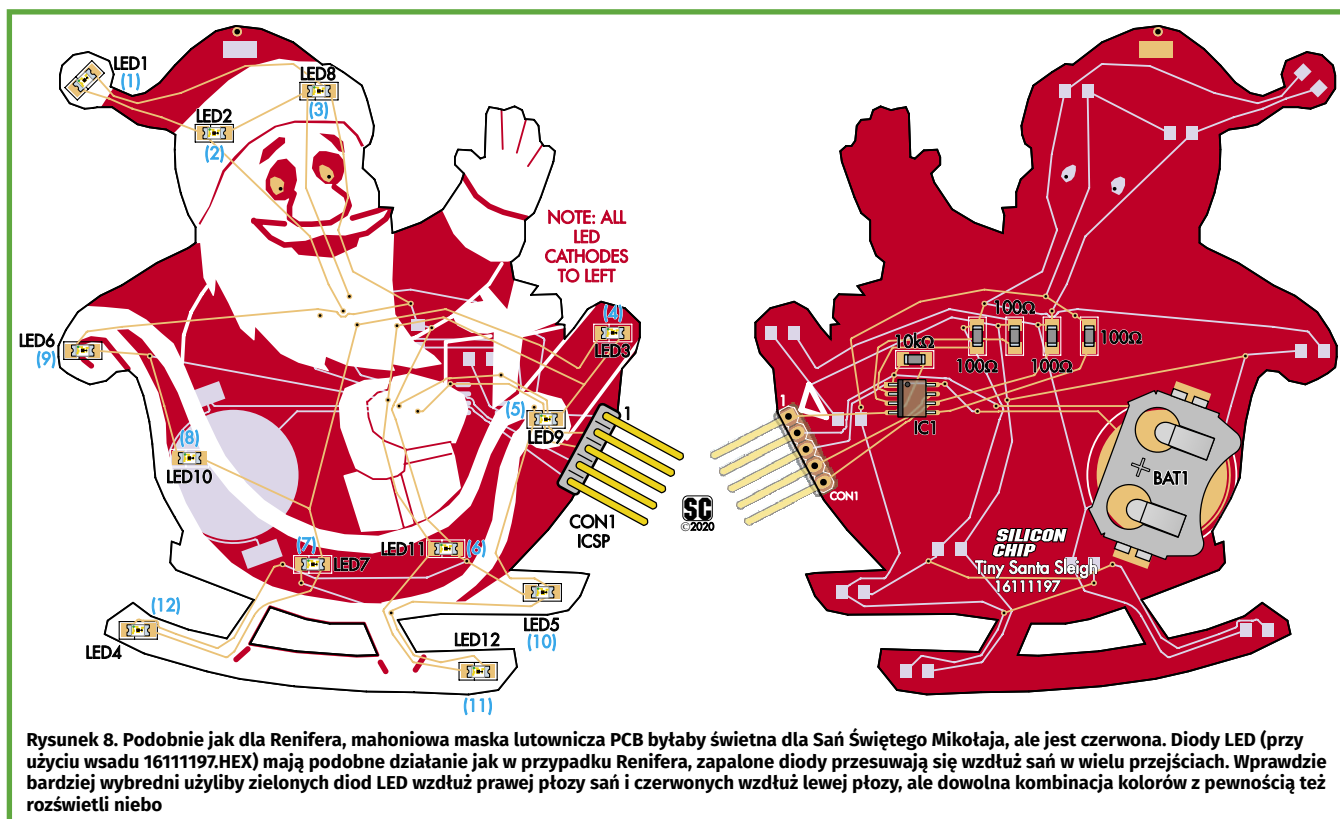
Pozostałe komponenty mają znacznie większe rozmiary i tylko po dwa wyprowadzenia. Dlatego są łatwiejsze do lutowania, a także znacznie mniej podatne na mostkowanie.

Następnie umieść rezystor 10 kΩ. Będzie on oznaczony jako „103” lub „1002” (choć do jego odczytania może być potrzebna lupa). Zarówno on, jak i pozostałe rezystory nie są spolaryzowane, więc można je przylutować w dowolny sposób. Sprawdź schemat montażowy PCB i sitodruk na płytce, aby zobaczyć, gdzie który element się znajduje.

Gdy rezystor 10 kΩ jest już na miejscu, przylutuj cztery rezystory po 100 Ω. Są one oznaczone jako «101» lub «1000» i pasują do pól oznaczonych 100.

Teraz odwróć płytkę drukowaną, aby przylutować diody LED. Są one spolaryzowane i ich katoda musi znajdować się najbliżej pola oznaczonego kreską. Zazwyczaj katoda diody LED jest oznaczona zieloną kropką lub strzałką, ale widzieliśmy też takie, których **anoda (!!!)** była oznaczona w ten sposób.

Dlatego najlepiej jest sprawdzić diody za pomocą DMM (jak opisano powyżej). Gdy dioda się świeci, czarna sonda znajduje się po stronie katody.



Rysunek 8. Podobieństwo dla Renifera, mahoniowa maska lutownicza PCB byłaby świetna dla Świąt Mikołaja, ale jest czerwona. Diody LED (przy użyciu wsadu 1611197.HEX) mają podobne działanie jak w przypadku Renifera, zapalone diody przesuwają się wzdłuż śnia w wielu przejściach. Wprowadź bardziej wybredni użyłoby zielonych diod LED wzdłuż prawej płozy śnia i czerwonych wzdłuż lewej płozy, ale dowolna kombinacja kolorów z pewnością też rozświetlił niebo

Zorientowaliśmy wszystkie diody LED na każdej płytce w ten sam sposób tak bardzo, jak to możliwe. Wszystkie katody powinny być skierowane w lewo i/lub w dół, gdy płytki są zorientowane tak, jak pokazano na rysunkach 2...8.

Strona katody diod LED jest wskazana na schematach montażowych PCB za pomocą ramki (dodatkowej poprzecznej kreski) wokół danej końcówki diody LED. Należy jednak pamiętać, że na rzeczywistych płytkach drukowanych niektóre dekoracyjne wzory sitodruku są naniesione nad konturami komponentów, więc nie zawsze wspomniane oznaczenia są widoczne.

Tak długo, jak pamiętasz o zasadzie lewo/dół i upewniasz się, że płytki są zorientowane tak, jak je pokazujemy, wszystkie diody LED powinny działać.

Użyj tej samej techniki, co poprzednio; przylutuj jedno wyprowadzenie elementu, upewnij się, że jest zgodny z zarysem i przylega płasko do PCB, a następnie przylutuj drugie wyprowadzenie i odśwież pierwszy lut.

Jest to podwójnie ważne dla diod LED, ponieważ jest to strona PCB, która będzie widoczna. Postaraj się więc o elegancki wygląd wszystkich spoin lutowniczych.

Po wykonaniu tej czynności, wyczyść pozostałości topnika na mniejszych komponentach za pomocą rozpuszczalnika, takiego jak alkohol izopropylowy czy spirytus „denaturat”.

Chociaż w przypadku większości topników nie jest to konieczne, pomaga to sprawić, że przód PCB będzie wyglądał schludniej, gdy ozdoba zostanie umieszczona na choince (lub gdziekolwiek planujesz ją wyeksponować).

W przypadku większości ozdób, ostatnim krokiem jest zamontowanie uchwytu na baterię pastylkową. Sprawdź poniższe uwagi dotyczące zestawu Renifera/Reniferów i Świąt Mikołaja, jeśli planujesz wykonać zestaw i podłączyć wiązkę przewodów.

W takim przypadku nie jest potrzebny uchwyt na baterię (choć nadal można go zamontować).

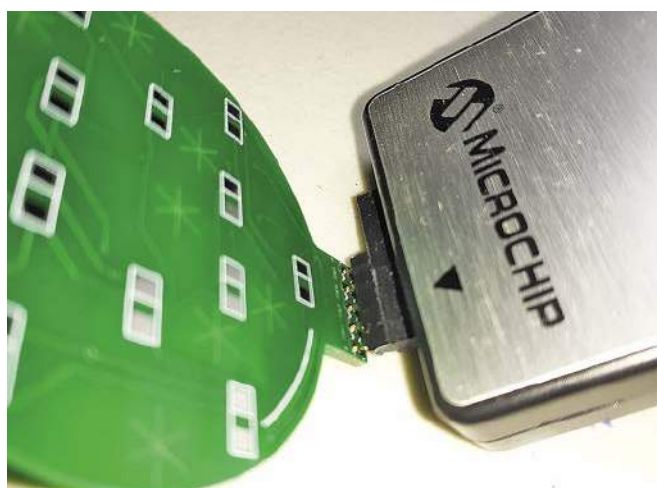
Orientacja pojemnika jest ważna, aby zapewnić możliwość wkładania i wyjmowania

ogniwa. Oba mniejsze pola lutownicze łączą się z dodatnią stroną baterii, a ujemny zacisk to duże okrągłe pole na płytce drukowanej.

W przypadku Cukierkowej Laski, Czapki i Skarpety można pojemnik dopasować w dowolny sposób. W przypadku pozostałych ozdób należy sprawdzić, czy małe wypustki na uchwycie baterii są skierowane w stronę środka płytki drukowanej. W ten sposób otwór uchwytu służący do wkładania/wyjmowania baterii będzie skierowany w stronę najbliższej krawędzi płytki.

Ponieważ uchwyt ogniwa jest większy niż inne komponenty i wykonany w całości z metalu, przed lutowaniem należy podnieść temperaturę grota lutownicy (jeśli to możliwe).

Proste umieszczenie pięciostykowego złącza szpilkowego w CON1 zapewnia wystarczający kontakt, aby zaprogramować PIC. Zastosuj delikatny nacisk, aby upewnić się, że styki lekko się potoczą podczas procesu programowania. Od Red.EdW: naszym zdaniem kątowna listwa kołkowa i najprostszy zacisk sprężysty pozwalają wykonać to tymczasowe połączenie znacznie wygodniej. Używamy PICKit 4, ale PICKit 3 również będzie działał



W przypadku innych komponentów, stosuj naszą sprawdzoną i zalecaną technikę montażu.

Jeśli masz zaprogramowany PIC, wszystko, co musisz zrobić, to zamontować ogniwo (dodatnią stroną do góry, zgodnie z oznaczeniem na uchwycie baterii), a Ornament powinien zacząć migać. Jeśli w ogóle nie miga, wyjmij baterię i sprawdź, czy nie ma zwarcia w uchwycie baterii lub na układzie PIC.

Jeśli działają tylko niektóre diody LED, sprawdź orientację diod LED oraz lutowanie diod LED, rezystorów i PIC.

Jeśli zamontowałeś pusty układ PIC, nie będzie on nic robił, dopóki go nie zaprogramujesz.

Programowanie w obwodzie

Chociaż złącze CON1 jest przeznaczone dla kąowego rzędu szpilek, nie trzeba go montować, nawet jeśli konieczne jest zaprogramowanie układu PIC na płytce. O ile nie planujesz wielokrotnego programowania PIC, samo przytrzymanie złącza programatora w tym miejscu, aby stykał się z kontaktami, jest zwykle wystarczające i daje schludniejszy efekt końcowy. Świetnie sprawdzi się w tym miejscu sprężysty klips do suszenia prania.

Nasze schematy i zdjęcia pokazują dołączoną kąową listwę kołkową, ponieważ pozwoliło nam to położyć programator i płytkę drukowaną płasko na stole, aby testować nasze oprogramowanie, choć oczywiście prosta listwa kołkowa również będzie działać. Możesz użyć kąowej listwy kołkowej, jeśli chcesz zaprogramować własne wzory.

Innym powodem, dla którego warto zamontować CON1, jest to, że styki 2 i 3 na CON1 mogą być używane do zasilania płytki, zamiast wbudowanego ogniwa.

Jeśli wolisz korzystać z zasilacza USB lub power banku, ozdoby będą z powodzeniem działać przy napięciu 5 V (i będą znacznie jaśniejsze). Doprowadź napięcie +5 V do kołka 2 złącza i podłącz masę do kołka 3. W ten sposób zasililiśmy naszego Mikołaja z reniferem, chociaż działa to również w przypadku innych ozdób.

Aby zamontować CON1, ułóż rząd kołków w otworach PCB, z rzędem szpilek umieszczonym od spodu, tak aby szpilki były mniej widoczne z przodu. Przylutuj jeden kołek i sprawdź, czy złącze jest proste, a następnie przylutuj pozostałe styki. **Od Red. EdW: Oczywiście na schematach montażowych złącza polutowano w odwrotny sposób, ułożone od góry.**

Po zakończeniu programowania można odłamać sekcję CON1 od płytki drukowanej. Jest to nieco bardziej kłopotliwe po zamontowaniu CON1, ale można to ostrożnie zrobić.

Wykaz elementów, kupuj w sklepie.avt.pl (W-wa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl):

- 1 uchwyt ogniwa pastylkowego CR2032 do montażu powierzchniowego [Digi-key BAT-HLD-001-ND, Mouser 712-BAT-HLD-001 lub podobny]
- 1 rezystor 10 kΩ SMD 1206 [np. Altronics R8188]
- 4 rezystory 100 Ω SMD 1206 [np. Altronics R8044]
- 12 diod LED SMD rozmiar 1206, dowolna kombinacja kolorów [np. Altronics Y1041, Y1056, Y1073, Y1079, Y1085]
- 1 litowa bateria pastylkowa CR2032 (CR2025 jest również odpowiednia, ale ma krótszą żywotność)
- 1 5-stykowy odcinek prostej lub kąowej jednorzędowej listwy kołkowej (CON1) (opcjonalnie; do programowania IC1)

Plus jeden z następujących elementów:

- Choinka: zielona, czerwona lub biała płytka PCB o wymiarach 54x41 mm i kodzie 16111191, plus PIC12F1572-I/SN zaprogramowany wsadem 16111191.HEX
- Czapka: czerwona płytka drukowana o wymiarach 54x56 mm i kodzie 16111193, plus PIC12F1572-I/SN zaprogramowany wsadem 16111193.HEX
- Skarpetka: czerwona lub zielona płytka drukowana o wymiarach 41x81 mm i kodzie 16111194, plus PIC12F1572-I/SN zaprogramowany wsadem 16111194.HEX
- Renifer: czerwona płytka drukowana o wymiarach 91x97 mm i kodzie 16111195, plus PIC12F1572-I/SN zaprogramowany wsadem 16111195.HEX
- Bombka: czerwona, żółta, zielona lub niebieska płytka PCB o wymiarach 53x46 mm i kodzie 16111196, plus PIC12F1572-I/SN zaprogramowany wsadem 16111196.HEX
- Sanie Świętego Mikołaja: czerwona płytka drukowana o wymiarach 79x91 mm i kodzie 16111197, plus PIC12F1572-I/SN zaprogramowany wsadem 16111197.HEX
- Gwiazda: biała płytka PCB o wymiarach 56x54 mm i kodzie 16111198, plus PIC12F1572-I/SN zaprogramowany wsadem 16111198.HEX
- Cukierkowa Laska: czerwona płytka drukowana o wymiarach 84x60 mm i kodzie 16111199, plus PIC12F1572-I/SN zaprogramowany wsadem 16111199.HEX

Dodatkowe części do uprząży reniferów (jeden zestaw dla każdego renifera) [brak w zestawie]

- 1 2-stykowe gniazdo żeńskie BLS02 raster 2,54 mm plus listwa ze stykami BLT-F ORAZ
- 2 przewody potężeniowe męsko-żeńskie LUB
- 2 odcinki emaliowanego przewodu miedzianego o średnicy 0,63 mm

Dodatkowe części do zasilania uprząży renifera z ogniw AA [brak w zestawie]

- 1 2-stykowe gniazdo żeńskie BLS02 raster 2,54 mm plus listwa ze stykami BLT-F
- 1 uchwyt na baterie 2xAA lub 3xAA, najlepiej z przetwornikiem (np. Jaycar PH9280)

Zestawy

Każdy zestaw zawiera wszystkie części wymagane do zbudowania jednej ozdoby (z wyjątkiem ogniwa pastylkowego) oraz 12 czerwonych, 12 zielonych i 12 białych diod LED, dzięki czemu można je dowolnie mieszać i dopasowywać. Dostępne są również inne kolory diod LED, wymienione poniżej. Wszystkie zestawy kosztują 14 AU\$ za sztukę (10% zniżki dla aktywnych subskrybentów SC) plus opłata pocztowa, która wynosi 10 AU\$ za zamówienie na terenie Australii. (W przypadku zamówienia do 50 zestawów koszt wysyłki wynosi 10 AU\$). Wszystkie zestawy mają ten sam kod katalogowy (SC5579) z opcjami dotyczącymi typu ozdoby i koloru płytki drukowanej (w przypadku ozdób dostępnych w więcej niż jednym kolorze). Na przykład: aby otrzymać zestaw z czerwoną Bombką, należy zamówić SC5579/bauble/R. W przypadku Śnia należy zamówić zestaw SC5579/sleigh (ponieważ dostępny jest tylko jeden kolor). Oryginalny zestaw Choinki Tiny LED Xmas Tree można nadal zamówić pod wcześniejszym kodem katalogowym SC5180.

Dodatkowe diody LED

- * 10 bursztynowych: Nr kat. SC3394, 0,70 AU\$
- * 10 żółtych: Cat SC3405, 0,70 AU\$
- * jedna różowa: Cat SC3406, 0,20 AU\$

- * 10 niebieskich: Nr kat. SC3396, 0,70 AU\$
- * 10 cyjanowych: Cat SC5199, 1,00 AU\$



Będziesz potrzebował PICKit 3 lub PICKit 4 (lub innego programatora, który może współpracować z PIC12F1572).

Używamy oprogramowania MPLAB X IPE (zintegrowane środowisko programowania), które można pobrać bezpłatnie jako część MPLAB X IDE (zintegrowane otoczenie środowiska programowania). Pliki do pobrania dla systemów Windows, Mac i Linux można znaleźć na stronie www.microchip.com/mplab/mplab-x-ide.

Najnowsza wersja (5.40) działa tylko z procesorami 64-bitowymi, więc jeśli masz procesor 32-bitowy, możesz potrzebować starszej wersji. Starsze wersje można znaleźć na stronie www.microchip.com/development-tools/pic-and-dspic-downloads-archive.

Podczas instalacji tego oprogramowania należy upewnić się, że włączona jest obsługa procesorów 8-bitowych (w tym PIC12F1572).

Przed podłączeniem programatora należy upewnić się, że do Ornametu nie jest włożona żadna bateria. Programatory PICKit mogą zasilать mikroprocesor napięciem 5 V i nie jest dobrym pomysłem podłączanie napięcia 5 V do ogniwa 3 V (programator jest prawdopodobnie wystarczająco inteligentny, aby tego uniknąć, ale lepiej być przezornym niż żałować... spalonego programatora, albo nawet tylko mikroprocesora).

Poniższy opis postępowania podczas programowania zakłada, że używasz PICKit-a i MPLAB X IPE, chociaż inne programatory będą działać podobnie. Zaczynaj



Rysunek 9. Zasililiśmy Sanie Świętego Mikołaja i dwa Renifery z pary baterii AA w obudowie Jaycar PH9280. Czerwone przewody odpowiadają napięciu +3 V, a szare 0 V; można użyć innych kolorów, tylko nie należy ich krzyżować! Ponieważ każda ozdoba pobiera mniej niż 1 mA prądu, w ten sposób można zasilać wiele innych ozdób. Układ IC1 działa w zakresie napięć od 2 V do 5,5 V, więc dobrze nadaje się do zasilania z dwóch lub trzech ogniw AA lub zasilacza USB. Jak pokazano na zdjęciu na stronie XX, podłączyliśmy go za pomocą wtyczek i gniazd, aby zapewnić elastyczność. Można jednak przylutować przewody bezpośrednio do pól lutowniczych. Najlepiej byłoby przetestować każdą ozdobę osobno przed połączeniem ich razem

od przeglądania katalogu z oprogramowaniem, aby otworzyć odpowiedni plik HEX (znajdujący się w pliku .zip oprogramowania do pobrania ze strony Silicon Chip, na stronie siliconchip.com.au/Shop/20/5579).

Istnieje plik HEX dla każdego projektu PCB; znajdź numer, który pasuje do programowanej płytki PCB. Alternatywnie, plik 16111190. HEX jest po prostu pół-losowym wzorem, który będzie działał z każdą z ozdób.

Podłącz programator do komputera, a następnie podłącz programator do ornamentu. Styk oznaczony strzałką (i czerwona linia na kablu taśmowym) na programatorach PICKit to styk 1, który łączy się z kołkiem 1 CON1.

Jeśli nie przylutowałeś CON1, umieść 5-cio-szpilkowy odcinek prostej (albo kątownej – jak Ci wygodnie) listwy kołkowej w złączu PICKit-a i oprzyj go na polach lutowniczych na płytce drukowanej, zamiast podłączać go do złącza programatora. Użyj delikatnie siły, aby zapewnić kontakt. Idealnie sprawdzi się sprężysty klips do suszenia bielizny.

Ustaw programator tak, aby dostarczał zasilanie do układu docelowego. MPLAB X IPE posiada przyciski „Apply” i „Connect”. Należy je kliknąć przed kliknięciem przycisku „Programm”.

Jeśli wszystko jest w porządku, diody LED powinny zacząć migotać po zakończeniu

programowania. Ponieważ jedna z końcówek LED jest również używana do programowania, niektóre diody LED mogą zapalić się poza sekwencją.

Odłącz programator i włóż ogniwo, aby sprawdzić działanie całości. Nie ma nic więcej do zrobienia.

Dekoracja

Ozdoby mają kilka opcji umieszczenia na choince. Większość z nich posiada przelotowy otwór na górze, do którego można przylutować pętlę z ocynowanego lub emaliowanego drutu miedzianego, dzięki czemu ozdobę można zawiesić na gałęzi drzewa.

Bombka nie ma tych otworów, ale można ją zawiesić na przewodzie przylutowanym do pól lutowniczych CON1 (środkowy styk jest najlepszy).

Choinkę Tiny LED Xmas Tree można postawić na płaskiej powierzchni poprzez przylutowanie dwóch lub więcej ocynowanych przewodów miedzianych do styków CON1 (np. końcówek 1 i 5) i wygięcie ich tak, aby stykały się z powierzchnią.

Drut można przykręcić do gałęzi choinki, aby go zabezpieczyć, chociaż w przypadku tradycyjnego świerka czy jodły (prawdziwej lub sztucznej) igły zazwyczaj dobrze utrzymują ozdobę na drzewku.

Większość ozdób ma również większe pole stykowe z tyłu płytki drukowanej. Można go użyć do przylutowania do ozdoby agrałki lub podobnego elementu, aby można było nosić ją na ubraniu lub w inny sposób przymocować do choinki.

Rysunek 9 i zdjęcie na stronie XX pokazują, jak można wzajemnie podłączyć PCB Renifera i Sań Świętego Mikołaja i wykorzystać je do stworzenia oszałamiającego centralnego elementu dekoracji.

Nie musisz poprzestawać na jednym czy dwóch reniferach; dodaj ich tyle, ile chcesz (osiem to tradycyjna liczba; i nie pomiń biednego Rudolfa z jego świecącym czerwonym nosem!).

Użyliśmy przewodów połączeniowych, ale można użyć emaliowanego drutu miedzianego o średnicy około 0,63 mm, który utrzyma cały zespół razem jako pojedynczą półsztywną jednostkę.

Po zasileniu ogniwami pastylkowymi powinny działać przez wiele miesięcy, zapewniając mnóstwo migających światełek na choince! ■

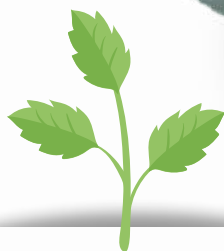
Tim Blythman

Adaptacja do wydania polskiego – Andrzej Nowicki

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

REKLAMA





Wyhoduj to sam

Sterowane cyfrowo, jednopojemnikowe rozwiązanie do upraw w pomieszczeniach

Co powiesz na małe pudełko zawierające „cyfrową farmę”? Przeczytaj poniższy artykuł, jeśli chcesz wiedzieć, jak zbudować mały pojemnik z czujnikami i mikrokontrolerem do uprawy roślin pod kontrolą cyfrową.

Bardzo zainspirowało mnie wystąpienie dyrektora Open Agriculture Initiative z MIT, Caleba Harpera, na temat rolnictwa cyfrowego pt. „Ten komputer będzie uprawiał twoją żywność w przyszłości” [1]. Najważniejszym pytaniem, na które odpowiedział w swoim przemówieniu, było: „A co by było, gdybyśmy mogli uprawiać pyszne, bogate w składniki odżywcze jedzenie w pomieszczeniach, w dowolnym miejscu na świecie?” I tak narodził się mój pomysł!

Cele projektu i rozważania

To, co próbowałem zbudować, było czymś w rodzaju pudełka lub inkubatora, który stworzyłby idealne warunki klimatyczne do wzrostu roślin

i dostarczyłby dokładnie tyle światła oraz składników odżywczych, tyle ile im potrzeba. Powinien zawierać symulator światła słonecznego, system nawadniania i system kontroli klimatu, a wszystko w eleganckiej i nowoczesnej obudowie.

Ponieważ chlorofil w roślinach reaguje głównie tylko na dwa pasma światła o długości fali około 450 i 650 nm (**rysunek 1**), to system oświetlenia powinien zawierać kombinację czerwonych i niebieskich diod LED, aby zapewnić idealną mieszankę wspierającą zarówno wzrost wegetatywny, jak i kwitnienie.

Chciałem zastosować nową metodę podlewania zwaną Aeroponics lub Fogponics [2], która wykorzystuje ultradźwiękowy atomizer

do nawilżania roślin przy pomocy mgły wodnej z nawozem. System powinien automatycznie dozować roślinom odpowiednią ilość składników odżywczych, dokładnie wtedy, gdy są one potrzebne.

Oczywiście mój system potrzebuje również czujników pH i TDS (Total Dissolved Solids – określa ilość rozpuszczonych w wodzie związków stałych). Pomagają one osiągnąć zrównoważone pH w zbiorniku wodnym, które jest najlepsze dla roślin oraz pomagają odpowiednio w dozowaniu składników odżywczych. Ponadto system powinien być wyposażony w przyłącze do automatycznej wymiany wody, które można by łatwo sterować naciśnięciem jednego przycisku.

System powinien posiadać obwody kontroli stanu powietrza, które pozwalają na precyzyjne utrzymanie temperatury i wilgotności wewnątrz naszej farmy, z dokładnością odpowiednio do 1°C i 1%. Do realizacji tego celu wymagany jest czujnik temperatury i wilgotności typu DHT22 lub DHT11 oraz wentylator z odpowiednim sterowaniem.

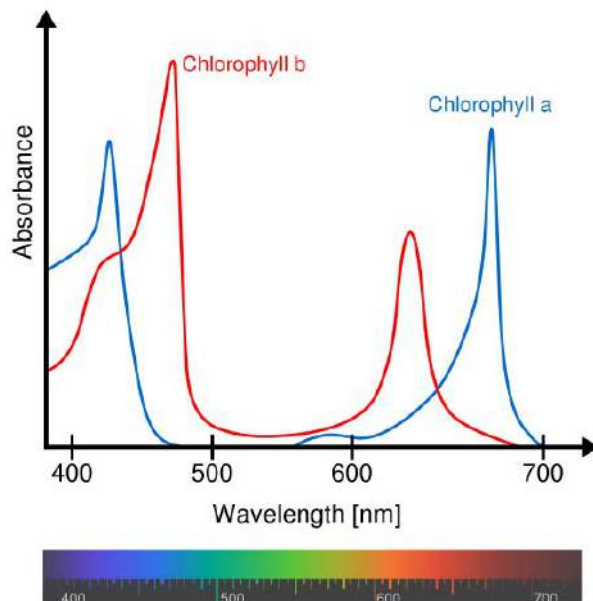
Wreszcie, system powinien być sterowany i monitorowany za pomocą aplikacji mobilnej. To byłby pierwszy raz w moim życiu, kiedy próbowałem stworzyć taką aplikację! Aplikacja powinna dostarczać w czasie rzeczywistym informacji o: pH, temperaturze, wilgotności, składnikach odżywczych itp. oraz dawać graficzną reprezentację ich zmienności w czasie, aby generować statystyki i dzielić się postępami we wzroście za pośrednictwem mediów społecznościowych. Chciałbym również zaimplementować inteligentne alarmy, które poinformują mnie, kiedy system będzie wymagał mojej interwencji. Oprócz wyświetlania wyników pomiarów, system powinien również oferować możliwość dostosowania jego parametrów pracy.

Jak możesz się domyślić, nie ja pierwszy wpadłem na taki pomysł, ale dobre wyniki często osiąga się przez ulepszanie tego, co już istnieje.

Na rynku istnieje kilka podobnych projektów, ale pomimo wszystkich zalet, są one także pewne wady: na przykład zajmują zbyt dużo miejsca, są za małe, za drogie lub są przeznaczone tylko dla jednej rośliny itp. W każdym razie chciałem opracować nowy, zaawansowany, otwarty system, który ma kilka wad i implementuje tylko zoptymalizowane funkcje. Brzmi to bardzo ambitnie i takie jest. Ponadto uwaga: projekt jest bardzo obszerny i wykracza daleko poza zakres artykułu Elektora. Dlatego prosimy o zapoznanie się ze stroną internetową projektu pod adresem [3], gdzie można znaleźć wiele szczegółów i informacji ogólnych, a także szczegółową instrukcję montażu.

Eksperymenty

Zajmowanie się roślinami jest bardzo czasochłonne! Zwykle potrzebują tygodni lub miesięcy, aby urosnąć – musisz być na to gotowy! Nawet jeśli podstawowa koncepcja jest jasna, nadal potrzeba kilku



Rysunek 1. Widma absorpcyjne chlorofilu a i b (Źródło: Daniele Pugliesi, CC BY-SA 3.0 [9])

wyprzedzających eksperymentów, aby zobaczyć, czy całość może się udać zgodnie z planem i jakie dostosowania do rzeczywistości są wymagane. Najpierw chciałem się dowiedzieć, czy uprawa roślin z użyciem mgły wzbogaconej w składniki odżywcze i oświetlenie LED jest lepsza niż konwencjonalny system z nawadnianą glebą i naturalnym światłem słonecznym. Teoretycznie tak powinno być, ale nie powinienemś niczego zakładać, dopóki tego nie spróbujesz!

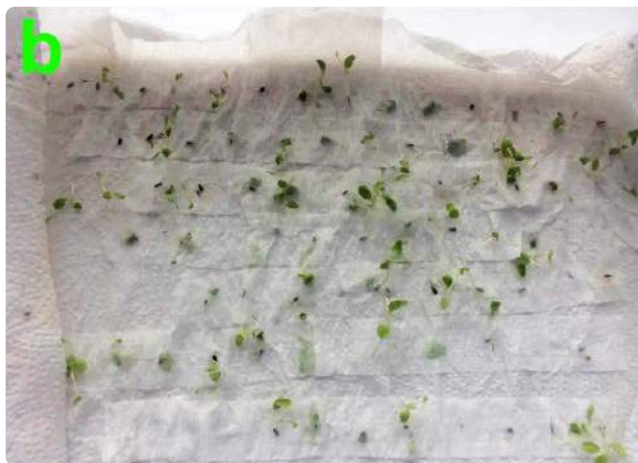
W swoich eksperymentach podzieliłem więc rośliny na kilka grup:

- Gleba + Światło słoneczne: Ta grupa jest uprawiana w warunkach naturalnych, sadzona w glebie i umieszczana na parapecie okna.
- Mgła + światło słoneczne: Ta grupa jest sadzona w pojemniku z bogatą w składniki odżywcze mgłą i umieszczona na tym samym parapecie.
- Gleba + oświetlenie LED: Ta grupa znajduje się w glebie, ale w sztucznym świetle LED zamiast światła słonecznego.
- Mgła + oświetlenie LED: Ten system ma być najlepszy, ponieważ łączy w sobie dwie główne cechy proponowanego systemu.

W dziale roślinnym lokalnego oddziału niemieckiej sieci supermarketów Kaufland, jako moje „króliki doświadczalne”, kupowałem nasioną sałaty mieszanej. W tym momencie proszę o trochę cierpliwości,



Rysunek 2. Pierwsze nasiona kietkują





Rysunek 3. Sadzonki przesadzone na inne podłoże przed (a) i po (b) tygodniu

ponieważ po raz pierwszy w życiu coś zasadziłem. Wcześniej przeprowadziłem badania i dowiedziałem się, że nasiona muszą najpierw wykiełkować. Wziąłem więc plastikowy pojemnik na nasiona i przykryłem go papierowym ręcznikiem. Na tym etapie nasiona potrzebują prawie 100% wilgoci, więc spryskałem papierowe ręczniki wodą i przykryłem całość plastikową torbą, aby woda nie odparowała (rysunek 2a). Zostawiłem nasiona na 10 dni, a kiedy otworzyłem pojemnik, byłem naprawdę zaskoczony: prawie wszystkie nasiona wykiełkowały (rysunek 2b)! Cieszyłem się jak dziecko.

Następnie musiałem wybrać najlepsze sadzonki, takie, które miały najgrubszą łodygę czyli po prostu wyglądały na większe i lepsze. Instruktaż na YouTube zalecał przesadzanie małych sadzonek na inne podłoże, do czasu aż uformują się tak zwane „drugie liście”. Na Amazonie zamówiłem granulację z włókna kokosowego. Ma on właściwości podobne do gleby, a szczególnie cenne jest to, że podłany zwiększa swą objętość 6-krotnie. Zamówiłem też, specjalną do sadzenia, skrzynkę z przekładkami. Dostałem go za jedyne 2 €, a to znacznie lepiej zorganizowało sprawę. W przegrodach pudełka umieściłem granulację kokosową, podlewałem je do całkowitego nawodnienia, potem położyłem małe roślinki na środku i przykryłem pudełko przezroczystym kawałkiem plastiku, który był dołączony do zestawu. Moja eksperymentalna szklarnia wyglądała jak na rysunku 3a. Pudełko zamknąłem i zostawiłem na tydzień. Kiedy go otworzyłem ponownie, byłem zaskoczony: rzeczywiście urosły. Rysunek 3b pokazuje, że już po tygodniu różnica

jest zauważalna! Teraz mogłem umieścić je w kulkach ziemi i gliny, i zacząć eksperymentować z mgłą!

Podczas kursu projektowania wspomaganego komputerowo zaprojektowałem tacki siatkowe, w które włożyłem rośliny (rysunek 4a). Kupiłem też plastikowy pojemnik odpowiedniej wielkości i zaznaczyłem miejsca pod otwory do wywiercenia. Na rysunku 4b i na moim filmie na YouTube [4] możesz zobaczyć mój zmontowany system testowy, z roślinami w granulach kokosowych oraz glinianymi kulkami i ziemią.

Szczerze mówiąc, wyniki mogłyby być lepsze. Tylko kilka roślin przeżyło mój system zamgławiania. Podejrzewam, że było to spowodowane nieodpowiednią pożywką lub mgłą, która nie wystarczała do utrzymania wilgoci w glinianych kulach i glebie. Co gorsza, nie mogłem nie robić w czasie ferii wielkanocnych, ponieważ laboratorium było zamknięte. W rezultacie wszystkie moje małe roślinki zginęły!

Po tym niepowodzeniu zdecydowałem się przetestować inny system, wykorzystujący wełnę mineralną jako alternatywne podłoże do upraw (rysunek 5a). Po tygodniu sprawdziłem swoje pudełko (rysunek 5b) i stwierdziłem, że jeden typ nasion odniósł sukces, ale inny nie! Aby przyspieszyć proces, kontynuowałem pracę z nasionami, które wykiełkowały. W międzyczasie ustawiłem inny system zamgławiania i próbowałem zoptymalizować czas jego działania, aby sprawdzić, czy uda mi się sprawić, by nasiona wykiełkowały, przy użyciu tylko jego! Nadal eksperymentuję z różnymi parametrami i warunkami, aby zobaczyć, jak daleko mogę się posunąć.



Rysunek 4. Tacki siatkowe wydrukowane w 3D (a) i gotowy zmontowany system testowy (b)



Rysunek 5. Następną próbą z nowym podłożem (a). Kiełkował tylko jeden typ nasion (b)

Projekt elektroniczny

Krótką historią: wybierając płytkę do sterowania całością chciałem wybrać coś pomiędzy satshakitem Daniela Ingrassia [5] i FabLeo Jonathana Grinhama [6]. Pierwsza jest w 100% kompatybilna z Arduino IDE i jego bibliotekami, doskonała i ulepszoną wersją Fabkit typu open source. Jest nie tylko tańszy, ale także szybszy (16 MHz) i łatwiejszy do polutowania. Arduino IDE rozpoznaje tę płytkę jako Arduino Uno.

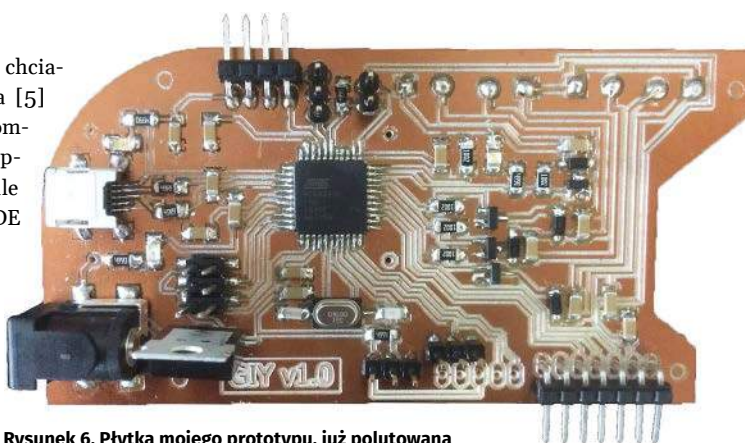
Z drugiej strony, FabLeo ma bardzo podobne funkcje, ale także sprzętowe USB. Ponieważ miałem już doświadczenie z ATmega328P, chciałem spróbować czegoś nowego i zdecydowałem się użyć FabLeo.

Do zaprojektowania płytki użyłem programu EAGLE. Darmowa wersja jest wystarczająca dla tego projektu, a pliki projektowe można pobrać z [3]. Sieć fab utrzymuje również stale aktualizowaną bibliotekę fab.lbr [7], z której korzystałem. Szczegóły dotyczące obwodu, jak również wykonania i lutowania płytki również można znaleźć w [3]. Cechą szczególną jest to, że potrzebowałem (impulsowej) przetwornicy podwyższającej napięcie z 12 V na 24 V do zasilania ultradźwiękowej wytwornicy mgły. Na spodzie płytki umieściłem płytkę Wi-Fi, którą zrobiłem podczas tygodnia sieci i komunikacji, na kursie, w którym uczestniczyłem. Powstała już wypełniona elementami płytkę można zobaczyć na rysunku 6.

Czujniki

Do programowania mojej płytki użyłem Arduino IDE. Podłączyłem płytkę Arduino do koncentratora USB i wgrałem kod [3]. Pierwszym podłączonym czujnikiem był DHT11 (rysunek 7a) do pomiaru temperatury i wilgotności. Te czujniki są bardzo proste i powolne, ale są idealne dla początkujących, którzy chcą wykonać proste rejestrowanie danych. DHT11 łączy w sobie pojemnościowy czujnik wilgotności i termistor. Wewnątrz znajduje się również bardzo prosty układ, który dokonuje konwersji analogowo-cyfrowej i wysyła cyfrowo dane dotyczące temperatury oraz wilgotności – łatwe do odczytania przez każdy mikrokontroler. DS18B20 (rysunek 7b), to cyfrowy czujnik temperatury firmy Maxim z interfejsem 1-wire, o dokładności od 9 do 12 bitów (zakres: od -55 do $125^{\circ}\text{C} \pm 0,5^{\circ}$). Do korzystania z tego czujnika, dostępne są w świecie Arduino, odpowiednie funkcje.

Dziwne jest to, że po pomyślnym zaprogramowaniu czujników płyta przestała działać, gdy połączyłem wszystkie czujniki razem. Dość długo zajęło mi zrozumienie, dlaczego tak się stało. Jako szybkie i brudne rozwiązanie zdecydowałem się użyć „nagiego” czujnika

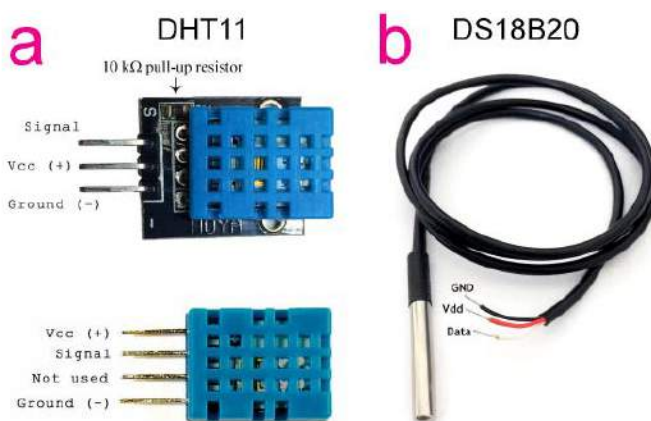


Rysunek 6. Płytkę mojego prototypu, już polutowaną

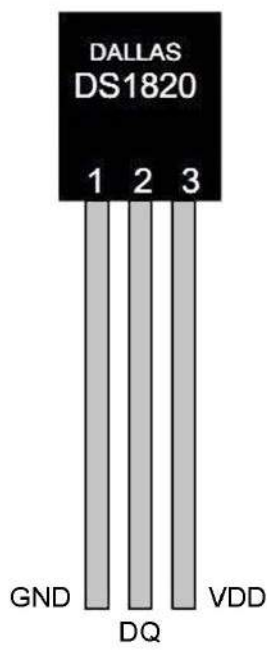
DS18B20 (rysunek 8), który zaizolowałem, w prawdziwym stylu „zrób to sam”, jednym palcem odciętym od gumowej rękawicy.

Następny czujnik to staromodny fotorezystor – raczej pasywny element elektroniczny. Jego rezystancja osiąga w ciemności około 1 M Ω przy oświetleniu $\approx 0,1$ lx, ale spada do około 1 k Ω przy oświetleniu ≈ 100 lx, (w zależności od modelu). Fotorezystor był włączony w układzie dzielnika napięcia z rezystorem i odczytywany przez wejście analogowe mikrokontrolera.

Kolejnym zadaniem był pomiar poziomu wody. Chcę mieć czujnik, który uruchomi alarm, gdy zbiornik na wodę będzie pusty i nadejdzie czas na jego uzupełnienie. Zbudowałem własny czujnik wody. Podstawową zasadą działania czujnika wody jest pomiar przewodności



Rysunek 7. Czujniki: DHT11 (a) i DS18B20 (b)



Rysunek 8. Samodzielna izolacja czujnika DS1820, palcem gumowej rękawicy

elektrycznej między dwiema elektrodami. Jest ona taka sama jak w przypadku czujnika wilgotności gleby. Zorientowałem się, jak mogę skalibrować czujnik, aby spełniał moje wymagania. Po kilku eksperymentach udało mi się to.

Podłączenia LCD

System wymaga wyświetlacza. Mój wyświetlacz LCD używał wielu końcówek – za dużo jak na możliwości mojej płyty. Musiałem więc znaleźć sposób na podłączenie ekranu 128×64 w inny sposób. Udało mi się podłączyć mój wyświetlacz za pomocą tylko trzech końcówek cyfrowych, pozostawiając resztę linii cyfrowych do innych celów. Drugą zaletą jest to, że potrzebuję mniej przewodów do połączenia wszystkich części razem, co pozwala uniknąć zbędnej plątaniny. Rysunek 9 pokazuje wyświetlacz, którego użyłem.



Rysunek 9. Zastosowany wyświetlacz LCD ma rozdzielczość 128×64 piksele i jest podłączony do płytki kontrolera za pomocą tylko trzech linii



Rysunek 10. Kalibracja „Czujnika pH v1.1” jest skomplikowana (Źródło: [10])

Czujnik pH

To najtrudniejszy czujnik, z jakim miałem do czynienia. Trudno było znaleźć informacje o „Czujniku pH v1.1” (rysunek 10). Postanowiłem więc sam go zbadać. Sonda działa jak (mała) bateria po umieszczeniu w wodnistej cieczy. W zależności od poziomu pH wytwarza dodatnie lub ujemne napięcie, o wartości kilku miliwoltów. Dlatego musiałem użyć wzmacniacza operacyjnego, aby wzmocnić ten słaby sygnał do rozsądnego zakresu 0...5 V przyswajalnego dla płyty Arduino. Proces kalibracji tego czujnika został szczegółowo opisany w [3].

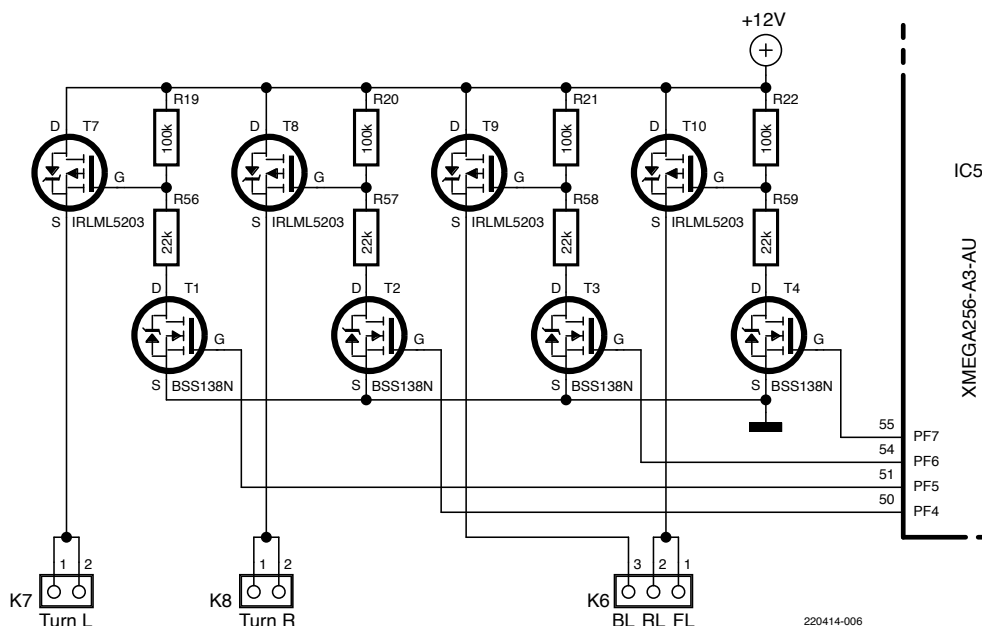
MOSFET-y

Elementy, które wymagają zasilania w moim systemie to taśma LED RGB i atomizer ultradźwiękowy. Do ich załączania użyłem tranzystorów MOSFET sterowanych z wyjść cyfrowych mikroprocesora. Rysunek 11 przedstawia część obwodu z wyjściami MOSFET.

Wygląd zewnętrzny

W moim ostatecznym systemie wykorzystałem wiele technik projektowania 3D i druku 3D. Musiałem używać dwóch różnych drukarek i dużo kombinowałem z ustawieniami, aż uzyskałem pożądane rezultaty. Wyzwanie polegało na tym, że chciałem, żeby wyglądało to naprawdę dobrze. Bardzo ważnym kryterium była estetyka, podobnie jak i funkcjonalność! Chciałem również, aby system można było złożyć, co sprawiło, że było to jeszcze trudniejsze. Chciałem również





Rysunek 11. Ta część obwodu (wykonana za pomocą programu EAGLE) pokazuje stopnie wyjściowe MOSFET

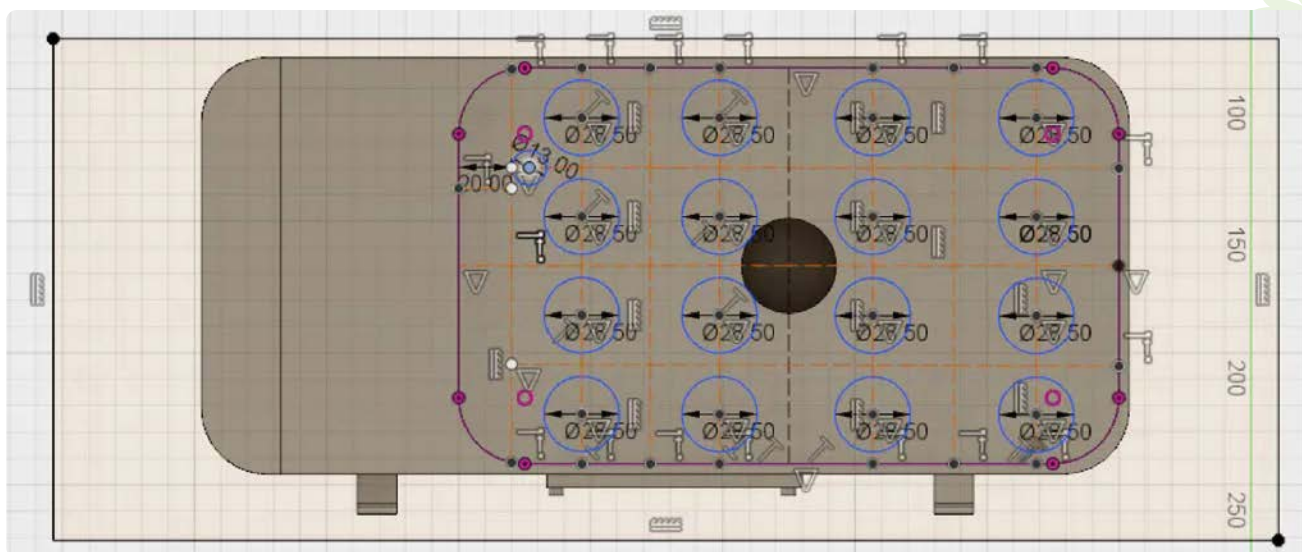


Rysunek 12. Kompletna konstrukcja obudowy i pojemnika na wodę

wykorzystać umiejętności, których się nauczyłem, takie jak drukowanie 3D, frezowanie CNC, cięcie laserowe itp.

Najpierw zrobiłem prosty szkic na papierze, a następnie zagłębiłem się w Autodesk Fusion 360. **Rysunek 12** przedstawia kompletny kontener, który powstał w wyniku tego całego wysiłku. Mój film na YouTube [8] pokazuje, jak to zostało wydrukowane w 3D. Kontener pełni kilka funkcji! Po pierwsze, jest wystarczająco dużo miejsca na rośliny, a ja wywierciłem w środku jego długości otwór na nebulizator ultradźwiękowy, ponieważ chciałem zrównoważyć nebulizator ze zbiornikiem. Rzecz w tym, aby poziom wody znajdował się 2 cm nad nebulizatorem, abyśmy nie wyrównali poziomu wody z wysokością nebulizatora! Gdybym umieścił nebulizator poniżej poziomu zbiornika, miałbym zero wody dokładnie w miejscu, w którym powinno być: zero dla zbiornika = 2 cm nad nebulizatorem. W ten sposób zużywana jest cała woda!

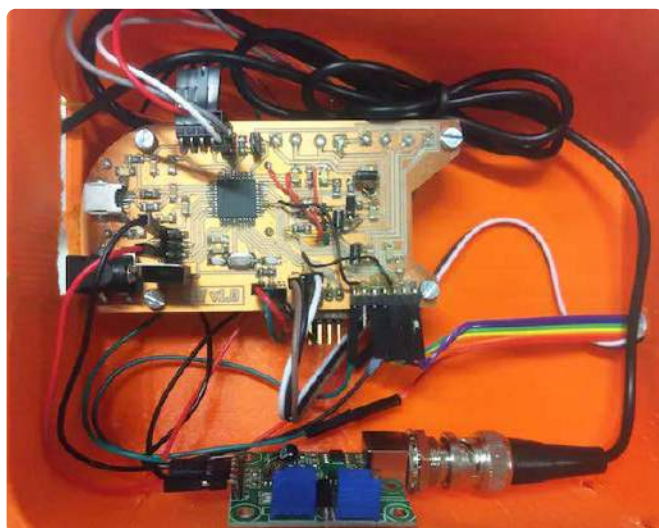
Było dużo więcej do zrobienia. Musiałem wydrukować w 3D tace z siatki dla roślin i wyciąć laserowo mocowanie dla nich (**rysunek 13**). Ogólnie rzecz biorąc, wiele arkuszy akrylowych wymaga cięcia



Rysunek 13. Plan wyrównania otworów dla roślin



Rysunek 14. Zastosowane paski LED RGB



Rysunek 15. Płytką i czujnik znajdują się w sekcji „elektronika”

i wierzenia. Więcej szczegółów, a zwłaszcza sposób, w jaki udało mi się uszczelnić próżniowo, plastikową płytką, pojemnik na wodę, można znaleźć w [3]. Na tej stronie jest więcej linków do filmów z YouTube.

Przystąp do hodowli!

Rysunek 14 pokazuje paski LED, których użyłem. Płytkę drukowaną znajduje się w sekcji „elektronika” (rysunek 15). Aby być jak najlepiej zrozumianym: rysunek 16 pokazuje, jak małe rośliny są trzymane w swoich siatkowych miseczkach, które znajdują się w dopasowanej pokrywie. A teraz proszę o brawa: rysunek 17 przedstawia cały system w akcji. Czy to nie piękne?

Jeśli zainspirowałeś się tym projektem i chcesz zbudować własny system lub jego zoptymalizowaną wersję, wiele informacji znajdziesz na wspomnianej kilka razy stronie projektu [3]. Zanim zaczniesz, pamiętaj jednak, że nie jest to projekt, który można ukończyć w jeden weekend!

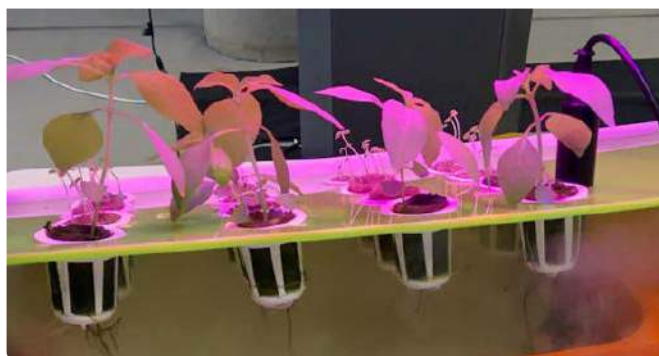
Z ostatniej chwili:

Rozważania autora na temat oświetlaczy LED naśladujących światło słoneczne: <https://shorturl.at/byDNY> ■

Dmitrij Albot (Moldawia)

O autorze

Dmitrij Albot jest absolwentem Fab Academy i byłym koordynatorem FabLab w Jordanii. Obecnie jest założycielem cityfarm (www.cityfarm.md) i ma misję wprowadzenia sztuki, nauki i biznesu AgriTech do miast.



Rysunek 16. Małe rośliny w siatkowych miseczkach, widziane z boku



Rysunek 17. Kompletny system w pełnej okazałości

Pytania lub komentarze?

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

PRZYPADNE LINKI SIECOWE

- [1] TED Talk: „Ten komputer wyhoduje w przyszłości twoje jedzenie”: <https://youtu.be/KJlrd3U1Kxk>
- [2] Fogponics: <https://en.wikipedia.org/wiki/Fogponics>
- [3] Strona projektu na: create.arduino.cc: <https://elektor.link/arduinoigiy>
- [4] Video pokazujące mój system testowy: <https://youtu.be/LF93Xjd8avk>
- [5] satshakit @ GitHub: <https://github.com/satshas/satshakit>
- [6] Dane płytki FabLeo: <https://elektor.link/fableoboard>
- [7] Biblioteka fab.lbr : <https://elektor.link/fablbr>
- [8] Druk 3D obudów (YouTube): https://youtu.be/938Yz_WegH8
- [9] Attribution-ShareAlike 3.0 license: <https://shorturl.at/exLMV>
- [10] Fotografia czujnika pH, v1.1: <https://elektor.link/phsensor11pic>

Podaruj ulubionej miękkiej przytulance bijące serce! Dzięki delikatnemu dźwiękowi i prawdziwemu odgłosowi bicia serca może zrelaksować do snu dziecko, szczeniaka lub kociaka, a nawet może pomóc Ci lepiej spać. Wszystko to jest możliwe dzięki Silicon Chip MiniHEART!

MiniHEART: miniaturowy symulator bicia serca

Wiele osesków – zarówno dzieci ludzkich, jak i zwierząt domowych – jest niespokojnych, gdy są pozostawione same do spania. Tęsknią za mamą i przeraża je samotność. Możliwość przytulenia się do zabawki wydającej odgłos bicia serca może pomóc złagodzić ich niepokój.

MiniHeart to małe urządzenie, które wytwarza kojący dźwięk bicia serca o niewielkim poziomie głośności, naśladujące bicie prawdziwego serca.

Częstotliwość uderzeń można regulować, aby dokładniej odpowiadała rytmowi serca, które ma naśladować, a minutnik wyłączy bicie serca po określonym czasie.

Urządzenie jest włączane i wyłączane za pomocą przełącznika, a dźwignia uruchamiająca tylko nieznacznie wystaje poza obudowę. Ma to zapobiec obrażeniom dziecka. Urządzenie jest w pełni zamknięte w plastikowej obudowie, która jest zatrzaskiwana, a Redakcja Silicon Chip-a dodała dodatkowe wsporniki śrub, aby upewnić się, że pudełko pozostanie zamknięte. W ten sposób dwa wewnętrzne ogniwa AAA nie będą łatwo dostępne, co zabezpiecza dziecko przed ryzykiem zadławienia bateriami wydłubanymi małymi paluszkami.

Zalecamy umieszczenie urządzenia w miękkiej, np. materiałowej lub pluszowej osłonie, która jest zszyta lub zapinana na zamek błyskawiczny. Zapewnia to dodatkowy margines bezpieczeństwa przed potencjalnym niebezpieczeństwem, który jest niezbędny, gdy urządzenie jest używane obok dziecka.

Powinniśmy podkreślić, że symulowane bicie serca nie jest głośnym dźwiękiem – nie ma takiej potrzeby.

Bardziej przypomina subtelny dźwięk prawdziwego bijącego serca; symulator musi być umieszczony blisko ucha i jest bardziej



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/cfw5m>

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

wyczuwalny niż słyszalny. Pomyśl o tym module jak o maleńkim sercu, ale w kształcie zaokrąglonego prostokątnego pudełka.

Głośny dźwięk bicia serca wymagałby dużego głośnika z odpowiednimi rezonatorami do wytwarzania basów oraz wzmacniacza o poważnej mocy.

Żaden z tych elementów nie jest cechą MiniHeart (ale można je dodać zewnętrznie).

Dźwięki serca

Podczas słuchania bicia serca usłyszysz dwa odrębne, oddzielne dźwięki, często nazywane „lub” i „dub”. Jeśli uważasz, że brzmią one raczej „pik” i „puk”, nie będziemy się z Tobą sprzeczać.

Te dwa dźwięki są wytwarzane przez zamykanie zastawek serca, kolejno przedsiłków i komór. Zastawki są niezbędne do wydajnego pompowania krwi – uniemożliwiają cofanie się krwi podczas skurczów serca..

Prawie na pewno widziałeś klasyczny kształt fali bicia serca pokazany na elektrokardiogramie (EKG).

Są to sygnały elektryczne wysyłane do mięśnia sercowego, monitorowane za pomocą elektrod umieszczonych na skórze. Są one przydatne w diagnozowaniu problemów z sercem. Odczyty elektrod nie reprezentują dźwięków i wibracji wytwarzanych przez serce;

dźwięki bicia serca są słyszalne za pomocą stetoskopu.

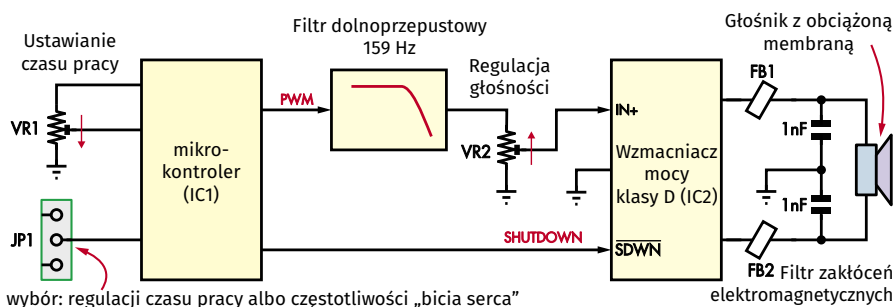
Schemat blokowy urządzenia MiniHeart przedstawiono na **rysunku 1**.

Mikroprocesor IC1 wytwarza przebieg imitujący bicie serca w postaci sygnału modulowanego szerokością impulsu (PWM). Częstotliwość impulsów wynosi 31,25 kHz, a szerokość impulsu jest zmieniana w celu uzyskania wygładzonego kształtu przebiegu o niższej częstotliwości. Usunięcie sygnałów o wysokiej częstotliwości następuje po przejściu przez filtr dolnoprzepustowy, dzięki czemu pozostaje tylko niskotonowy przebieg bicia serca.

Rysunek 2 pokazuje, w jaki sposób sygnał PWM jest wykorzystywany do wytworzenia wygładzonego sygnału o niższej częstotliwości. Czerwony przebieg jest sygnałem wyjściowym PWM z mikroprocesora IC1, podczas gdy zielony przebieg jest jego średnią wartością, po odfiltrowaniu częstotliwości impulsów PWM. Dla wygody pokazujemy sygnał sinusoidalny, chociaż można wygenerować dowolny kształt fali.

Jeśli sygnał PWM ma 50% wypełnienia, tj. równy czas w stanie wysokim i niskim, wówczas przefiltrowane napięcie będzie znajdować się na poziomie połowy potencjału wysokiego poziomu napięcia.





Rysunek 1. Ten schemat blokowy pokazuje, że układ symulatora MiniHeart jest dość prosty, wykorzystując do generowania dźwięku tylko mikroprocesor i wzmacniacz klasy D. Podstawowy filtr dolnoprzepustowy RC zamienia sygnał na wyjściu PWM mikroprocesora na sygnał analogowy dla wzmacniacza, podczas gdy koraliki ferrytowe i kondensatory redukują przenikanie zakłóceń EMI ze wzmacniacza klasy D do głośnika

Aby wytworzyć wyższe napięcie, cykl pracy sygnału PWM jest zmieniany tak, aby czas, gdy jest na wysokim poziomie, był dłuższy niż czas, gdy jest na niskim (tj. cykl pracy z wypełnieniem >50%).

I odwrotnie, dla niższego napięcia, stan PWM jest utrzymywany na niskim poziomie dłużej niż na wysokim (cykl pracy z wypełnieniem <50%).

Zielona krzywa przedstawia sygnał, który pojawia się po usunięciu wszystkich wyższych częstotliwości przez filtr dolnoprzepustowy. Należy pamiętać, że ten sygnał PWM jest tylko prezentacją – w rzeczywistości częstotliwość sygnału PWM jest znacznie wyższa (około 700 razy wyższą!) niż częstotliwość pokazanej fali sinusoidalnej i nie można jej odtworzyć na wykresie w rzeczywistej skali.

Na następnej ilustracji pokazujemy różne przebiegi wytwarzane przez MiniHeart.

Oscylogram 1...oscylogram 3 pokazują ogólne działanie. Oscylogram 1 pokazuje kilka okresów sygnału PWM o częstotliwości około 31 kHz (podstawa czasu 25 μs). **Oscylogram 2 i oscylogram 3** (podstawa czasu 10 ms) to sygnały „lub” i „dub” powstałe po przefiltrowaniu sygnału PWM.

Oscylogram 4 pokazuje pojedyncze uderzenie serca z przebiegami „lub” i „dub”, podczas gdy **oscylogram 5** pokazuje dwa uderzenia serca, z widoczną przerwą między każdym uderzeniem serca.

Okres między każdym uderzeniem serca, częstotliwość przebiegów „lub” i „dub” oraz okres między przebiegami „lub” i „dub” mają dodany niewielki zakres losowości. Ma to zapobiec zbyt sztucznemu brzmieniu bicia serca. Symuluje to zmienność tętna i czasu bicia prawdziwego serca.

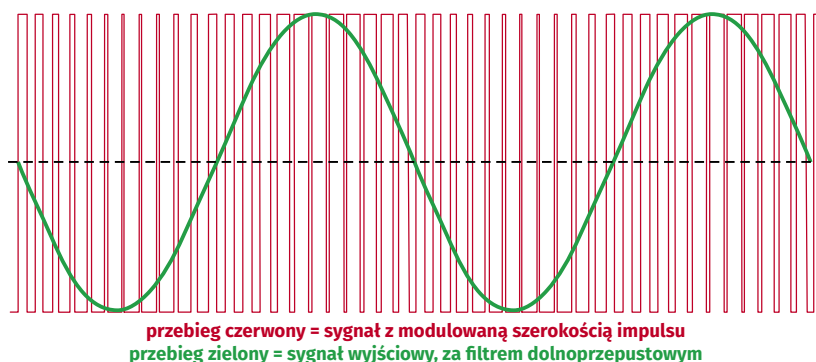
Te przebiegi są podawane do małego wzmacniacza klasy D (tj. przełączającego), który jest zwykle używany w telefonach komórkowych. Został on zaprojektowany tak, aby miał bardzo wysoką sprawność. Zasila on w trybie mostkowym mały głośnik, aby zmaksymalizować moc

wyjściową ograniczoną zasilaniem 3 V DC. Głośnik ma ciężką membranę, tzn. membrana głośnika jest obciążona dodatkową masą. Dzięki temu vibracje o niskiej częstotliwości będą słyszalne i odczuwalne.

Szczegóły obwodu

Pełny schemat ideowy obwodu pokazano na **rysunku 3**. Jego kluczowym elementem (nomen omen „sercem”) jest mikroprocesor PIC12F617, IC1. Wejście IC1, Master Clear (MCLR), styk 4, jest podłączone do szyny zasilającej 3 V za pomocą rezystora 10 kΩ, aby zapewnić funkcję resetowania po włączeniu zasilania.

Układ IC1 podaje, za pośrednictwem wyjścia cyfrowego GP5, napięcie 3 V na potencjometr regulacyjny VR1. Jest ono w stanie wysokim tylko wtedy, gdy pozycja potencjometru jest monitorowana przez wejście analogowe AN3 układu IC1 (styk 3). Po doprowadzeniu wyjścia GP5 do poziomu 3 V, napięcie na AN3 jest konwertowane na wartość cyfrową za pośrednictwem wewnętrznego przetwornika analogowo-cyfrowego (ADC) układu IC1. Po odczytaniu wartości, wyjście GP5 ponownie przechodzi w stan niski (0 V) w celu oszczędzania energii.



Rysunek 2. Pokazuje, w jaki sposób sygnał prostokątny z modulacją szerokości impulsu o wysokiej częstotliwości może być zamieniony przez filtr dolnoprzepustowy na płynnie zmieniający się przebieg o pożądanym kształcie i o niższej częstotliwości (pokazany na zielono jako sinusoida). Chwilowe napięcie zielonego sygnału jest równe średniemu napięciu czerwonych impulsów. W rzeczywistości częstotliwość impulsów prostokątnych jest znacznie wyższa w porównaniu do wynikowego impulsu niskiej częstotliwości

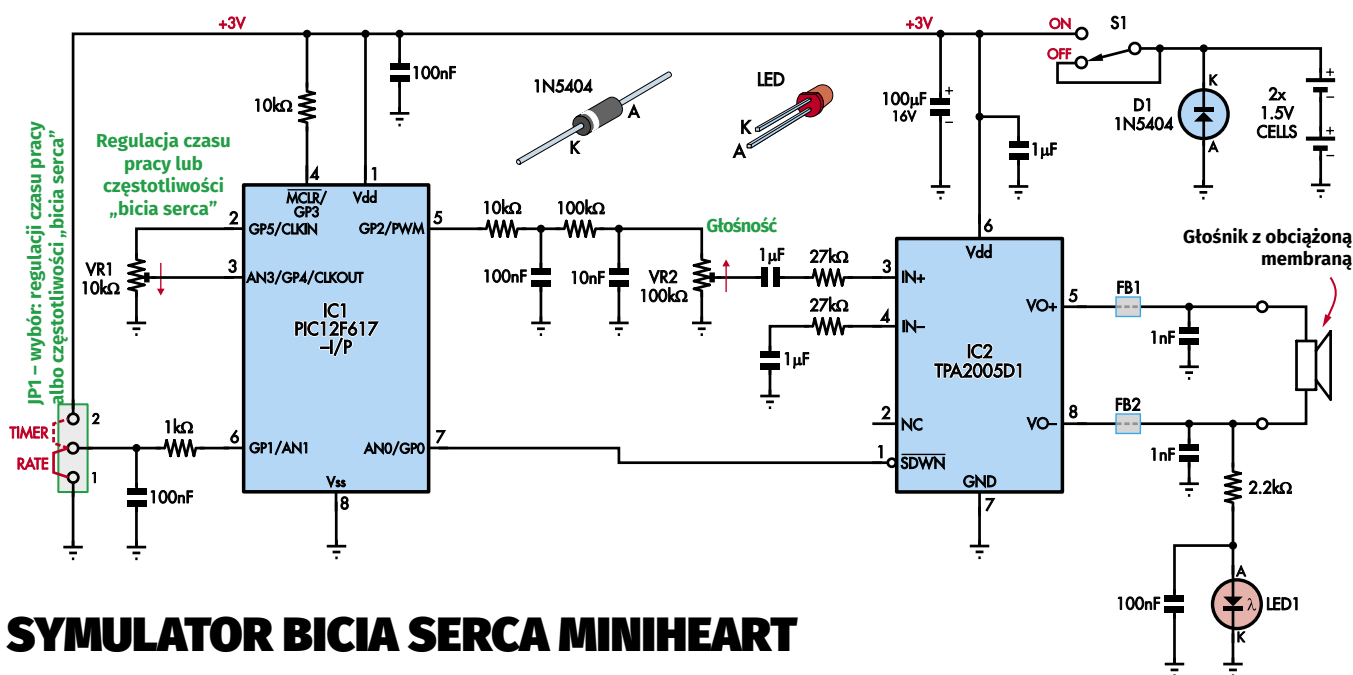
Właściwości i parametry

- Kompaktywny rozmiar
- Regulowana głośność
- Regulowany czas pracy i „tętno”
- Migająca dioda LED zsynchronizowana z „biciem serca”
- Włącznik/wyłącznik zasilania
- Zasilanie: dwa ogniwa AAA (nominalnie 3 V), działające aż do napięcia poniżej 2,5 V
- Pobór prądu: średnio 10 mA podczas pracy, 500 nA w trybie gotowości (typowo)
- Czas pracy: regulowany od dwóch minut do czterech godzin
- Częstotliwość „bicia serca”: od 42 do 114 bpm (uderzeń na minutę)
- Losowość „tętna”: około 15% zmienności
- Częstotliwość dźwięku: 45 Hz...51 Hz (z losowością 2 Hz)
- Metoda generowania kształtu przebiegu: PWM przy 31,25 kHz
- Częstotliwość próbkowania kształtu przebiegu: około 1 kHz

Zworka JP1 może być ustawiona w jednym z dwóch położen: położenie 1, w którym wejście GP1 jest polaryzowane napięciem 0 V (do masy), lub położenie 2, w którym GP1 jest polaryzowane napięciem zasilania 3 V. W pozycji 1, potencjometr nastawny VR1 reguluje częstotliwość bicia serca. Gdy JP1 znajduje się w pozycji 2, VR1 reguluje czas pracy urządzenia.

Częstotliwość bicia serca można ustawić w zakresie od 42 do 114 uderzeń na minutę (BPM). Limit czasu pracy można ustawić w zakresie od dwóch minut do czterech godzin.

Częstotliwość bicia serca można regulować podczas generowania bicia serca, ale limit czasu pracy jest sprawdzany tylko po włączeniu zasilania. Dlatego po zmianie wartości limitu czasu działania za pomocą VR1 należy wyłączyć i ponownie włączyć zasilanie, aby nowy limit czasu zaczął obowiązywać.



SYMULATOR BICIA SERCA MINIHEART

Rysunek 3. Pełny schemat ideowy symulatora MiniHeart nie jest dużo bardziej skomplikowany niż schemat blokowy. Tutaj można zobaczyć szczegóły filtra dolnoprzepustowego drugiego rzędu, kondensatory sprzęgające AC do wejść IC2 i rezystory szeregowo, które ustawiają jego wzmocnienie. Dioda LED1 reaguje na średnie napięcie dostarczane do głośnika, więc zaczyna świecić po wytworzeniu dźwięku

Generowanie tętna wyłącza się po upływie ustawionego czasu. Oszczędza to energię w przypadku pozostawienia włączonego urządzenia.

Jeśli zworka JP1 jest zdjęta, styk 6 wejścia GP1 jest w stanie nieokreślonym. Napięcie na nim może wahać się w zakresie od 0 V do 3 V. Może to prowadzić do wysokiego poboru prądu przez układ IC1, skracając żywotność baterii, ponieważ wejścia cyfrowe powinny znajdować się w jednym lub drugim stanie logicznym.

Dlatego układ IC1 sprawdza ten stan, zamieniając GP1 na wyjście i ustawiając je na wysokim poziomie przez 1 ms. Rezystor 1 kΩ ładuje wtedy kondensator 100 nF do napięcia

3 V. Następnie GP1 jest zamieniany z powrotem na wejście, a jego poziom jest sprawdzany. Jeśli napięcie wejściowe pozostaje wysokie, oznacza to, że albo zworka w pozycji 2 podnosi poziom wejściowy, albo nie ma zworki, a wejście jest utrzymywane w stanie wysokim przez naładowany kondensator 100 nF. Jeśli napięcie wejściowe zmienia się na niskie, oznacza to, że zworka JP1 jest założona w pozycji 1, co kończy test.

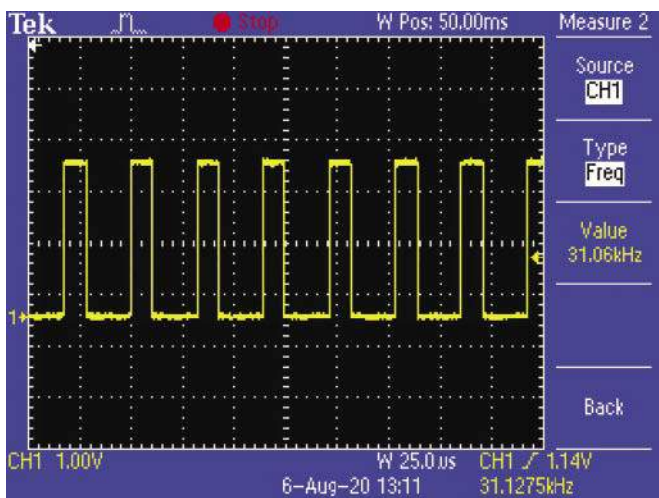
Dla pierwszej możliwości (poziom na GP1 pozostaje wysoki) test ten jest powtarzany przy niskim poziomie wyjściowym. Jeśli poziom uległ zmianie, należy założyć, że zworka JP1 jest założona (w pozycji 2). Aby zapobiec

plywającemu stanowi wejścia przy braku zworki JP1, GP1 jest zamieniane na wyjście o niskim poziomie (0 V) i pozostawiane w tym stanie, minimalizując zużycie energii.

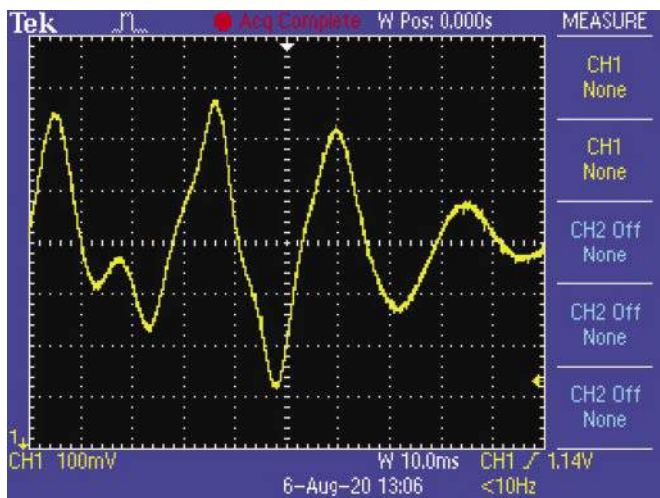
Generowanie bicia serca

Układ IC1 wykorzystuje do generowania sygnału PWM 31,25 kHz wewnętrzny oscylator 8 MHz z wyjściem na styku 5. Sygnał jest podawany do dwustopniowego filtra dolnoprzepustowego RC. Pierwszy stopień składa się z rezystora 10 kΩ i kondensatora 100 nF, aby zapewnić tłumienie -3 dB przy 159 Hz.

Drugi stopień ma taką samą częstotliwość tłumienia, ale wykorzystuje rezystor 100 kΩ



Oscylogram 1. Pokazuje nieco ponad siedem okresów sygnału PWM ~32 kHz, który jest wytwarzany na wyjściu 5 układu IC1. Wartość szczytowa sygnału wynosi 3 V, a podstawa czasu 25 μs



Oscylogram 2. Ten sygnał „lub” symulujący prawdziwy dźwięk bicia serca, wytwarzany jest przez filtrowanie sygnału PWM. Pomiar na styku potencjometru VR2. Należy zwrócić uwagę na dłuższą podstawę czasu (10 ms/działkę)

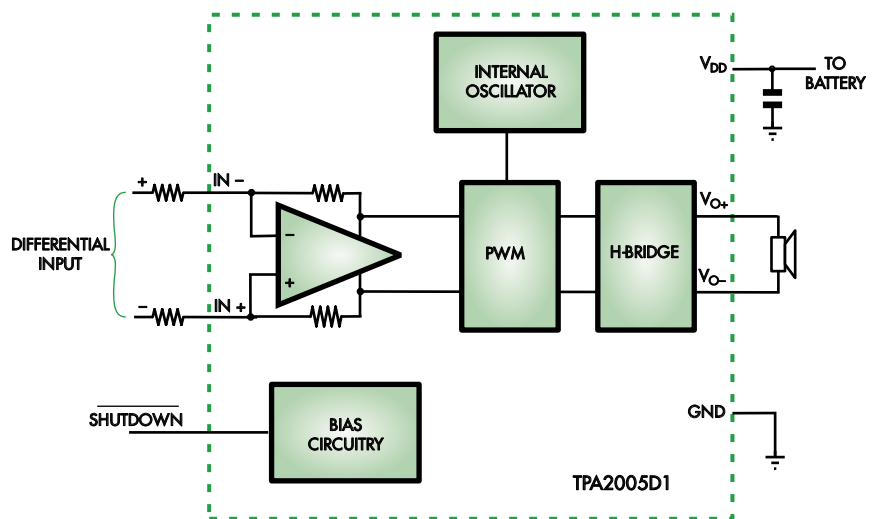
z kondensatorem 10 nF. Komponenty te zapewniają impedancję, która jest 10 razy większa od impedancji filtra pierwszego stopnia, minimalizując obciążenie pierwszego stopnia przez drugi stopień. Przetworzony sygnał jest podawany do potencjometru nastawnego regulacji głośności VR2, a następnie do wejścia nieodwracającego, styk 3, wzmacniacza IC2 poprzez kondensator 1 μ F i rezystor 27 k Ω .

IC2 to wzmacniacz klasy D (tj. przełączający) TPA2005D1 w niewielkiej obudowie SMD o wymiarach zaledwie 3×5 mm. Został on specjalnie zaprojektowany do użytku w telefonach komórkowych, gdzie kluczowe znaczenie ma jego wysoka sprawność. Schemat blokowy wzmacniacza TPA2005D1 pokazano na rysunku 4.

Układ posiada wejścia różnicowe wewnętrznego wzmacniacza. Wzmacniacz zasilają sekcje PWM z częstotliwością przełączania 250 kHz, ustawioną przez wewnętrzny oscylator. Sekcja PWM zasilają obwód mostka H do sterowania zewnętrznym głośnikiem.

Arkusz danych dla TPA2005 podkreśla dwa interesujące punkty. Pierwszym z nich jest wysoki CMRR (common-mode rejection ratio = współczynnik tłumienia sygnału wspólnego), który rzekomo eliminuje potrzebę stosowania wejściowych kondensatorów sprzęgających. Jednak z tym wysokim CMRR mamy do czynienia tylko wtedy, gdy wzmacniacz jest używany w trybie zbalansowanym, z obydwoma wejściami na tym samym poziomie DC.

W naszym układzie używamy go w trybie niezbalansowanym, z wejściem odwracającym uziemionym dla sygnałów akustycznych (przez kondensator 1 μ F), więc musimy użyć dwóch kondensatorów wejściowych. Rezystor 27 k Ω dla wejścia nieodwracającego, w połączeniu z wewnętrznym



Rysunek 4. Wewnętrzny schemat blokowy układu wzmacniacza audio TPA2005 klasy D. Sygnał na jego wejściach różnicowych trafiają do zbalansowanego wzmacniacza analogowego, a następnie do modulatora PWM, który zasilają mostek H na MOSFET-ach, a ten z kolei zasilają głośnik. Zapewnia to wysoką sprawność i dużą moc przy niskim napięciu zasilania. Jak pokazano, układ w trybie klasy D może zasilają głośnik bez filtra

rezystorem sprzężenia zwrotnego 150 k Ω , ustawia wzmacnienie wzmacniacza na około 5,5 raza. Ponieważ wzmacniacz jest typu mostkowego, całkowite wzmacnienie jest dwukrotnie większe, tj. 11 razy.

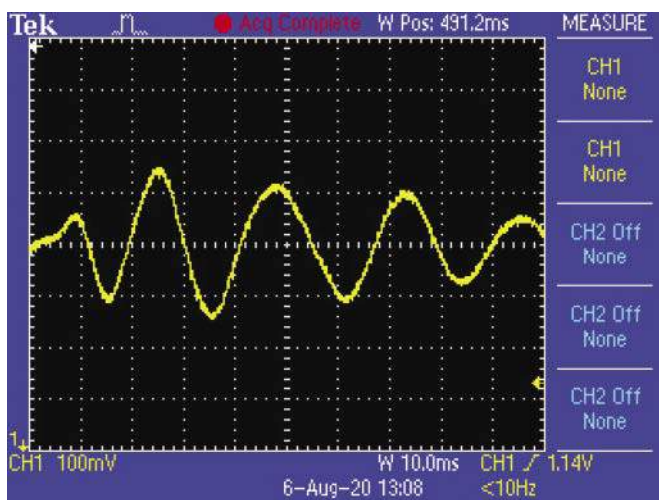
Drugą interesującą kwestią jest to, że TPA2005 może działać bez filtra wyjściowego, który zwykle byłby wymagany do usunięcia sygnału przełączającego 250 kHz. To znaczy może, pod warunkiem, że przewody wyjściowe są krótkie. Mimo to używamy koralików ferrytowych (FB1 i FB2) oraz kondensatorów bocznikujących 1 nF w celu zmniejszenia zakłóceń elektromagnetycznych (EMI).

Zasilanie

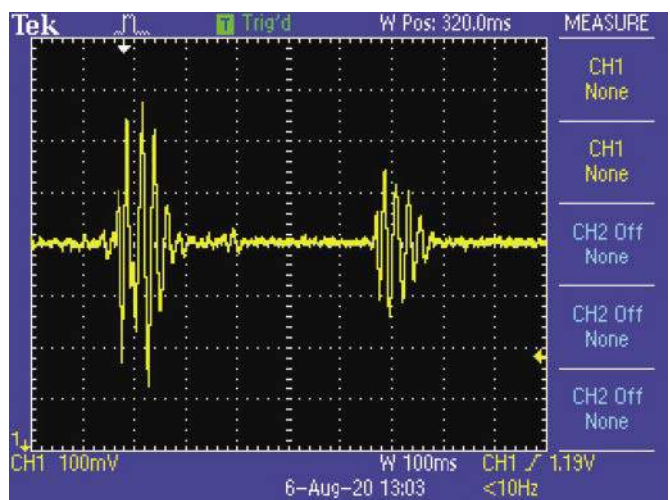
Zasilanie pochodzi z dwóch połączonych szeregowo ogniw AAA, aby zapewnić nominalne

napięcie 3 V, włączane lub wyłączane przez przełącznik zasilania S1. Kondensator 100 μ F łądodzi impulsy włączania/wyłączania zasilania. Razem z kondensatorem ceramicznym 1 μ F w pobliżu wejść zasilających IC2; i kondensatorem 100 nF przy stykach zasilających IC1 obniża też wewnętrzną impedancję źródła zasilania. Obecność tych kondensatorów wynika z impulsowego charakteru pracy IC1 oraz IC2 i odpręża zasilanie.

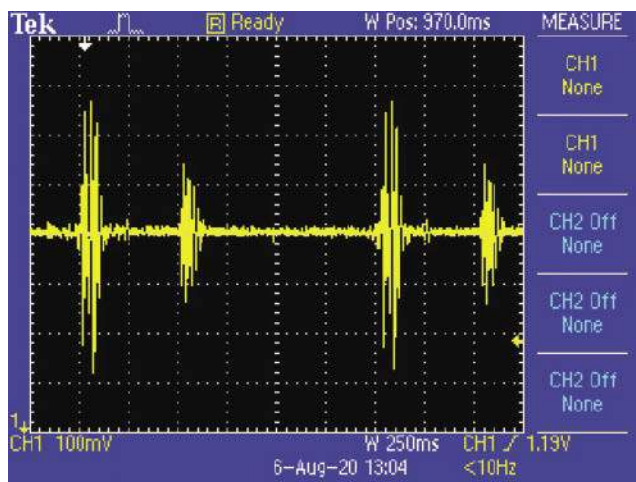
Dioda D1 jest dołączona w celu ochrony przed uszkodzeniem komponentów, jeśli ogniwa zostaną włożone z odwrotną polaryzacją. W takim przypadku dioda będzie przewodzić i ograniczać ujemne napięcie w obwodzie. Wadą tego rozwiązania jest szybkie rozładowanie ogniw (praktycznie przez zwarcie diodą), ale prawdopodobnie użytkownik



Oscylogram 3. Jest to sygnał „dub” mierzony identycznie jak sygnał „lub” pokazany na Oscylogramie 2. Ponownie, jest to symulacja prawdziwego dźwięku bicia serca



Oscylogram 4. Pojedynczy dźwięk bicia serca z przebiegiem „lub” i „dub”. Widać ich nieco inne kształty i amplitudy, a także opóźnienie między nimi.



Oscylogram 5. Dwa uderzenia serca pokazane na oscylogramie 4. Przy wolniejszej podstawie czasu można również zobaczyć opóźnienie między uderzeniami

zauważy, że urządzenie nie działa i natychmiast usunie błąd.

Alternatywna metoda zabezpieczenia, z diodą połączoną szeregowo z zasilaniem, zbyt mocno obniża napięcie dla tego zastosowania. Nawet dioda Schottky'ego, z jej niższym napięciem przewodzenia, nie byłaby odpowiednia. **Od Red. EdW: jak widać, nikt już nie pamięta o archaicznych diodach germanowych, np. DZG6, mających przy 100 mA napięcie przewodzenia <200 mV.** Autor nie mógł też uzasadnić w tej roli MOSFET-a z jego kosztami zakupu (miałby on spadek napięcia przewodzenia rzędu kilku mV). **Od Red. EdW: i słusznie, stosowanie MOSFET-a w tym miejscu nie miałoby też sensu ze względu na pasożytniczą przeciwsobną diodę w jego strukturze.**

Sygnalizacja

Dioda LED1 zapala się jednocześnie z dźwiękami „lub”/„dub” i jest sterowana przez wyjście VO- układu IC2. Bez sygnału napięcie na tym wyjściu wynosi średnio 1,5 V. Jest ono tworzone przez filtr dolnoprzepustowy RC (2,2 kΩ/100 nF) z sygnału fali prostokątnej 250 kHz na styku 8 układu IC2. Waha się między 0 V a 3 V z 50% cyklem wypełnienia w stanie beczynności.

Dioda LED zapala się, gdy napięcie to wzrośnie powyżej zwykłego napięcia przewodzenia czerwonej diody LED wynoszącego około 1,8 V, a dzieje się tak, gdy cykl wypełnienia na wyjściu, styku 8, wzrasta powyżej 60%.

Oszczędzanie energii

Ponieważ urządzenie jest zasilane z baterii AAA, musimy zminimalizować zużycie energii, aby zachować żywotność ogniw. Zazwyczaj obwód pobiera średnio 10 mA podczas generowania bicia serca. Jednak po zakończeniu okresu działania pobór prądu musi spaść do bardzo

niskiego poziomu, aż do wyłączenia urządzenia.

Osiąga się to na kilka sposobów. Po pierwsze, jak już wspomniano, przez większość czasu nie ma napięcia na VR1. Ponadto po upływie limitu czasu mikroprocesor IC1 przechodzi w tryb uśpienia i pobiera tylko około 150 nA.

Wzmacniacz IC2 jest również wyłączany przez IC1, który wystawia niski

poziom na wyjściu GPO, połączonym z wejściem SDWN (wyłączenie) układu IC2. IC2 pobiera wtedy tylko około 500 nA. W naszym prototypie w stanie wyłączenia zmierzaliśmy pobór prądu 500 nA dla całego symulatora bicia serca (pół mikroampera!). Przy tak małym poborze prądu ogniwa powinny wystarczyć na cały okres użytkowania.

Budowa

MiniHeart Simulator jest zbudowany na dwustronnej, przelotowej płytce drukowanej o kodzie 01109201 i wymiarach 70×73 mm. Jest ona umieszczona w wenty-

lowanej obudowie z tworzywa sztucznego o wymiarach 80×80×20 mm.

Rysunek 5 przedstawia schemat montażowy elementów na PCB. Rozpocznij od zamontowania układu wzmacniacza SMD klasy D, IC2. Wymaga on bardzo cienkiego grotu lutownicy i, najlepiej, oświetlonej lupy na gęsiej szyi lub lupy biurkowej (dobra lupa z opaską LED również działa dobrze).

Pod powiększeniem zidentyfikuj jego styk 1, a następnie ustaw go tak, jak pokazano na rysunku 5, ze stykiem 1 w kierunku otworu głośnika. Dodaj trochę pasty topnikowej na środek centralnego pola pod układem (lub płynnego topnika, jeśli nie masz pasty), ostrożnie umieść IC2 na jego polach lutowniczych, a następnie przylutuj styk 4 do jego podłączenia.

Sprawdź, czy układ scalony jest nadal wyrównany z polami lutowniczymi PCB po obu swoich stronach; w razie potrzeby ponownie użyj lutownicy, aby wyrównać układ. Jeśli wszystko jest w porządku, przylutuj pozostałe styki narożne, a następnie styki 2, 3, 6 i 7. Użyj miedzianej plecionki lutowniczej i dobrej pasty, aby usunąć wszelkie mostki lutownicze między końcówkami układu scalonego.

IC2 ma na spodzie podkładkę uziemiającą, którą należy przylutować do płytki drukowanej. Można to zrobić, podając lut od spodu płytki drukowanej przez otwór umieszczony pod układem scalonym. Użyj minimalnej ilości lutu, aby zapobiec szerokiemu rozplątaniu

Wykaz elementów, kupuj w sklepie [avt.pl](http://www.avt.pl) (W-wa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl):

- 1 dwustronna, przelotowa płytka drukowana o wymiarach 70×73 mm i kodzie 01109201
- 1 wentylowana obudowa Hammond 1151V4, 80×80×20 mm [Jaycar HB6118]
- 2 uchwyty na ogniwa AAA do montażu na płytce drukowanej
- 2 ogniwa alkaliczne AAA
- 1 głośnik z membraną Mylar o średnicy 40 mm [Jaycar AS3004] i impedancji 8 Ω lub
- 1 głośnik z membraną Mylar o średnicy 40 mm i impedancji 18-32 Ω ze starych słuchawek nausznych
- 1 przełącznik SPDT (S1) do montażu na płytce drukowanej [Altronics S1421]
- 1 podstawa DIL-8 dla IC
- 2 koralki ferrytowe o średnicy 4 mm i długości 5 mm (FB1, FB2) [Altronics L5250A, Jaycar LF1250]
- 13-szpilekowy odcinek prostej listwy kołkowej, raster 2,54 mm plus zworka (JP1)
- 2 gwintowane kotki dystansowe M3 o długości 9 mm
- 2 śruby z łbem walcowym M3x6
- 4 wkręty samogwintujące nr 4
- 2 poliamidowe śruby M3x6 (preferowany też stopżkowy)
- 1 nakrętka M8 ze stali nierdzewnej niemagnetycznej 316 (klasy morskiej), odpowiednik ISO: stal A4 (wysokość 6,35 mm)
- 1 40 mm pocynowanego drutu miedzianego o średnicy 0,7 mm (dla FB1 i FB2)
- 1 przewód potężnościowy o długości 100 mm (lub 2-drożny kabel taśmowy albo przewód o przekroju „8”)
- 1 mała tubka neutralnie utwardzanego uszczelnacza silikonowego (np. silikon dachowy i rynnowy)

Półprzewodniki:

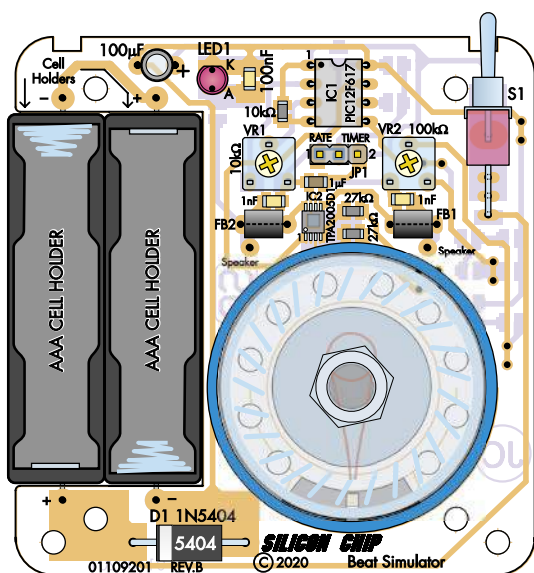
- 1 mikrokontroler PIC12F617-I/P zaprogramowany kodem 0110920A.hex (IC1)
- 1 monofoniczny wzmacniacz klasy D bez filtra TPA2005D1DGNRQ1 1,4 W (IC2)
- 1 dioda 1N5404 3 A (D1)
- 1 czerwona dioda LED 3 mm o wysokiej jasności (LED1)

Kondensatory:

- 1 kondensator elektrolityczny 100 μF 16 V PC
 - 3 kondensatory ceramiczne 1 μF 6.3 V SMD 1206 X7R*
 - 4 kondensatory ceramiczne 100 nF 50 V SMD 1206 X7R
 - 1 kondensator ceramiczny 10 nF 50 V SMD 1206 X7R
 - 2 kondensatory ceramiczne 1 nF 50 V SMD 1206 X7R
- * typ Y5V działał w prototypie SC, ale lepszym wyborem będzie X5R lub X7R.

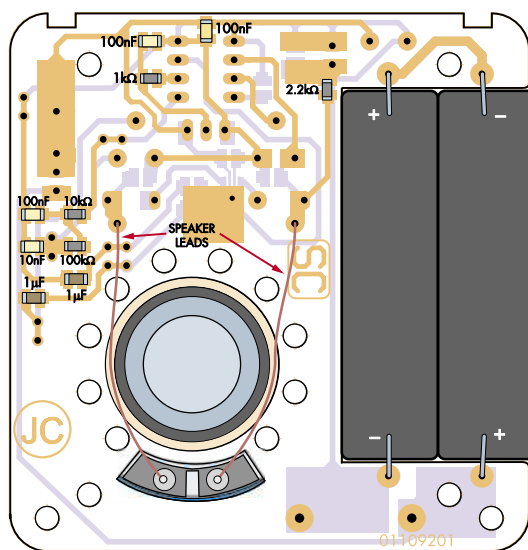
Rezystory: (wszystkie 1% SMD 1206)

- 1 szt. 100 kΩ
- 2 szt. 27 kΩ
- 2 szt. 10 kΩ
- 1 szt. 2,2 kΩ
- 1 szt. 1 kΩ
- 1 mini potencjometr nastawny poziomy 10 kΩ (VR1)
- 1 mini potencjometr nastawny poziomy 100 kΩ (VR2)

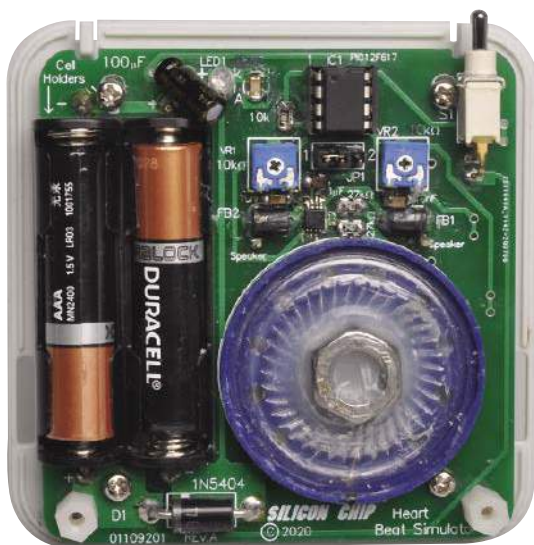


TOP VIEW

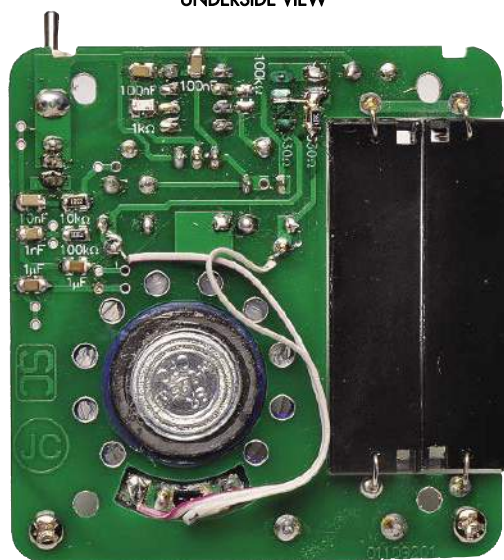
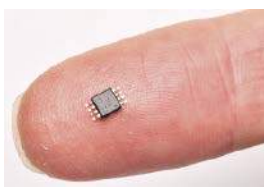
Rysunek 5. Te (i pasujące zdjęcia poniżej) pokazują, gdzie są zamontowane komponenty po obu stronach płytki drukowanej. Ogólnie rzecz biorąc, przed lutowaniem elementów przewlekanych najlepiej jest zamontować wszystkie elementy SMD na górnej stronie (i ewentualnie również na dolnej), ze względu na ich mały rozmiar i niewielką wysokość. Zwróć uwagę, w jaki sposób zorientowany jest głośnik, aby jego zaciski pasowały do wycięcia w płytce, a także jak są wygięte przewody pojemników baterii, aby pasowały do pól lutowniczych PCB. Są one wygięte od spodu do góry, przeprowadzone od spodu przez otwory w PCB i przylutowane na górze płytki



UNDERSIDE VIEW



IC1 to zwykły IC w obu-
dowie DIP-8 ... ale IC2
(TPA2005D1DGNRQ1) jest
małutki (pokazano go poni-
żej w naturalnej wielkości).
Słowo ostrzeżenia: nie
kichaj ani nie włączaj wen-
tylatora, jeśli jeszcze kie-
dykolwiek chcesz ten układ
zobaczyć...



się lutowia i zwarcia do wyprowadzeń układu scalonego.

Dodany wcześniej topnik pomoże lutowi spłynąć na podkładkę na spodzie układu scalonego.

Teraz zainstaluj rezystory i kondensatory do montażu powierzchniowego. Elementy te znajdują się po obu stronach płytki drukowanej. Kondensatory są zwykle nieoznaczone, za wyjątkiem opakowania zbiorczego.

Rezystory będą prawdopodobnie oznaczone skróconym kodem. Pierwsze kilka cyfr wskazuje wartość rezystancji, po której na ostatniej pozycji następuje liczba dodatkowych zer. Na przykład rezystor 1 kΩ będzie miał kod 102 lub 1001. Jest to 10 z dwoma zerami lub 100 z jednym zerem. Dla 10 kΩ kod będzie wynosił 103 lub 1002 itd.

Następnie zamontuj diodę D1, zwracając uwagę na jej prawidłową orientację. Potem przylutuj koraliki ferrytowe FB1 i FB2,

najpierw przeprowadzając pocynowany drut miedziany przez środkowy otwór, a następnie wkładając i lutując drut do otworów w PCB. Podczas lutowania należy drut przytrzymać, aby zapobiec poluzowaniu się koralików.

Użyliśmy podstawki dla IC1 na wypadek, gdybyś kiedykolwiek chciał go wyjąć w celu przeprogramowania. Należy zwrócić uwagę na prawidłową orientację gniazda (wycięcie w kierunku krawędzi płytki drukowanej).

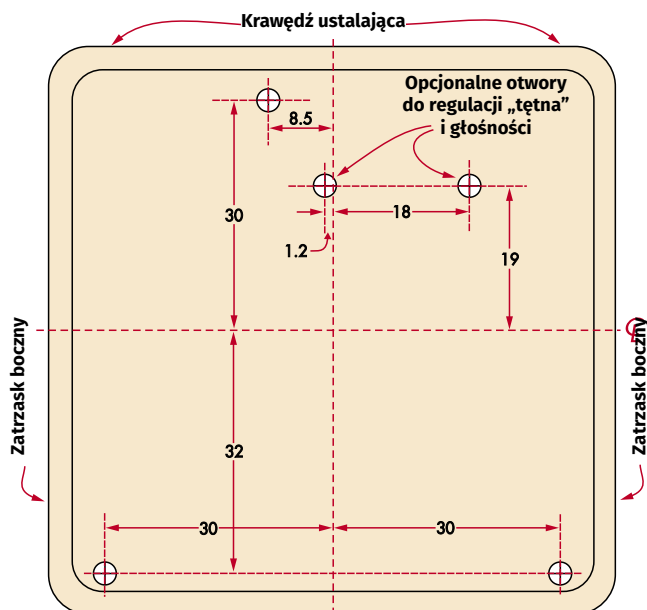
Teraz można zamontować potencjometry nastawne VR1 i VR2. Należy uważać, aby umieścić ten o rezystancji 10 kΩ w pozycji VR1, a drugi, o rezystancji 100 kΩ, w pozycji VR2. Następnie należy zamontować trójszpilkowy odcinek prostej listwy kołkowej pod złącze JP1 z krótszymi końcami szpilek lutowanymi do otworów w płytce drukowanej.

Przełącznik zasilania (S1) jest instalowany w pokazanej pozycji.

Użyty przez SC w prototypie przełącznik różni się nieco od tego z listy części, ponieważ dźwignia w zalecanym przełączniku jest dłuższa. Dlatego mocowanie przełącznika zostało przesunięte dalej od krawędzi płytki drukowanej. W ten sposób dźwignia przełącznika będzie wystawać z obudowy mniej więcej tyle, ile w przypadku pierwotnego prototypu.

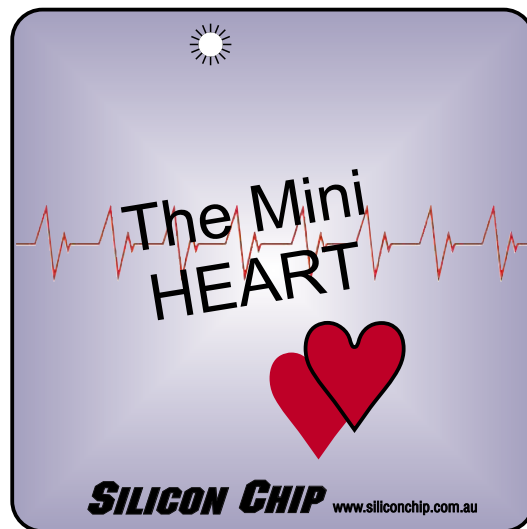
Dioda LED1 jest montowana anodą (dłuższym wyprowadzeniem) w otworze oznaczonym „A”. Przylutuj ją tak, aby górna część soczewki znajdowała się 11 mm nad górną powierzchnią płytki drukowanej

W przypadku uchwytów na ogniwa AAA należy wygiąć końcówki przewodów tak, aby wystawały z krótszych boków uchwytów, a następnie wygiąć je w górę, aby poprowadzić przewody przez otwory w płytce drukowanej od spodu i przylutować je na górze. Uchwyty ogniw muszą być prawidłowo zorientowane, jak pokazano na schemacie montażowym.



Wszystkie otwory o średnicy 3 mm  Wszystkie wymiary w mm

Rysunek 6. Ten schemat wiercenia pokrywki pokazuje lokalizację 3 mm otworu na diodę LED, dwóch 3 mm otworów do mocowania pokrywki (na dole) i opcjonalnych otworów umożliwiających dostęp do potencjometrów nastawnych bez konieczności zdejmowania pokrywki



Rysunek 7. Grafika „panelu przedniego” z otworem na diodę LED. Zobacz link do naszej strony internetowej w tekście, aby dowiedzieć się, jak ją wydrukować i przymocować do pokrywki. Plik PDF z grafiką można pobrać ze strony Silicon Chip-a.

Podstawy pojemników baterii powinny być ustawione tak, aby znajdowały się na podstawie obudowy, gdy płytka drukowana jest osadzona na czterech wspornikach dystansowych spodu obudowy. Oznacza to, że denka obudów baterii będą niżej niż dolna krawędź płytki drukowanej.

Następnie zamontuj kondensator 100 μ F. Włóż jego przewody, z dłuższym przewodem przez otwór oznaczony +, a następnie zegnij go tak, aby korpus kondensatora znajdował się między diodą LED a uchwytem ogniwa AAA. Kondensator nie może znajdować się wyżej niż 11 mm nad górną krawędzią płytki drukowanej. Pozwoli to na dopasowanie pokrywki.

Można teraz wltować dwa kołki PC do podłączenia głośnika, krótszym końcem włożonym od góry do płytki drukowanej.

Na tym etapie nie należy jeszcze wkładać mikroprocesora PIC (IC1). Jeśli zakupiłeś PIC12F617-I/P do tego projektu w sklepie Silicon Chip online shop, będzie on już miał załadowane oprogramowanie układowe (0110920A.hex). Jeśli chcesz go zaprogramować samodzielnie, plik można pobrać ze strony internetowej Silicon Chip-a.

Obudowa

Wciśnij boczne zatraski pokrywki obudowy, aby oddzielić ją od podstawy. Pokrywa jest zabezpieczona na miejscu w podstawie przy pomocy kołnierzy ustalających podstawy pasujących do jednej krawędzi pokrywki.

Płytkę PCB została zaprojektowana do montażu na zintegrowanych wspornikach podstawy

obudowy. Istnieje tylko jedna prawidłowa orientacja, a mianowicie z dwoma wycięciami wzdłuż górnej krawędzi PCB pasującymi do kołnierzy blokujących pokrywę obudowy w podstawie. Płytkę PCB jest przymocowana do wsporników za pomocą małych wkrętów samogwintujących.

Do płytki PCB dołączamy dwa gwintowane poliamidowe kołki dystansowe M3 o długości 9 mm, aby umożliwić przykręcenie pokrywki (dobrze widać je na fotografii na przedostatniej stronie). Jest to dodatek do bocznych klipsów pokrywki, które blokują ją na miejscu. Następnie dwie śruby są wkręcone w kołki dystansowe od zewnętrznej strony pokrywki. Przymocuj te elementy dystansowe, wprowadzając przez dwa narożne otwory od spodu płytki drukowanej krótkie gwintowane śruby M3, a następnie dokręć kołki dystansowe na trzpieniach śrub – widać je na prawej fotografii pod rysunkami 5.

Szablon wiercenia pokrywki (rysunek 6) pokazuje położenie dwóch otworów wymaganych dla śrub mocujących. Pokazuje również lokalizację otworu na diodę LED i dwóch opcjonalnych otworów dostępowych do regulacji potencjometrów nastawnych. Te ostatnie możesz pominąć, jeśli zamierzasz otwierać obudowę w celu dokonania jakichkolwiek

regulacji. Zresztą bez otwarcia obudowy nie ma dostępu do złącza JP1!

Grafika etykiety na pokrywę (rysunek 7) jest również dostępna do pobrania ze strony internetowej SC. Szczegóły dotyczące drukowania i dołączania grafiki panelu można znaleźć na stronie <https://tiny.pl/c967z>.

Testowanie

Umieść zwórkę JP1w pozycji 1 i podłącz dwa przewody o długości około 80 mm do dwóch kołków PC pod płytką drukowaną w celu przylutowania do miniaturowego głośnika 8-omowego. Użyliśmy dwóch odizolowanych



Ten widok pokazuje, w jaki sposób płytka drukowana jest przymocowana do pokrywki podstawy obudowy, ale co ważniejsze, pokazuje „tłumik” (obciążnik) przyklejony do poliestrowej membrany głośnika (w tym przypadku nakrętka ze stali nierdzewnej). Nie ulegaj pokusie użycia nakrętki ze zwykłej stali: są one magnetyczne i natychmiast przylgną do magnesu głośnika blokując membranę. Od Red. EdW: możesz oczywiście użyć nakrętki z też niemagnetycznego mosiądzu



na końcach przewodów z tężowego kabla taśmowego; wszelkie podobne przewody również działałyby dobrze.

Głośnik montuje się na górze płytki drukowanej z zaciskami głośnikowymi w wyciętym obszarze. Przewody podłącza się do zacisków głośnika od spodu płytki drukowanej. Na razie głośnik nie będzie zablokowany w PCB.

Włóż dwie baterie AAA i włącz zasilanie. Sprawdź, czy między stykami 1 i 8 gniazda IC1 jest napięcie około 3 V.

Odłącz zasilanie i włącz zaprogramowany układ PIC do podstawki, upewniając się, że jest prawidłowo ustawiony (wycięciem w kierunku krawędzi płytki drukowanej). Ponownie podłącz zasilanie, a głośnik powinien zacząć się poruszać w odpowiedzi na generowany dźwięk „lub” – „dub”. Jeśli nie, upewnij się, że VR2 jest ustawiony przynajmniej częściowo zgodnie z ruchem wskazówek zegara. Dalsze skrócenie w prawo zapewni większą siłę dźwięku.

Należy pamiętać, że dźwięk będzie miał ton tła o częstotliwości około 1 kHz. Dzieje się tak dlatego, że chociaż ton ten jest odfiltrowywany w obwodzie, głośnik jest znacznie bardziej wydajny w wytwarzaniu tonu 1 kHz w porównaniu z niskotonowymi dźwiękami „lub” – „dub”, o częstotliwości około 47 Hz. Należy również pamiętać, że dźwięk „lub” – „dub” tak naprawdę nie będzie słyszalny, ale będzie wyczuwalny po umieszczeniu palca na środku membrany głośnika.

Membrana głośnika musi zostać obciążona, aby bicie serca było słyszalne i aby zapobiec odtwarzaniu tonów o wyższej częstotliwości. W tym celu używamy nakrętki M8 ze stali nierdzewnej (niemagnetycznej A4 wg ISO) jako obciążenia membrany głośnika. Należy użyć nakrętki niemagnetycznej; w przeciwnym razie membrana głośnika

byłaby na stałe docięnięta do magnesu głośnika przez nakrętkę.

Udało nam się tego uniknąć, ponieważ membrana głośnika jest wykonana z poliestru Mylaru, a więc jest dość mocna i sztywna. Oznacza to, że cewka głośnika jest nadal wyśrodkowana w szczeliny magnesu, nawet przy dodatkowej masie obciążnika.

Aby przymocować nakrętkę, należy nałożyć niewielką ilość neutralnie utwardzanego silikonu (idealny jest silikon do dachów i rynien) na jedną stronę nakrętki i przycisnąć ją centralnie do membrany głośnika. Dodatkowa porcja silikonu jest wymagana do wypełnienia wnętrza nakrętki. Upewnij się, że wewnątrz nakrętki jest wypełnione silikonem aż do membrany. Na górze nakrętki silikon powinien znajdować się równo z jej krawędzią. Nałóż również ciekłą warstwę silikonu na membranę głośnika wokół nakrętki.

Podczas pracy warto zabezpieczyć koraliki ferrytowe (FB1 i FB2) za pomocą silikonu, aby unieruchomić je na płycie drukowanej. Konieczna jest tylko niewielka ilość. Zapobiegnie to grzechotaniu i dodawaniu nieokreślonych dźwięków do „bicia serca”.

Głośnik jest przymocowany do płytki drukowanej również za pomocą silikonu wokół centralnego magnesu, który pasuje do otworu w płycie drukowanej.

Należy pamiętać, że głośnik musi być prawidłowo umieszczony, z punktami lutowania przewodów zasilających umieszczonymi nad wycięciem w płycie drukowanej i z tylną częścią magnesu głośnika spoczywającą na podstawie obudowy.

Podczas utwardzania silikonu płytka drukowana powinna być tymczasowo umieszczona na zintegrowanych wspornikach w obudowie. W ten sposób głośnik znajdzie się na odpowiedniej wysokości nad płytką drukowaną.

Korzystanie z symulatora

Dostosuj czas działania urządzenia, aby dźwięk bicia serca trwał przez wymagany czas. Odbywa się to za pomocą zworki JP1 w pozycji 2. Aby to zrobić, ustaw JP1 w pozycji 2 przy wyłączonym zasilaniu i ustaw wymagany czas. Pełne skrócenie potencjometru VR1 w prawo daje 4-godzinny limit czasu. Położenie środkowe to dwie godziny, a położenie ok. 25% skrócenia to około jedna godzina.

Ustaw limit czasu działania, a następnie włącz zasilanie. Limit czasu działania zostanie zapisany. Wszelkie dalsze regulacje VR1 przy włączonym zasilaniu zostaną zignorowane. Rejestrowane jest tylko ustawienie VR1 po włączeniu zasilania, gdy zworka JP1 znajduje się w pozycji 2. Ustawienie jest przechowywane w nieulotnej pamięci flash i zapamiętywane do użycia następnym razem.

Gdy zworka 1 znajduje się w pozycji 1, można regulować częstotliwość bicia serca. Można ją zmienić po włączeniu zasilania, w zakresie od 42 do 114 uderzeń na minutę. Ustawienie jest również przechowywane w pamięci flash, a zostanie użyte ostatnie ustawienie, jeśli urządzenie zostanie włączone ze zworką JP1 w pozycji 2.

Głośność ustawia się za pomocą potencjometru nastawnego VR2. Jednak dźwięk głośnika zostanie zniekształcony, jeśli VR2 zostanie obrócony zbyt mocno w prawo, więc należy użyć pozycji mniejszej niż połowa skrócenia. ■

John Clarke

Adaptacja do wydania
polskiego – Andrzej Nowicki

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

REKLAMA

Publikujemy dla projektantów i programistów elektroniki. Odwiedź

ELPORTAL.pl

Znajdziesz nas również na Facebooku: facebook.com/ElportalPL

Gdy chcesz, aby Ci nie przeszkadzać. Jestem zajęty – dla wszystkich!

OK, to wygląda trochę żartobliwe... ale może mieć inne, poważniejsze zastosowania. Busy Dunny Door Warning (Nie wchodź tutaj!) miga jasnym światłem LED na drzwiach, gdy, hmmm, nie chcesz, aby ktoś wkraczał do środka. Po wyjściu i otwarciu drzwi światło gaśnie! To prosty pomysł z naprawdę prostym obwodem – a przy tym jest to świetny projekt dla początkujących...

Pomysł na ten mały projekt zrodził się, gdy zapalony czytelnik Silicon Chip-a, John Chappell, siedział i czytał swój najnowszy egzemplarz czasopisma z pasjonującym projektem... i drzwi do, hmmm, krępującego pomieszczenia otworzyły się, z oczywistym zażenowaniem otoczenia.

Być może zajęło mu to trochę więcej czasu niż zwykle; być może był tak pochłonięty magazynem, że nie usłyszał, jak ktoś dobija się i krzyczy... ale zaczęło go to zastanawiać, jak uniknąć takiej delikatnej sytuacji w przyszłości.

Jednym z problemów było to, że zamki w drzwiach, hmmm, nie działały.

Więc jak bez wymiany zamka poinformować innych, że najlepsze miejsce w domu jest zajęte – bez zażenowania?!

Pomysł z lampką LED

Oczywiście, to była odpowiedź: jasna, migająca dioda LED, która dawałaby innym znać, by się nie wtrącać.

Gdyby była ona w pewnym stopniu automatyczna – tj. wyłączała się, gdy drzwi do WC otwierają się, aby go wypuścić, tym lepiej. Rezultatem jest ten naprawdę prosty obwód.

Po naciśnięciu przycisku (S1), zarówno dioda LED zamontowana na drzwiach, jak i wewnętrzna dioda LED zaczynały migać.

Dlaczego dwie diody LED? Jedna z nich to bardzo jasna ostrzegawcza dioda LED zamontowana na (lub przez) drzwiach toalety, aby ostrzec innych, że pomieszczenie jest zajęte. Druga (wewnętrzna) dioda LED jedynie potwierdza, że obwód działa.

Od Red. EdW: zapobiega to też w pewnym sensie wpadnięciu w trans letargiczny...

Przesada? Być może – ale koszt diody LED – 0,10\$ i rezystora – 0,05\$ nie zwiększa kosztów projektu.

Po otwarciu wspomnianych drzwi, magnetyczny zestyk z kontaktronem resetuje obwód,



a diody LED wyłączają się. To naprawdę takie proste!

Jak powiedzieliśmy wcześniej, jest to świetny projekt dla początkujących. Części są tanie jak chipsy; jest zasilany bateryjnie (a bateria wystarcza na wiele lat) i nie wykorzystuje żadnego z tych nieznoszących podzespołów do montażu powierzchniowego, z którymi początkujący mają tak wiele trudności podczas lutowania. I oczywiście nie wymaga programowania mikroprocesora.

Całkowity czas montażu nie powinien przekroczyć godziny.

Obwód

Schemat ideowy jest pokazany na **rysunku 1** – i jak widać, nie zawiera zbyt wielu części!

Wykorzystuje on układ 4093B – cztery 2-wejściowe bramki NAND CMOS z przerzutnikami Schmitta (IC1). Jeśli jesteś przerażony, nie martw się: zobacz ramkę „Co to jest bramka NAND?”, a wszystko się wyjaśni.

Każda z czterech bramek NAND jest skonfigurowana w odmienny sposób.

IC1a jest inwerterem: gdy jej wejścia są na niskim poziomie, wyjście jest na wysokim (i odwrotnie).

Gdy drzwi są zamknięte, magnes powoduje zamknięcie „normalnie otwartego” kontaktronu, co z kolei oznacza, że oba wejścia IC1a są w stanie niskim, a wyjście jest w stanie wysokim.

Bramki IC1b i IC1d tworzą zatrask, z wejściami IC1d normalnie (gdy drzwi są otwarte), w stanie wysokim. Zatem wyjście 11 IC1d normalnie jest w stanie niskim. Pomyśl o zatrasku jak o bistabilnym przełączniku dwupozycyjnym: normalnie jest w stanie otwartym, ale gdy ktoś go uruchomi, przechodzi w stan zamknięty, i pozostaje w nim tak długo, aż nie zostanie ponownie otwarty.

W tym przypadku, po naciśnięciu przycisku („Start”), zatrask jest resetowany poprzez wymuszenie niskiego poziomu wejścia 12 układu IC1d, co wymusza wysoki poziom na wyjściu 11. Czyli mamy kontaktron zwarty przez zamknięcie drzwi, więc wyjście 3 IC1a i wejście 5 IC1b są w stanie wysokim. Przyciskiem „Start” wymuszamy niski poziom na wejściu 12 IC1d, czyli wyjście 11 IC1d przechodzi w stan wysoki, wymuszając stan wysoki na wejściu 6 IC1b. Wyjście 4 IC1b przechodzi w stan niski, analogicznie jak wejście 13 IC1d. Ponieważ wejście 12 IC1d, po zwolnieniu przycisku „Start”, jest na wysokim poziomie dzięki rezystorowi 100 kΩ, wyjście 11 IC1d



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/cfw5g>
Materiały dodatkowe są również dostępne na stronie elportal.pl/ do-pobrania



Tak, patrzysz na cały projekt! Po lewej stronie znajduje się kontaktron i magnes, które wyłączają diodę LED po otwarciu drzwi. Po prawej stronie znajduje się zamontowana na drzwiach bardzo jasna dioda LED, podczas gdy wewnętrzna dioda LED jest w tym przypadku zintegrowana z przyciskiem „Start”.

jest wtedy nadal w stanie wysokim, i w takim stanie układ się „zatrzaskuje”.

Umożliwia to również układowi IC1c, z rezystorem 47 kΩ i kondensatorem 10 μF, rozpoczęcie oscylacji, z wyjściem przechodzącym na przemian w stan wysoki (kondensator rozładowany) i niski (kondensator naładowany) z częstotliwością określoną przez czas ładowania i rozładowywania kondensatora tantalowego – w tym przypadku częstotliwość wynosi około jednej sekundy.

Gdy wyjście 10 IC1c przechodzi w stan niski, dwie diody LED połączone szeregowo między jego wyjściem 10 a zasilaniem +9 V zostają spolaryzowane w kierunku przewodzenia i świecą.

Nawiasem mówiąc, szybkość błyskania można zmienić, zmieniając rezystor i/lub kondensator. Zwiększenie jednego z nich (lub obu) spowolni szybkość, jakiej można się spodziewać, zmniejszenie przyspieszy szybkość migania.

Gdy drzwi zostaną otwarte, kontaktron otworzy się (gdy magnes się odsunie), wejścia IC1a przejdą w stan wysoki dzięki rezystorowi spolaryzującemu 100 kΩ podłączonemu do zasilania 9 V i obwód powróci do stanu uśpienia.

Dla lepszego zrozumienia, zamieszczamy obok tabelkę stanów logicznych wszystkich końcówek układu IC1 w 4-ch możliwych sytuacjach. Warto zauważyć, że diody LED będą migać również w sytuacji, gdy przycisk „Start” zostanie zwarty na stałe.

Zasilanie

Obwód jest zasilany pojedynczą baterią 9 V, która ze względu na sporadyczne rozładowywanie powinna wytrzymać prawie tak długo, jak jej okres przydatności do użycia. Z tego samego powodu włącznik/wyłącznik nie jest przewidziany ani potrzebny. (Oczywiście, jeśli zdecydujesz się czytać „Wojnę i Pokój” podczas swoich „wizyt”, możesz nie uzyskać takiej żywotności baterii).

Przewody z klipsem zatrzaskowym do baterii można podłączyć do gniazda złącza 402-2,

z wtykiem 403-2 wlotowym na PCB, lub poprowadzić pod płytką i w górę przez otwór w lewym dolnym rogu, a następnie przylutować do odpowiednich pól lutowniczych na płycie. Zapewnia to jakieś zabezpieczenie, aby zapobiec zerwaniu raczej cienkich przewodów.

Dioda krzemowa 1N4004 włączona szeregowo z baterią zapobiega uszkodzeniu, jeśli spróbujesz podłączyć baterię odwrotnie (zaskakująco łatwe do zrobienia!).

Pojedynczy kondensator 10 μF filtruje zasilanie 9 V. Kondensator jest na liście części wyszczególniony jako tantalowy, ale na zdjęciach widać, że użyto standardowego kondensatora elektrolitycznego 10 μF 16 V. Każdy z nich będzie w porządku – ale drugi kondensator 10 μF (na wejściu 8 IC1c) powinien być tantalowy.

Budowa

Do przylutowania na płytce drukowanej jest tylko dziesięć elementów; i tylko pięć z nich jest spolaryzowanych: oczywiście układ scalony 4093B, wbudowana dioda LED, dioda 1N4004 i dwa kondensatory. W pierwszej kolejności należy zamontować rezystory – oprócz

ew. odczytania kolorowych kodów na częściach, wygodniej jest użyć multimetru, aby potwierdzić ich wartość.

W przypadku kondensatorów tantalowych, znak „+” nadrukowany na ich korpusie idzie do znaku „+” na płycie drukowanej („zwykle” elektrolity aluminiowe mają zaznaczoną końcówkę „-” i to oczywiście ona trafia do pola „-” na płycie drukowanej).

Podobnie, upewnij się, że pasek na diodzie pokrywa się z paskiem oznaczonym na płycie drukowanej. Na koniec zwróć uwagę na wycięcie na końcu układu scalonego: znajduje się ono najbliżej prawej krawędzi płytki.

Anoda wewnętrznej diody LED ma dłuższe z dwóch wyprowadzeń – trafia do pola oznaczonego „A” na płycie drukowanej.

S1, przełącznik „Start”, powinien być przylutowany bezpośrednio do płytki drukowanej.

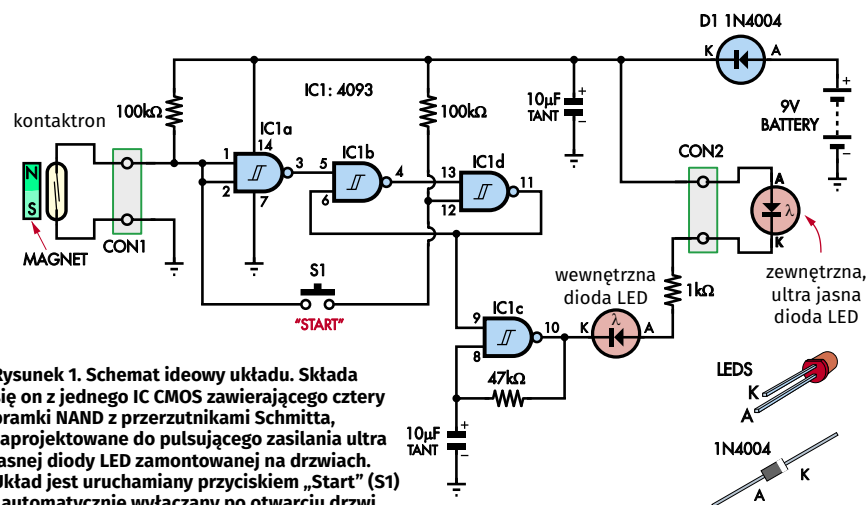
Kontaktron i zewnętrzna dioda LED są podłączone cienkimi izolowanymi przewodami do odpowiednich złączy 402-2/403-2 na płycie drukowanej (kontaktron do CON1; dioda LED do CON2). Zwróć uwagę na polaryzację diody LED – upewnij się, że anoda łączy się z oznaczeniem „A” na CON2. W przypadku prototypu – widać to na **rysunku 2** – użyliśmy miniaturowych dwuzaciskowych złączy śrubowych.

Przed wywierceniem otworów w obudowie i zamontowaniem gotowej płytki drukowanej podłącz baterię 9 V i sprawdź działanie. Przytrzymaj magnes drzwi blisko kontaktronu, a następnie naciśnij S1. Obie diody LED powinny zacząć migać; odsuń magnes od kontaktronu a wtedy powinny przestać migać.

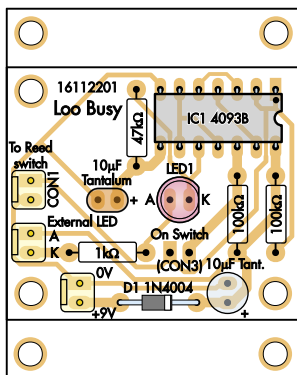
Jeśli tak się nie stanie, sprawdź rozmieszczenie komponentów, orientację i lutowanie. Przy tak małej liczbie części niewiele rzeczy może pójść źle. Jeśli wszystko inne zawiedzie, zmierz napięcie baterii, gdy obwód powinien

Tabela stanów logicznych układu IC1

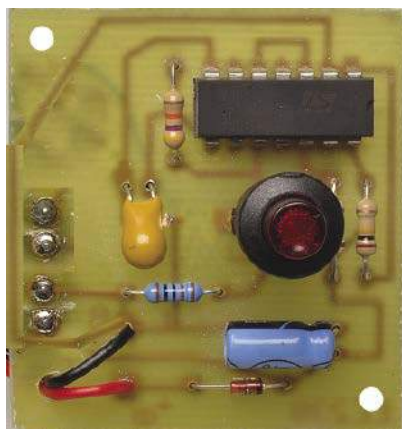
	IC1a			IC1b			IC1d			IC1c		
Drzwi otwarte	1	2	3	5	6	4	13	12	11	9	8	10
	1	1	0	0	0	1	1	1	0	0	1	1
Drzwi zamknięte	IC1a			IC1b			IC1d			IC1c		
	1	2	3	5	6	4	13	12	11	9	8	10
	0	0	1	1	0	1	1	1	0	0	1	1
Drzwi zamknięte „Start”	IC1a			IC1b			IC1d			IC1c		
	1	2	3	5	6	4	13	12	11	9	8	10
	0	0	1	1	1	0	0	0	1	1	0/1	1/0
Drzwi zamknięte „Start” zwolniony	IC1a			IC1b			IC1d			IC1c		
	1	2	3	5	6	4	13	12	11	9	8	10
	0	0	1	1	1	0	0	1	1	1	0/1	1/0



Rysunek 1. Schemat ideowy układu. Składa się on z jednego IC CMOS zawierającego cztery bramki NAND z przerzutnikami Schmitta, zaprojektowane do pulsującego zasilania ultra jasnej diody LED zamontowanej na drzwiach. Układ jest uruchamiany przyciskiem „Start” (S1) i automatycznie wyłączany po otwarciu drzwi



Rysunek 2. Schemat montażowy i fotografia gotowego modułu pomogą umieścić części we właściwych pozycjach. Zwróć uwagę na polaryzację IC1, diody i LED oraz obu kondensatorów. Schemat płytki drukowanej różni się od zdjęcia poniżej stronię tym, że ma «przedłużenia», które umożliwiają jej zatrzasknięcie w pudełku Jiffy. Można je odciąć, jeśli nie są potrzebne



Zdjęcie płytki drukowanej jest odwzorowane w rozmiarze większym niż rzeczywisty. Jest to nowoczesny prototyp i istnieją pewne różnice między schematem montażowym a tą płytką – na przykład S1 i LED1 są umieszczone w tej samej obudowie jako zestaw zespółony (można użyć tego typu lub oddzielnej diody LED i przycisku). Również w tym przypadku złącze baterii jest „na stałe” wlutowane do pól na płycie i wykorzystuje otwór w lewym dolnym rogu do zabezpieczenia przed wyrwaniem

być włączony. Powinno ono wynosić 9 V lub być bardzo bliskie tej wartości.

Montaż płytki drukowanej

Jak widać na Rysunku 2, płytka znajduje się w pudełku Jiffy elementami do spodu – płytka jest zaprojektowana tak, aby zatrzasknąć się w prowadnicach po bokach pudełka. Będziesz musiał wywiercić otwory w dolnej części obudowy (która staje się górną!) dla przełącznika „Start” (i wewnętrznej diody LED).

Jeśli przełącznik startowy jest przyłutowany bezpośrednio do płytki drukowanej, trzeba być dość dokładnym w rozmieszczeniu otworów.

Kolejny otwór jest potrzebny w górnej części obudowy (która staje się dolną!), aby wyprowadzić przewody do kontaktronu i diody LED na drzwiach.

Montaż oku drzwi

Dokładna lokalizacja ostrzegawczej diody LED zależy wyłącznie od Ciebie – niezależnie od tego, co zapewnia najlepszą widoczność.

Może to być montaż przelotowy przez drzwi lub na ościeżnicy. Dostępna jest szeroka gama ramek LED, z których niektóre są zaprojektowane do pracy w otworach przelotowych lub do montażu na ościeżnicy.

Można też po prostu przykleić płaską podstawę bardzo jasnej diody LED od zewnętrznej strony drzwi, z kilkoma drobnymi otworami na przewody.

Kontaktron i jego magnes muszą być umieszczone w taki sposób, aby po zamknięciu drzwi magnes zbliżał się do kontaktronu (nie uderzając w niego!). Jeśli moduł montowany jest na wewnętrznej stronie drzwi, najlepiej jest umieścić kontaktron na drzwiach, a magnes na ościeżnicy drzwi.

Wykaz elementów, kupuj w sklepie avt.pl (W-wa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl):

- 1 płytka drukowana o wymiarach 38,5×49 mm i kodzie 16112201
- 1 obudowa UB5 Jiffy, 83×54×31 mm [np. Jaycar HB6025]
- 1 zestaw przełącznika kontaktronowego (kontaktron i magnes – często sprzedawany do systemów alarmowych – np. Jaycar LA5027)
- 1 mały chwilowy przełącznik przyciskowy (S1) #
- 2 mini złącza do montażu na PCB (np. 402-2/403-2) lub:
- 2 dwupozycyjne złącza śrubowe z rastrem 2,54 mm
- 1 poczwórna bramka NAND CMOS 4093B z przerzutnikami Schmitta (IC1)
- 1 dioda 1N4004 (D1)
- 1 ultra jasna czerwona dioda LED [np. Jaycar ZD0102]
- 1 standardowa czerwona dioda LED #
- Przewody do montażu wewnętrznej i zewnętrznej diody LED
- 1 zatrzask do baterii 9 V
- 1 bateria 6LR61 9 V

Kondensatory:

2 kondensatory tantalowe 10 µF 16 V

Rezystory: (0,25 W, 1%)

2 szt. 100 kΩ 1 szt. 47 kΩ 1 szt. 1 kΩ

Użyliśmy przełącznika przyciskowego ze zintegrowaną diodą LED; na płytce drukowanej przewidziano miejsce na to lub na oddzielny przełącznik i diodę LED.

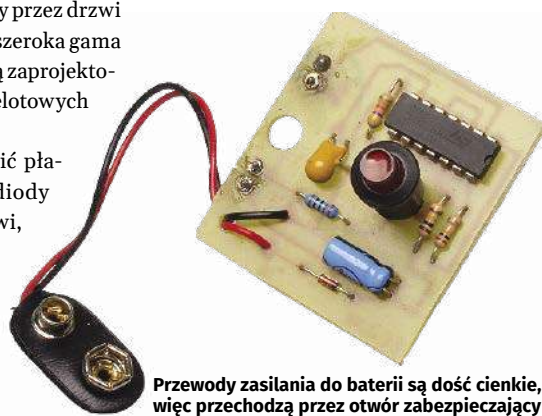
Istnieją poręczne zestawy kontaktronów, które są dostarczane w plastikowych uchwytach z otworami na śruby, przeznaczone do systemów alarmowych (np. Jaycar LA5027). Inne są przeznaczone do całkowicie ukrytego montażu – kontaktron jest wpuszczany w ościeżnicę, a magnes montuje się wewnątrz drzwi – lub odwrotnie (np. Jaycar LA5075).

Używanie

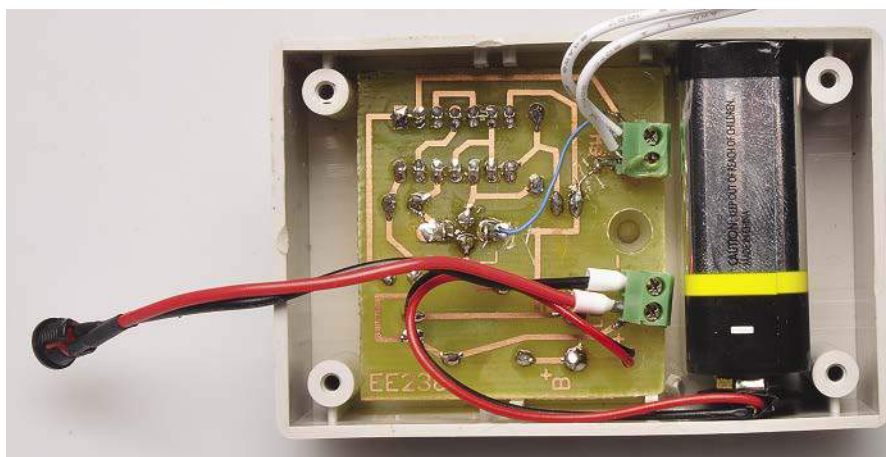
Prostota sama w sobie!

Po wejściu do toalety należy nacisnąć chwilowy (tj. normalnie otwarty) przełącznik „Start” (S1). Spowoduje to miganie obu diod LED (wewnętrznej diody LED po to, aby upewnić się, że nie masz rozładowanej baterii).

Pozostanie tak aż do momentu otwarcia drzwi. Gdy magnes oddali się od kontaktronu



Przewody zasilania do baterii są dość cienkie, więc przechodzą przez otwór zabezpieczający w płytce drukowanej przed przyłutowaniem do odpowiednich pól. Jak wspomniano w tekście, kondensator w prawym dolnym rogu jest na liście części określony jako tantalowy, ale tutaj wystarczą standardowe elektrolity aluminiowe. Drugi kondensator (pomarańczowa pastylka) powinien być tantalowy ze względu na mniejszą upływność



Rysunek 3. Płytkę drukowaną jest montowana w obudowie elementami do spodu, utrzymywana na miejscu przez wycięcia w krawędzi obudowy. Element po lewej stronie (na czerwono-czarnych przewodach) to ultra jasna dioda LED montowana na drzwiach

(S2), przy otwieraniu drzwi, następuje wyłączenie obwodu. Jest on teraz gotowy na następnego użytkownika.

„Automatyczne” wyłączanie kontaktronu zostało uwzględnione ze względu na wysokie prawdopodobieństwo, że ktoś zapomni wyłączyć go ręcznie, co spowoduje kolejkę do drzwi pustej toalety!

Mogliśmy zrobić wszystko w pełni automatycznie (tj. diody LED zaczynają migać zaraz po wejściu), ale uznaliśmy, że dodatkowa komplikacja nie jest warta zachodu. Ale dla eksperymentatorów nie byłoby to trudne.

Od Red. EdW:

- drzwi do toalety, kiedy jest pusta, mogą być zamknięte; nie ma to żadnego wpływu na działanie układu,
- przy modyfikacji układu warto rozważyć możliwość automatycznego uruchamiania zestyku „Start”, nawet przy tak prymitywnym zamknięciu, jak haczyk. Jak zaznaczono wcześniej, zwarcie styku „Start” na stałe nie wpływa na pracę układu, diody LED migają,
- w numerze EdW z sierpnia 2020 r. zaproponowano „Zadanie Główne Szkoły



Dwa typy kontaktronów, oba odpowiednie do opisanego zastosowania. Typ po lewej stronie (Jaycar LA5072) przeznaczony jest do montażu powierzchniowego (stąd otwory montażowe), podczas gdy typ po prawej (Jaycar LA5075) jest w pełni ukryty, montowany w otworach wywierconych w drewnianej ramie drzwi (lub okna). Przetątnik składa się z dwóch części – samego kontaktronu i magnesu przetaczającego. Przetątnik kontaktronowy jest normalnie otwarty i zwiiera się po zbliżeniu magnesu

Konstruktorów nr 294 – Elektroniczna sygnalizacja otwarcia furtki”. Nadesłane rozwiązanie opublikowano w numerze styczniowym EdW z 2021 r. Mogą one stanowić inspirację do własnej modyfikacji rozwiązania z SC. ■

John Chappell

Adaptacja do wydania polskiego – Andrzej Nowicki

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

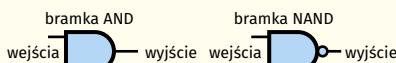
Czym jest bramka NAND?

Wewnątrz układu 4093B znajdują się cztery identyczne bramki, z których każda działa niezależnie (ale ze wspólnym zasilaniem). Stąd nazwa „poczwórny”.

Najpierw przyjrzymy się bramce AND. Pomyśl o bramce jak o furtce w ogrodzeniu. Może być otwarta lub zamknięta. W przypadku dwóch furtek OBA wejścia muszą być w stanie zamkniętym, aby zablokować wejście, lub przynajmniej jedno (albo oba) musi być otwarte, aby wejście było możliwe. W przypadku bramki AND, jeśli oba wejścia są w stanie wysokim, wyjście będzie w stanie wysokim. Jeśli jedno z nich (lub oba) jest w stanie niskim, wyjście będzie w stanie niskim. Dlatego ten układ nazywa się to bramką AND. Realizuje ona iloczyn logiczny (operację mnożenia).

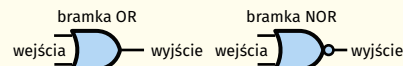
Ale układ 4093 ma dodatkowy obwód w każdej bramce, który „odwraca” stan wyjścia. Zatem, gdy na obu wejściach pojawia się stan „wysoki”, na wyjściu pojawia się stan „niski” (i odwrotnie). To sprawia, że jest to bramka NAND, skrót od NOT AND. Małe kółko na wyjściu bramki NAND wskazuje, że jest to bramka z inwerterem (bramka AND nie ma takiego kółka). Bramka NAND realizuje iloczyn logiczny (operację mnożenia), z inwersją wyniku. Najlepiej wyjaśnia to tabela:

Wejście 1	Wejście 2	Wyjście
0	0	1
0	1	1
1	0	1
1	1	0



Zanim opuścimy bramkę AND/NAND, często zobaczysz inny typ prostej bramki, OR/NOR. W przypadku tej bramki, jak sama nazwa wskazuje, wystarczy aby jedno LUB drugie wejście było w stanie wysokim, aby otrzymać na wyjściu stan wysoki. Bramka OR realizuje sumę logiczną (operację alternatywy).

Ale jeśli jest to bramka NOR, w odróżnieniu od bramki OR, stan wyjścia będzie odwrócony (podobnie jak różnica między bramkami NAND i AND). Bramka NOR realizuje sumę logiczną (operację alternatywy) z inwersją wyniku.

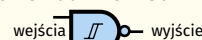


Wreszcie, skąd się wzięło określenie „Przerzutnik Schmitta” („Schmitt Trigger”)?

W większości bramek przejście między napięciami dla stanu wysokiego i niskiego jest dość szerokie – napięcie musi być poniżej pewnego napięcia progowego, aby było określone jako niskie (blisko 0 V) i powyżej innego napięcia progowego, aby było określone jako wysokie (znacznie bliżej napięcia zasilania). Napięcia progowe między napięciami progowymi stanu niskiego i wysokiego nie odpowiadają zdefiniowanym stanom logicznym.

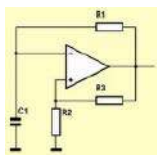
Jest to jednak często niepożądane, więc wewnątrz bramki znajdują się obwody, które sprawiają, że przejście ze stanu niskiego do wysokiego lub wysokiego do niskiego jest znacznie dokładniej określone dzięki histerezie. Nazywa się to przerzutnikiem Schmitta i oznacza na symbolu bramki małym znakiem histerazy.

bramka NAND z przerzutnikiem Schmitta



Praktyczny kurs op-ampów

24. Generator fali prostokątnej

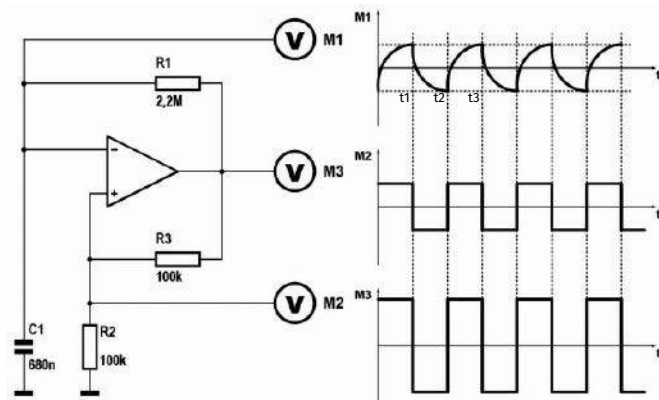


Z jednym op-ampem, trzema rezystorami i kondensatorem możesz zbudować układ, który generuje piękne przebiegi prostokątne. Dzięki niewielkiej rozbudowie możesz przekształcić ten układ w generator impulsów z regulacją ich szerokości i częstotliwości.

Schemat generatora

To może być prostsze!

Z generatora funkcyjnego omówionego w odcinku tego kursu „Generator funkcyjny” można uzyskać przebieg prostokątny. Jeśli nie potrzebujesz przebiegu trójkątnego, można zbudować prostsze układy do generowania samej fali prostokątnej. W końcu w większości przypadków potrzebujesz takich sygnałów tylko do sterowania układami cyfrowymi i wtedy nie korzystasz z przebiegu trójkątnego. Dlatego w tym artykule omawiamy najprostszą metodę generowania sygnału prostokątnego za pomocą op-ampa. Schemat jest pokazany na poniższym rysunku.



Podstawowy schemat op-ampa jako generatora przebiegu prostokątnego (© 2018 Jos Verstraten)

Schemat

W układzie widać podwójne sprzężenie zwrotne: rezystancyjne między wyjściem a wejściem dodatnim, obwód RC między wyjściem a wejściem ujemnym. Działanie układu ilustrują wykresy po prawej stronie.

Przyjmij, że po włączeniu napięcia zasilania napięcie wyjściowe op-ampa jest równe dodatniemu napięciu zasilania. Za pomocą dzielnika napięcia R2-R3 połowa tego napięcia wyjściowego dotrze do wejścia dodatniego. Przyjmijmy, że kondensator C1 został całkowicie rozładowany. Napięcie na wejściu ujemnym jest bardziej ujemne niż na wejściu dodatnim, napięcie wyjściowe jest równe dodatniemu napięciu zasilania.

Kondensator zaczyna się ładować poprzez rezystor R1. Rośnie więc napięcie na wejściu ujemnym i po pewnym czasie staje się ono większe od napięcia na wejściu dodatnim. Stan na wyjściu Op-ampa zmienia się, napięcie wyjściowe osiąga wartość ujemnego napięcia zasilania. Wejście dodatnie jest ustawione na połowę tego napięcia. Kondensator zaczyna się rozładowywać, ponieważ R1 jest

podłączony do bardzo ujemnego napięcia. Po pewnym czasie napięcie na wejściu ujemnym staje się mniejsze od napięcia na wejściu dodatnim, stan na wyjściu Op-ampa zmienia się ponownie.

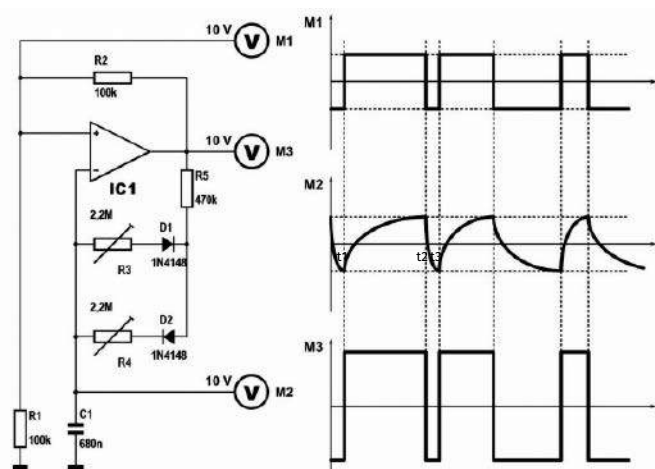
Zmieniając wartości C1 i R1, możesz zmieniać częstotliwość. W punkcie pomiarowym M3 możesz to poznać po szybkości, z jaką wskazówka miernika przesuwają się z jednego końca skali do drugiego. Na wyjściu układu pojawia się przebieg prostokątny, którego częstotliwość zależy od szybkości ładowania i rozładowywania kondensatora C1 przez rezystor R1. Wystarczy więc zmieniać wartość któregoś z tych elementów, aby sterować częstotliwością. Zazwyczaj przełącza się kondensator, aby wybrać zakres częstotliwości i używa potencjometru w miejscu R1 do sterowania częstotliwością w wybranym zakresie.

Symetryczny czy asymetryczny?

Rozbudowa obwodu o regulację szerokości impulsu

Układ pokazany na podstawowym schemacie generuje w przybliżeniu symetryczną falę prostokątną. Przedział czasowy t1-t2 jest w przybliżeniu równy przedziałowi czasowemu t2-t3. Przy niewielkich rozszerzeniach można przekształcić ten układ w generator impulsów, który może generować wąskie impulsy dodatnie lub ujemne. Uniwersalny schemat generatora impulsów z op-ampem został pokazany na schemacie poniżej. Za pomocą dwóch diod ładowanie i rozładowywanie kondensatora jest kontrolowane oddzielnie przez dwa rezystory R3 i R4. Jeśli napięcie wyjściowe op-ampa jest dodatnie, to D1 nie przewodzi, a D2 pracuje w kierunku przewodzenia. Prąd ładowania popłynie przez D2, jego wielkość jest określona przez wartość rezystora R4. Jeśli wyjście jest ujemne, D2 nie przewodzi, a D1 pracuje w kierunku przewodzenia. Kondensator C1 jest wtedy rozładowywany przez rezystor R3.

Na wykresach po prawej stronie, napięcia w różnych punktach obwodu są wykreślone dla dwóch różnych ustawień potencjometrów. W jednym przypadku układ generuje wąskie impulsy ujemne, w drugim nieco węższe impulsy dodatnie.

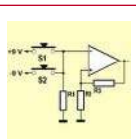


Rozbudowa obwodu o regulację szerokości impulsu (© 2018 Jos Verstraten)

Ten eksperyment pokazuje, że można włączyć op-ampa do systemu cyfrowego zbudowanego z układów CMOS. Często w systemie cyfrowym potrzebny jest oscylator zegarowy, a wszystkie bramki dostępne w stosowanych cyfrowych układach scalonych są zajęte przez inne funkcje. Wtedy taniej jest zastosować tani op-amp, niż wprowadzać nowy cyfrowy układ scalony, którego tylko połowa lub jedna czwarta jest wykorzystywana.

Jedno ostrzeżenie: zwykle tanie op-ampy, takie jak 741 i 3140, są w zasadzie układami małej częstotliwości. Nie jest możliwe za pomocą tych układów generowanie impulsów o czasach narastania rzędu kilkudziesięciu nanosekund, jak to jest możliwe w przypadku układów cyfrowych. Ta właściwość op-ampów ogranicza przydatność tych układów w systemach cyfrowych. ■

25. Przerzutnik (ang. flip-flop)



Op-amp jest typowym układem analogowym. Jednak w tym i kolejnych odcinkach tego kursu pokażemy, że można również użyć op-ampy do budowy podstawowych układów elektroniki cyfrowej. Zaczniemy od przerzutnika.

Zasada działania układu

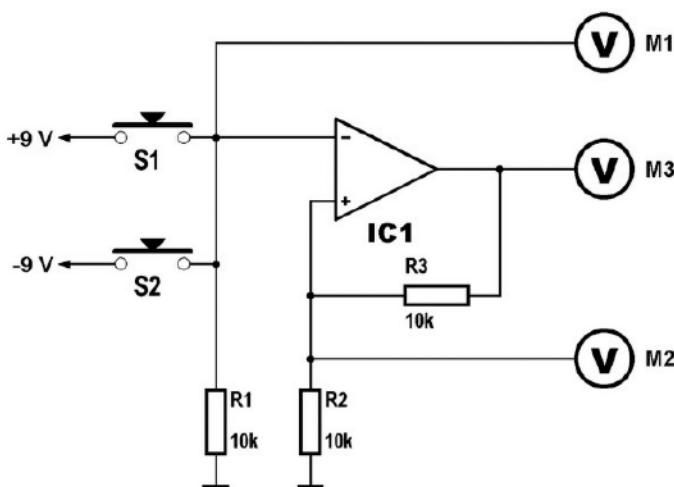
Wstęp

Przerzutniki są podstawowymi układami elektroniki cyfrowej. Za pomocą przerzutnika można „zapamiętać” obecność chwilowego impulsu napięcia. Przerzutnik jest więc elektroniczną pamięcią zdolną do przekształcania pojawienia się chwilowego impulsu w „wieczny” stan na wyjściu. Istnieją setki cyfrowych układów scalonych zawierających różne wersje przerzutników. Używanie op-ampa do tej funkcji nie jest więc takie oczywiste.

Niemniej jednak, czasami może być bardziej ekonomicznie zbudować przerzutnik z op-ampa, zamiast stosować jeden z wielu cyfrowych układów scalonych. Szczególnie teraz, gdy na rynku jest wiele układów scalonych zawierających cztery identyczne op-ampy, często pojawia się praktyczna sytuacja, w której zostaje nam op-amp i jednocześnie potrzebujemy przerzutnika. Wtedy oczywiście bardziej opłaca się zastosować niewykorzystany op-amp.

Schemat podstawowy

Podstawowa realizacja op-ampa jako przerzutnika została pokazana na schemacie poniżej. Wejście dodatnie op-ampa jest połączone z wyjściem za pomocą dzielnika rezystorowego. Jest to ta sama konfiguracja, jak w przypadku generatora przebiegu prostokątnego! Nawiasem mówiąc, w tym samym celu: aby ustalić napięcie na wyjściu na poziomie jednego z napięć zasilania, dopóki coś nie stanie się na wejściu



Podstawowy schemat przerzutnika z op-ampem (© 2018 Jos Verstraten)

ujemnym. Tym „czymś” w naszym przypadku jest pojawienie się wąskiego impulsu dodatniego lub ujemnego.

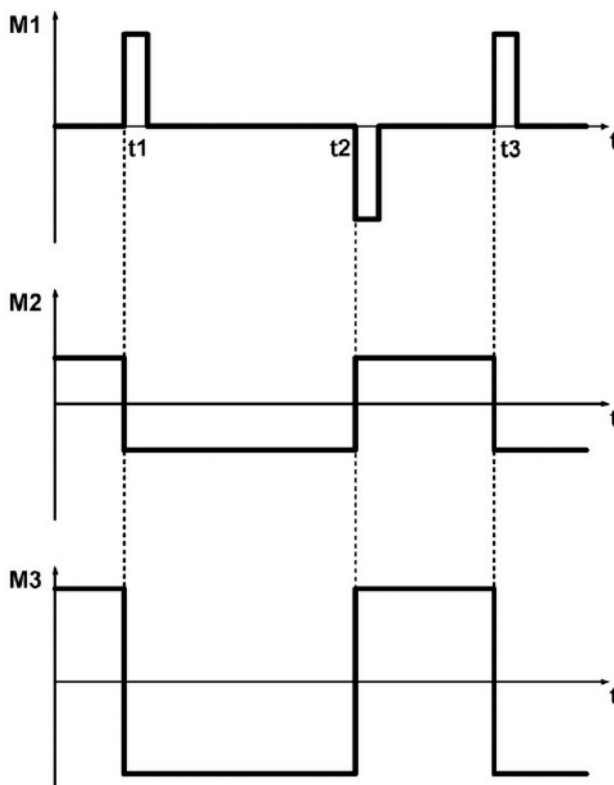
Możesz to zasymulować używając ponownie dwóch baterii 9 V połączonych szeregowo, które w tym kursie były już kilkakrotnie używane do podawania dodatnich lub ujemnych impulsów. Podłącz plus jednej baterii do minusa drugiej baterii. Połącz ten punkt z masą twojej płytki eksperymentalnej. Przylutuj dwa pozostałe bieguny baterii do dwóch przycisków. Drugie końcówki przycisków idą razem do wejścia odwracającego op-ampa. To wejście jest również połączone z masą przez rezystor R1. Jeśli krótko naciśniesz przycisk S1, podłączony do dodatniego bieguna baterii, pojawi się impuls +9 V. Gdy naciśniesz drugi przycisk, pojawi się impuls -9 V.

Zwróć uwagę!

Nigdy, przenigdy nie wolno naciskać obu przycisków jednocześnie! Spowodujesz wtedy zwarcie obu baterii, co spowoduje przepływ dość dużego prądu zwarcowego i nagrzanie się baterii.

Działanie układu

Przyjmij, że napięcie wyjściowe op-ampa jest równe dodatniemu napięciu zasilania, gdy napięcie zasilania jest włączone. Sprzężenie



Przebiegi napięć w układzie (© 2018 Jos Verstraten)

zwrotne powoduje, że wejście dodatnie staje się dodatnie w stosunku do wejścia ujemnego. To sprzężenie zwrotne stabilizuje więc sytuację włączenia. Teraz, naciskając S1, przyłoż do wejścia ujemnego impuls dodatni. Napięcie +9 V na tym wejściu jest bardziej dodatnie niż +5 V na wejściu dodatnim, op-amp zmienia stan na wyjściu. Wyjście przechodzi do -10 V, a sprzężenie zwrotne powoduje, że na wejściu dodatnim jest napięcie -5 V. Wejście ujemne jest nadal bardziej dodatnie niż wejście dodatnie, ta sytuacja jest stabilna.

Teraz zwolnij przycisk. Ujemne wejście op-ampa ma potencjał masy. Ale nawet teraz to wejście jest nadal bardziej dodatnie niż napięcie -5 V na wejściu nieodwracającym. Tak więc na wyjściu pozostaje napięcie bliskie ujemnemu napięciu zasilania!

Wniosek

Podanie krótkiego impulsu dodatniego na wejście ujemne powoduje przerzucenie napięcia wyjściowego z +10 V na -10 V. Układ działa jak element pamięci, przerzutnik. Oczywiście można przywrócić układ do stanu początkowego poprzez przyłożenie do wejścia ujemnego impulsu ujemnego.

Chcąc określić ten układ w terminologii układów cyfrowych, można powiedzieć, że op-amp jako przerzutnik odpowiada najprostszej

wersji przerzutnika: typu S-R, gdzie funkcje set i reset działają na jedno i to samo wejście, ale żądają przeciwnych polaryzacji. W porównaniu z szeroką gamą dostępnych wariantów przerzutników cyfrowych (R-S, D, J-K, itd.), op-amp ma więc tylko ograniczone zastosowania.

Praca z tabelą prawdy

W układach cyfrowych, takich jak ten przerzutnik, można graficznie pokazać związek pomiędzy wielkościami wejściowymi i wyjściowymi w postaci tzw. tabeli prawdy. W takiej tabeli wpisujesz działanie na wejściu (wejściach) w pierwszej kolumnie (kolumnach) i odnotowujesz odpowiedź wyjścia w ostatniej kolumnie. Nasz przerzutnik jest doskonałym narzędziem do zapoznania się z pojęciem tabeli prawdy. Na rysunku poniżej znajduje się tabela prawdy tego typu przerzutnika. ■

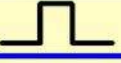
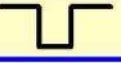
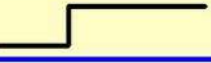
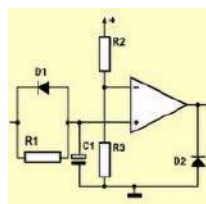
Przycisk S1	Przycisk S2	Reakcja na wyjściu
		
		

Tabela prawdy op-ampa jako przerzutnika R-S (© 2018 Jos Verstraten)

26. Opóźnienie

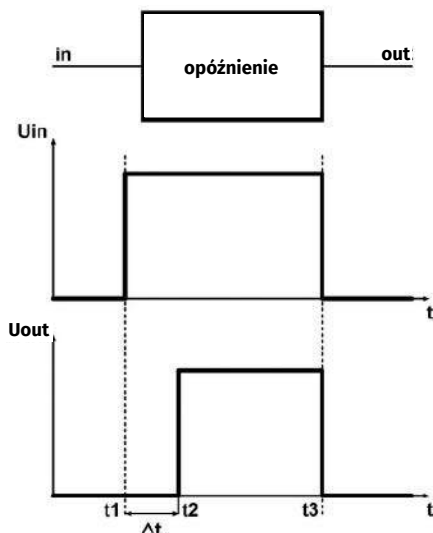


Opóźnienia są często stosowane w elektronice cyfrowej. Alarm przeciw włamaniom ma przycisk aktywacji. Jeśli go naciśniesz, masz jeszcze dwadzieścia sekund na opuszczenie budynku, zanim włączy się alarm. To opóźnienie jest zapewnione przez układ opóźniający, który możesz zaprojektować na wzmacniaczu operacyjnym.

Zasada działania opóźnienia

Najprostsze opóźnienie

Istnieje kilka rodzajów układów opóźniających. Najprostszy z nich ilustruje poniższy rysunek. W tym opóźnieniu tylko zbocze narastające



Działanie opóźnienia: zbocze narastające impulsu jest opóźnione o czas Δt (© 2018 Jos Verstraten)

impulsu jest opóźnione o czas Δt . Symbol Δ , grecka litera delta, jest powszechnie stosowany w inżynierii do oznaczenia małego odstępu czasu. W stanie spoczynku wejście i wyjście znajdują się na potencjale masy. Dodatni impuls na wejściu w chwili t_1 pojawia się na wyjściu z opóźnieniem w chwili t_2 . Jeśli jednak w chwili t_3 impuls wejściowy zaniknie, to impuls wyjściowy również ustanie.

Istnieją układy opóźniające, które pozwalają opóźnić nie tylko zbocze narastające, ale także zbocze opadające impulsu (szerokość impulsu pozostaje wtedy stała), a nawet istnieją układy pozwalające na przesunięcie w czasie krótkiego impulsu. Innymi słowy, impuls wejściowy znika zanim opóźnienie wygeneruje impuls wyjściowy.

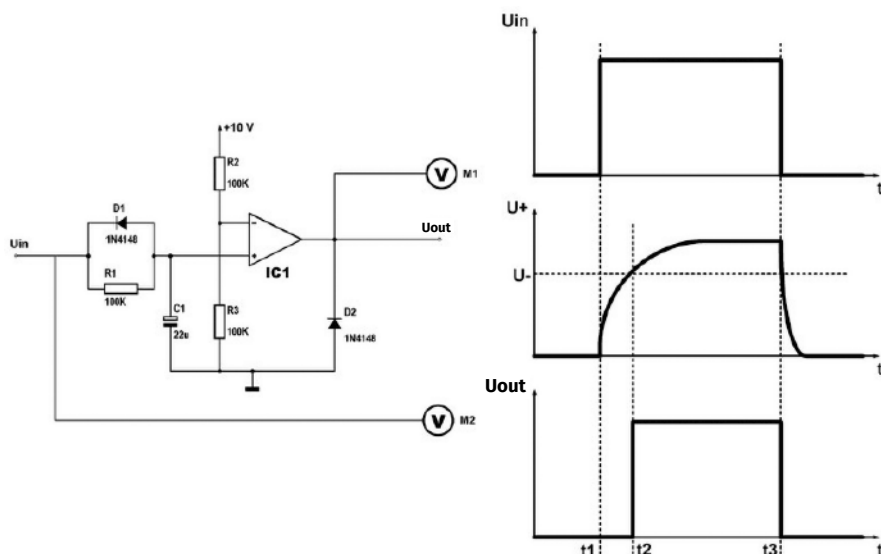
Podążając za klasyczną metodą cyfrową, tego typu układy buduje się przy użyciu scalonych multiwibratorów monostabilnych, takich jak 74121 dla TTL lub 4047 dla CMOS. Jeśli nie masz zbyt dużych wymagań co do szybkości przełączania (tj. pracujesz w środowisku o niskiej częstotliwości), możesz zbudować takie opóźnienia równie łatwo za pomocą op-ampów.

Opóźnienie z op-ampem

Schemat prezentuje przykład obwodu, który opóźnia zbocze narastające dodatniego impulsu o regulowany czas. Wykresy po prawej stronie ilustrują szczegóły działania. Op-amp jest w rzeczywistości przełączany jako komparator. Napięcie progowe jest ustawione na połowę dodatniego napięcia zasilania za pomocą dzielnika napięcia R2-R3.

Jeśli napięcie wejściowe jest zerowe, to ujemne wejście układu scalonego ma potencjał bardziej dodatni niż wejście nieodwracające. Wyjście w tej sytuacji dąży do ujemnego napięcia zasilania. Jednak dioda D2 zapewnia, że wyjście op-ampa nie może zejść poniżej poziomu -0,6 V. W elektronice cyfrowej pracuje się ze standardowymi poziomami, a tutaj zwyczajowo przyjmuje się, że „0” (L) odpowiada wartości zero woltów.

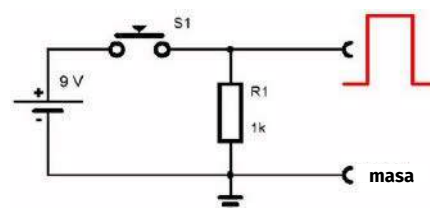
Uwaga! Nie każdy typ op-ampa pozwala na zwarcie jego wyjścia do masy za pomocą diody. W przypadku 741 jest to możliwe, można zwierzać bez ograniczeń.



Opóźnienie analogowe złożone z integratora RC i komparatora (© 2018 Jos Verstraten)

Działanie układu

Założmy, że na wejście w chwili t_1 podajemy impuls dodatni. Przez R_1 płynie prąd, który zaczyna ładować kondensator C_1 . Napięcie na wejściu dodatnim wzrasta. W chwili t_2 napięcie na tym wejściu staje się równe napięciu progowemu na wejściu odwracającym. Wyjście op-ampa zmienia stan, staje się dodatnie. Opóźnienie czasowe Δt



Generowanie dodatniego impulsu (© 2018 Jos Verstraten)

jest określone przez parametry obwodu RC włączonego pomiędzy wejściem układu a nie-odwracającym wejściem op-ampa.

W chwili t_3 impuls wejściowy zanika. Naładowany kondensator rozładowuje się teraz przez przewodzącą diodę D_1 . Napięcie na $+IN$ staje się mniejsze od progu na $-IN$, komparator zmienia stan – wyjście przechodzi do zera.

Generowanie impulsu wejściowego

Do eksperymentów z układem musimy wygenerować krótki, dodatni impuls. Oczywiście już wiesz jak to zrobić – potrzebujesz baterii 9 V, przycisku i rezystora 1 k Ω . Elementy łączysz jak na rysunku powyżej. Po naciśnięciu przycisku, generowany jest ładny impuls o napięciu 10 V, a po zwolnieniu przycisku, poprzez rezystor ustalany jest stan 0 V na wyjściu. ■

Jos Verstraten



Kurs op-ampów

Rozwiązanie znajdziesz na www.elportal.pl/quizy

Kiedy op-amp działa jako generator?

- ☐ Gdy ma dodatnie sprzężenie zwrotne i obwód RC w jako ujemne sprzężenie zwrotne
- ☐ Gdy ma ujemne sprzężenie zwrotne i obwód RC w jako dodatnie sprzężenie zwrotne
- ☐ Gdy ma dodatnie sprzężenie zwrotne i obwód RC na wyjściu

Rozbudowa generatora z op-ampem o regulację szerokości impulsu wymaga:

- ☐ Dodania dwóch potencjometrów w obwodzie ujemnego sprzężenia zwrotnego
- ☐ Dodania dwóch potencjometrów i dwóch diod w obwodzie ujemnego sprzężenia zwrotnego
- ☐ Dodania dwóch diod w obwodzie ujemnego sprzężenia zwrotnego

Jak działa przerzutnik?

- ☐ Generuje przebieg prostokątny o regulowanym wypełnieniu
- ☐ Zamienia chwilowy impuls na wejściu w wieczny stan na wyjściu
- ☐ Zamienia chwilowy impuls na wejściu w impuls o określonym czasie na wyjściu

Op-amp jako przerzutnik odpowiada przerzutnikowi:

- ☐ S-R
- ☐ D
- ☐ J-K

Czym jest tabela prawdy?

- ☐ Pokazuje przy jakim napięciu układ cyfrowy działa prawidłowo
- ☐ Pokazuje wartości elementów RC i odpowiadające im częstotliwości generatora cyfrowego
- ☐ Pokazuje związek pomiędzy wielkościami wejściowymi i wyjściowymi w układach cyfrowych



Symbol Δ , grecka litera delta, połączony z literą t , czyli Δt , oznacza:

- ☐ Jest stosowany w inżynierii do oznaczenia małego odstępu czasu
- ☐ Jest stosowany w inżynierii do oznaczenia małej zmiany temperatury
- ☐ Jest stosowany w inżynierii do oznaczenia punktu pomiaru temperatury

Co to znaczy, że zasilanie jest symetryczne:

- ☐ Dostarcza dwóch napięć – dodatniego i ujemnego względem wspólnej masy
- ☐ Dostarcza dwóch napięć o takim samym prądzie maksymalnym
- ☐ Dostarcza dwóch napięć z możliwością regulacji

Co oznacza określenie flip-flop

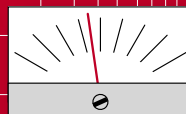
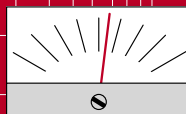
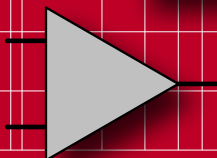
- ☐ Jest to generator
- ☐ Jest to układ opóźniający z op-ampem
- ☐ Jest to przerzutnik

Czy op-ampy:

- ☐ Przeważnie są tak szybkie, jak układy cyfrowe
- ☐ Przeważnie są układami o małej częstotliwości maksymalnej
- ☐ Przeważnie są szybsze, niż układy cyfrowe

Układ typu uA741:

- ☐ To wzmacniacz operacyjny ogólnego przeznaczenia
- ☐ To super szybki wzmacniacz operacyjny
- ☐ To dwa wzmacniacze operacyjne w jednej obudowie

AUDIO
OUT

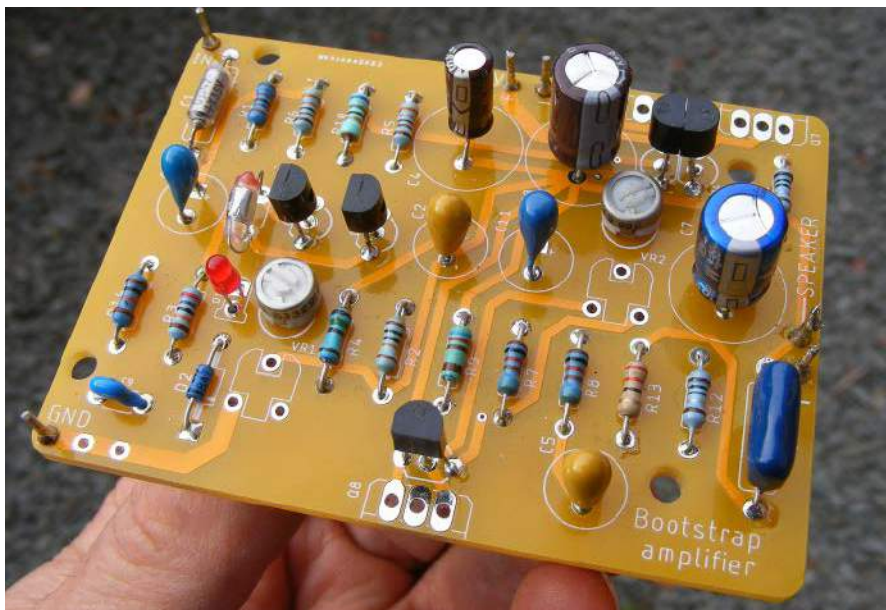
Wzmacniacz audio do Theremina, część 1

Opisany projekt był już prezentowany w Practical Electronics w sierpniu 2019 roku, jako część Theremina. Wówczas był to tylko schemat obwodu (rysunek 34 na stronie 54), który został zmontowany na płytce prototypowej. Teraz powstała dla niego płytka drukowana, dzięki czemu stał się samodzielnym modułem, który sprawdzi się w wielu zastosowaniach.

Zaprezentowany wzmacniacz charakteryzuje się niewielkim prądem spoczynkowym, wynoszącym ok. 3 mA. Może być zasilany tanią cynkowo-węglową baterią 9 V typu 6F22 (lub alkaliczną 6LR61), jeżeli użyjemy głośnika o wysokiej impedancji, np. 25 Ω (rysunek 1). Maksymalna moc wyjściowa wynosi 275 mW, a pobór prądu dochodzi do 50 mA przy pełnym wysterowaniu. Jest idealny do przenośnego Theremina, radia i różnych eksperymentów, dlatego zostaną omówione wszystkie szczegóły jego budowy i działania.

Projekt na elementach dyskretnych

Większość komercyjnych projektów Theremina, dla uproszczenia produkcji, zawierało scalony wzmacniacz mocy typu LM386, TB820 lub LM384 (rysunek 2). Mają one kilka wad, a najistotniejszą jest wysoki prąd spoczynkowy, co powoduje krótki czas działania przy zasilaniu z baterii. Mają one również wysokie wzmocnienie napięciowe wynikające z rozbudowanych obwodów wejściowych, co przy ustawionej minimalnej głośności skutkuje znacznymi szumami i stanowi problem np. przy nagrywaniu. Zaletą rozwiązania dyskretnego (czyli wykonanego z użyciem pojedynczych tranzystorów) jest to, że każdy parametr może być zoptymalizowany. Takie podejście zwykle nie jest rekomendowane w projektach komercyjnych, ale jest idealne do nauki i eksperymentów.



Wzmacniacz z układem bootstrap – idealny partner dla projektu theremina opublikowanego w PE

Te nieznosne elementy

Projekty układów scalonych mają tę przewagę nad rozwiązaniami dyskretnymi, że wszystkie użyte w nich tranzystory są wykonane na tym samym kawałku krzemu i w jednym procesie produkcyjnym, co przekłada się na ich właściwe dopasowanie. Projektanci mogą również zastosować tyle tranzystorów, ile potrzebują np. w źródłach i lustrach prądowych oraz obwodach polaryzacji. Ilość parametrów wymagających regulacji jest

ograniczona lub nawet jest brak konieczności regulacji.

W większości obwodów dyskretnych konieczne strojenia są wykonywane za pomocą rezystorów nastawnych – potencjometrów i helitrimów (ang. *Trimpots*). Elementy te są słabym punktem konstrukcji. Nie ma wątpliwości, że otwarte potencjometry, takie jak pokazane na rysunku 3) stają się niepewne, gdy do ich wnętrza dostanie się brud lub wilgoć i zacznie się proces utleniania. Należy unikać takich



Rysunek 1. Dla wydłużenia żywotności baterii 9 V zastosowano głośnik o wysokiej impedancji (25 Ω)

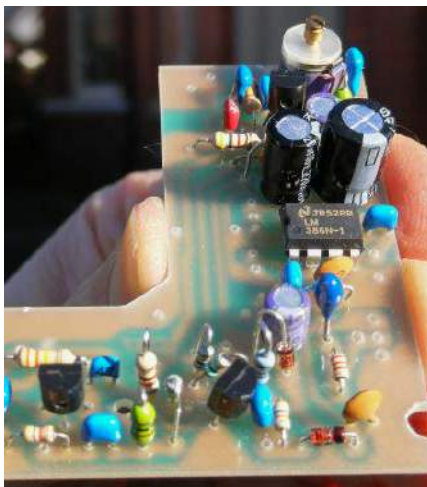
rozwiązań i projektować swoje urządzenia lepiej. Mój wykładowca, w 1983, powiedział mi: „używaj analizy układowej” a nie tylko wstawiaj potencjometry dostrojcze. Niestety, zawsze uwielbiałem „kręcić pokrętkami” i dostrajać te idealne punkty minimalnych zniekształceń lub symetrii obcinania sygnału. Dla mnie technologia elektroniki muzycznej zawsze była sprzężeniem zwrotnym ludzkich zmysłów kontrolujących elektroniczne zmienne.

Sporym problemem jest wzrost cen komponentów (który w pewnym stopniu omija tylko elementy i układy SMD) szczególnie tam, gdzie był spadek popytu na części analogowe, np. tranzystory JFET lub przełączniki mechaniczne. Trymery nie są wyjątkiem, standardowy potencjometr hermetyczny Bournsa 3329H TO5, pokazany na **rysunku 4**, zdrożał kilkukrotnie w ciągu ostatnich lat. Zauważ, że na płytce są dwa różne fizyczne kontury dla różnych typów potencjometrów (montujemy tylko jeden w danym momencie). W każdym razie dość narzekania, opisany układ jest bazą do zdobycia doświadczenia w strojeniu, które jest tego warte.

Bootstrap – stara, ale przydatna sztuczka

W czasach kiedy tranzystory były drogie, a ich parametry niezbyt imponujące układ pompy ładunkowej był szeroko stosowany, dla uzyskania dodatkowego wzmocnienia napięciowego w otwartej pętli i większego zakresu napięć wyjściowych. Proponowany wzmacniacz używa tego rozwiązania w dwóch miejscach. Płytkę drukowaną jest opisana jako „Wzmacniacz bootstrap” (ang. *Bootstrap amplifier*), ponieważ taka była jego robocza nazwa w trakcie projektowania.

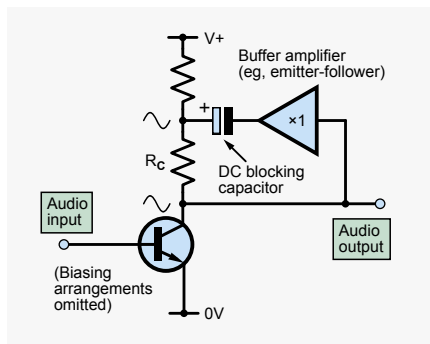
Nazwa bootstrap pochodzi od wyrażenia „podnosić kogoś za paski od butów”. W elektronice oznacza to podnoszenie napięcia jednego końca elementu (zwykle rezystora), podczas gdy drugi koniec jestysterowany w tym samym kierunku. Efektem tego jest zmniejszenie prądu płynącego przez rezystor, dzięki czemu jego pozorną rezystancja jest znacznie wyższa. Górny



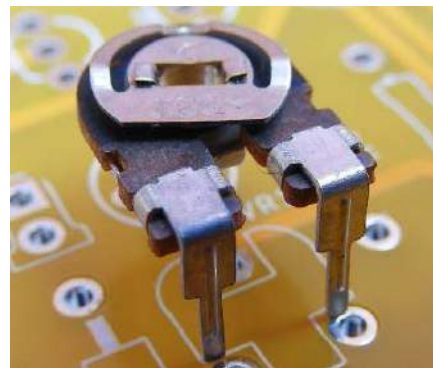
Rysunek 2. Scalone wzmacniacze małej mocy, jak ten 8-pinowy LM386, już dawno zastąpiły rozwiązania dyskretnie w elektronice komercyjnej. Jednak każdy powinien uruchomić mały układ na elementach dyskretnych zanim przejdzie na drogę konstrukcji Hi-Fi. Pokazany klasyczny LM386 firmy National Semiconductor jest używany w moim Thereminie Eclipse (<http://theremin.co.uk>)

koniec musi być sterowany przez bufor, np. wtórnik emiterowy, ponieważ musi tam być dodatkowe źródło energii. Pojedynczy tranzystor nie może „podnieść” sam siebie. Ta technika zastosowana do wzmacniacza w układzie wspólnego emitera, jest pokazana na **rysunku 5**. Może być także zastosowana do wtórnika emiterowego, jak pokazano na **rysunku 6**. Efekt zwiększenia impedancji, można również uzyskać, zastępując rezystor źródłem prądowym, rozwiązanie to jest stosowane w większości wzmacniaczy scalonych.

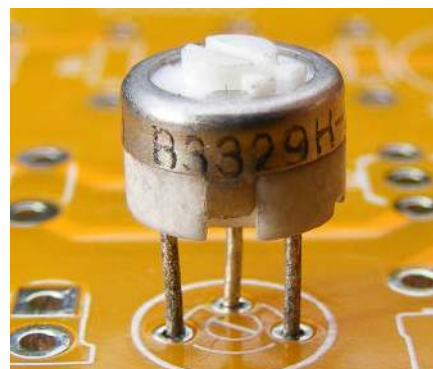
Pompa ładunkowa daje dodatkową zaletę – zwiększa efektywne napięcie zasilania, zwiększając zakres napięć wyjściowych, co jest bardzo przydatne przy baterijnym zasilaniu obwodów. Odwrotnie, źródła prądowe zmniejszają dostępny zakres zmienności Uwy o około 1,5 V, które jest potrzebne do ich poprawnej pracy (zobacz **rysunek 7**). Wzmacniacz z pompą ładunkową daje kilka



Rysunek 5. Bootstrap zastosowany we wzmacniaczu ze wspólnym emiterem znacznie zwiększa wzmocnienie napięciowe przez to, że R_c wydaje się mieć dużo większą wartość niż ma w rzeczywistości



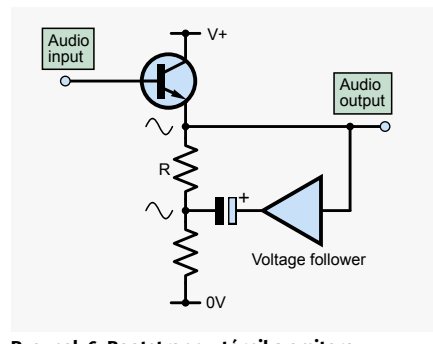
Rysunek 3. Powinno się unikać elementów nastawczych o otwartej konstrukcji, jeżeli wymagana jest niezawodność długoterminowa



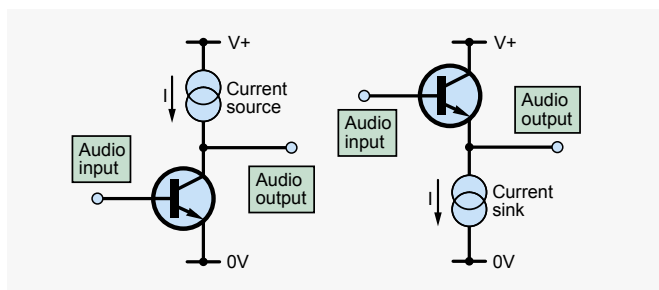
Rysunek 4. Uszczelnione elementy nastawcze takie jak pokazano na fotografii, są całkiem niezawodne, ale mogą kosztować więcej niż układ scalony wzmacniacza mocy

woltów więcej na wyjściu, niż rozwiązanie ze źródłem prądowym.

Wadą układu bootstrap jest to, że wymaga on sprzężenia zmiennoprądowego przez kondensator o dużej wartości, od 1 do 220 μF , aby uzyskać efekt podbicia napięcia. Te kondensatory działają podobnie jak te w układzie podwajacza napięcia, gdzie napięcie z kondensatora jest dodawane do napięcia zasilania. Niestety trzeba tu użyć kondensatorów elektrolitycznych, które są i owszem tanie, ale mają duże wymiary i z czasem wysychają. Jeżeli wybierzemy kondensatory SAL czy



Rysunek 6. Bootstrapp wtórnika emiterowego zwiększa wydajność wyjściową poprzez wyeliminowanie efektu obciążenia R. Wiele bootstrapów może także zwiększyć zakres zmienności napięcia wyjściowego w małych wzmacniaczach mocy o kilka woltów



Rysunek 7. Źródła prądowe zapewniają taki sam efekt wzmacnienia napięciowego i prądu jak pompa ładunkowa, utrzymując działanie bez ograniczeń dolnej częstotliwości pracy wzmacniacza oraz unikają potrzeby stosowania kondensatorów elektrolitycznych. Występuje jednak na nich spadek napięcia od 1 do 2 V

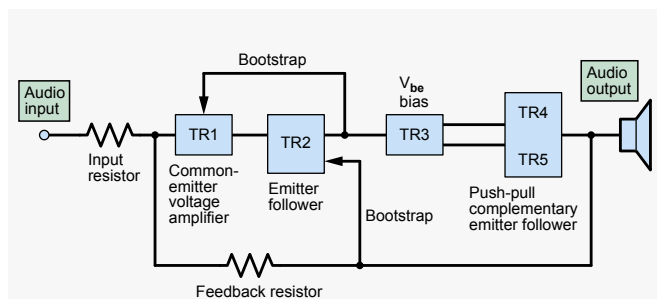
tantalowe, to trzeba uwzględnić że ich cena jest ok. pięciokrotnie większa (a do tego kondensatory SAL, ze stałym elektrolitem, nie są już produkowane).

Bootstrapowanie nie jest też do przyjęcia w nowoczesnych wzmacniaczach Hi-Fi, bo efekt pompy ładunkowej nie działa przy bardzo małych częstotliwościach, dając wzrost zniekształceń w tym zakresie. Innym problemem jest powrót po przesterowaniu, gdzie punkt pracy może się przesuwać, aczkolwiek każdy kto doprowadza wzmacniacz Hi-Fi do przycinania sygnału, wysterowuje go po prostu zbyt mocno. Oczywiście, możliwe jest też zastosowanie pompy ładunkowej do źródła prądowego dla połączenia zalet obu rozwiązań. Zrobiono tak w układzie wzmacniacza TBA820. Obciążenie z cewką indukcyjną jest jeszcze innym sposobem, najskuteczniejszym ze wszystkich, ale takie cewki audio są drogie i trudne do zdobycia.

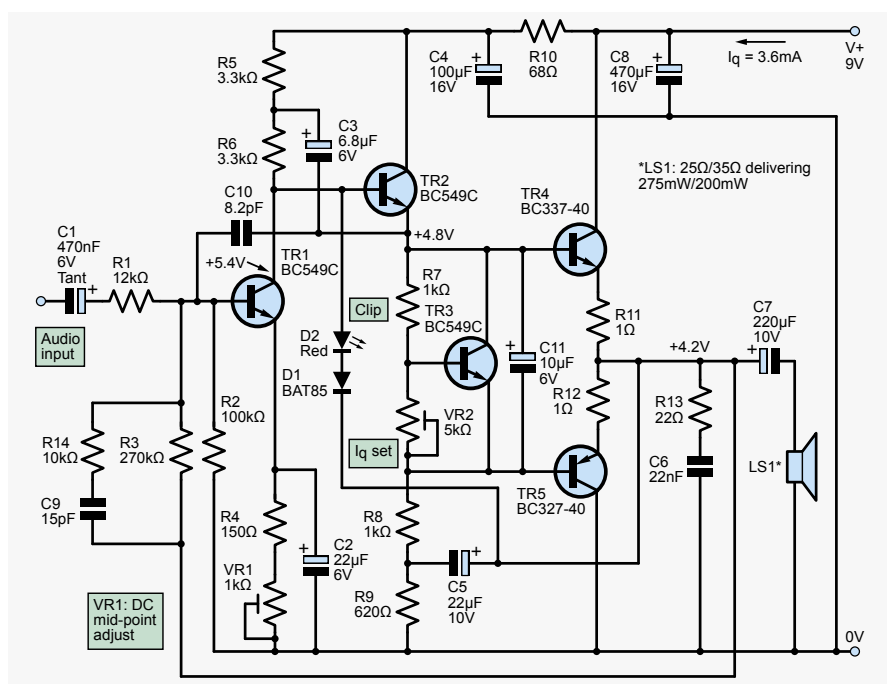
Projekt obwodu

Obwód składa się z połączonych ze sobą: wzmacniacza napięciowego (VAS) w układzie wspólnego emitera (TR1) i stopnia wyjściowego w układzie przeciwobnego wtórnika emiterowego push-pull (TR4 i TR5), jak pokazano na **rysunku 8**. Dodatkowo, między te stopnie włączono wtórnik emiterowy (TR2). Zwiększa on prąd sterujący końcówką i zmniejsza obciążenie stopnia wejściowego VAS. Wtórnik ten zapewnia również sterowanie obwodu bootstrap (C3) do zasilania rezystora obciążenia pierwszego stopnia (R6) dla zwiększenia wzmocnienia wzmacniacza w otwartej pętli. Wtórnik pracuje w klasie A, ponieważ nie ma tu sensu używać wtórnika takiego jak wyjściowy, jak to ma miejsce w wielu projektach, ponieważ pracuje on w klasie B. Nie byłoby wtedy działania bootstrap wokół punktu przejścia (połowa napięcia wyjściowego), który jest właśnie tam, gdzie jest potrzebny, aby zredukować to paskudne zniekształcenie brzmienia.

Aby poprawić wydajność dla pracy baterijnej, rezystor obciążający TR2 jest pompowany



Rysunek 8. Schemat blokowy wzmacniacza z obwodem bootstrap. Zwróć uwagę na wtórnik emiterowy TR2 pracujący w klasie A, sterujący stopniem wyjściowym i bootstrap dla TR1



Rysunek 9. Pełny obwód wzmacniacza bootstrap. Pięć tranzystorów to minimum zapewniające dźwiękowo akceptowalną wydajność przy niskim zużyciu prądu

(bootstrap) przez C5. TR3 to standardowy mnożnik napięcia U_{be} do polaryzacji wyjścia, do pracy w klasie AB. Kondensatory C6, C9 i C10 (zwiększa efekt Millera dla TR1) oraz rezystor R14 mają zapobiegać wzbudzeniu się wzmacniacza na wysokich częstotliwościach. Wzmocnienie napięciowe wynosi 22 V/V (ustalone przez rezystory R1 i R3), połowę tego co mają wzmacniacze scalone. Rezystory R11 i R12 tworzą ujemne sprzężenie zwrotne zapobiegające wzrostowi prądu spoczynkowego tranzystorów końcowych przy wzroście temperatury. R13 zapewnia rezystancyjne obciążenie wzmacniacza dla częstotliwości ponadakustycznych. VR1 pozwala ustawić napięcia wyjściowe (na dodatniej okładce C7). VR2 służy do regulacji prądu spoczynkowego tranzystorów końcowych. C4 i R10 filtrują zasilanie części sterującej wzmacniacza. Pełny schemat układu pokazano na **rysunku 9**.

Wskaźnik przesterowania

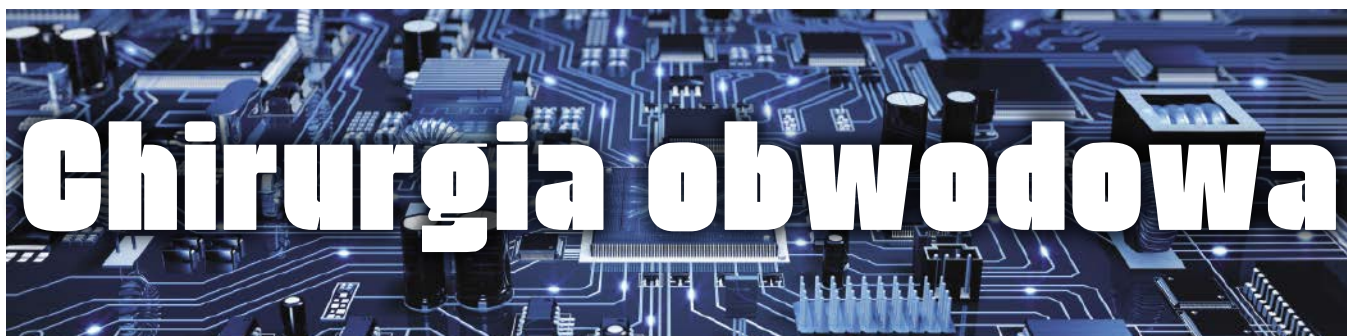
W układzie przewidziano wskaźnik przesterowania, który sygnalizuje obcinanie sygnału przez wzmacniacz. Dioda LED1 zaświeca się, gdy napięcie na kolektorze TR1 przekracza wyjściowe o 1,85 V, co występuje w sytuacji przesterowania. Zapewnia również miękkie obcinanie dodatnich szczytów sygnału jeśli szeregowo dioda Schottky'ego (D1) jest zwarta. Zaokrąglone szczyty fali brzmią przyjemnie, jak wzmacniacz gitarowy z efektem fuzz.

W następnym miesiącu

To był niezły przegląd wzmacniacza z układem bootstrap, za miesiąc go zbudujemy. ■

Jake Rothman

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, listopad 2020 (www.epemag3.com)



Chirurgia obwodowa

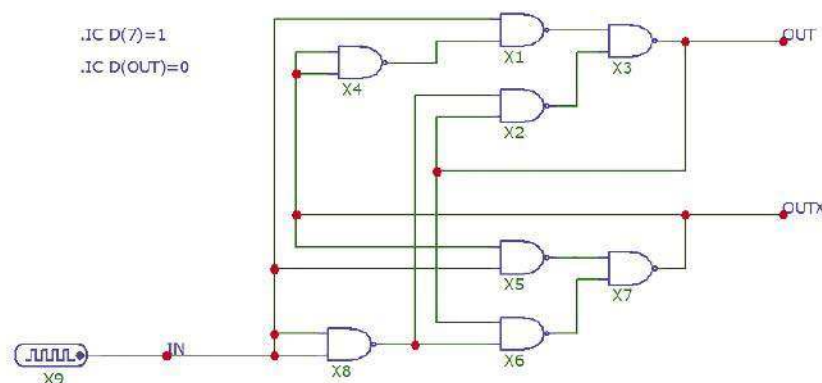
Problemy z symulacją cyfrowego układu dzielenia przez dwa

W odpowiedzi na artykuł Circuit Surgery na temat oprogramowania do rysowania i symulacji obwodów elektrycznych – Micro-Cap 12 (PE grudzień 2020), Ken Wood postanowił spróbować zasymulować obwód, z którym miał problemy wiele lat temu. Píše: „Byłem zachwycony, gdy dowiedziałem się o dostępności Micro-Cap 12 i szybko go pobrałem. Ostatnim razem korzystałem z podobnego narzędzia około 1988 roku, kiedy to mieliśmy w pracy ogromną maszynę uniksową. Z tamtych czasów zapamiętałem pewien problem...”

Zapamiętany problem dotyczył „odtworzenia standardowych funkcji rejestru logicznego przy użyciu samych bramek. Nie udało mi się wtedy zasymulować tych obwodów, więc uznałem, że warto spróbować jeszcze raz. Zaprojektowałem prostą funkcję dzielenia przez dwa (wyjście przełącza się z prędkością równą połowie wejścia) przy użyciu ośmiu bramek i wprowadziłem ją do MC12 (patrz rysunek 1). Uruchomienie symulatora powodowało niestabilność (patrz rysunek 2). W obwodzie jest kilka bramek sprzężonych krzyżowo, a wraz z innymi sygnałami w odpowiednim stanie tworzy to pętlę dodatniego sprzężenia zwrotnego z opóźnieniem, która oscyluje w symulatorze”.

Z doświadczenia wiem, że to nie będzie oscylować, będzie zatrząskiwać się w jednym lub drugim stanie, więc zbudowałem obwód i udowodniłem, że działa. Zdjęcie pokazuje obwód na płytce prototypowej zbudowany z dwóch układów scalonych CD4011 (quad-NAND), zasilany napięciem 5 V i taktowany sygnałem zegarowym kalibracyjnym – 1 kHz, 5 V z oscyloskopu (rysunek 3). Zadziałało za pierwszym razem, dokładnie tak, jak zostało zaprojektowane”.

„Czy jest jakaś odpowiedź na ten problem z symulacją? Myślę, że może to być spowodowane tym, że modele są idealne i tworzą zjawisko wyścigu czasowego, podczas gdy prawdziwe urządzenia nigdy nie są idealne. Próbowalem zmienić niektóre opóźnienia bramek, ale nie udało mi się znaleźć działającej



Rysunek 1. Schemat w symulatorze Micro-Cap 12 Kena dla jego obwodu dzielenia przez dwa

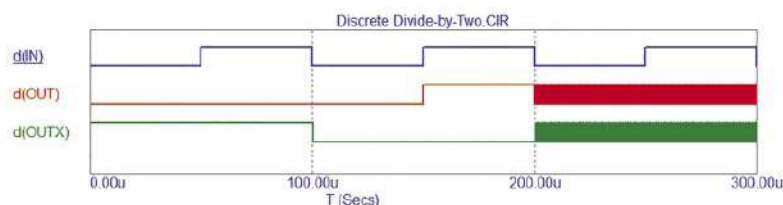
kombinacji. Teraz już wiem, dlaczego moje symulacje nie działały w 1988 roku. Symulator był uważany za „be all and end all”, i dobrze, że nigdy nie potrzebowałem takiego obwodu w profesjonalnym projekcie, ponieważ trudno byłoby przejść przegląd przedprodukcyjny bez potwierdzenia symulatora, że zadziała!”

Zanim przyjrzymy się układowi Kena i temu, co może być nie tak w symulacji, wrócimy do podstaw i zobaczymy, jak są zbudowane układy flip-flop (przerzutniki) z bramek

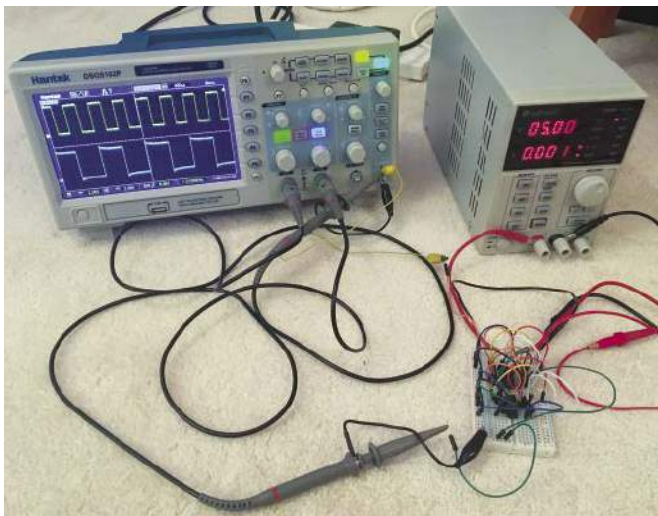
logicznych i jakie problemy mogą wystąpić podczas ich działania.

Z tego powstają wspomnienia

Wyobraź sobie dwa inwertery (bramki NOT) połączone szeregowo (patrz rysunek 4). Stan 1 na wejściu A daje stan 1 na wyjściu C, a 0 na wejściu daje 0 na wyjściu. Zastanów się, co się stanie, gdy podłączymy wyjście z powrotem do wejścia (patrz rysunek 5). Nie ma konfliktu, ponieważ wejście i wyjście



Rysunek 2. Wyniki symulacji Micro-Cap 12 obwodu z rysunku 1 – ciągłe kolorowe bloki na wyjściach to oscylacje, które są zbyt szybkie, aby zobaczyć je w tej skali czasowej



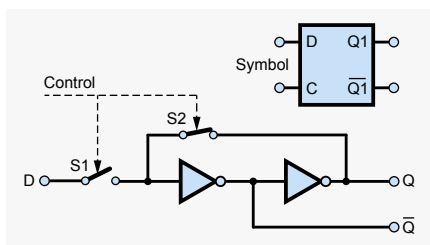
Rysunek 3. Obwód z rysunku 1 zbudowany na płytce drukowanej, działający poprawnie

są na tym samym poziomie logicznym (tak naprawdę nie ma już wejścia). Jeśli punkt A (teraz taki sam jak C) zostanie ustawiony na 1, pozostanie tam w nieskończoność. Jeśli ustawimy A na 0, pozostanie na 0. Te dwa warunki są określane jako „stany stabilne”, a obwód jest opisywany jako bistabilny.

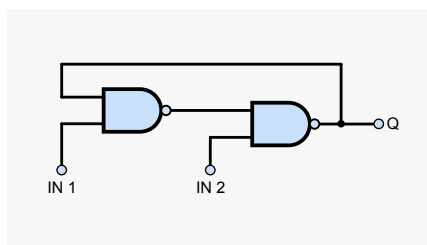
Zdolność do utrzymywania w nieskończoność jednego z dwóch możliwych stanów jest podstawą statycznej pamięci cyfrowej, gdzie termin „styczna” odnosi się do faktu, że stan jest utrzymywany na stałe (tak długo, jak zasilanie jest podłączone) i nie zanika ani nie wymaga odświeżania.

Funkcja pamięci realizowana przez obwód z **rysunku 5** nie jest szczególnie przydatna, ponieważ nie ma wejścia, za pomocą którego można by podać wartość do zapamiętania. Możemy jednak zmodyfikować obwód na kilka sposobów, aby to osiągnąć. Po pierwsze, możemy przerwać pętlę za pomocą przełącznika, jednocześnie używając innego przełącznika do ustawienia wejścia. Po drugie, można zastosować dodatkowe bramki do modyfikacji logiki pętli i ustawienia stanu.

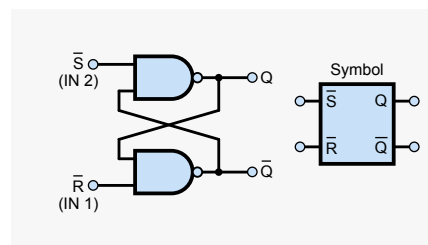
Sposób, w jaki przełączniki mogą być używane do ustawiania stanu pętli z **rysunku 5**, pokazano na **rysunku 6**. Obwód ten może być zaimplementowany w układach CMOS, gdzie MOSFETy mogą być użyte jako przełączniki. Dwa przełączniki (S1, S2) są sterowane razem i są ustawione tak, że gdy S1 jest otwarty, S2 jest zamknięty i odwrotnie. Gdy S1 jest otwarty, a S2 zamknięty, wejście jest izolowane i pętla jest zamknięta (jak na **rysunku 5**). Pętla pozostaje w swoim aktualnym stanie i nie ma na nią wpływu wejście. Gdy S1 jest zamknięty, a S2 jest otwarty, wejście jest podłączone do wyjścia za pośrednictwem dwóch inwerterów (jak na **rysunku 4**), a wyjście będzie podążać za wejściem. Gdy przełączniki zmieniają się na S1 otwarty i S2 zamknięty, tworzona jest pętla, przechowująca dane. Pojemność bramki inwertera będzie utrzymywać wartość podczas przełączania, dopóki pętla nie zablokuje zapisanego stanu.



Rysunek 6. Obwód zatrasku i symbol



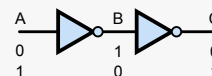
Rysunek 7. Przerzutnik set-reset



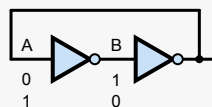
Rysunek 8. Obwód i symbol przerzutnika set-reset

Przerzutnik set-reset

Obwód na **rysunku 7** zastępuje inwertery z **rysunku 5** dwiema bramkami NAND. Jeśli $IN1 = IN2 = 1$, obwód jest równoważny pętli z dwoma inwerterami. Jednak, jeśli jedno z wejść wynosi zero, wyjście Q jest



Rysunek 4. Dwa inwertery połączone szeregowo



Rysunek 5. Dwa inwertery w pętli – obwód bistabilny, podstawa przerzutnika

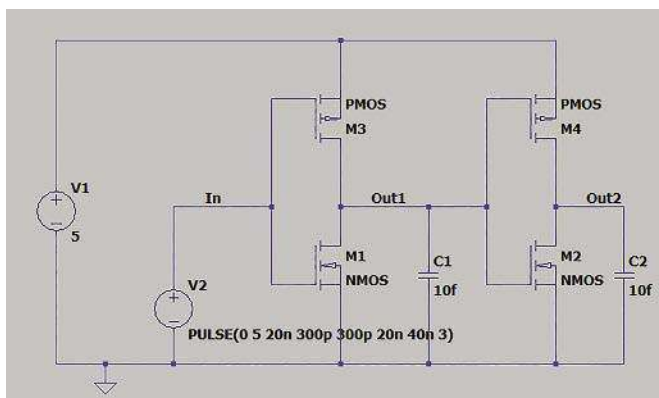
wymuszane na 1 lub 0, niezależnie od wcześniej zapisanej wartości. Jeśli $IN1 = 0$ i $IN2 = 1$, na wyjściu generowana jest wartość $Q = 0$. Wartość ta zostanie zachowana, gdy $IN1$ powróci do 1. Jeśli $IN1 = 1$ i $IN2 = 0$, na wyjściu pojawi się $Q = 1$. Wartość ta zostanie zachowana, gdy $IN2$ powróci do 1.

Niski stan na wejściu $IN1$ wymusza zapisanie 0, więc wejście to nazywane jest resetem (R). Niski stan na wejściu $IN2$ wymusza zapisanie wartości 1, więc wejście to nazywane jest set (S). Ten typ obwodu jest często rysowany w formie pokazanej na **rysunku 8** i może być reprezentowany przez symbol blokowy, również pokazany na **rysunku 8**. Wejścia set i reset są aktywne w stanie niskim, stąd paski narysowane nad nimi na schemacie. Jeśli połączysz dwie bramki NOR w tej samej konfiguracji, otrzymasz przerzutnik SR z aktywnymi wejściami w stanie wysokim. W obwodzie Kena znajdują się dwa przerzutniki SR, jeden utworzony przez X2 i X3, a drugi przez X5 i X7.

Wyjścia obu bramek NAND na **rysunku 8** mogą być używane i są komplementarne (przeciwnie poziomy logicznie względem siebie), a zatem są oznaczone jako Q i !Q. Jeśli jednak ustawimy oba wejścia na zero, oba wyjścia przejdą na 1. Jest to „nielegalny” warunek, być może sugerujący, że nasz przerzutnik nie jest używany prawidłowo – w rzeczywistości prosimy go o ustawienie i zresetowanie w tym samym czasie, co tak naprawdę nie ma sensu. Jeśli jednak usuniemy jeden z aktywnych warunków wejściowych, pozostały zostanie ustawiony poprawnie. Na przykład, jeśli zastosujemy $R=0$ i $S=0$ do przerzutnika na **rysunku 8**, a następnie zmienimy na $S=1$, zachowując $R=0$, przerzutnik zostanie zresetowany. Jeśli warunek, w którym oba wyjścia mają wartość 1, nie powoduje problemów w innych częściach obwodu, użycie przerzutnika w ten sposób może być w porządku.

Nieprzewidywalny

Znacznie poważniejszy problem pojawia się, gdy oba wejścia są jednocześnie ustawione na 1. Przerzutnik skutecznie zapamiętuje ostatnią otrzymaną instrukcję (set lub reset). Jeśli zastosujemy obie instrukcje jednocześnie i usuniemy je jednocześnie, to który stan powinien zostać zapamiętany, biorąc pod uwagę, że może on przechowywać tylko jeden stan? Obwód jest symetryczny i nic nie wskazuje na to, co powinno się stać – odpowiedź nie jest zdefiniowana. Stosowanie obwodu z niezdefiniowanym zachowaniem nie jest dobrym pomysłem, ale przerzutnik



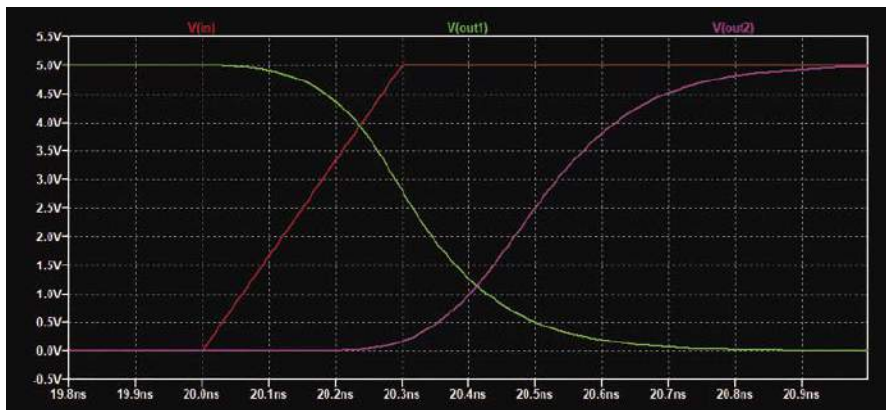
Rysunek 9. Schemat w LTspice dwóch inwerterów CMOS połączonych szeregowo

SR jest przydatny w innym przypadku, więc możemy przyjąć zasadę, aby nigdy tego nie robić i spróbować zaprojektować go tak, aby nigdy się nie pojawił. Jednak jeśli usuniemy set i reset jednocześnie, to rzeczywisty obwód zrobi coś, co warto zrozumieć (na wypadek, gdyby wystąpiło w jakimś naszym projekcie).

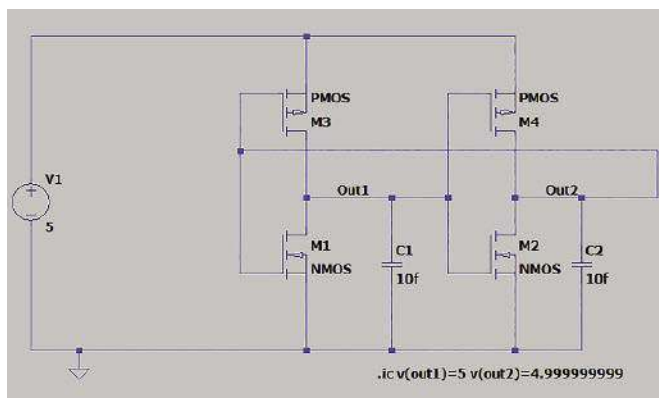
Gdy jednocześnie zmieniamy oba wejścia przerzutnika NAND SR z 0 na 1, skutecznie tworzymy sytuację podobną do zapętlenia inwerterów na rysunku 5, gdzie oba mają wyjście 1. Oba inwertery zareagują na swoje wejścia (również oba 1), zmieniając swoje wyjścia w kierunku logicznego 0. Jeśli uznamy, że bramka zachowuje się jak funkcja Boole'a NOT z opóźnieniem, to oba inwertery wyprowadzą 0 po upływie czasu opóźnienia. Jeśli mają identyczne opóźnienia, to oba zmienią stan na 0 jednocześnie. Mamy teraz zapętlenie inwerterów z rysunku 5, oba wysyłające 0, oba zmieniają się na 1 w tym samym czasie i wrócimy do miejsca, w którym zaczęliśmy. Proces ten będzie kontynuowany w nieskończoność – zapętlenie inwerterów będą oscylować, tak jak dzieje się to w symulacji Kena.

W rzeczywistym obwodzie jest mało prawdopodobne (prawie niemożliwe), aby dwa inwertery (w rzeczywistości bramki NAND w przerzutnikach SR z rysunku 8) miały dokładnie takie samo opóźnienie. W takiej sytuacji najszybszy inwerter przełączy się najpierw na zero, w którym to momencie pętla inwertera znajdzie się w jednym ze stabilnych stanów pokazanych na rysunku 5. Stan ten jest stabilny, więc pozostanie tam – nie wystąpią żadne oscylacje. Problem polega na tym, że generalnie nie wiemy, która bramka jest najszybsza, więc wynik jest nieprzewidywalny – przerzutnik może się ustawić lub zresetować, tego nie wiemy.

Ta nieprzewidywalność może oznaczać, że fizyczny obwód zachowuje się inaczej w różnych momentach lub różne kopie obwodu robią różne rzeczy – oba przypadki mogą mieć bardzo poważne konsekwencje.



Rysunek 10. Wyniki symulacji dla układu z rysunku 9 pokazujące zmianę wejścia z 0 na 1 i odpowiedź dwóch bramek



Rysunek 11. Schemat LTspice dwóch inwerterów CMOS w pętli. Zwróć uwagę na dyrektywę .ic SPICE, aby ustawić początkowe napięcie na wyjściach

Możliwe jest również, że (prawie) każda kopia zbudowanego przez nas obwodu zachowuje się spójnie, na przykład dlatego, że konfiguracja obwodu zapewnia, że jedna z zapętlenych bramek jest (prawie) zawsze najwolniejsza. Jeśli takie zachowanie jest zgodne z intencją projektową, może to prowadzić do fałszywego poczucia bezpieczeństwa – obwód może zawieść w określonych warunkach.

Jeśli spojrzymy nieco głębiej na to, co dzieje się, gdy zmieniamy oba wejścia bramki NAND SR z 0 na 1 jednocześnie, stwierdzimy, że widok bramki logicznej bazujący na funkcji Boole'a plus opóźnienie nie opisuje tego, co się dzieje. Może to oznaczać, że symulacja cyfrowa (bazuje na funkcji logicznej i opóźnieniu) nie jest w stanie przewidzieć zachowania obwodu w omawianych przez nas warunkach. Niektórzy mogą powiedzieć, że oznacza to, że nie należy ufać symulatorowi – zamiast tego należy zbudować obwód. Ale zbudowanie prototypu, w którym projekt jest z natury nieprzewidywalny, może być równie mylące. Jeśli ta kopia zachowuje się w określony sposób, nie oznacza to, że inne kopie będą robić to samo.

Symulacja analogowa

Symulacja cyfrowa upraszcza zachowanie podstawowych obwodów bramki logicznej, umożliwiając symulację bardzo dużych i złożonych obwodów w rozsądnym czasie. Możemy oczywiście dokładniej symulować stosunkowo małe obwody cyfrowe za pomocą symulatora analogowego i pełnego schematu obwodu tranzystora. Przyjrzymy się symulacji LTspice obwodów z rysunku 4 i rysunku 5, aby uzyskać wgląd w to, co dzieje się, gdy zmieniamy oba wejścia NAND SR z 0 na 1 jednocześnie. Użyjemy tylko inwerterów zbudowanych z generycznych tranzystorów MOSFET, zamiast próbować modelować pełny układ przerzutnika SR używany przez Kena – będzie to wystarczające do zilustrowania typowego zachowania.

Obwód z **rysunku 9** to dwa inwertery CMOS połączone szeregowo, które wykorzystamy jako punkt odniesienia podczas analizy zachowania pętli. Kondensatory 10fF (10×10^{-15} F) reprezentują pojemność okablowania łączącego inwertery (tak jak może to być w układzie scalonym). Wyniki symulacji na **rysunku 10** pokazują, że średnie opóźnienie propagacji dwóch inwerterów wynosi około 175 ps (175×10^{-12} s).

Rysunek 11 to obwód z rysunku 9 przekształcony w pętlę inwerterów, jak na rysunku 5. Nie mamy już sygnału wejściowego z V2 i musimy określić napięcia na wejściach/wyjściach inwertera na początku symulacji za pomocą dyrektywy SPICE.ic (warunek początkowy). Ustawiając warunki początkowe, w których oba wejścia/wyjścia

mają wartość logiczną 1 (w tym przypadku 5 V), symulujemy to, co dzieje się bezpośrednio po zmianie obu wejść NAND SR z 0 na 1 jednocześnie.

Rysunki 12 i 13 pokazują wyniki dwóch symulacji z dwoma wyjściami przy początkowym napięciu 5 V oraz o 1 nV mniejszym niż 5 V (4,999999999 V) – próbując w obie strony. Ta niewielka różnica jest wystarczająca do określenia ostatecznego stanu obwodu (które wyjście przechodzi do logiki 1 (5 V), a które do logiki 0 (0 V)). Wyjście, które zaczyna się przy bardzo nieznacznie niższym napięciu, kończy się na 0.

Jest kilka kluczowych rzeczy, na które należy zwrócić uwagę. Po pierwsze, nie ma oscylacji, w przeciwieństwie do tej przewidywanej przez model bramki cyfrowej z równymi opóźnieniami (i symulację Kena). Po drugie, oba napięcia pozostają bardzo blisko siebie przez długi czas w porównaniu do opóźnienia bramek widocznego na rysunku 10. Jest to ponad dziesięć razy więcej niż 175 ps opóźnienia propagacji inwerterów, zanim pojawi się jakkolwiek znacząca różnica napięcia widoczna na rysunku 12 i rysunku 13.

Metastabilność

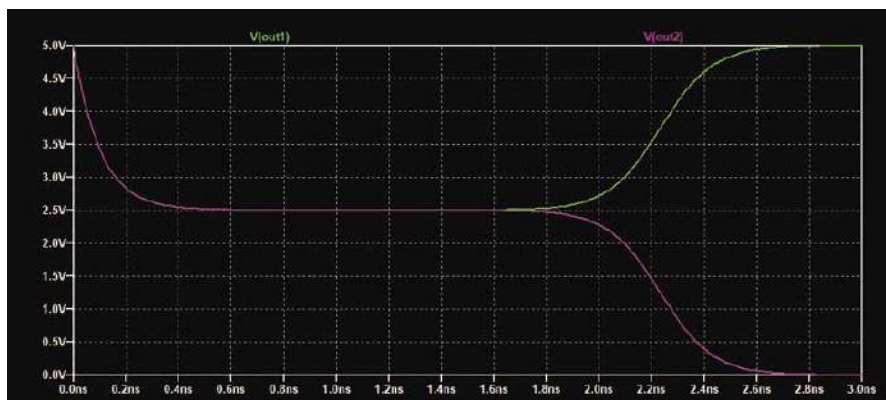
Zachowanie widoczne na rysunku 12 i rysunku 13 jest formą tak zwanej „metastabilności”. Jak wspomniano wcześniej, obwód ma dwa stabilne stany, ale ma również stan metastabilny, w którym napięcia wejściowe/wyjściowe obu inwerterów są równe. Jest to punkt idealnej równowagi, jak balansowanie kulką na czubku igły – teoretycznie mogłaby tam pozostać na zawsze, ale sytuacja jest z natury niestabilna – niewielki ruch powietrza spowoduje, że kulka spadnie z igły, i podobnie każda niewielka perturbacja napięcia spowoduje, że obwód spadnie do jednego z dwóch stabilnych stanów.

Odpowiedź metastabilnych układów przerzutników może być analizowana poprzez rozwiązywanie równań różniczkowych opisujących zachowanie obwodu. Wyniki pokazują, że im doskonalsza jest równowaga początkowa, tym dłużej trwa wyjście ze stanu metastabilnego (znane jako czas rozdzielczości), a także, że w niektórych obwodach mogą wystąpić krótkie oscylacje, zanim wyjście się ustabilizuje.

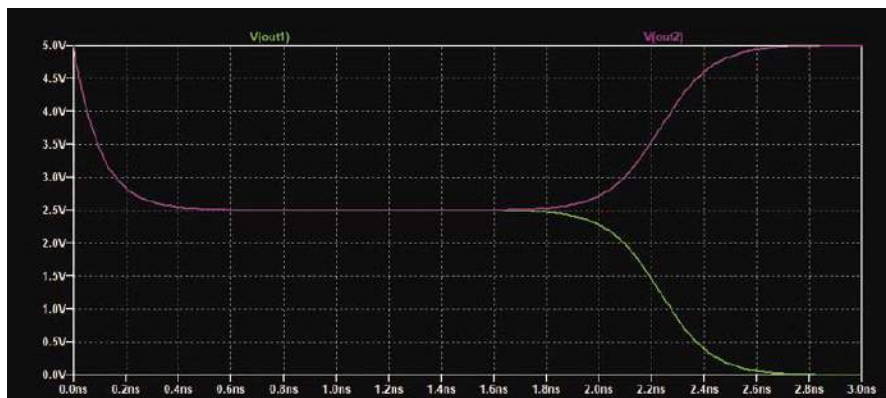
Metastabilność może powodować awarie systemów cyfrowych. Nieprzewidywalne wyniki i utrzymujące się pośrednie poziomy napięcia mogą powodować propagację nieprawidłowych wartości logicznych w obwodzie. Nawet jeśli poziom logiczny z metastabilnego przerzutnika ostatecznie okaże się prawidłowy, jeśli czas rozdzielczości jest zbyt długi (np. dłuższy niż cykl zegara), wystąpią błędy (np. nieprawidłowa wartość jest próbkowana przez układ podłączony do wyjścia metastabilnego).

Śledzenie problemu

Powodem, dla którego symulacja w programie Micro-Cap 12 obwodu Kena pokazuje oscylację, jest to, że wejścia przerzutnika SR mogą jednocześnie zmieniać się z 0 na 1, a domyślnie symulacja zawiera dokładnie równe opóźnienia dla wszystkich bramek. Rozważmy układ przerzutnika SR utworzony przez X2 i X3. Gdy zegar ma wartość 0, wejścia set i reset mają wartość 1 poprzez inwerter X8 i bramkę NAND X1, które są sterowane przez zegar. Gdy zegar wynosi 1, wejścia S=OUTX (przez



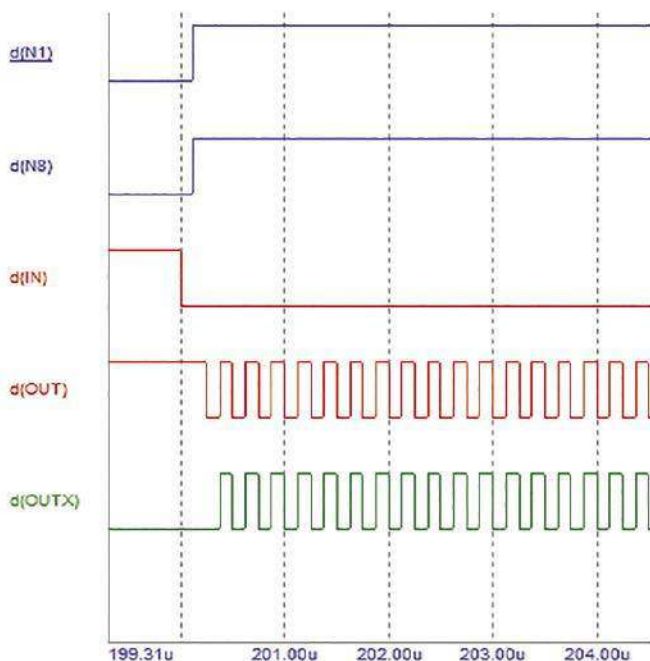
Rysunek 12. Symulacja obwodu z rysunku 11 z warunkami początkowymi $v(out1)=5$ i $v(out2)=4.999999999$



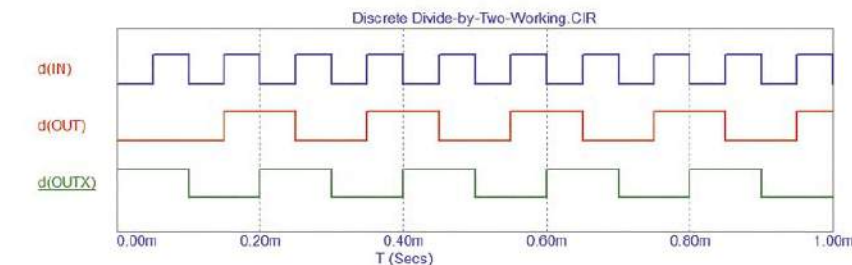
Rysunek 13. Symulacja obwodu z rysunku 11 z warunkami początkowymi $v(out1)=4.999999999$ i $v(out2)=5$

X4 i X1) i R=0 (przez X8). Jeśli OUTX wynosi 0, gdy zegar zmienia się z 1 na 0, wejścia przerzutnika zmieniają się jednocześnie z 0 na 1. Inicjuje to oscylację z powodu równych opóźnień dla X2 i X3, jak opisano powyżej.

Sytuację tę pokazano na **rysunku 14**. Jest to powiększenie wyników pokazanych na rysunku 2 z dodatkowymi śladami pokazującymi wyjścia bramek X1 (N1) i X8 (N8). Widzimy, jak zegar (IN) zmienia się



Rysunek 14. Szczegółowe spojrzenie na oscylacje w obwodzie Kena



Rysunek 15. Wyniki symulacji Micro-Cap 12 obwodu z rysunku 1 z opóźnieniem X8 dostosowanym tak, aby nie odpowiadało opóźnieniu X1 (porównaj z oryginalną symulacją z rysunku 2 i działającym obwodem z rysunku 3)

z 1 na 0, a następnie N1 i N8 zmieniają się z 0 na 1 jednocześnie ale minimalnie później. Oscylacja rozpoczyna się w przerzutniku utworzonym przez X2 i X3 po czasie opóźnienia bramki (przebieg dla OUT).

W fizycznej implementacji obwodu Kena, stan metastabilny w X2/X3 może nie zostać wyzwolony, ponieważ zależy on od tego, czy opóźnienia X1 i X8 są prawie równe. Ponadto, jak widzieliśmy, trwałe oscylacje prawdopodobnie nie wystąpią w rzeczywistym obwodzie, jeśli wystąpi stan metastabilny – więc nie spodziewalibyśmy się zobaczyć tego na oscyloskopie, nawet gdyby istniała metastabilność.

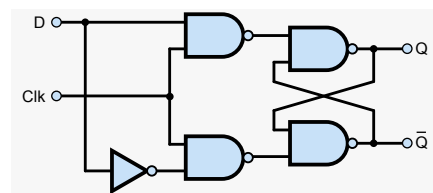
Jeśli założymy, że opóźnienia X1 i X8 są wystarczająco różne, aby nie wywoływać problemu, możemy zasymulować obwód w tych warunkach. Możemy to zrobić w Micro-Cap 12, ustawiając PARAMS:=MNTYMXDLY=1 dla X8, aby wybrać model minimalnego opóźnienia czasowego dla bramki, zamiast wartości domyślnej. Kliknij dwukrotnie X8 na schemacie, aby zobaczyć i ustawić ustawienia parametrów. Wyniki tego pokazano na rysunku 15 – obwód działa zgodnie z zamierzeniami Kena i odpowiada wynikom uzyskanym z rzeczywistego obwodu.

Pytania

Wyniki te prowadzą do pytań dotyczących symulacji. Po pierwsze, czy symulacja pokazana na rysunku 2 jest błędna? Możemy powiedzieć „nie” – symulator wykonał dokładnie to, o co został poproszony i zasymulował obwód ze wszystkimi bramkami o dokładnie równych opóźnieniach, a wyniki są poprawne na zastosowanym poziomie abstrakcji. Nie odwzorowuje on dokładnie analogowego zachowania metastabilnego przerzutników, ale nie spodziewalibyśmy się, że będzie to możliwe, biorąc pod uwagę modele bramek logicznych używane do zapewnienia szybkiej symulacji cyfrowej. Po drugie, konkretna sytuacja pokazana na rysunku 2 jest mało prawdopodobna w rzeczywistym obwodzie, więc możemy również zapytać, czy wynik jest przydatny? Tutaj możemy powiedzieć „tak” – ostrzega nas o tym, że obwód może zawierać ukryty problem metastabilności. Wiedząc to, możemy zdecydować, czy ryzyko wystąpienia tego problemu jest akceptowalne, czy też nie, co zależy od kontekstu, w jakim obwód będzie używany.

Symulatory nie są najlepszymi narzędziami do wykrywania problemów z synchronizacją, ale pakiety oprogramowania przeznaczone do zaawansowanego projektowania cyfrowego mogą zawierać analizatory synchronizacji, które wykrywają możliwe problemy z synchronizacją, w tym metastabilność.

Kolejnym pytaniem, które się nasuwa, jest to, czy możemy zbudować podobny obwód, który



Rysunek 16. Zatrzask danych zawierający przerzutnik SR NAND

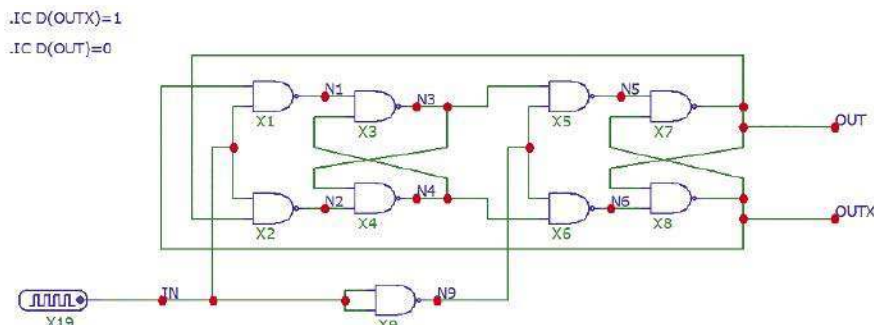
nie zawiera problemów omówionych powyżej? Jedną z możliwości jest zastosowanie zatrzasku danych pokazanego na **rysunku 16**. Zawiera on również przerzutnik NAND SR do przecho-

wywania danych, ale robi to pod kontrolą zegara. Jest on skonfigurowany w taki sposób, że wejścia S i R przerzutnika SR nie mogą być aktywne w tym samym czasie. Gdy zegar (Clk) jest niski, dwie bramki NAND podłączone do zegara wymuszają na wejściach S i R stan nieaktywny (1). Gdy zegar jest wysoki, wejście D jest przekazywane do wejść S i R, ale bramka NOT zapewnia, że tylko jedno z wejść S lub R jest aktywne. Przerzutnik ustawia się lub resetuje w zależności od wartości D.

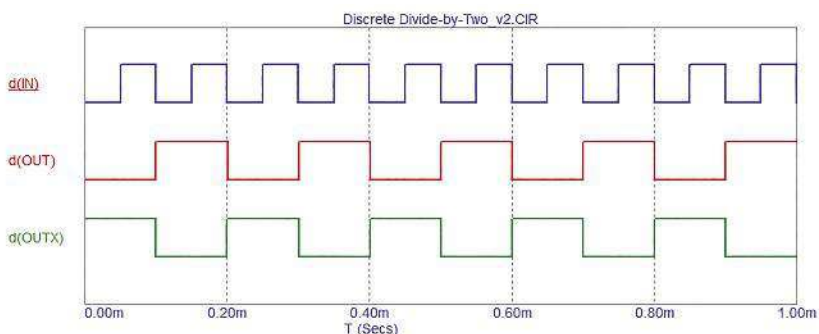
Użycie dwóch takich zatrzasków w szeregu z wyjściem drugiego zatrzasku odwróconym i doprowadzonym z powrotem do pierwszego tworzy obwód dzielenia przez dwa, co jest równoważne z podłączeniem wyjścia Q wyzwalanego zboczem ujemnym przerzutnika typu D z powrotem do jego wejścia D. Schemat proponowanego obwodu pokazano na **rysunku 17**. Nie obejmuje on bramek NOT z rysunku 16, ponieważ możemy użyć komplementarnych wyjść z przerzutników SR. Wyniki symulacji pokazano na **rysunku 18**, gdzie można zobaczyć operację dzielenia przez dwa. Obwód nie jest odporny na metastabilność – może to być spowodowane zmianą zegara zatrzasku i danych w tym samym czasie, ale aby to osiągnąć, musielibyśmy taktować obwód wystarczająco szybko, aby opóźnienie przez zatrzaski odpowiadało okresowi zegara. ■

Ian Bell

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, luty 2021 (www.epemag3.com)



Rysunek 17. Obwód dzielenia przez dwa bazujący na dwóch zatrzaskach D



Rysunek 18. Symulacja Micro-Cap 12 układu z rysunku 17 (wszystkie opóźnienia bramek są równe)



Fazowa regulacja mocy



Regulacja fazowa jest najprostszą zasadą kontrolowania mocy prądu przemiennego dostarczanego do obciążeń rezystancyjnych i indukcyjnych.

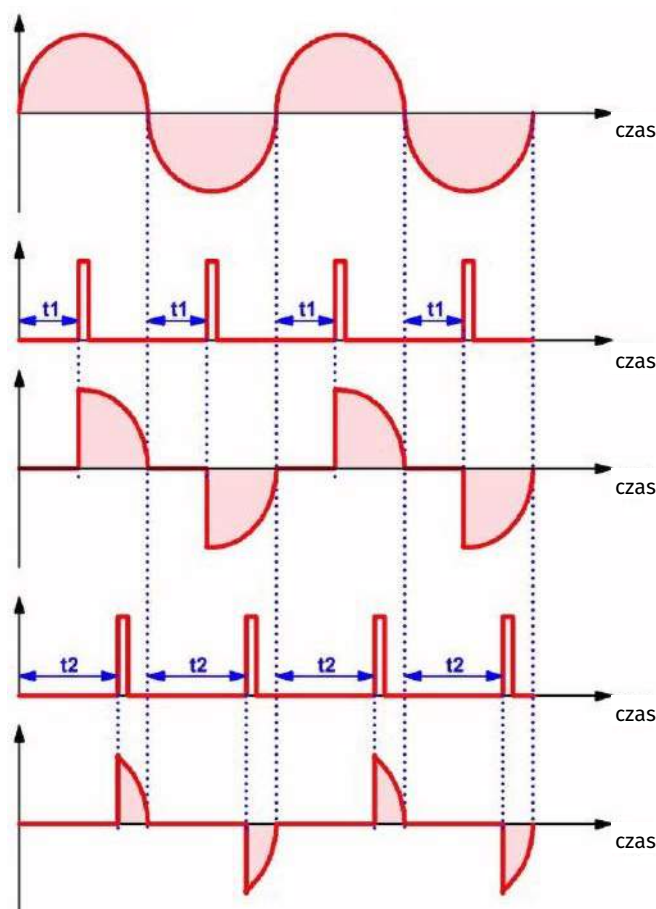
Zasada kontroli przebiegu napięcia sieciowego

Regulacja mocy dostępnej z napięcia sieciowego

Jeśli chcesz regulować moc dostarczaną do obciążenia poprzez napięcie sieciowe, możesz zastosować regulację fazową. Wartość skuteczną napięcia sieciowego można zmieniać w zakresie od zera do maksimum, niejako przesyłając napięcie sieciowe przez elektroniczny przełącznik, który włącza i wyłącza napięcie sieciowe synchronicznie z półokresami napięcia. Wpływając na moment załączenia przełącznika określamy, czy odbiornikowi oddawana jest mniejsza, czy większa część z całego półokresu.

Graficzne wyjaśnienie zasady

Zasadę wyjaśniono graficznie za pomocą wykresów na **rysunku 1**. Górny wykres pokazuje dwa okresy napięcia sieciowego. Drugi i trzeci wykres pokazują, co się stanie, gdy przełącznik elektroniczny zostanie zamknięty na okres czasu t_1 po przejściu przebiegu napięcia



1. Zasada regulacji fazowej

przez zero. Pierwsza część półokresu jest zablokowana, tylko ta część półokresu, która znajduje się za t_1 jest doprowadzona do obciążenia. Jeśli zamkniesz przełącznik w dalszej części okresu, na przykład w chwili t_2 , to do obciążenia zostanie doprowadzona jeszcze mniej-sza część półokresu.

Przełączanie za pomocą tyrystorów lub triaków

W ten sposób można bardzo płynnie regulować moc, jaką napięcie sieciowe może wygenerować w obciążeniu, w zakresie od zera do wartości maksymalnej. Ważnym warunkiem jest oczywiście ponowne otwarcie przełącznika przy następnym przejściu napięcia przez zero.

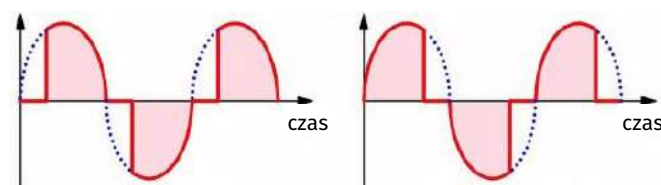
Tyrystor lub triak jest idealnym przełącznikiem elektronicznym do tego zastosowania. Jeśli przyłożysz napięcie do tego elementu to prąd nie będzie płynął przez niego. Jeśli jednak przyłożysz wąski impuls do bramki – elektrody sterującej, element załączy się i pozostanie w tym stanie, dopóki prąd płynący przez tyrystor lub triak nie spadnie poniżej wartości podtrzymania.

Załączony element można porównać do zamkniętego przełącznika o bardzo małej rezystancji wewnętrznej. Ze względu na właściwość fizyczną, która polega na tym, że półprzewodnik automatycznie wyłącza się, gdy prąd spadnie poniżej wartości podtrzymania, nie trzeba podejmować działań, aby wyłączyć element dokładnie przy przejściu fali sinusoidalnej przez zero. W momencie przejścia przez zero napięcie sieciowe przez krótki czas wynosi zero i przez obwód nie będzie przepływał żaden prąd. Dlatego prąd płynący przez tyrystor lub triak spada poniżej wartości podtrzymania i element wyłącza się automatycznie. Aby zaprojektować uniwersalny regulator mocy, wystarczy zaprojektować obwód generujący wąskie impulsy zsynchronizowane z przejściem fali sinusoidalnej przez zero i którego moment pojawienia się można przesuwając w sposób ciągły przez cały półokres napięcia sieci zasilającej.

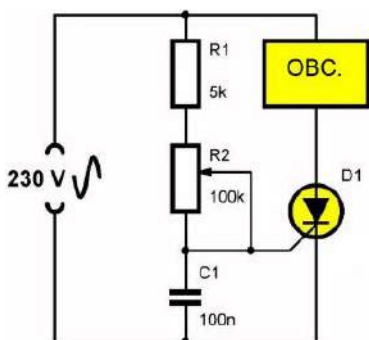
Wyprowadzenie fazy a odcięcie fazy

Istnieje również inna zasada regulacji mocy, którą może dostarczyć napięcie sieciowe. Zasada ta nazywa się „regulacją z odcięciem fazy” (**rysunek 2**). Dzięki tej zasadzie przełącznik elektroniczny zaczyna przewodzić, gdy napięcie sieciowe przekroczy zero i w pewnym momencie półokresu przełącznik elektroniczny zostaje ponownie wyłączony.

Kontrolery odcięcia fazy nadają się do wszystkich obciążeń rezystancyjnych i wszystkich obciążeń pojemnościowych, takich jak ściemniacze stateczniki elektroniczne. Jednakże ściemniacze z odcięciem fazy absolutnie nie mogą być używane z obciążeniami indukcyjnymi, takimi jak silniki, ponieważ prawie zawsze będzie to niszczące dla ściemniacza. Tylko najnowocześniejsze sterowniki nadają się do wszystkich obciążeń rezystancyjnych, a ponadto można bez problemu sterować wszystkimi obciążeniami indukcyjnymi, takimi jak



2. Porównanie regulacji fazowej i regulacji z odcięciem fazy



3. Najprostszy układ z regulacją fazową na bazie tyrystora

silniki i lampy halogenowe 12 V zasilane za pomocą konwencjonalnych transformatorów.

Regulacja fazowa za pomocą tyrystora

Najprostszy obwód

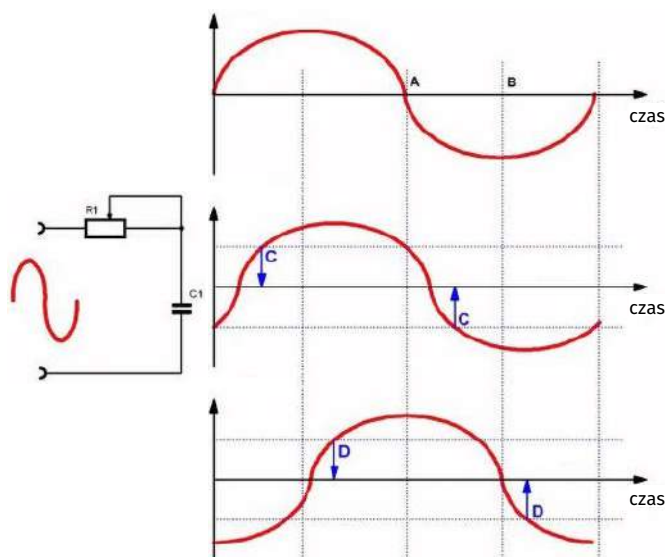
Najprostszy obwód realizujący w praktyce zasadę regulacji fazowej pokazano na rysunku 3. Załóżmy, że po włączeniu napięcia sieciowego rozpoczyna się ono od dodatnio skierowanej fali sinusoidalnej. Tyrystor D1 jest zablokowany, a napięcie na anodzie tyrystora wzrasta pod wpływem obciążenia. W tym samym czasie przez R1 i R2 przepływa prąd, który zaczyna ładować kondensator C1. W pewnym momencie napięcie kondensatora wzrosło tak bardzo, że do bramki wpływa tak duży prąd, że element się załącza. W tym momencie tyrystor zaczyna przewodzić. Dodatnia połowa sinusoidy na anodzie tyrystora jest wtedy częściowo przepuszczana. Na przykład, jeśli ustawisz potencjometr R2 w pozycji środkowej, dodatnia połowa sinusoidy jest już w połowie, zanim tyrystor zacznie przewodzić. Obciążenie również otrzymuje moc tylko przez połowę dodatniej połowy sinusoidalnej i otrzymuje niewielką moc. Przy ujemnej połowie sinusoidalnej tyrystor nie może się otworzyć, w końcu jest to dioda, która pozwala na przepływ prądu tylko w jednym kierunku.

Kiedy napięcie sieciowe osiągnie zero, prąd płynący przez tyrystor spada do zera, innymi słowy tyrystor zostaje odcięty. Dopiero podczas następnej dodatniej połowy okresu napięcia tyrystor może ponownie zacząć przewodzić. Jeśli ustawisz R2 na minimalną rezystancję, C1 zostanie bardzo szybko naładowany i zapali tyrystor na samym początku dodatniej połowy sinusoidy. Przy maksymalnej rezystancji R2 ładowanie C1 jest opóźnione w taki sposób, że tyrystor zapala się dopiero w ostatniej części połowy sinusoidy, a obciążenie otrzymuje minimalną moc.

Jak powstaje przesunięcie fazowe?

Obwód działa, ponieważ napięcie na kondensatorze C1 jest opóźnione w stosunku do napięcia sieciowego. Nazywa się to wówczas „przesunięciem fazowym”. Być może zastanawiasz się, co to oznacza. Z pomocą przychodzi teoria prądu przemiennego. Pełny okres napięcia sieciowego jest podzielony na 360 stopni kątowych (°). Obwód RC, jak pokazano na rysunku 4, powoduje przesunięcie fazowe do 90°. Dlatego napięcie na kondensatorze jest opóźnione w stosunku do napięcia sieciowego o kilka stopni kątowych.

Fizycznie można to wyjaśnić w następujący sposób. Ładowanie kondensatora za pomocą rezystora zajmuje określoną ilość czasu. Jeśli to napięcie jest napięciem stałym, ładowanie następuje raz. Po pewnym czasie napięcie kondensatora zrównuje się z napięciem prądu stałego i układ znajduje się w stanie spoczynku. Jednakże w pokazanym przykładzie napięcie ładowania jest napięciem przemiennym, którego wartość zmienia się sinusoidalnie. Kondensator C1 chce w sposób ciągły dostosowywać swoje napięcie do zmieniającej się



4. Układ RC zapewnia napięcie, którego przesunięcie fazowe jest regulowane

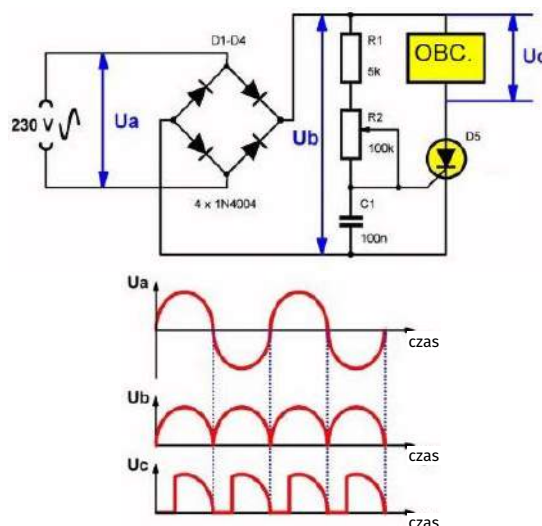
fali sinusoidalnej, ale rezystor ładujący R1 uniemożliwia to. W rezultacie napięcie kondensatora jest stale opóźnione w stosunku do napięcia sieciowego. Nazywa się to „przesunięciem fazowym”.

W chwili A napięcie sieciowe wynosi już zero, podczas gdy napięcie na kondensatorze jest nadal dodatnie. W punkcie B napięcie sieciowe jest maksymalnie ujemne, ale kondensator wciąż jest na drodze do wartości maksymalnej. Oczywiście jest, że wartość rezystora określa stopień opóźnienia zmian napięcia kondensatora.

Podsumowując, można powiedzieć, że im większy R1, tym dłużej kondensator osiąga określone napięcie, powiedzmy 40 V. Przy małej wartości R dzieje się to np. w chwili C, a przy dużej wartości w chwili D. Regulując rezystorem R1 regulujemy czas zapłonu.

Regulacja fazy w układzie z mostkiem prostowniczym

Wadą opisanego wcześniej układu jest prostownicze działanie tyrystora. W rezultacie ujemnie skierowane połówki sinusoidy nie są przepuszczane, a obciążenie może otrzymać maksymalnie połowę napięcia sieciowego. Można przezwyciężyć tę wadę, prostując najpierw napięcie sieciowe za pomocą mostka diodowego. Oznacza to, że „odwracasz ujemne połówki przebiegu w górę”, tak aby stały się również dodatnie. Pokazano to na rysunku 5. Za pomocą tego obwodu możliwa jest regulacja obciążenia od 0% do 100% maksymalnej dostępnej mocy. Metoda



5. Rozbudowa obwodu o mostek prostowniczy

ta jest jednak rzadko stosowana w praktyce, ponieważ problem prostowania można rozwiązać w prostszy sposób za pomocą triaka.

Kontrola fazy za pomocą triaka

Najprostszy obwód

Praktyczny obwód pokazano na **rysunku 6**. Od razu widać duże podobieństwo do obwodu tyrystorowego, a działanie jest również identyczne, z tym wyjątkiem, że w tym przypadku ujemne połówki sinusoidy są również przepuszczane do obciążenia. Jedyną znaczącą różnicą jest dioda wyzwalająca D2 – diak.

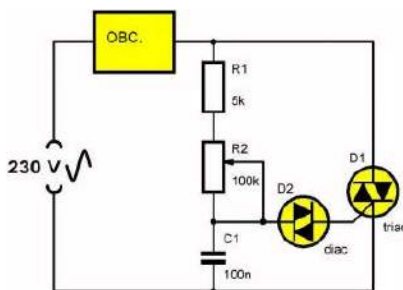
Diak jest niezbędny w tym sterowaniu triakiem, gdyż triak wymaga dość dużego prądu zapłonowego. Diak blokuje i dlatego nie pozwala na przepływ prądu, dopóki napięcie na elemencie nie wzrośnie do około 25 V. Kondensator C1 może zatem najpierw naładować się do około 25 V, zanim diak się otworzy. W tym momencie w kondensatorze jest wystarczający ładunek, aby dostarczyć duży prąd zapłonowy niezbędny dla triaka.

Ponadto należy pamiętać, że rezystory ładujące są teraz podłączone między obciążeniem a triakiem. Jest to absolutnie konieczne przy regulacji fazowej z zastosowaniem diaka!

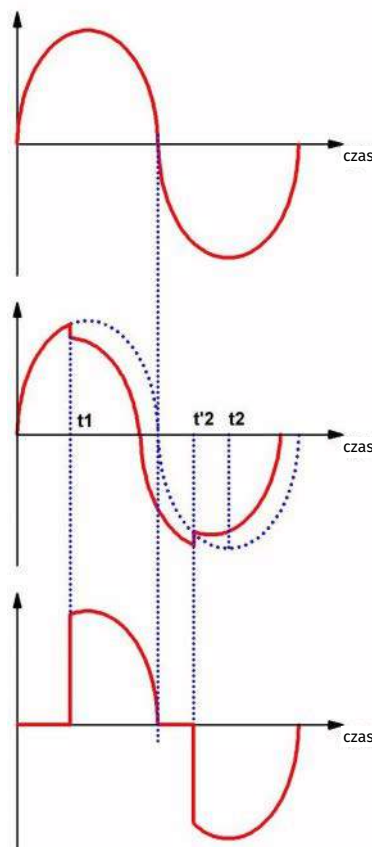
Zjawisko „histerezy ściemniacza”

Omawiany schemat działa dobrze i jest właściwie standardowym obwodem we wszystkich tanich ściemniaczach światła, ale jest jeszcze wiele do poprawienia. Do tej pory zakładano, że dodanie diaka nie ma wpływu na napięcie na kondensatorze. Jednak tak właśnie jest i powoduje to zjawisko zwane „histerezą ściemniacza”. Zjawisko histerezy wyjaśniono graficznie za pomocą wykresów pokazanych na **rysunku 7**. Górny wykres pokazuje napięcie na kondensatorze, gdy wartość rezystorów R1+R2 jest duża. Kondensator nie może wówczas naładować się do napięcia przebicia diaka. Napięcie kondensatora jest czysto sinusoidalne.

Drugi wykres pokazuje sytuację, gdy zmniejszysz wartość rezystancji tak bardzo, że kondensator będzie mógł naładować się tylko do napięcia zapłonu diaka. Następuje to w chwili t1. Zadziałanie



6. Powszechnie stosowany układ regulacji fazowej z triakiem



7. „Histereza ściemniacza” zaprezentowana graficznie

diaka powoduje przepływ prądu. Prąd ten może być dostarczony tylko przez kondensator. W rezultacie napięcie kondensatora nagle nieznacznie spada.

Liniami przerywanymi pokazano, co by się stało, gdyby spadek napięcia nie nastąpił. W następnej połowie cyklu napięcia sieciowego napięcie kondensatora w chwili t2 ponownie stałoby się równe napięciu zadziałania diaka. Jednak ze względu na zmniejszenie napięcia w wyniku przepływu prądu przez diak, ten drugi moment nastąpi nieco wcześniej, a mianowicie w chwili t'2.

Konsekwencje

Co to oznacza w praktyce? Bez opisanego efektu obciążenie otrzymałoby napięcie sieciowe w pierwszej połowie okresu w chwili t1 i w drugiej połowie okresu w chwili t2. Oba punkty są przesunięte w czasie o tę samą wartość od momentu przejścia napięcia sieciowego przez zero. W obu półokresach występowałoby zatem identyczne przesunięcie fazowe pomiędzy przejściem sieci przez zero a włączeniem obciążenia.

Załóżmy, że obciążeniem jest żarówka. Lampa będzie delikatnie świecić. Ten nagły spadek napięcia na kondensatorze spowoduje wcześniejsze zaświecenie lampy w drugiej połowie okresu. W praktyce lampa nie zapala się łagodnie, lecz nagle, z dość dużą intensywnością. Zjawisko to nazywane jest „histerezą ściemniacza”. Każdy, kto kiedykolwiek instalował ściemniacz oświetlenia, wie z praktyki, że większość tanich ściemniaczy światła ma taki efekt.

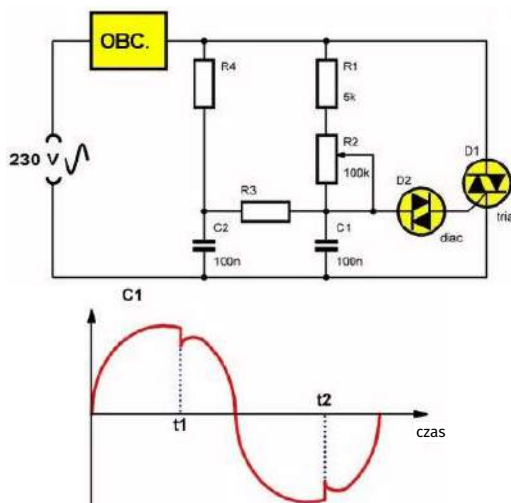
Kompensacja histerezy

Zjawisko histerezy można skompensować w dość prosty sposób. Jak to zrobić pokazano na **rysunku 8**. Istnieją teraz dwa obwody RC, które są połączone ze sobą za pomocą rezystora R3. Kiedy napięcie na C1 wzrośnie do napięcia przebicia diaka i diak zacznie przewodzić prąd z kondensatora, kondensator C2 przyjdzie na ratunek swojemu odpowiednikowi C1. Spadek napięcia na C1 jest szybko uzupełniany prądem dostarczanym z C2 przez R4. Pozostaje teraz tylko tak obliczyć składowe, aby w krytycznym momencie regulacji, czyli w momencie, gdy napięcie na C1 zrówna się z napięciem przebicia diaka, napięcie na C2 było na tyle duże, aby skompensować redukcję napięcia na C1.

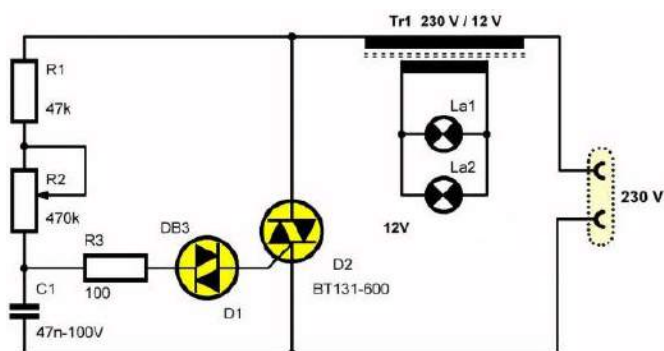
Regulacja fazowa przy niskich napięciach

Czy można stosować również do lamp halogenowych 12 V?

Wszystkie omówione do tej pory obwody pracują z napięciem sieciowym 230 V AC i nadają się do ściemniania żarówek 230 V i reflektorów halogenowych. Jednak prawdopodobnie nadal masz też małe żarówki halogenowe 12 V, które często są używane pod blatami i lampami biurkowymi. Czy można je przyciemnić za pomocą regulacji fazowej?



8. Kompensacja histerezy ściemniacza



9. Schemat ściemniacza do lamp halogenowych 12 V

Opisanego układu nie można używać jako źródła zasilania o napięciu przemiennym 12 V lub 24 V. Można jednak zastosować klasyczny transformator z uzwojeniem pierwotnym na 230 V i wtórnym na 12 V lub 24 V. Wystarczy podłączyć lampy halogenowe do uzwojenia wtórnego i przyciemnić uzwojenie 230 V.

Praktyczny obwód

Schemat na rysunku 9 pokazuje praktycznie przydatny obwód, którego można użyć do przyciemnienia lamp halogenowych 12 V. Jako triak zastosowano BT131-600 firmy NXP (dawniej Philips Semiconductors). Ten półprzewodnik może przełączać napięcie 600 V i ma maksymalną obciążalność prądową 1 A. Triak jest dostarczany w standardowej plastikowej obudowie TO-92. Możesz użyć dowolnego dostępnego diaka. Szybka wycieczka po sklepach z elektroniką pokazuje, że prawdopodobnie, podobnie jak my, skończysz z DB3 firmy STMicroelectronics.

Jedyny kondensator musi mieć napięcie robocze 100 V. Potencjometr jest typu liniowego i można wybierać pomiędzy 470 kΩ i 1 MΩ. Za pomocą pierwszej wartości można ustawić punkt początkowy regulacji – aż żarówka zacznie się delikatnie ściemniać. Za pomocą drugiej wartości możesz zmniejszać jasność, aż do momentu, w którym nie będzie już więcej światła. W systemie wymagane jest umieszczenie wyłącznika/wyłącznika. Dzięki temu możesz mieć pewność, że po wyłączeniu lampy nie zużywa żadnej energii, ponieważ nawet po wyregulowaniu potencjometrem do zaniku światła nadal płynie niewielki prąd.

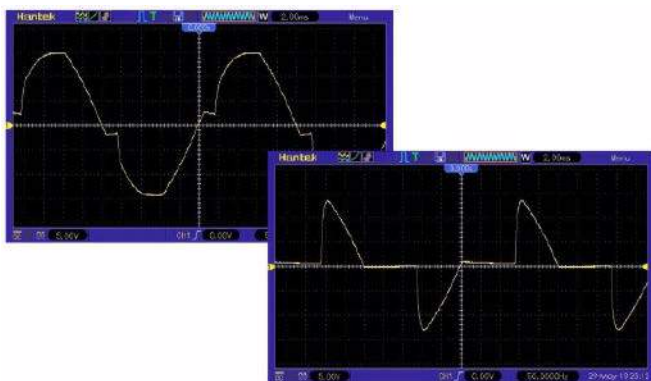
Pomiary

Na rysunku 10 widać, jak wygląda napięcie wtórne transformatora, jeśli regulujemy uzwojenie 230 V w opisany sposób. Oscylogramy te rejestrowano przy obciążeniu wtórnym dwiema lampami 12 V G4 o mocy 20 W.

Kontrola fazy za pomocą napięcia sterującego

Napięcie sterujące zamiast potencjometru

Czasami bardzo niewygodne jest, aby potencjometr, za pomocą którego regulujesz moc, był bezpośrednio podłączony do napięcia



10. Napięcie wtórne 12 V za transformatorem

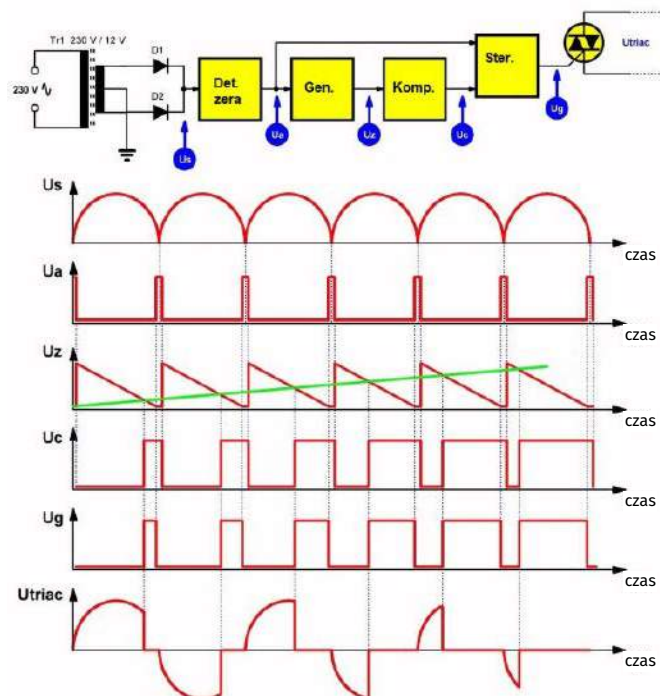
sieciowego. Przy dużej liczbie elementów sterujących mocą możesz potrzebować możliwości regulowania mocy za pomocą napięcia sterującego, na przykład napięcia stałego od 0 V do 10 V. Przy 0 V obciążenie może nie pobierać żadnej energii, przy 10 V obciążenie musi być w pełni podłączone do sieci i pobierać maksymalną moc. Takie systemy istnieją i nawet nie są tak trudne do samodzielnego zaprojektowania. Chociaż istnieją różne systemy sterowania triakiem z napięcia stałego, opisana tutaj metoda jest najbardziej niezawodna.

Zasada działania

Działanie układu wynika ze schematu blokowego pokazanego na poniższym rysunku. Impuls U_a jest wyprowadzany z napięcia sieciowego 230 V poprzez transformator (może to być transformator mocy) i pojawia się w okolicach przejścia przez zero napięcia sieci. Impuls ten jest używany do uruchomienia generatora przebiegu piłokształtnego, który generuje ząb o ujemnym nachyleniu. Na początku cyklu napięcie jest zatem maksymalne i będzie spadać liniowo do zera. Dzięki impulsom ząb piłokształtny U_z jest synchronizowany z półokresami napięcia sieciowego.

Ten ząb piłokształtny U_z jest porównywany w komparatorze z napięciem sterującym U_{ster} w zakresie od 0 V do +10 V. Wynikiem jest impuls U_c w chwili, gdy napięcie sterujące jest większe niż napięcie piłokształtne. Napięcie wyjściowe komparatora U_c jest podawane wraz z pierwszym impulsem U_a do bramki, co zapewnia, że napięcie wyjściowe U_g będzie dodatnie, gdy U_c stanie się dodatnie i powróci do zera, gdy pojawi się impuls U_a . „Ogon” jest jakby odcięty od impulsów U_c .

Jeśli porównać impulsy U_g z napięciem sieciowym, można zauważyć, że zbocze opadające impulsów występuje tuż przed przejściem napięcia sieciowego przez zero. Zbocze impulsu wyzwalającego przypada gdzieś w okresie i jest to określone przez wielkość napięcia sterującego U_{ster} . Im większe napięcie, tym wcześniej w okresie pojawia się impuls. Im mniejsze napięcie, tym bliżej następnego przejścia przez zero pojawia się impuls. Jest jasne, że za pomocą tego impulsu możesz sterować triakiem. Wykres napięcia triaka U_{triac} pokazuje, że im szerszy impuls, tym mniejsze napięcie pozostaje na triaku i tym więcej napięcia, a tym samym mocy, jest dostępne dla obciążenia.



11. Schemat blokowy sterowania fazowego za pomocą napięcia stałego

Zalety systemu

Ta modulacja szerokości impulsu jest najbardziej niezawodną metodą sterowania regulacją fazową dla napięcia przemiennego sieci z następujących powodów:

- Przez cały okres przewodzenia triaka do bramki przesyłany jest prąd zapłonowy. Dlatego triak nigdy nie może nieoczekiwanie wyłączyć się z powodu napięć zakłócających i nadal będzie przewodził, nawet przy bardzo małych obciążeniach.
- Prąd bramki zanika tuż przed przekroczeniem przez falę sinusoidalną zera. Dlatego też nie ma absolutnie żadnego niebezpieczeństwa, że triak „przeskoczy” – przypadkowo pozostanie w przewodzeniu, ponieważ na bramce wokół przejścia napięcia przez zero powstaje sygnał załączający.
- Impulsy zapłonowe bramki mogą być przesyłane w separacji galwanicznej z jednego obwodu do drugiego bez żadnych problemów za pomocą sprzęgacza optycznego.
- Zanim system synchronizacji ulegnie dezorientacji, muszą pojawić się bardzo duże impulsy zakłócające na napięciu sieciowym.

Od schematu blokowego do praktyki

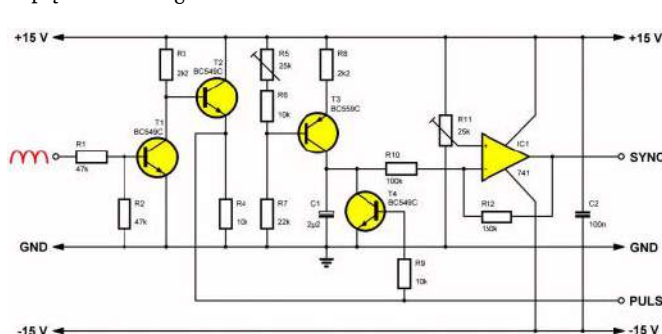
Sterowanie triakiem za pomocą napięcia stałego występuje w praktyce projektowej tak często, że uznałem, że dobrym pomysłem będzie skonkretyzowanie schematu blokowego z **rysunku 11**. Obwód zaproponowany na trzech poniższych schematach był używany dziesiątki razy w ściemniaczach o mocy 2 kW do reflektorów teatralnych i gwarantuje, że będzie działał.

Generator zębaty

Zsynchronizowany z siecią generator piłokształtny jest najtrudniejszą częścią obwodu. Większość generatorów przebiegów wytwarza ząb, który zaczyna się od 0 V, a następnie wzrasta liniowo do pewnej wartości maksymalnej. Jednak ten generator piłokształtny zaczyna przebieg od pewnej wartości dodatniej, a następnie spada liniowo do 0 V. Schemat generatora pokazano na **rysunku 12**. Napięcie U_s pochodzące z transformatora sieciowego przykładane jest do bazy tranzystora T1. Następnie dzielnik napięcia R1/R2 zapewnia, że napięcie baza/emiter staje się mniejsze niż napięcie przewodzenia 0,65 V. W związku z tym na kolektorze tego tranzystora wokół przejścia sieci zasilającej przez zero powstaje wąski dodatni impuls U_a . Napięcie to jest przykładane do wtórnika emiterowego T2 i ten sam sygnał U_a znajduje się na rezystorze R4, ale teraz jest w stanie dostarczyć prądu niezbędnego do napędzania wszystkich obwodów.

Tranzystor T3 jest podłączony jako proste źródło prądowe. Prąd stały dostarczany przez ten półprzewodnik jest określony przez położenie suwaka potencjometru regulacyjnego R5. Prąd stały dostarczany przez kolektor ładuje kondensator C1. Dlatego w tej części powstaje liniowo rosnące napięcie.

Tranzystor T4 zwiera ten kondensator za każdym razem, gdy pojawia się impuls U_a . W rezultacie na kondensatorze tworzy się rosnący ząb piłokształtny, który jest ładnie zsynchronizowany z półokresami napięcia sieciowego.



12. Praktyczny obwód generatora piłokształtnego

Omówienie podstaw działania regulatora w poprzedniej sekcji pokazało, że potrzebny jest ząb piłokształtny opadający a nie narastający. Stąd wzmacniacz operacyjny IC1. Stopień ten jest podłączony jako odwracający wzmacniacz różnicowy. Przebieg jest podawany na wejście odwracające. Dzielnik rezystancyjny R10/R12 ustawia wzmacniacz operacyjny jako wzmacniacz odwracający. Gdyby wejście nieodwracające było podłączone do masy, to na wyjściu pojawiłby się ząb piłokształtny, który jest odwrócony w stosunku do napięcia wejściowego, ale opadający i wahałby się od 0 V do około -8 V. Sygnał ten trzeba niejako „wzmocnić” o 8 V, aby zrównał się z sygnałem wymaganym.

Potencjometr R11 umożliwia ustawienie dodatniego wejścia wzmacniacza operacyjnego na określone stałe napięcie dodatnie. To właśnie to napięcie zapewnia, że przebieg piłokształtny na wyjściu wzmacniacza operacyjnego jest taki, jak U_z na wykresie. Obracając potencjometr można regulować napięcie wyjściowe tak, aby najniższy poziom był dokładnie taki sam jak oś 0 V.

Przebieg piłokształtny, sygnał SYNC i dodatni impuls przechodzący przez zero, PULS, są używane w modulatorze szerokości impulsu, który generuje napięcie sterujące dla triaka.

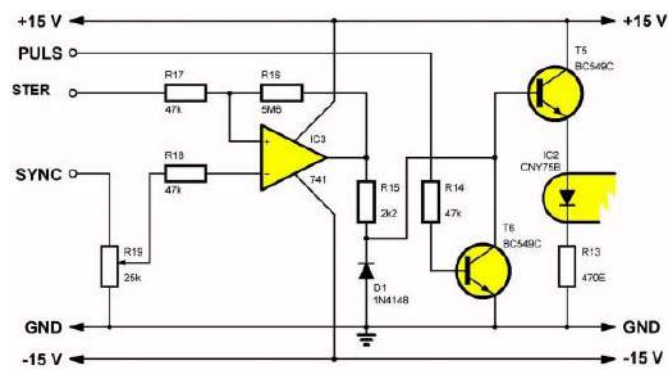
Modulator szerokości impulsu

Tak naprawdę, co widać na **rysunku 13**, jest to wzmacniacz operacyjny IC3 i przetwornik napięcie-prąd podłączony w charakterze komparatora. Ząb piłokształtny SYNC jest podłączony do wejścia odwracającego wzmacniacza operacyjnego IC3 za pomocą potencjometru regulacyjnego R19. Napięcie sterujące STER z zakresu od 0 V do +10 V podawane jest na wejście nieodwracające. Komparator dostarcza impuls dodatni, gdy napięcie sterujące jest większe niż poziom napięcia przebiegu piłokształtnego. Powstaje dodatni impuls na wyjściu, którego krawędź natarcia przesuwają się w przód i w tył w miarę zmiany wielkości napięcia sterującego.

Komparator IC3 zapewnia małą histerezę poprzez zainstalowanie dużego rezystora sprzężenia zwrotnego R16 pomiędzy wyjściem a wejściem dodatnim. Dzięki tej dodatkowej rezystancji na punkt przełączania komparatora nie będą miały wpływu żadne tąpnięcia ani szumy napięcia sterującego.

Wzmacniacz operacyjny IC3 jest zasilany symetrycznie. Impuls wyjściowy przeskakuje zatem tam i z powrotem pomiędzy -15 V, a +15 V. Jednak przy napięciu ujemnym nie można nic zrobić, dlatego jest ono zwierane do masy przez diodę D1. Impuls o modulowanej szerokości trafia do wtórnika emiterowego T5. Gdy pojawi się impuls, tranzystor włącza się i wysyła prąd o natężeniu około 25 mA przez podczerwoną diodę LED transoptora IC2.

Teraz musisz upewnić się, że impuls zapłonowy triaka zniknął w okolicach przejścia przez zero sinusa sieci. Stąd tranzystor T6, który zwiera sygnał na bazie T5 do masy, gdy pojawi się impuls przejścia przez zero na linii PULS.



13. Praktyczny obwód sterowania transoptorem

Obwód triaka

Schemat obwodu z triakiem pokazano na **rysunku 14**. Fototranzystor IC2 jest połączony szeregowo z rezystorem R20 pomiędzy niezależnym napięciem zasilania +15 V a punktem neutralnym (N) sieci. Kiedy dioda LED transoptora emituje światło (występuje impuls o modulowanej szerokości), fototranzystor zaczyna przewodzić, a na rezystorze R20 pojawia się pełne +15 V. Sygnał ten jest podawany do wtórnika emitera T7, a następnie wysyła prąd o natężeniu około 35 mA do bramki triaka D5 poprzez rezystor szeregowy R22.

Bardzo ważna uwaga do zasilania obwodu triaka – można używać TYLKO zasilania +15 V pokazanego na rysunku 14. Przecież ten zasilacz jest podłączony bezpośrednio do napięcia sieciowego jednym biegunem. Dotykanie zasilacza i zasilanego obwodu może zatem zagrażać życiu, dlatego należy go używać wyłącznie w tym obwodzie i nigdzie indziej.

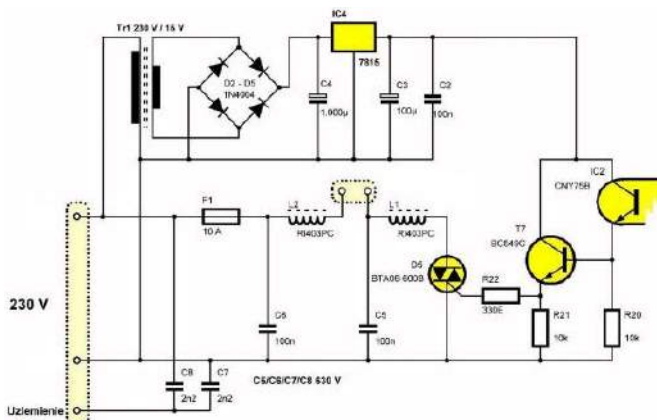
Tłumienie zakłóceń jest absolutnie konieczne

Skąd biorą się zakłócenia?

Artykuł na temat działania regulacji fazowej nie jest kompletny bez zwrócenia uwagi na niezbędne techniki tłumienia zakłóceń. W obwodach, które współpracują z regulacją fazową, z definicji powstają duże, szybko narastające prądy. Stronie zbocza tych impulsów prądu zawierają wiele wyższych harmonicznych, dlatego w obwodzie, który w zasadzie działa tylko przy częstotliwości 50 Hz, nadal można wykryć wiele wysokich częstotliwości! Te sygnały o wysokiej częstotliwości przepływają głównie przez przewody łączące pomiędzy tyrystorem lub triakiem a obciążeniem. W praktyce mają one często dziesiątki metrów długości, więc przewody te są naprawdę doskonałymi antenami nadawczych i obficie rozpraszają w otoczeniu fale elektromagnetyczne o wysokiej częstotliwości.

Bądź odpowiedzialny!

Na ostatnim rysunku widać kilka części, które nie mają nic wspólnego z działaniem triaka, ale wszystkie mają związek z jak najlepszym



14. Praktyczny obwód sterowania triakiem

tłumieniem sygnałów zakłócających o wysokiej częstotliwości: L1, L2, C5, C6, C7 i C8. Te obwody LC tworzą filtry dolnoprzepustowe, które tłumią sygnały zakłócające o wysokiej częstotliwości o dziesiątki dB i umożliwiają przejście częstotliwości 50 Hz sieci praktycznie bez zakłóceń.

Sześć drogich komponentów do tłumienia zakłóceń? Nie za dużo tego dobrego? Nie, absolutnie nie! Dobrze znany układ, który można znaleźć w tanich ściemniaczach i który składa się z jednego kondensatora i jednej małej pierścieniowej cewki tłumiącej, jest naprawdę niewystarczający.

Zakłócanie zgodnie z zasadami

Dobre tłumienie zakłóceń jest sztuką i dlatego poświęcimy tej sztuce jeden z kolejnych cykli. Tymczasem zdecydowanie zalecamy uważne przeczytanie tego artykułu, gdy rozpoczynasz prace z układami regulacji fazowej. ■

Jos Verstraten



Fazowa regulacja mocy

Rozwiązanie znajdziesz na www.elportal.pl/quizy

Kiedy wyłączy się tyrystor?

- ☐ Gdy zaniknie impuls sterujący na jego bramkę
- ☐ Gdy prąd spadnie poniżej wartości podtrzymania
- ☐ Po określonym czasie

W jakim zakresie mocy możliwa jest regulacja fazowa:

- ☐ 0...50%
- ☐ 0...100%
- ☐ 1%...100%

Z jaką częstotliwością napięcie sieci energetycznej w Polsce przechodzi przez zero?

- ☐ 60 razy na sekundę
- ☐ 100 razy na sekundę
- ☐ 50 razy na sekundę

Fazowa regulacja mocy polega na:

- ☐ Odcięciu czoła każdego półokresu przebiegu napięcia sieci
- ☐ Odcięciu ogona każdego półokresu przebiegu napięcia sieci
- ☐ Odcięciu wierzchołka amplitudy każdego półokresu przebiegu napięcia sieci

Czym różni się triak od tyrystora?

- ☐ Ogranicza przepływ prądu z sieci
- ☐ Może przewodzić zarówno przy dodatniej, jak i ujemnej półokresie sinusoidy napięcia sieci
- ☐ Działa jak diodowy mostek prostowniczy

Jak skompensować zjawisko histerezy ściemniacza?

- ☐ Poprzez dodatkowy obwód RC
- ☐ Poprzez zastosowanie diaka DB3
- ☐ Poprzez zmniejszenie prądu obciążenia

W układzie regulacji fazowej terowanym napięciem potrzebny jest generator sygnału:

- ☐ Sinusoidalnego
- ☐ Trójkątnego
- ☐ Piłokształtnego z opadającym zboczem

Przyczyna zakłóceń w obwodach z regulacją fazową to:

- ☐ Strome zbocza impulsów prądu
- ☐ Wytwarzanie sygnału piłokształtnego
- ☐ Zbyt długie przewody połączeniowe

Jak przeciwdziała się powstawaniu zakłóceń w obwodach z regulacją fazową?

- ☐ Stosuje się transformator na wyjściu regulatora
- ☐ Stosuje się dolnoprzepustowe filtry LC w obwodzie triaka
- ☐ Stosuje się transoptor w obwodzie bramki triaka

Jak działa diak?

- ☐ Powoduje przesunięcie fazowe do 90°
- ☐ Blokuje przepływ prądu, aż do osiągnięcia określonego napięcia
- ☐ Generuje dodatkowy impuls prądu na bramkę triaka

Punkt pracy tranzystora bipolarnego



Punkt pracy tranzystora bipolarnego oznacza ustalenie prądów stałych przepływających przez tranzystor w przypadku braku sygnału na wejściu obwodu.

Pojęcia i definicje

Na czym polega ustawienie punktu pracy tranzystora bipolarnego?

Ustawiając punkt pracy tranzystora określasz wartości:

- prądu bazy,
- prądu kolektora,
- prądu emitera.

W większości przypadków dokonuje się tego poprzez obliczenie rezystancji w obwodach bazy, kolektora i emitera. Po dołączeniu tych rezystorów i podaniu napięcia zasilania tranzystor bipolarny, w wyniku swoich fizycznych cech, zachowuje się podobnie jak rezystor zmienny i powoduje przepływ pożądanych prądów.

Ogólna zasada konfiguracji

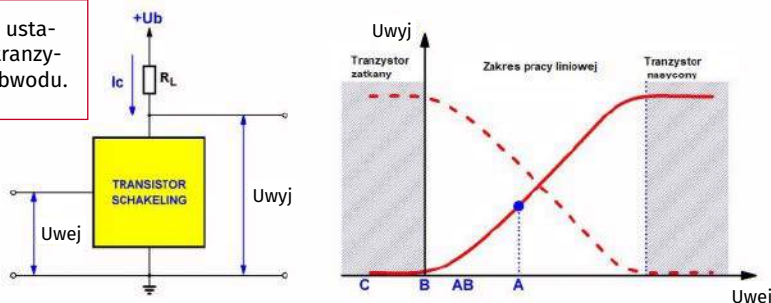
Ogólną zasadę konfiguracji tranzystora bipolarnego pokazano na **rysunku 1**. Tranzystor wraz z elementami nastawczymi jest reprezentowany przez jeden blok. Jest on połączony szeregowo z rezystorem obciążenia R_L , który jest podłączony pomiędzy kolektorem a zasilaczem. W rezultacie prąd polaryzacji I_c popłynie przez połączenie szeregowo rezystora R_L i tranzystora. Ten prąd polaryzacji nazywany jest także „prądem spoczynkowym”.

Konsekwencją prądu I_c jest spadek napięcia na rezystorze R_L i tranzystorze. Napięcie na styku obu elementów nazywa się napięciem wyjściowym U_{out} . Wykres po prawej stronie rysunku pokazuje zależność pomiędzy prądem polaryzacji I_c , napięciem wyjściowym U_{out} i napięciem wejściowym U_{in} .

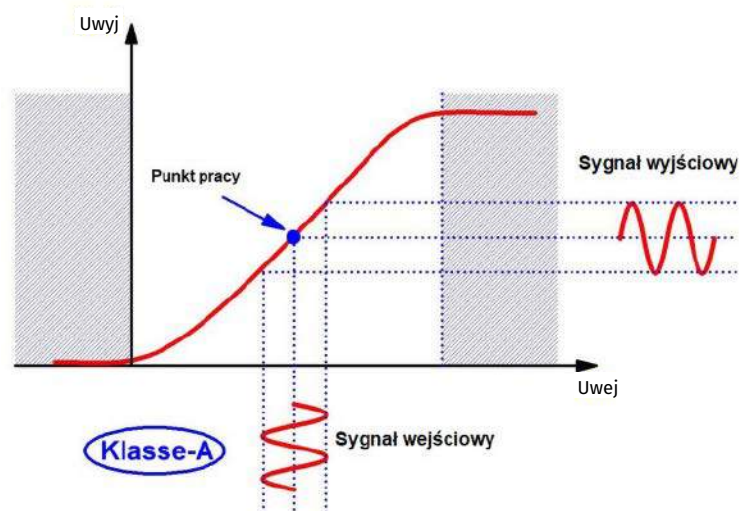
Zależność pomiędzy prądem i napięciem jest pokazana linią ciągłą, stosunek pomiędzy dwoma napięciami linią przerywaną. Jeżeli napięcie U_{in} jest mniejsze od określonej wartości, przez obwód nie będzie przepływał żaden prąd. Obwód zachowuje się wtedy jak przerwa. Jest to pokazane w lewym zaciemnionym obszarze. Napięcie wyjściowe jest wówczas maksymalne i równe wartości napięcia zasilania. Mówi się, że tranzystor jest „zatkany”. Jeśli napięcie wejściowe przekroczy określoną wartość, prąd polaryzacji I_c stanie się równy maksymalnej wartości, która może przepływać przez obwód. Wartość ta jest równa $+U_b / R_L$, czyli napięcie zasilania podzielone przez rezystancję obciążenia. W tym momencie napięcie wyjściowe wynosi zero, a tranzystor zachowuje się jak zwarcie. Jest to reprezentowane przez prawy zaciemniony obszar na wykresie. Mówi się, że tranzystor jest „w stanie nasycenia”.

Zakres pracy

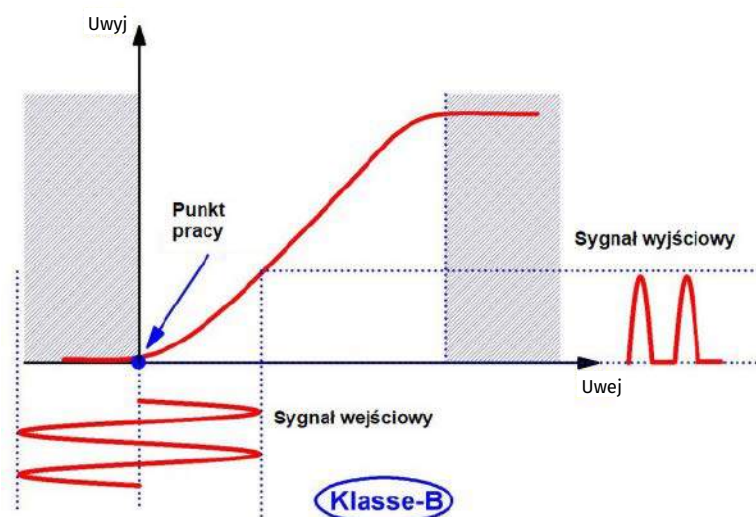
Gdy tranzystor pracuje, tj. gdy obecny jest sygnał wejściowy, ten sygnał wejściowy zapewnia wzrost i spadek napięcia



1. Ogólna idea działania wzmacniacza na tranzystorze bipolarnym i zobrazowanie położenia punktu pracy



2. Ustawienie pracy tranzystora bipolarnego w klasie A



3. Ustawienie pracy tranzystora bipolarnego w klasie B

polaryzacji Uout. Wartość tego napięcia przebiega zatem w tę i z powrotem po poziomej osi wykresu. W wyniku tej modulacji natężenie prądu Ic również będzie się zwiększać i zmniejszać. Istnieją dwa podstawowe systemy określające obszar roboczy, w którym będzie płynął prąd:

- **Regulacja cyfrowa.** Jeśli tranzystor ma działać jako element cyfrowy, należy upewnić się, że półprzewodnik jest regulowany tylko w zacięzionych obszarach. Prąd jest wtedy maksymalny lub zerowy. Napięcie wyjściowe staje się zatem zerowe, co cyfrowo odpowiada „L”, lub maksimum, co cyfrowo odpowiada „H”.
- **Ustawienie liniowe.** Jeżeli tranzystor ma pracować jako wzmacniacz, należy zwrócić uwagę, aby półprzewodnik był ustawiony pomiędzy zacięzionymi obszarami. Prąd nie może nigdy kończyć się w żadnym z zacięzionych obszarów. Tranzystor działa wówczas jako rezystor zmienny i dzięki temu działaniu na wyjściu pojawi się napięcie proporcjonalne do napięcia wejściowego.

Klasy ustawień

Na arenie międzynarodowej uzgodniono kodowanie literowe opisujące obszar, w którym ustawiony jest tranzystor bipolarny. To kodowanie pokazano również na wcześniejszym wykresie. Wyróżnia się następujące ustawienia – klasy:

- klasa A,
- klasa B,
- klasa AB,
- klasa C,
- klasa D.

Ustawienie klasy A

Przy tym ustawieniu tranzystor pracuje całkowicie w niezacięzionej części wykresu, dlatego zawsze działa jako rezystor zmienny. Wartość zadana A znajduje się w środku niezacięzionego obszaru. Takie ustawienie jest domyślnym ustawieniem dla przedwzmacniacza audio, regulatorów tonu, mikserów i tym podobnych. To ustawienie pokazano schematycznie na **rysunku 2**.

Ustawienie klasy B

Przy tym ustawieniu tranzystor jest ustawiony tak, że prąd polaryzacji wynosi dokładnie zero. Wartość zadana B znajduje się dokładnie w miejscu wyłączenia tranzystora. To ustawienie jest często stosowane we wzmacniaczach mocy, gdzie dwa tranzystory są połączone w specjalny sposób tak, że każdy półprzewodnik odpowiada tylko za połowę sygnału. To ustawienie jest schematycznie zaprezentowane na **rysunku 3**.

Ustawienie klasy AB

To ustawienie mieści się pomiędzy wartościami zadanymi A i B. W stanie spoczynku przez tranzystor przepływa dość mały prąd polaryzacji. Ale jeśli nałożysz sygnał na prąd spoczynkowy, tranzystor czasami się wyłączy. To ustawienie jest również często używane w przypadku wzmacniaczy mocy.

Ustawienie klasy C

Przy tym ustawieniu tranzystor jest całkowicie ustawiony w obszarze blokującym. Jeśli modulujesz prąd spoczynkowy małym sygnałem, tranzystor nadal pozostanie wyłączony. Tylko jeśli modulujesz prąd spoczynkowy bardzo dużym sygnałem, dodatnie szczyty tego prądu sprawią, że tranzystor będzie przewodził. Wartość zadana C znajduje się zatem całkowicie w lewej zacięzionej strefie. To ustawienie jest często używane we wzmacniaczach HF, gdzie duże zniekształcenia sygnału związane z tym ustawieniem nie mają znaczenia.

Ustawienie klasy D

Nazywa się to czasem „ustawieniem cyfrowym” tranzystora. W tej klasie nie ma prądu spoczynkowego, a prąd sygnałowy ma tylko dwie wartości. Zapewniają one,

że tranzystor jest albo zablokowany, albo w stanie nasycenia. Napięcie wyjściowe wynosi zatem zero lub maksimum, co odpowiada sygnałom cyfrowym „L” i „H”. To ustawienie stosowane jest np. w cyfrowych wzmacniaczach mocy, gdzie sygnał dźwiękowy prezentowany jest w formie cyfrowej (sygnał PWM) i w takiej postaci trafia do głośnika. Pomiedzy wyjście wzmacniacza a głośnik podłączony jest filtr dolnoprzepustowy, który zapewnia konwersję sygnałów cyfrowych z powrotem na sygnał analogowy.

Podstawowe obwody tranzystorów bipolarnych

Wprowadzenie

Mówiąc o stopniu tranzystorowym często przyjmuje się za pewnik, że sygnał wejściowy jest podłączony do bazy, a sygnał wyjściowy pobierany jest z kolektora. To rzeczywiście najczęściej używany obwód. Można jednak ustawić tranzystor w trzech różnych podstawowych obwodach:

- układ wspólnego emitera,
- układ wspólnej bazy,
- układ wspólnego kolektora.

Każdy z tych trzech obwodów ma swoje specyficzne zastosowania i właściwości.

Układ wspólnego emitera

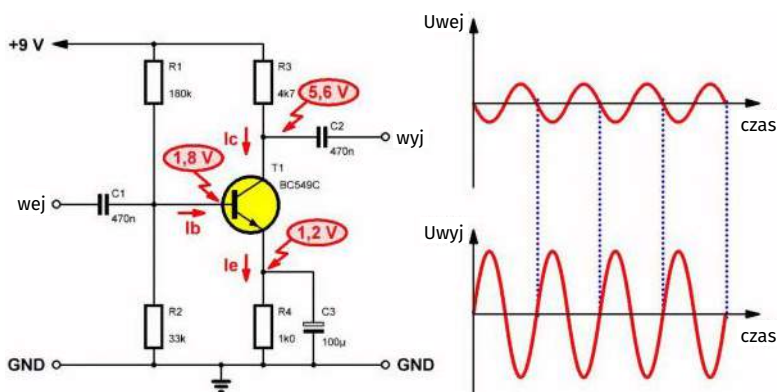
Schemat

Jest to, jak pokazano na **rysunku 4**, najczęściej używany obwód do ustawiania tranzystora. Emiter jest połączony z masą poprzez rezystor R4 i kondensator C3. Duży kondensator tworzy zwarcie dla napięć przemiennych i zapewnia połączenie emitera sygnału z masą. W końcu duży kondensator elektrolityczny ma impedancję (rezystancję przy prądzie przemiennym), która jest praktycznie nieskończenie wysoka dla napięcia stałego i praktycznie zerowa dla napięcia przemiennego. Stąd nazwa „układ wspólnego emitera”. Bazę ustala się za pomocą dzielnika rezystancji R1 i R2. Ten dzielnik rezystancji określa napięcie na bazie, a tym samym prąd spoczynkowy Ic płynący przez tranzystor. Kolektor jest podłączony do dodatniego napięcia zasilania poprzez rezystor obciążenia R3. Należy eksperymentalnie dobrać dwa rezystory do bazy tak, aby w stanie spoczynku na kolektorze było około połowy napięcia zasilania.

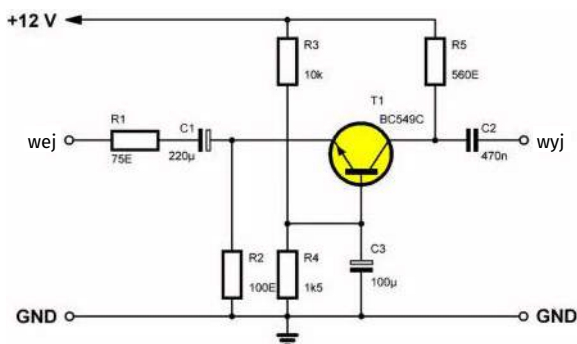
Sygnał wejściowy jest podawany do bazy poprzez kondensator izolujący C1. Sygnał napięcia przemiennego na wejściu moduluje napięcie polaryzacji bazy, a tym samym prąd bazy. Te zmiany prądu są wzmacniane przez tranzystor, w wyniku czego prąd spoczynkowy Ic również się zmienia. Ten zmienny prąd powoduje powstanie sygnału napięcia wyjściowego na kolektorze.

Właściwości układu wspólnego emitera

Dzięki uziemionemu obwodowi emitera można w jednym stopniu osiągnąć wzmocnienie napięcia od 50 do 2000. Ponieważ obwód



4. Układ ze wspólnym emiterem jest najbardziej znanym układem pracy tranzystora



5. Układ ze wspólną bazą

wzmacnia zarówno prąd wejściowy, jak i napięcie wejściowe, obwód ma duży zysk mocy. Obwód działa odwracająco. Wzrost prądu bazy pod wpływem dodatniego sygnału na wejściu powoduje wzrost prądu kolektora. Powoduje to jednak spadek napięcia na kolektorze. Zależność pomiędzy sygnałem na wejściu i sygnałem na wyjściu pokazano na wykresie na powyższym rysunku. Dodatni półcykl na wejściu powoduje zatem ujemny półcykl na wyjściu.

Impedancja wejściowa obwodu zależy głównie od równoległej wartości zastępczej dwóch rezystorów w bazie. W większości przypadków impedancja ta nie będzie bardzo wysoka. W pokazanym przykładzie impedancja wejściowa jest w każdym przypadku niższa niż 33 kΩ. Stąd kondensator izolujący C1 na wejściu musi mieć dość dużą wartość. Jeśli podasz temu kondensatorowi zbyt małą wartość, impedancja (rezystancja prądu przemiennego) kondensatora dla napięć przemiennych utworzy dzielnik napięcia z rezystorami bazowymi. Osłabi to sygnały o niskiej częstotliwości.

Szerokość pasma tego obwodu jest dość niska. Wynika to z katastrofalnego wpływu pojemności pasożytniczych w tranzystorze na działanie obwodu przy wysokiej częstotliwości (patrz dalej). Dlatego układ wspólnego emitera można znaleźć głównie w obwodach niskiej częstotliwości, takich jak przedwzmacniacze audio, regulatory tonów, miksery itp.

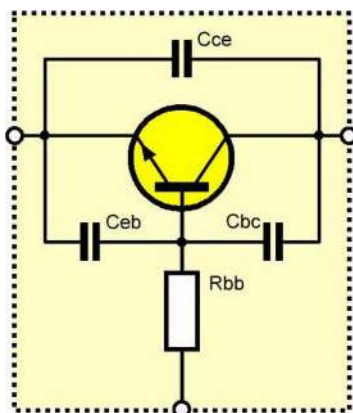
Układ wspólnej bazy

Schemat

W układzie wspólnej bazy, baza jest połączona z masą. Jak pokazano na rysunku 5, osiąga się to poprzez podłączenie za pomocą dużego kondensatora C3. Napięcie bazy ustala się za pomocą dzielnika napięcia R3 i R4. Sygnał wejściowy jest teraz doprowadzany do emitera poprzez kondensator izolujący C1 i rezystor szeregowy R1. Sygnał wyjściowy pobierany jest z kolektora. Ponownie rezystor obciążenia R5 jest połączony szeregowo z kolektorem. Naturalnie emiter musi być podłączony do masy poprzez rezystor R2. Bez tego R2 żaden prąd nie mógłby płynąć przez półprzewodnik.

O pojemnościach pasożytniczych

Możesz zadać sobie pytanie, jaka jest zaleta tego układu. Zaleta ta jest widoczna głównie przy wysokich częstotliwościach. Każdy tranzystor ma pasożytnicze pojemności pomiędzy trzema stykami, co pokazuje rysunek 6. Jednak pomiędzy



6. Rozmieszczenie pasożytniczych komponentów Ceb i Rbb w układzie ze wspólną bazą

połączeniem bazy a pojemnością pomiędzy bazą a emiterem (Ceb) występuje wewnętrzna rezystancja bazy Rbb. W układzie wspólnego emitera te dwa elementy tworzą filtr dolnoprzepustowy, który tłumi wysokie częstotliwości. Dlatego szerokość pasma takiego obwodu jest ograniczona.

W układzie wspólnej bazy wygląda to zupełnie inaczej. Jak widać na rysunku, Rbb i Ceb są teraz podłączone odwrotnie między wejściem a masą. W rezultacie obwód ten nie tworzy teraz filtra dolnoprzepustowego, ale filtr górnoprzepustowy. Pasożytnicza pojemność nie osłabi wysokich częstotliwości sygnału, ale nawet nieznacznie je wzmocni.

Właściwości układu ze wspólną bazą

Jak widać na schemacie, układ ze wspólną bazą ma bardzo niską impedancję wejściową. Przecież między wejściem a masą jest tylko rezystor R1+R2=175 Ω. Dlatego wartość kondensatora izolującego C1 musi być bardzo duża, aby zapobiec tłumieniu niskich częstotliwości. Dlatego niezbędny jest do tego duży kondensator elektrolityczny. Jednak impedancja wyjściowa obwodu jest bardzo duża. Uziemiony obwód podstawowy nie działa odwracająco. Jeśli sygnał wejściowy wzrośnie, napięcie wyjściowe również wzrośnie.

W porównaniu z szerokością pasma B_E układu ze wspólnym emiterem, szerokość pasma układu ze wspólną bazą B_B jest równa:

$$B_B = B_E \cdot \beta$$

gdzie β oznacza wzmocnienie prądowe tranzystora.

Będzie zatem jasne, że wzmocnienie układu ze wspólną bazą pozostaje stałe aż do bardzo wysokich częstotliwości sygnału. Dlatego w obwodach HF często można znaleźć układ ze wspólną bazą. Co więcej, bardzo niska impedancja wejściowa nie powoduje praktycznie żadnych wad. Wręcz przeciwnie, często można dobrze wykorzystać tę niską impedancję, aby prawidłowo podłączyć wejście obwodu do kabla o niskiej impedancji własnej, na przykład 50 Ω lub 75 Ω.

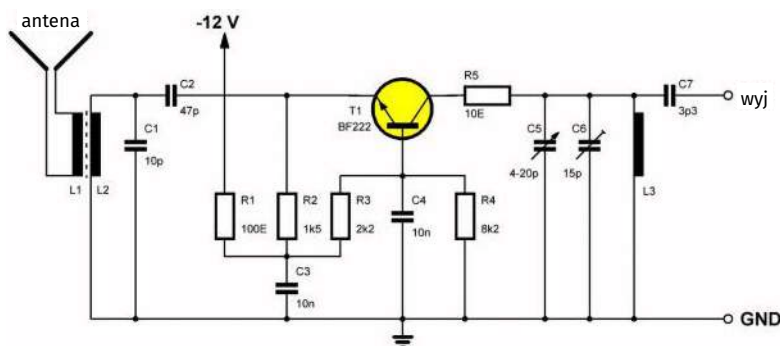
Effekt Millera

Drugim aspektem odgrywającym rolę we wzmocnieniu HF jest tak zwany „efekt Millera”. Pomiędzy bazą a kolektorem występuje pojemność pasożytnicza Cbc. W układzie ze wspólnym emiterem pojemność ta jest połączona między wyjściem a wejściem. Ta pojemność tworzy zatem sprzężenie zwrotne z wyjścia do wejścia. W takich warunkach oddziaływanie tej pojemności na wejściu wzmacniacza będzie pomnożone przez wzmocnienie napięciowe tranzystora. Zatem jeśli stopień ma wzmocnienie napięciowe 1000, a wartość pojemności pasożytniczej wynosi 0,5 pF, to w obwodzie podłączonym do wejścia wydaje się, że między wejściem a masą znajduje się kondensator o wartości 0,5 nF. Nie jest to oczywiście zbyt dobre, jeśli chcesz przetwarzać wysokie częstotliwości za pomocą tego układu.

W układzie ze wspólną bazą pasożytnicza pojemność Cbc jest połączona pomiędzy wyjściem obwodu a bazą połączoną z masą. Nie ma już sprzężenia zwrotnego w przeciwnej fazie pomiędzy wyjściem i wejściem, w związku z czym wpływ tej pojemności na pracę układu jest znacznie mniejszy.

Praktyczny obwód

Poniższy rysunek 7 pokazuje praktyczny przykład układu ze wspólną bazą w tunerze odbiornika FM. Baza jest połączona bezpośrednio z masą za pomocą kondensatora C4, dzięki czemu spełniony jest warunek podstawowy tego obwodu. Sygnał wejściowy pobierany jest z transformatora antenowego L1/L2 poprzez kondensator C2 47 pF. Obciążenie składa się z rezystora 10 Ω R5 i dostrojonego filtra C5, C6, L3. Wzmocniony sygnał jest przekazywany do następnego obwodu poprzez kondensator C7 o pojemności 3,3 pF. Kolektor jest podłączony do masy poprzez rezystor R5 i cewkę L3. Aby tranzystor zadziałał, emiter należy podłączyć do ujemnego napięcia zasilania -12 V prądu stałego.



7. Praktyczny przykład układu ze wspólną bazą

Układ ze wspólnym kolektorem

Schemat

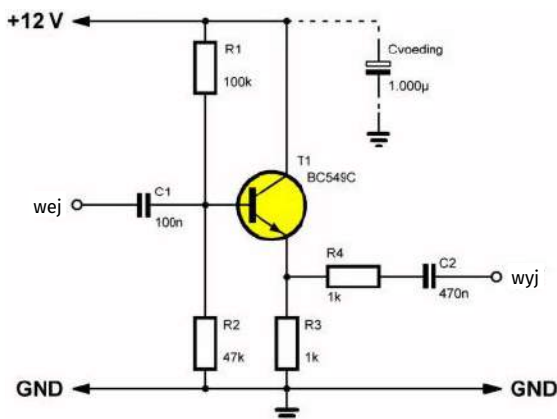
Podstawowy schemat układu ze wspólnym kolektorem został pokazane na **rysunku 8**. Kolektor jest połączony bezpośrednio z napięciem zasilania. Zasilacz zawiera duży kondensator filtrujący, co zapewnia, że na kolektorze nie ma sygnałów napięcia przemiennego. Bazę ustawia się w znany sposób za pomocą rezystorów R1 i R2. Sygnał wejściowy jest dostarczany poprzez kondensator separujący C1. W emiterze rezystor R3 jest podłączony do masy, sygnał wyjściowy pobierany jest z emitera. Rezystor szeregowy R4 nie jest konieczny, ale chroni tranzystor przed nadmiernym obciążeniem. Jeżeli tego rezystora nie ma i przypadkowo zwieramy wyjście do masy, tranzystor (patrząc na poziom sygnału) jest podłączony bezpośrednio pomiędzy zasilaniem a masą. W rezultacie półprzewodnik uległby dość szybkiemu uszkodzeniu z powodu zbyt dużego prądu.

Dobrze znany wtórnik emiterowy

Układ ze wspólnym kolektorem jest również znany jako „wtórnik emiterowy”. Sygnał na emiterze podąża za sygnałem na bazie. Oznacza to, że jedną z głównych cech obwodu kolektora z uziemieniem jest to, że obwód ten nie wytwarza wzmocnienia napięcia! Sygnał na emiterze jest dokładną kopią sygnału wejściowego na bazie. Jednak obwód ma inne bardzo przydatne właściwości.

Właściwości układu ze wspólnym kolektorem

Jedną z najważniejszych właściwości jest bardzo wysoka impedancja wejściowa. Jest to spowodowane ekstremalnym sprzężeniem zwrotnym pomiędzy emiterem a bazą, które powoduje bardzo małe obciążenie źródła. Drugą ważną właściwością jest bardzo niska impedancja wyjściowa. Wartości 50Ω są łatwe do osiągnięcia. Dlatego wtórnik emiterowy jest idealnym obwodem do budowania stopni buforowych. Taki stopień buforowy musi oddzielić źródło, które nie powinno być przeciążane i zapewnia sygnał wyjściowy źródła przy bardzo niskiej impedancji wyjściowej.



8. Schemat obwodu z tranzystorem w układzie wspólnego kolektora

Układ ze wspólnym kolektorem nie jest, jak już napisano, wzmacniaczem napięcia. Wzmacniane są prądy. Ze względu na bardzo wysoką impedancję wejściową, prąd wejściowy jest oczywiście bardzo niski. Jednak ze względu na niską impedancję wyjściową obwód może dostarczać duży prąd do obciążenia. Układ nie działa odwracająco, dzięki czemu sygnał wyjściowy na emiterze jest w fazie z sygnałem wejściowym na bazie.

Obwody z tranzystorem bipolarnym w praktyce

Wprowadzenie

Samo ustawienie tranzystora nie jest trudną pracą.

Można postępować na dwa sposoby:

- sposób „oficjalny” – obliczanie za pomocą kart katalogowych i kalkulatora,
- metoda „zmień i sprawdź” – w której pracujesz wyłącznie eksperymentalnie.

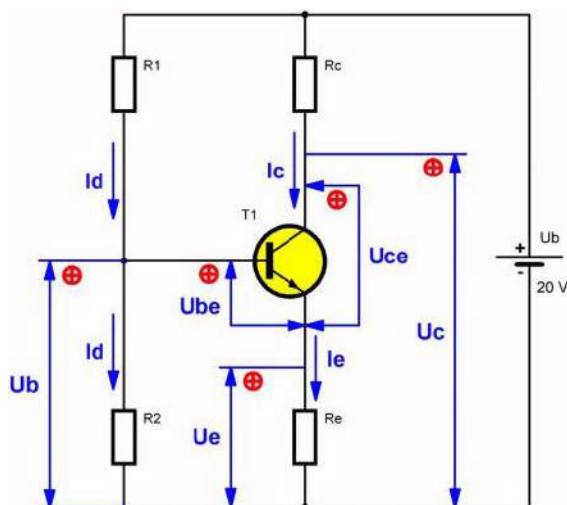
Obliczanie parametrów pracy tranzystora

Problem z tą metodą polega na tym, że potrzebujesz charakterystyki tranzystora. Ale na szczęście możesz to szybko znaleźć za pomocą Google, więc nie stanowi to problemu. Jako przykład omówiono ustawienie stopnia tranzystorowego w klasie A i w układzie wspólnego emitera. Konfigurowanie pozostałych podstawowych obwodów i pozostałych klas odbywa się w zasadzie w podobny sposób. Jednak nigdy nie zapominaj, że obliczenia dotyczące tranzystora są zawsze bardzo przybliżone. Tranzystory tego samego typu mają bardzo duże rozrzuty, a publikowane charakterystyki są niczym więcej niż średnią z dużej liczby tranzystorów. Mogą jednak wystąpić indywidualne odchylenia przekraczające 25%!

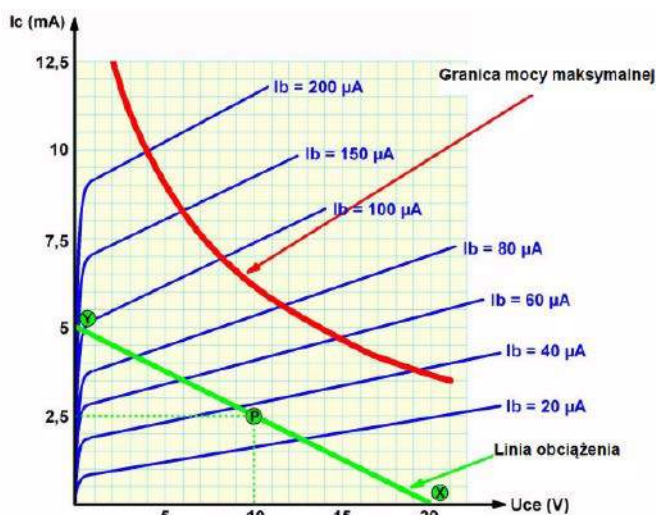
Podstawowy schemat dla przykładu został pokazany na **rysunku 9**. Ten rysunek pokazuje:

- U_c dla napięcia spoczynkowego na kolektorze,
- U_b dla napięcia spoczynkowego na bazie,
- U_e dla napięcia spoczynkowego na emiterze,
- I_c dla prądu spoczynkowego płynącego przez kolektor,
- I_b dla prądu spoczynkowego płynącego przez bazę,
- I_e dla prądu spoczynkowego płynącego przez emiter,
- U_{be} oznacza napięcie przewodzenia tranzystora,
- I_d dla prądu płynącego przez rezystory potencjału bazy.

Napięcie przewodzenia tranzystora wynosi w przybliżeniu 0,6 V dla krzemu i około 0,2 V dla germanu. Z tej wiedzy można zatem



9. Schemat pomocniczy do obliczenia parametrów pracy tranzystora



10. Linia obciążenia narysowana w charakterystyce $I_c=f(U_{ce})$ tranzystora

wywnioskować związek pomiędzy U_b i U_e . Napięcie spoczynkowe emitera jest zawsze o 0,6 V lub 0,2 V niższe niż napięcie spoczynkowe na bazie, przynajmniej w przypadku tranzystorów NPN.

Układ zasilany jest z napięcia stałego U_b o napięciu 20 V.

Wartość zadana

Wartość zadana, zwana P, jest najważniejszą informacją potrzebną do skonfigurowania stopnia tranzystorowego. Punkt ten określa prąd spoczynkowy płynący przez tranzystor i napięcie spoczynkowe na wyjściu obwodu. Jeśli tak jak w opisanym przykładzie chcesz ustawić pracę w klasie A, to jasne będzie, że wartość zadana musi znajdować

się pośrodku niezacienionej części omawianego wcześniej wykresu. Oznacza to, że napięcie spoczynkowe na kolektorze musi być w przybliżeniu równe połowie napięcia zasilania.

Jest to możliwe tylko wtedy, gdy prąd spoczynkowy płynący przez tranzystor może wzrosnąć tak bardzo, jak i może spaść tak, że nie znajdzie się w jednym z zacienionych obszarów. Jeśli więc stopień tranzystora zostanie podłączony do napięcia zasilania 20 V, napięcie spoczynkowe w klasie A na kolektorze U_c musi być równe 10 V. Uzyskuje się to poprzez obliczenie określonej wartości rezystancji kolektora R_c i rezystancji emitera R_e .

Linia obciążenia

Linia obciążenia określa zależność pomiędzy prądem kolektora I_c a napięciem kolektora/emitera U_{ce} tranzystora dla określonej wartości rezystorów w kolektorze i emiterze. Teraz pojęcia I_c i U_{ce} coś znaczą. Tworzą one osie charakterystyki wyjściowej $I_c=f(U_{ce})$ tranzystora (wymawiaj: „ I_c w funkcji U_{ce} ”). Ta cecha pokazuje związek między U_{ce} i I_c dla różnych wartości prądu bazowego I_b (niebieskie linie). Dlatego musisz mieć pod ręką tę charakterystykę, aby narysować linię obciążenia. Na rysunku 10 linia obciążenia (zielona) jest narysowana na wyidealizowanej charakterystyce $I_c=f(U_{ce})$ tranzystora.

Jak ustalana jest ta linia obciążenia?

Wartość zadana P musi w każdym przypadku znajdować się na linii pionowej odpowiadającej U_{ce} wynoszącemu 10 V. Następnie należy wybrać prąd spoczynkowy I_c , upewniając się, że pozostaje on znacznie poniżej „linii mocy maksymalnej”. Ta linia (czerwona) jest rysowana przez producenta tranzystora dla większości charakterystyk $I_c=f(U_{ce})$. W tym przykładzie wybrano wartość spoczynkową prądu kolektora 2,5 mA. W tym momencie wartość zadana P jest już określona. Punkt ten znajduje się na przecięciu linii pionowej w punkcie



Punkt pracy tranzystora bipolarnego

Rozwiązanie znajdziesz na www.elportal.pl/quizy

Co oznacza, że tranzystor jest „w stanie nasycenia”?

- ☐ Prąd polaryzacji wynosi dokładnie zero
- ☐ Tranzystor nie przewodzi
- ☐ Tranzystor zachowuje się jak zwarcie

W jakim ustawieniu przez tranzystor przepływa dość mały prąd polaryzacji w stanie spoczynku?

- ☐ W klasie A
- ☐ W klasie AB
- ☐ W klasie B

Które ustawienie czasem nazywa się „ustawieniem cyfrowym”?

- ☐ Praca w klasie B
- ☐ Praca w klasie C
- ☐ Praca w klasie D

Układ wspólnego emitera:

- ☐ Ma duże wzmocnienie napięciowe i nie odwraca sygnału
- ☐ To „wtórnik emiterowy” i nie ma wzmocnienia napięciowego
- ☐ Ma duże wzmocnienie napięciowe i odwraca sygnał

Układ wspólnej bazy:

- ☐ Wzmacnia sygnały nawet do bardzo wysokich częstotliwości
- ☐ Ze względu na wpływ pojemności pasozytniczych ma dość wąskie pasmo
- ☐ Ma dużą impedancję wejściową

Układ ze wspólnym kolektorem:

- ☐ Działa tak, że sygnał na emiterze jest dokładną kopią sygnału wejściowego na bazie
- ☐ Ma bardzo niską impedancję wejściową
- ☐ Jest wzmacniaczem instrumentalnym

Jakie jest napięcie przewodzenia tranzystora krzemowego?

- ☐ W przybliżeniu 0,2 V
- ☐ W przybliżeniu 0,6 V
- ☐ W przybliżeniu 1,0 V

Co wyraża linia obciążenia?

- ☐ Określa zależność pomiędzy prądem kolektora I_c , a prądem bazy I_b
- ☐ Określa maksymalne obciążenie wyjścia wzmacniacza na bazie tranzystora
- ☐ Określa zależność pomiędzy prądem kolektora I_c , a napięciem kolektora/emitera U_{ce}

Jakie powinno być napięcie na emiterze w układzie ze wspólnym emiterem?

- ☐ W przybliżeniu równe połowie napięcia zasilania
- ☐ W przybliżeniu równe 1/10 napięcia zasilania
- ☐ W przybliżeniu równe 2 V

Jakie powinno być napięcie na kolektorze w układzie ze wspólnym emiterem dla tranzystora pracującego w klasie A, przy zasilaniu 12 V?

- ☐ W przybliżeniu równe 0,6 V
- ☐ W przybliżeniu równe 6 V
- ☐ W przybliżeniu równe 2 V

Uce równej 10 V i linii poziomej w punkcie Ic równy 2,5 mA (zielone, kropkowane linie). W każdym razie punkt P leży na linii obciążenia.

Wymagany jest jednak również drugi punkt, który musi leżeć na prostej. Jeżeli prąd kolektora wynosił zero, to jasne jest, że w przykładzie napięcie kolektora/emitera jest równe napięciu zasilania, czyli 20 V. W ten sposób można wyznaczyć punkt X linii obciążenia. Następnie łączysz punkty P i X linią prostą i rozciągasz ją do osi pionowej. Linia obciążenia przecina tę oś w punkcie Y. Z tego widać, że prąd kolektora jest równy 5 mA, gdy między kolektorem a emiterem nie ma napięcia.

Jeśli przy $U_{ce}=0$ V przez łańcuch przepływa prąd 5 mA, a napięcie zasilania wynosi 20 V, to automatycznie wynika z tego, że suma trzech rezystorów R_e , R_c i R_{ce} musi być równa 4 kΩ. Ponieważ R_{ce} w tym przypadku wynosi zero (nie ma napięcia na nim – $U_{ce}=0$ V), suma R_e i R_c musi być równa 4 kΩ. Korzystając z prawa Ohma (podziel napięcie 20 V przez prąd 5 mA) możesz obliczyć, że suma wartości tych rezystorów musi być równa 4 kΩ.

Obliczanie rezystancji

Aby obliczyć obie rezystancje, musimy przejść z punktu Y z powrotem do punktu P. Praktyczna zasada mówi, że napięcie spoczynkowe na emiterze powinno być w przybliżeniu równe jednej dziesiątej napięcia zasilania. W tym przykładzie odpowiada to napięciu 2 V. Prąd spoczynkowy I_e jest nieznan, ale ze względu na duże wzmocnienie prądowe tranzystora można go porównać z prądem spoczynkowym płynącym przez kolektor, w tym przypadku 2,5 mA.

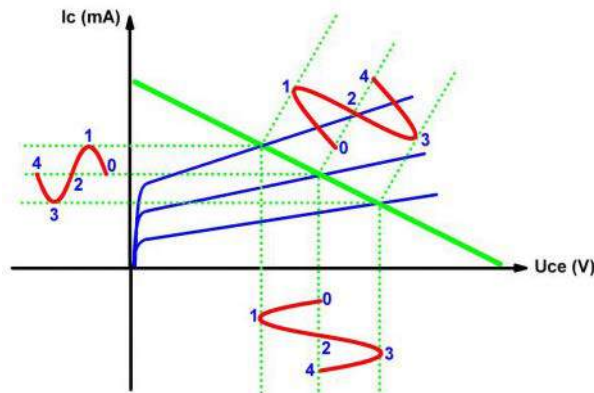
Można wówczas obliczyć wartość rezystora emitera: 2 V podzielone przez 2,5 mA daje wartość 800 Ω. Od razu znasz wartość rezystancji kolektora: 4 kΩ – 0,8 kΩ równa się 3,2 kΩ. Na ten rezystor spada napięcie 8 V. W tym przypadku na rezystor R_{ce} spada napięcie 10,0 V. Zatem rezystancja dynamiczna jest równa 4 kΩ w punkcie P. Kolektor ma napięcie 12 V.

Taki wybór wartości rezystancji spowoduje lekkie przesunięcie wartości zadanej w lewo, ale nie jest to katastrofa. W końcu obliczenia są bardzo przybliżone!

Nastawa podstawowa

Z charakterystyki $I_c=f(U_{ce})$ widać, że wartość zadana P leży w przybliżeniu na linii odpowiadającej prądowi bazy 40 μA. Następnie oszacuj prąd bazy w punkcie P, w tym przypadku 35 μA nie będzie daleko. Staje się to zatem wartością spoczynkową I_b prądu bazy. Ten prąd bazy jest dostarczany przez rezystory R_1 i R_2 . Druga praktyczna zasada mówi, że należy upewnić się, że prąd I_d płynący przez te rezystory jest około dziesięciokrotnie większy niż prąd spoczynkowy płynący do bazy.

Dlatego przez oba rezystory musi przepływać prąd I_d o natężeniu 350 μA. Rezystory są połączone szeregowo przez napięcie zasilania, więc wartość szeregową można obliczyć, korzystając z prawa Ohma: 20 V podzielone przez 350 μA daje wartość 57 kΩ. Pozostaje teraz tylko kwestia określenia stosunku pomiędzy dwoma oporami. Znane jest napięcie emitera, czyli 2 V. Ponadto wiadomo, że napięcie przewodzenia tranzystora S_i wynosi 0,6 V. Baza musi zatem znajdować się pod napięciem spoczynkowym U_b wynoszącym 2,6 V. Jest to również napięcie na rezystorze R_2 . Znasz już napięcie na rezystorze (2,6 V)



11. Modulowanie wartości spoczynkowej prądu bazy za pomocą sygnału wejściowego i konsekwencje dla napięcia i prądu kolektora

i przepływający przez niego prąd (350 μA). Wartość rezystancji obliczasz za pomocą prawa Ohma. 2,6 V podzielone przez 350 μA równa się 7,4 kΩ. Oczywiście, ostatni opór R_1 jest teraz również łatwy do określenia: 50 kΩ minus 7,4 kΩ równa się 42,6 kΩ.

W ten sposób w pełni określiliśmy parametry punktu pracy tranzystora.

Najnowsze obliczenia

Wzmocnienie prądowe dla napięcia stałego obwodu jest łatwe do obliczenia. Przecież znasz wartości spoczynkowe prądu bazy i prądów kolektora: 35 μA i 2,5 mA. Stosunek obu wartości daje wzmocnienie DC ok. 71,4. Wzmocnienie prądu przemiennego będzie znacznie większe, jeśli zmostkujesz rezystor emitera dużym kondensatorem. Na wykresie $I_c=f(U_{ce})$ możesz narysować, co się stanie, jeśli zmienisz prąd bazy symetrycznie wokół ustalonej wartości 35 μA, na przykład od 15 μA do 55 μA. Dzieje się tak, gdy do podstawy przykładany jest sygnał napięcia przemiennego poprzez kondensator izolujący i rezystor szeregowy.

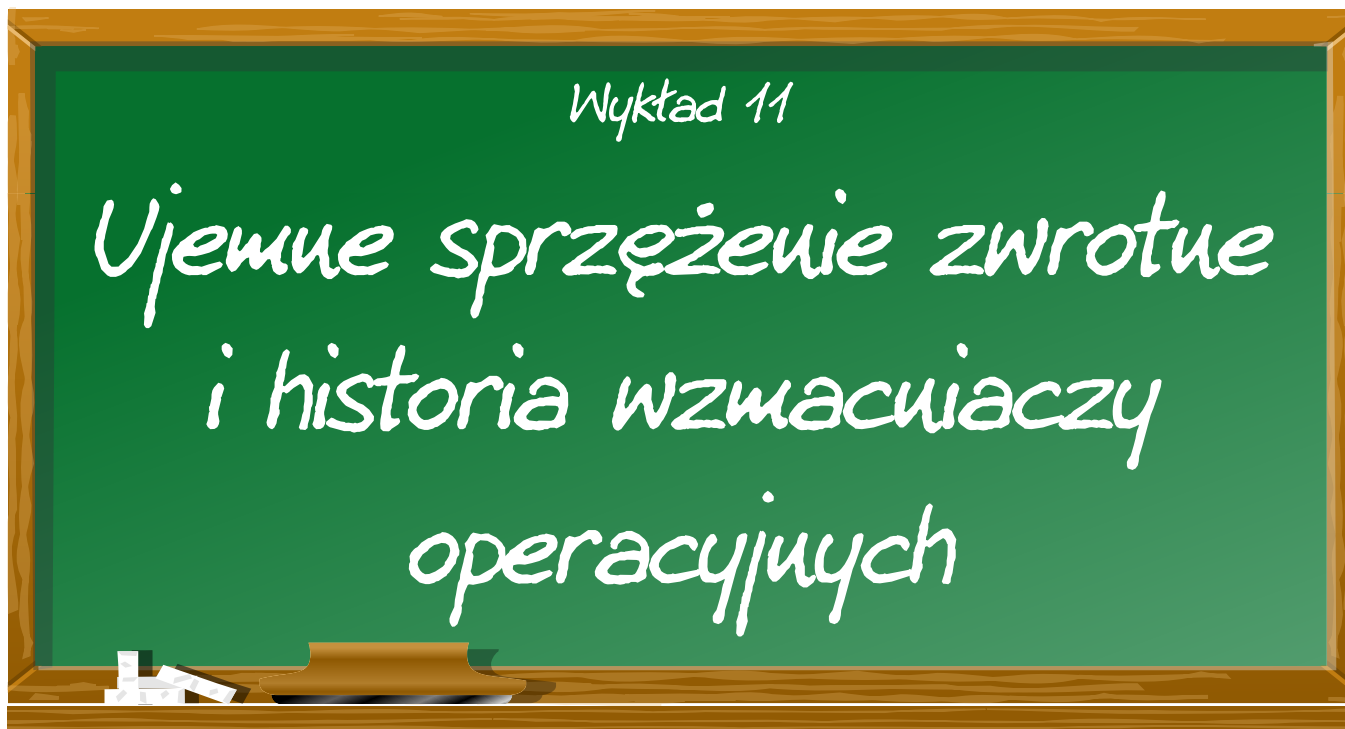
Punkt pracy będzie wówczas przesunął się w przód i w tył wzdłuż linii obciążenia. Zostało pokazane to na powyższym rysunku 11. Następnie modulowane jest również napięcie i prąd kolektora, dzięki czemu wzmocnienie prądu i napięcia stopnia tranzystorowego jest natychmiast widoczne. Jeśli prąd bazy wzrośnie od punktu 0 do punktu 1, wówczas U_{ce} zmniejszy się od punktu 0 do punktu 1, a I_c wzrośnie od punktu 0 do punktu 1. Od $I_c=f(U_{ce})$ można wywnioskować, że przy wspomnianej modulacji podstawowej prąd kolektora będzie się wahał od 3,75 mA do 1,25 mA, a napięcie kolektora/emitera będzie się zmieniać od 5 V do 15 V. Wartość szczytowa wzmocnionego sygnału wyjściowego wynosi zatem równe 10 V. Wystarczy obliczyć sterowanie podstawowe w taki sposób, aby maksymalne dodatnie i ujemne odchylenia od prądu spoczynkowego nie przekroczyły podanego 20 μA. Można to zrobić na przykład poprzez włączenie rezystora szeregowego pomiędzy poprzednim stopniem a wejściem obwodu. Możesz oczywiście obliczyć tę wartość, korzystając z prawa Ohma. ■

Jos Verstraten

REKLAMA



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.



Wzmacniacze operacyjne i układy z ujemnym sprzężeniem zwrotnym są dziś wszechobecne i można by pomyśleć, że istnieją od zawsze. Był jednak czas, gdy elektronika wciąż się rozwijała, a takie układy nie zostały jeszcze wynalezione. Zmieniło się to w 1927 roku wraz z pomysłem pewnego sprytnego człowieka...

Jednym z najbardziej znaczących wczesnych pomysłów w dziedzinie układów elektronicznych było wynalezienie przez Harolda Stevensa Blacka ujemnego sprzężenia zwrotnego. W 1927 roku Harold S. Black (1898...1983) płynął promem w kierunku swojego biura w West Street Labs firmy Western Electric, prekursora Bell Telephone Laboratories w Nowym Jorku. W jego głowie pojawił się pomysł, który diametralnie zmienił komunikację elektroniczną i utworzył nowe perspektywy dla rozwoju układów elektronicznych.

Jego pomysł dotyczył wzmacniacza z ujemnym sprzężeniem zwrotnym, w którym wzmocnienie jest dokładnie ustawione, a zniekształcenia ograniczone poprzez doprowadzenie części sygnału wyjściowego z powrotem do wejścia wzmacniacza. Black naszkicował swój pomysł na błędnie wydrukowanej stronie swojego egzemplarza New York Timesa, jedynej gazety, jaką miał przy sobie. Kiedy Black dotarł do swojego biura, poprosił kolegę o poświadczenie i podpisanie szkicu – patrz **rysunek 1**.

Przedtem, przez ostatnie sześć lat Black pracował nad ulepszeniem trzy- i czterokanałowych wzmacniaczy telefonicznych. W przypadku długodystansowych połączeń telefonicznych konieczne było dodanie repeaterów w celu odtworzenia sygnału słabnącego w miarę pokonywania odległości. Jednak wzmacniacze te miały zbyt duże zniekształcenia, więc zanim sygnał audio dotarł do miejsca docelowego, był niezrozumiały. Black zdał sobie sprawę, że zniekształcenia i szumy wzmacniacza można zmniejszyć za pomocą ujemnego sprzężenia zwrotnego, kosztem zmniejszonego wzmocnienia wzmacniacza. Później powiedział, że nie wie, co sprawiło, że jego pomysł pojawił się w jego głowie – po prostu przyszedł.

Black użył swojego nowego pomysłu do zaprojektowania szerokopasmowych wzmacniaczy o niskich zniekształceniach, które w końcu nadawały się do długodystansowych połączeń telefonicznych. Umożliwiło to przesyłanie większej liczby kanałów przez parę przewodów.



Harold Steven Black

Jego patenty

Harold S. Black otrzymał w swojej karierze 62 patenty, z których 18 dotyczyło ujemnego sprzężenia zwrotnego. Jego najbardziej znanym patentem jest numer 2,102,671, który można zobaczyć na stronie <https://patents.google.com/patent/US2102671A> (jeśli zastąpisz numer w tym linku innymi numerami patentów plus prefiks „US”, możesz wyświetlić odpowiedni plik PDF).

Patent ten, zatytułowany „WAVE TRANSLATION SYSTEM”, został złożony w 1932 roku i przyznany w 1937 roku. Liczy on 87 stron i zawiera wiele szczegółowych rysunków (w tym obwodów i wykresów) oraz mnóstwo tekstu objaśniającego. Jeden z najważniejszych zestawów schematów obwodów (ale nie jedyny!) w tym patencie, pojawiający się na stronie czwartej, został pokazany na **rysunku 2**. Pokazuje on cztery różne sposoby realizacji jego pomysłu przy użyciu lamp elektronowych. Inne ważne wątki w patencie obejmują charakterystyki wzmocnienia, kryteria stabilności, obwody równoważne i kilka praktycznych implementacji wynalazku.

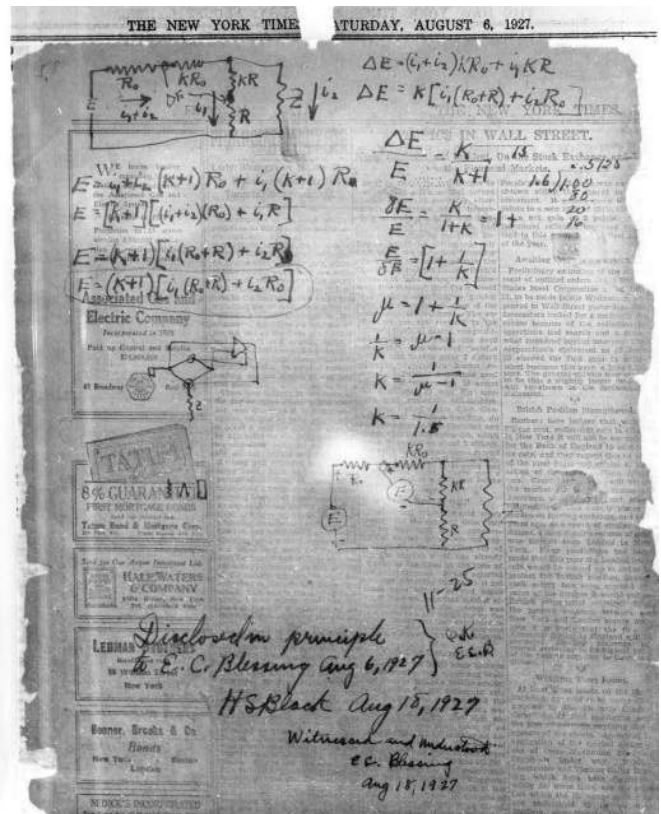
Znaczenie negatywnych informacji zwrotnych

Niemal wszystkie produkowane obecnie urządzenia analogowe korzystają z ujemnego sprzężenia zwrotnego. Obejmuje to układy obsługujące sygnały audio, analogowe wideo, sterowanie silnikami, monitorowanie baterii, oprzyrządowanie, a czasem także RF.

Rozwiązanie to jest stosowane w telewizorach, radiach, komputerach, sprzęcie medycznym, układach sterujących, przyrządach pomiarowych i telefonach komórkowych. Bardzo trudno jest znaleźć urządzenie elektroniczne, które nie korzysta z ujemnego sprzężenia zwrotnego.

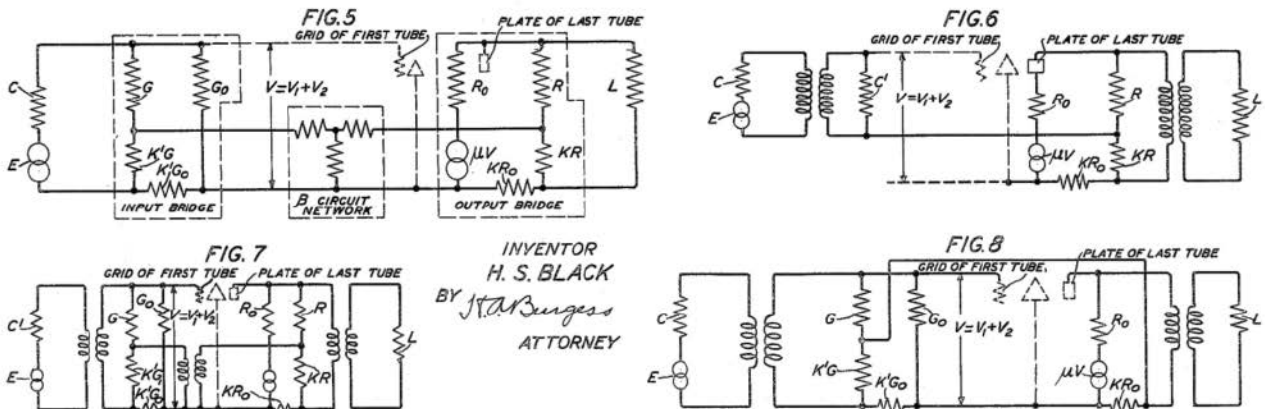
Wzmacniacze operacyjne

Utorowało to drogę do rozwoju wzmacniaczy operacyjnych (op amps); zasadniczo, monolitycznej implementacji układu, który zawiera ujemne sprzężenie zwrotne. Dostępne są tysiące



Rysunek 1. Oryginalne odręczne notatki Harolda Blacka dotyczące zasady zastosowania ujemnego sprzężenia zwrotnego do eliminacji zniekształceń

Rysunek 2. Strona z jednego z wielu patentów Harolda Blacka dotyczących ujemnego sprzężenia zwrotnego. Ten pochodzi z patentu 2,102,671 i pokazuje kilka możliwych sposobów budowy wzmacniacza z ujemnym sprzężeniem zwrotnym przy użyciu lamp elektronowych



różnych typów wzmacniaczy operacyjnych, które pasują do niemal każdego zastosowania – typy o niskim poborze mocy, typy o wysokiej szybkości działania, typy o wysokim wzmocnieniu, typy precyzyjne, pojedyncze, podwójne, poczwórne itp.

Termin „wzmacniacz operacyjny” pochodzi z około 1943 roku, kiedy to nazwa ta została wymieniona w artykule napisanym przez R. Ragazinniego pod tytułem „Analiza problemów w dynamice”. Artykuł był dziełem amerykańskiej Narodowej Rady Badań Obronnych (1940), został opublikowany przez IRE w maju 1947 roku i jest uważany za klasyczne dzieło w literaturze elektronicznej.

Firma George A. Philbrick Researches wprowadziła bazujący na lampach elektronowych wzmacniacz operacyjny ogólnego zastosowania K2-W w 1952 roku, ponad dekadę przed pojawieniem się pierwszej wersji tranzystorowej (**rysunki 3 i 4**).

Pierwszy tranzystor półprzewodnikowy został pomyślnie zademonstrowany 23 grudnia 1947 roku, ale minęło trochę czasu, zanim tranzystory znalazły się w powszechnym



Rysunek 3. Popularny wczesny lampowy wzmacniacz operacyjny, Philbrick Researches K2-W

użyciu. Pierwsza seria półprzewodnikowych wzmacniaczy operacyjnych została wprowadzona przez Burr-Brown Research Corporation i GA Philbrick Researches Inc. w 1962 roku.

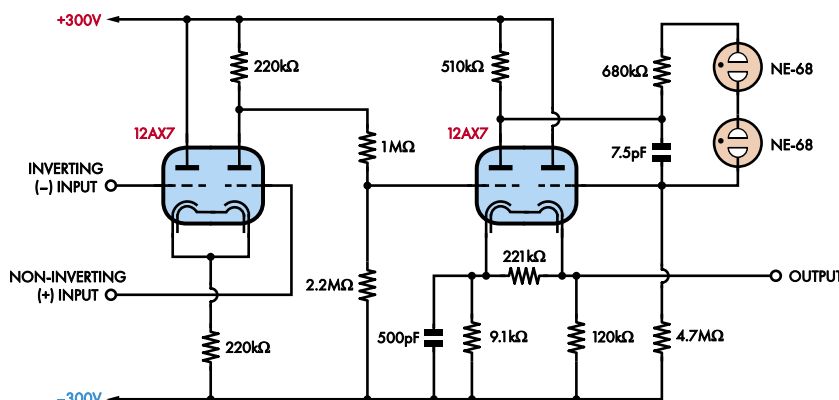
Pierwszym półprzewodnikowym monolitycznym wzmacniaczem operacyjnym, zaprojektowanym przez Boba Widlara w 1963 roku i oferowanym publicznie, był uA702 produkowany przez Fairchild Semiconductors. Wymagał jednak dziwnych napięć zasilania, takich jak +12 V i -6 V, i miał tendencję do przepalania się. Mimo to był najlepszy w swoich czasach i sprzedawany za około 300 USD (dziś to fortuna!).

Ze względu na wysoką cenę był używany głównie przez amerykańskie wojsko.

Następnie uA709 od Fairchild Semiconductors pojawił się w 1965 roku. Został on wprowadzony na rynek w cenie około 70 USD i jako pierwszy przełamał barierę ceny 10 USD, a niewiele później barierę 5 USD. W 1969 roku wzmacniacze operacyjne były sprzedawane za około 2 dolary za sztukę. Od tego momentu wielu producentów produkuje wzmacniacze operacyjne w wielu odmianach, aż do dnia dzisiejszego.

Jednym ze szczególnie popularnych modeli był uA741, który został ulepszony od czasu jego pierwszego wprowadzenia w 1968 roku. Niektóre jego warianty, takie jak LM741, są produkowane do dziś! Jego równoważny obwód pokazano na **rysunku 5**. Nowoczesne wzmacniacze operacyjne działają na tych samych zasadach, ale różnią się niektórymi szczegółami implementacji, takimi jak metoda wewnętrznej kompensacji częstotliwościowej.

Dużą zaletą wzmacniacza operacyjnego jest jego elastyczność. Może on wykonywać szeroki zakres „funkcji” analogowych po dodaniu kilku elementów pasywnych. Funkcje te obejmują mieszanie sygnałów, wzmacnianie, filtrowanie (dolnoprzepustowe, górnoprzepustowe, pasmowoprzepustowe, wąskopasmowe itp.), całkowanie, różniczkowanie, mnożenie, symulowaną indukcyjność i inne. O wzmacniaczach operacyjnych można myśleć jako o „koniach pociągowych” współczesnej elektroniki układów analogowych.



Rysunek 4. K2-W to konstrukcja podobna do tranzystorowych wzmacniaczy operacyjnych, z parą wejściową (jedna podwójna trioda 12AX7), po której następuje stopień wzmacnienia/buforowania napięcia wykonany z innej podwójnej triody 12AX7 i dwóch lamp neonowych

Ujemne sprzężenie zwrotne

W jaki sposób ujemne sprzężenie zwrotne jest używane do sterowania wzmacniaczem operacyjnym w celu zmniejszenia zniekształceń i ustalenia stałego wzmocnienia?



Ujemne sprzężenie zwrotne

Rozwiązanie znajdziesz na www.elportal.pl/quizy

Pierwszy monolityczny wzmacniacz operacyjny, opracowany w 1962 roku, to:

- ☐ uA741
- ☐ uA702
- ☐ K2-W

Jakie będzie wzmocnienie wzmacniacza operacyjnego, gdy wejście odwracające połączymy bezpośrednio z jego wyjściem?

- ☐ Duże, w niektórych przypadkach ponad milion
- ☐ -1 V/V
- ☐ Powstanie bufor o jednostkowym wzmocnieniu

Do skonfigurowania wzmacniacza operacyjnego jako wzmacniacza napięciowego o stałym wzmocnieniu większym od 1, potrzebne są:

- ☐ Dwa rezystory
- ☐ Dwa rezystory i kondensator
- ☐ Tylko jeden rezystor

Czy wzmocnienie może być ujemne?

- ☐ Nie, wzmacniacz zwiększa amplitudę sygnału
- ☐ Tak, wzmacniacz odwracający zawsze ma wzmocnienie ujemne
- ☐ Tak, wzmacniacz odwracający może mieć wzmocnienie ujemne

Wzmacniacz różnicowy:

- ☐ Daje na wyjściu wartość różnicy napięć pomiędzy jego wejściami, pomnożoną przez wartość wzmocnienia
- ☐ Daje na wyjściu wartość różnicy napięć pomiędzy jego wejściami
- ☐ Daje na wyjściu wartość różnicy napięć pomiędzy jego wejściem i napięciem zasilania

Co to jest żyrator?

- ☐ Obwód, który symuluje cewkę indukcyjną przy niskich wartościach prądu
- ☐ Filtr pasmowoprzepustowy drugiego rzędu
- ☐ Wzmacniacz instrumentalny o wysokim współczynniku CMRR

Do czego służy układ „Twin-T” (podwójne T)?

- ☐ Jest to prostownik precyzyjny
- ☐ Jest to aktywny filtr wycinający
- ☐ Jest to aktywny filtr pasmowoprzepustowy

Czym różni się komparator od zwykłego wzmacniacza operacyjnego?

- ☐ Ma usunięty obwód kompensacji częstotliwości w celu uzyskania szybszej reakcji na wyjściu
- ☐ Ma silne dodatnie sprzężenie zwrotne powodujące histerezę
- ☐ Ma tranzystory wyjściowe, które mogą przewodzić znaczne prądy

Czym wyróżniają się wzmacniacze operacyjne typu rail-to-rail?

- ☐ Przystosowane są do sterowania stosunkowo niskich impedancji obciążenia
- ☐ Napięcie wyjściowe może przyjmować wartości praktycznie w całym zakresie napięcia zasilania
- ☐ Mogą pracować przy bardzo niskim napięciu zasilania, np. 1,8 V

Czym wyróżniają się wzmacniacze operacyjne CMOS?

- ☐ Pracują jako oscylatory do generowania przebiegów sinusoidalnych
- ☐ Mają dużą impedancję wejściową, nawet w zakresie teraomów
- ☐ Są zbudowane z bramek logicznych CMOS

Napięcie wyjściowe wzmacniacza operacyjnego to nieodwracające napięcie wejściowe minus odwracające napięcie wejściowe pomnożone przez duży współczynnik wzmocnienia (w niektórych przypadkach ponad milion). Jeśli mówimy, że wzmocnienie wynosi dokładnie milion, oznacza to, że:

Jeśli wejście „+” wynosi $100\ \mu\text{V}$, a wejście „-” $99\ \mu\text{V}$, na wyjściu będzie $+1\ \text{V}$;

Jeśli wejście „+” wynosi $100\ \mu\text{V}$, a wejście „-” wynosi $100\ \mu\text{V}$, na wyjściu będzie $0\ \text{V}$;

Jeśli wejście „+” wynosi $100\ \mu\text{V}$, a wejście „-” $101\ \mu\text{V}$, na wyjściu będzie $-1\ \text{V}$.

Z tego widać, że jeśli różnica między napięciami wejściowymi jest większa niż kilka mikrowoltów, napięcie wyjściowe zostanie „nasycone” do poziomu jednej lub drugiej szyny zasilającej. Tak więc, o ile nie używamy wzmacniacza operacyjnego jako komparatora (jedna z funkcji wzmacniacza operacyjnego), wejścia prawie zawsze będą miały bardzo podobne napięcie. Ujemne sprzężenie zwrotne jest zwykle skonfigurowane tak, aby zapewnić taki stan.

Załóżmy, że podajemy 10% napięcia wyjściowego z powrotem na wejście odwracające i przykładamy $1\ \text{V}$ do wejścia nieodwracającego. Gdy napięcie wyjściowe jest mniejsze niż $10\ \text{V}$, różnica napięć między wejściami będzie dodatnia, więc napięcie wyjściowe wzrośnie. Gdy napięcie wyjściowe jest większe niż $10\ \text{V}$, różnica napięć między wejściami będzie ujemna, więc napięcie wyjściowe spadnie. W ten sposób napięcie wyjściowe będzie dążyć do $10\ \text{V}$.

Jedynymi rzeczywistymi źródłami błędów w kontekście prądu stałego są: napięcie offsetu wejściowego (wyjście nie wynosi dokładnie $0\ \text{V}$ przy obu wejściach o tym samym napięciu) i skończone wzmocnienie, które dodaje kilka dodatkowych mikrowoltów błędu. Ale to tylko jedna część na milion.

Jest to sytuacja bardzo zbliżona do idealnego wzmacniacza o stałym wzmocnieniu. Z pewnością nie jest tak w przypadku typowego wzmacniacza z pojedynczym tranzystorem! Ze względu na tolerancje produkcyjne, trudno jest skonfigurować (polaryzować) pojedynczy tranzystor tak, aby zapewnić dokładnie ustalone wzmocnienie. Nawet jeśli uda się je osiągnąć (np. poprzez wyregulowanie), prawdopodobnie będzie się ono zmieniać wraz z temperaturą i upływem czasu.

Zwróć uwagę, że dokładna znajomość wzmocnienia wzmacniacza operacyjnego bez sprzężenia zwrotnego nie ma praktycznego znaczenia. Ogólne wzmocnienie jest ustawiane przez dzielnik sprzężenia zwrotnego, zwykle wykonany z rezystorów (i czasami kondensatorów), więc łatwo jest ustawić go blisko żądanej wartości. W razie potrzeby można go odpowiednio dobrać i jest mało prawdopodobne, aby dryfował.

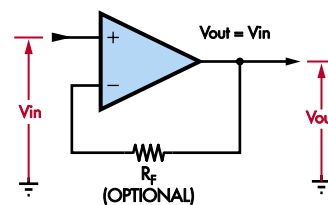
Ujemne sprzężenie zwrotne daje również prawie idealne wyniki dla sygnałów AC, o ile ich częstotliwość jest znacznie poniżej szerokości pasma wzmacniacza operacyjnego (zwykle określanej jako szerokość pasma wzmocnienia, którą należy podzielić przez skonfigurowane wzmocnienie). W ten sposób wzmacniacz bazujący na wzmacniaczu operacyjnym może zapewnić zasadniczo płaską krzywą wzmocnienia w całym zakresie częstotliwości, podczas gdy tranzystor zazwyczaj nie będzie miał płaskiej charakterystyki, chyba że jest to specjalny model.

Podstawowe układy ze wzmacniaczami operacyjnymi

1) Bufor o jednostkowym wzmocnieniu

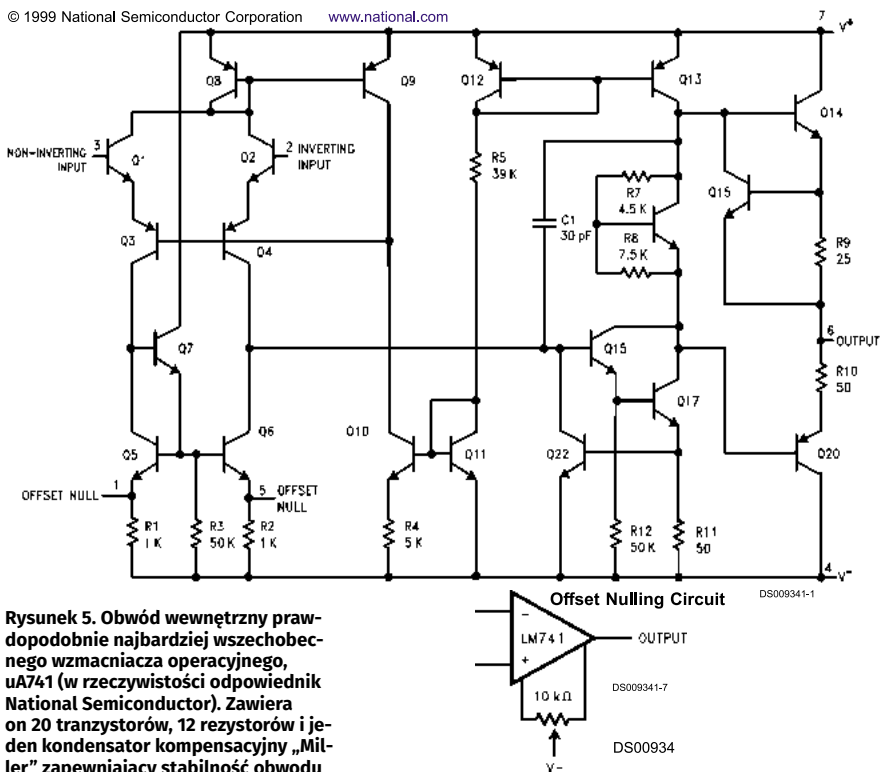
Na rysunku 6 został pokazany wzmacniacz operacyjny zaaranżowany jako bufor o jednostkowym wzmocnieniu. Wyjście jest podawane z powrotem na wejście odwracające, więc napięcie wyjściowe podąża za wejściem nieodwracającym. Ponieważ wyjście wzmacniacza operacyjnego ma impedancję bliską zero, ale wejście ma stosunkowo wysoką impedancję, ta konfiguracja jest przydatna, aby uniknąć obciążenia źródła sygnału podawanego na wejście przez układ sterowany z wyjścia wzmacniacza operacyjnego.

Często wyjście jest podłączane bezpośrednio do wejścia odwracającego. Jednak w niektórych przypadkach niedopasowanie impedancji źródła między wejściami może powodować dryft temperaturowy i inne problemy. Aby tego uniknąć, rezystor R_f można dobrać tak, aby pasował do impedancji wejścia nieodwracającego.



Rysunek 6. Użycie wzmacniacza operacyjnego do buforowania sygnału może być tak proste, jak podłączenie jego wyjścia do wejścia odwracającego. Jednak rezystor R_f jest dobrym pomysłem, aby zrównoważyć prądy wejściowe, jeśli impedancja źródła dla wejścia odwracającego jest stosunkowo wysoka

© 1999 National Semiconductor Corporation www.national.com



Rysunek 5. Obwód wewnętrzny prawdopodobnie najbardziej wszechobecnego wzmacniacza operacyjnego, uA741 (w rzeczywistości odpowiednik National Semiconductor). Zawiera on 20 tranzystorów, 12 rezystorów i jeden kondensator kompensacyjny „Miller” zapewniający stabilność obwodu

2) Wzmacniacz nieodwracający

Na **rysunku 7** pokazano wzmacniacz operacyjny zapewniający wzmocnienie nieodwracające.

Napięcie wyjściowe jest sygnałem prądu przemiennego o takim samym kształcie jak sygnał wejściowy, ale o zwiększonej wielkości, o współczynnik . Podobnie jak w przypadku bufora, układ ten może być podłączony do źródła sygnału o wysokiej impedancji, zapewniając wyjście o niskiej impedancji.

Kondensator C1 można pominąć, ale zazwyczaj warto go zachować. Zmniejsza on wzmocnienie układu przy wyższych częstotliwościach, zwiększając tym samym stabilność i zapobiegając wzmacnianiu niepożądanych sygnałów o wysokiej częstotliwości.

W dolnej części dzielnika sprzężenia zwrotnego, między dolną końcówką R1 a masą, można zobaczyć kondensator o dużej wartości, pokazany jako alternatywne połączenie dla R1 na **rysunku 7**. Ustawia on wzmocnienie obwodu dla napięcia stałego (DC) na jednostkę niezależnie od wzmocnienia dla sygnałów zmiennych (AC). Dlatego jest często używany w obwodach wzmacniania sygnałów AC (patrz także **rysunek 19**). Zmniejszanie wzmocnienia DC układu zapobiega zablokowaniu wyjścia na szynie dodatniej przy dodatnich skokach sygnału, a także zmniejsza wzmocnienie wejściowego napięcia błędu offsetu.

Praktyczny limit wzmocnienia zależy od szerokości pasma wzmocnienia wzmacniacza operacyjnego i maksymalnej częstotliwości sygnału. Na przykład, wzmacniacz operacyjny o szerokości pasma wzmocnienia 3 MHz ma maksymalne praktyczne wzmocnienie 30 razy dla sygnałów do 100 kHz (3 MHz : 100 kHz). Szumy i zniekształcenia na wyjściu zwiększają się wraz ze wzrostem wzmocnienia, ponieważ wzmacniacz operacyjny ma mniejsze sprzężenie zwrotne (szerokość pasma zamkniętej pętli).

3) Wzmacniacz odwracający

Dzięki doprowadzeniu sygnału do wejścia odwracającego, a nie nieodwracającego, poprzez rezystor, sygnał jest odwracany i nadal można uzyskać wzmocnienie, jak pokazano na **rysunku 8**. Wzmocnienie wynosi, więc w przeciwieństwie do wersji nieodwracającej, wartości wzmocnienia mogą być mniejsze niż jednostka (tj. tłumienie) bez oddzielnego tłumika wejściowego.

Niefortunną konsekwencją tej konfiguracji jest to, że typowo wysoka impedancja wejściowa wzmacniacza operacyjnego jest zredukowana do wartości R_{in} , więc źródło sygnału podawanego na wejście jest bardziej obciążone. Problem ten można rozwiązać poprzez dodanie bufora o jednostkowym wzmocnieniu pomiędzy źródłem sygnału a wzmacniaczem odwracającym.

Jedną z zalet tej konfiguracji jest to, że oba wejścia wzmacniacza operacyjnego są utrzymywane na stałym napięciu (V_{bias}), więc nie ma sygnału w trybie wspólnym, a zatem nie ma zniekształceń w trybie wspólnym (często dominujący mechanizm zniekształceń).

W tym układzie kondensator C1 pełni podobną rolę jak na **rysunku 7**, chociaż jest prawdopodobnie bardziej skuteczny, ponieważ zmniejsza wzmocnienie przy bardzo wysokich częstotliwościach do zera, a nie do jednostki.

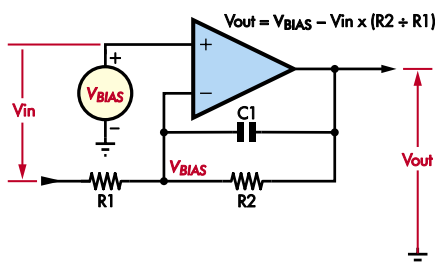
4) Wirtualny mieszacz

Na **rysunku 9** został pokazany układ, który jest w zasadzie wzmacniaczem odwracającym z wieloma rezystorami doprowadzającymi różne sygnały do wejścia odwracającego. Ponieważ wejście odwracające jest utrzymywane na stałym napięciu stałym przez ujemne sprzężenie zwrotne, nie ma możliwości przesłuchu między sygnałami (co może być istotne w konsoli miksującej, gdzie się one podawane z wielu źródeł).

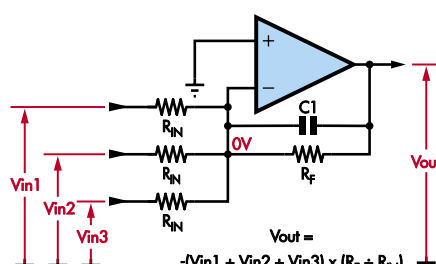
5) Wzmacniacz różnicowy

Jest to bardzo przydatny układ stosowany w wielu różnych formach. Chociaż można go zbudować przy użyciu zwykłych wzmacniaczy operacyjnych, jest on prawdopodobnie szerzej stosowany w monolitycznych wzmacniaczach instrumentalnych (choć w zmodyfikowanej formie), wzmacniaczach różnicowych i monitorach boczników prądowych.

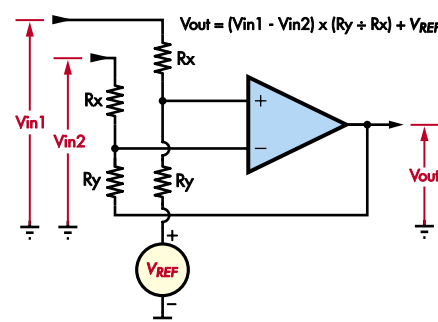
Na **rysunku 10** pokazano podstawową wersję tego układu. Zapewnia on niezwykle przydatną funkcjonalność – pobiera różnicę między dwoma napięciami, mnoży ją przez stałą wartość wzmocnienia (określoną przez wartości rezystorów), a następnie ewentualnie dodaje do dodatnie lub ujemne napięcie przesunięcia. Jednak V_{ref} jest często ustawiane na 0 V, więc napięcie wyjściowe jest odniesione do masy.



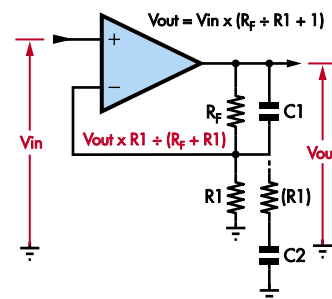
Rysunek 8. Konfiguracja wzmacniacza odwracającego również zawiera dwa rezystory i jeden opcjonalny kondensator. Jego zaletą jest to, że wzmocnienie może być mniejsze niż jednostka, ale wadą jest to, że impedancja wejściowa jest równa R_{in} , a nie zwykle znacznie wyższa wartość dla wejść wzmacniacza operacyjnego



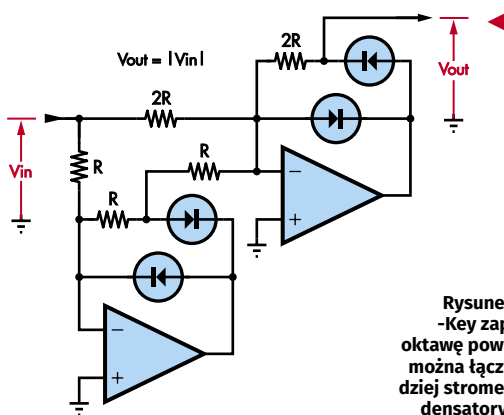
Rysunek 9. Mieszacz z wirtualną masą to wzmacniacz odwracający z wieloma źródłami sygnału. Ponieważ oba wejścia wzmacniacza operacyjnego są utrzymywane bardzo blisko 0 V, nie ma możliwości, aby podawane sygnały oddziaływały na siebie nawzajem, z wyjątkiem wyjścia mieszacza



Rysunek 10. Podstawowy wzmacniacz różnicowy oblicza różnicę między dwoma napięciami, pomnożoną przez stałą wartość wzmocnienia plus przesunięcie. Wymaga dobrego dopasowania rezystorów

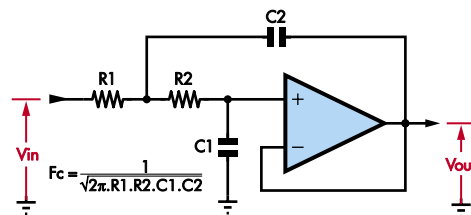


Rysunek 7. Do skonfigurowania wzmacniacza operacyjnego jako wzmacniacza napięciowego o stałym wzmocnieniu potrzebne są tylko dwa rezystory. Ponieważ sygnał jest podawany bezpośrednio na wejście nieodwracające, impedancja wejściowa jest wysoka. Opcjonalny kondensator C1 ogranicza szerokość pasma w celu zapewnienia stabilności, podczas gdy C2 może być użyty do zmniejszenia wzmocnienia DC do jednostki przy jednoczesnym wyższym wzmocnieniu AC



Rysunek 11. Ten układ prostownika pełnokresowego zawiera wzmacniacze operacyjne, aby skutecznie zniwelować napięcie przewodzenia diod. W rezultacie, dla dodatnich napięć V_{in} , V_{out} działa bardzo dokładnie (w zakresie mikrowoltów, biorąc pod uwagę wystarczająco precyzyjne wzmacniacze operacyjne), a dla ujemnych napięć V_{in} , $V_{out} = -V_{in}$ (ponownie, w zakresie mikrowoltów). Jest to idealne rozwiązanie dla układów, które muszą wykrywać szczytowe poziomy sygnału, takie jak mierniki obcinania sygnału audio

Rysunek 12. Ten filtr dolnoprzepustowy Sallen-Key zapewnia redukcję amplitudy przy -12 dB/oktawę powyżej częstotliwości -3 dB, a wiele stopni można łączyć kaskadowo, aby uzyskać jeszcze bardziej strome nachylenie. Zmiana rezystorów na kondensatory i kondensatorów na rezystory sprawia, że jest to filtr górnoprzepustowy



Układ ten wymaga precyzyjnego dopasowania rezystorów w celu uzyskania dobrego współczynnika tłumienia sygnału wspólnego (CMRR). Nawet w przypadku rezystorów o tolerancji 0,1%, trudno jest zagwarantować CMRR powyżej 60 dB. Kalibrowanie może dać dobre wyniki, chociaż procedura może być trudna. Ogólnie rzecz biorąc, lepiej jest używać laserowo kalibrowanych układów monolitycznych, takich jak wzmacniacze instrumentalne, które mogą mieć współczynnik CMRR powyżej 100 dB.

Większość wzmacniaczy instrumentalnych zawiera nieco inny obwód wewnętrzny, który zawiera trzy wzmacniacze operacyjne; oprócz bardzo dobrego współczynnika CMRR, ma to tę zaletę, że wzmocnienie można ustawić za pomocą jednego zewnętrznego rezystora. Podstawowa zasada działania jest jednak taka sama.

Wzmacniacz różnicowy jest zasadniczo wzmacniaczem instrumentalnym, w którym napięcia wejściowe mogą znajdować się znacznie poza (zwykle powyżej) zakresem zasilania układu. Monitor bocznika prądowego jest wyspecjalizowaną wersją wzmacniacza instrumentalnego. Wszystkie wewnętrznie bazują na wzmacniaczach operacyjnych lub podobnych układach.

Monitor bocznikowy umożliwia umieszczenie rezystora bocznikowego o niskiej wartości w dodatnim zasilaniu układu i uzyskanie napięcia odniesienia do masy w celu zasilania przetwornika analogowo-cyfrowego (ADC) lub podobnego. Mają one wysoki współczynnik CMRR, aby stłumić tętnienia zasilania.

6) Prostowniki precyzyjne

Precyzyjny prostownik działa jak dioda lub mostek prostowniczy, ale bez spadku napięcia przewodzenia. Jest to ważne w przypadku prostowania sygnałów o niskim poziomie (zbyt niskim, aby uzyskać polaryzację diody w kierunku przewodzenia) lub w przypadku dokładnego prostowania sygnałów prądu przemiennego w celu pomiaru ich wielkości itp. Są one powszechnie stosowane w urządzeniach takich jak mierniki VU lub monitory prądu przemiennego.

Na **rysunku 11** pokazano wersję pełnofalową, podobną do prostownika mostkowego. Wersja półfalowa to w zasadzie tylko jedna z sekcji wzmacniacza operacyjnego.

Wzmacniacze operacyjne zmniejszają efektywne napięcie przewodzenia diod o współczynnik ich wzmocnienia w otwartej pętli, co oznacza, że spadek napięcia $\sim 0,7$ V na standardowej diodzie krzemowej jest efektywnie mniejszy niż $1 \mu\text{V}$ dla wzmocnienia w otwartej pętli wynoszącego około miliona.

Pokazane wartości rezystorów dają wzmocnienie równe jedności. Ten układ pierwotnie pochodzi od National Semiconductor, który określił $R=100$ k Ω , chociaż można użyć innych wartości. W razie potrzeby wartości można zmieniać, aby uzyskać stałe wzmocnienie.

7) Aktywny filtr dolnoprzepustowy

Najprostszym sposobem implementacji filtra dolnoprzepustowego za pomocą wzmacniacza operacyjnego jest połączenie podstawowego filtra dolnoprzepustowego RC z buforem o jednostkowym wzmocnieniu. Jednak bardziej ekonomicznym układem jest filtr dolnoprzepustowy Sallen-Key pokazany na **rysunku 12**. Ma on nachylenie -12 dB/oktawę, w porównaniu do -6 dB/oktawę dla filtra RC, przy użyciu tylko jednego wzmacniacza operacyjnego. Umożliwia on również uzyskanie wzmocnienia.

Na **rysunku 13** został pokazany filtr dolnoprzepustowy z wielokrotnym sprzężeniem zwrotnym. Zapewnia on dokładnie taką samą funkcję jak filtr Sallen-Key, ale jest mniej podatny na przenikanie sygnału, co oznacza, że działa znacznie bliżej idealnego filtra przy częstotliwościach zbliżonych do szerokości pasma wzmacniacza operacyjnego. Jediną wadą jest użycie jednego rezystora więcej.

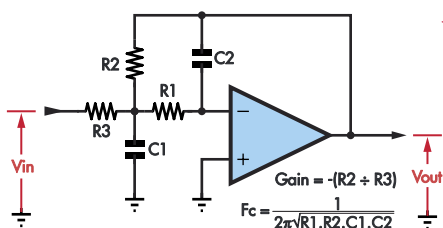
Aby obliczyć wymagane wartości rezystorów i kondensatorów dla danej częstotliwości odcięcia, odwiedź stronę siliconchip.com.au/link/aajq. Należy pamiętać, że możliwe jest zbudowanie aktywnego filtra dolnoprzepustowego trzeciego rzędu Sallen-Key przy użyciu pojedynczego wzmacniacza operacyjnego. Pozwoli to uzyskać 18 dB/oktawę przy jednym wzmacniaczu operacyjnym, 30 dB/oktawę przy dwóch itd. Pokazano to na stronie siliconchip.com.au/link/ab8v.

8) Aktywny filtr górnoprzepustowy

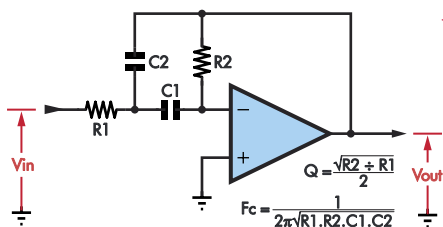
Aby przekształcić filtry dolnoprzepustowe pokazane na rysunkach 12 i 13 w filtry górnoprzepustowe, wystarczy przetransponować rezystory i kondensatory. Podobnie jak w przypadku filtrów dolnoprzepustowych, zapewnią one nachylenie 12 dB/oktawę na wzmacniacz operacyjny.

Zarówno w przypadku filtrów dolnoprzepustowych, jak i górnoprzepustowych, poprzez dostosowanie rezystancji i pojemności, możliwe jest zaprojektowanie filtrów o charakterystyce innej niż Butterworth. Butterworth ma minimalne (zasadniczo prawie żadne) tętnienie w paśmie przepustowym, ale różne typy filtrów, takie jak Chebyshev, kompensują zwiększone tętnienie pasma przepustowego w zamian za bardziej stromy spadek poza nim.

Aby obliczyć wymagane wartości komponentów, zobacz siliconchip.com.au/link/ab8w



Rysunek 13. Ten filtr z wielokrotnym sprzężeniem zwrotnym wykonuje to samo zadanie, co filtr Sallen-Key, ale jest bardziej skuteczny przy wyższych częstotliwościach. Jest to ważne w przypadku filtrów dolnoprzepustowych, ponieważ w przeciwnym razie może przepuszczać sygnały, które filtr ma blokować. Ponieważ potrzebny jest tylko jeden dodatkowy rezystor, jest to optyczne ulepszenie, a wzmacnienie można ustawić bez dodatkowych rezystorów (choć odwraca sygnał)



Rysunek 14. Ten aktywny filtr pasmowoprzepustowy blokuje sygnały poza danym zakresem częstotliwości, chociaż nachylenie wynosi tylko -6 dB/oktawę. W przypadku bardziej stromych zboczy (np. -12 dB/oktawę), jeden z opisanych powyżej aktywnych filtrów dolnoprzepustowych można połączyć szeregowo z podobnym filtrem górnoprzepustowym

Rysunek 15. Ten aktywny filtr wycinający Twin-T (podwójne T) tłumi sygnały o określonej częstotliwości. Zarówno częstotliwość, jak i stromość/głębokość wycięcia można kontrolować poprzez staranny dobór wartości elementów pasywnych

9) Aktywny filtr pasmowoprzepustowy

Filtr pasmowoprzepustowy drugiego rzędu można utworzyć poprzez połączenie aktywnych filtrów dolnoprzepustowego i górnoprzepustowego drugiego rzędu. Alternatywnie, można użyć konfiguracji pokazanej na **rysunku 14**, gdzie pojedynczy wzmacniacz operacyjny może działać jako filtr pasmowoprzepustowy pierwszego rzędu z regulowanym wzmacnieniem i Q (dobroć) do 25. Ta konfiguracja odwraca fazę sygnału, jednak w przypadku łączenia łańcuchowego wielu filtrów, może ona zostać ponownie odwrócona przez inny stopień.

10) Aktywny filtr wycinający

Na **rysunku 15** został pokazany aktywny filtr wycinający „Twin-T” (podwójne T). Jednym z interesujących aspektów tego projektu jest to, że Q, a tym samym głębokość wycięcia, zmienia się w zależności od wybranych wartości rezystora i kondensatora. Zobacz kalkulator online na stronie siliconchip.com.au/link/ab8x

11) Żyrator

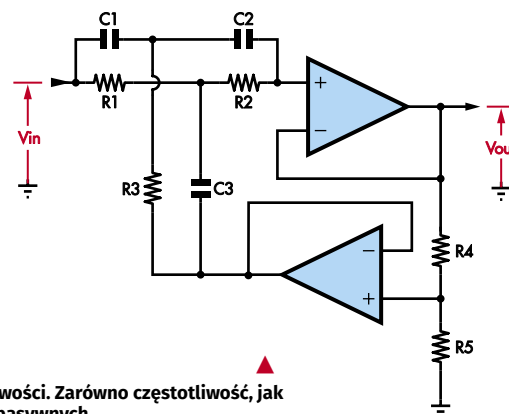
Na **rysunku 16** został pokazany „żyrator”, aktywny element, który zachowuje się podobnie do idealnej cewki indukcyjnej przy niskich wartościach prądu. Jest to możliwe dzięki zastosowaniu ujemnego sprzężenia zwrotnego wzmacniacza operacyjnego do skutecznego odwrócenia zachowania kondensatora C. Jest to przydatne w obwodach takich jak korektory graficzne, gdzie potrzebne są elementy rezonansowe (LC) o dokładnych częstotliwościach rezonansowych, niskich zniekształceniach i niewielkich rozmiarach. Tolerancje cewek indukcyjnych są zwykle znacznie szersze niż kondensatorów, a cewki o wysokiej wartości mogą być bardzo nieporęczne, więc w obwodach przetwarzania sygnału żyrator jest prawie zawsze lepszy niż obwód rezonansowy bazujący na rzeczywistej cewce indukcyjnej.

12) Filtr aktywny Baxandall

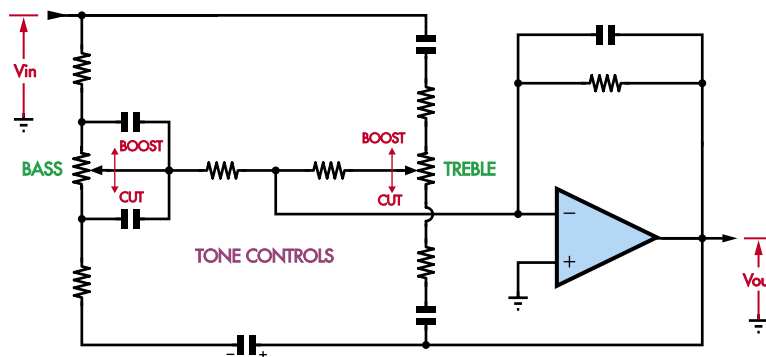
Na **rysunku 17** pokazano podstawową wersję szeroko stosowanego aktywnego regulatora barwy dźwięku Baxandall. To rozwiązanie ma wiele dobrych właściwości, takich jak możliwość ustalania dowolnej liczby pasm, bez interakcji między regulatorami i bez specjalnych wymagań dotyczących potencjometrów. Na zdjęciu pokazano potencjometry tonów niskich i wysokich, ale z łatwością można dodać jeden lub dwa potencjometry tonów średnich.

Na **rysunku 18** został pokazany aktywny regulator głośności Baxandall. Tradycyjną metodą regulacji głośności jest potencjometr logarytmiczny, ale wersje podwójne (stereo) potencjometrów zwykle mają kiepskie parametry, co jest słyszalne przy niskim poziomie sygnałów, więc nie są zalecane do obwodów stereo.

Aktywny obwód Baxandall zapewnia logarytmiczną kontrolę z zastosowaniem liniowego potencjometru



Rysunek 16. Żyrator jest szczególnie sprytnym obwodem. Zawiera on ujemne sprzężenie zwrotne, aby kondensator zachowywał się jak cewka indukcyjna. W wielu zastosowaniach związanych z przetwarzaniem sygnałów jest on lepszy od rzeczywistej cewki indukcyjnej



Rysunek 17. Regulacja tonów Baxandall została początkowo zaprojektowana z lampą lub tranzystorem jako elementem aktywnym, ale działa jeszcze lepiej ze wzmacniaczem operacyjnym. Jest to układ elegancki i rozszerzalny, praktycznie bez interakcji między stopniami (w tym przypadku dwoma – regulacją tonów niskich i wysokich). Bez względu na liczbę pasm, wymagany jest tylko jeden wzmacniacz operacyjny na kanał (tj. dwa dla stereo)

67

Również inne funkcje matematyczne mogą być stosowane do napięć przetwarzanych przez wzmacniacz operacyjny, w tym dodawanie, odejmowanie, dzielenie i odwrotność logarytmu (wspomniane wyżej potęgowanie).

Wzmacniacze operacyjne mogą być również stosowane do budowy kontrolowanych źródeł prądu, w tym stałych obciążeń, poprzez łączenie wzmacniaczy operacyjnych z dużymi tranzystorami, które mogą obsługiwać dużą moc przy wystarczającym chłodzeniu.

Uogólniony konwerter impedancji zawiera dwa wzmacniacze operacyjne do przedstawienia impedancji obciążenia proporcjonalnej do innej impedancji. Współczynnik można ustawić za pomocą stałych lub zmiennych rezystorów (lub nawet innych impedancji!).

Wiele wzmacniaczy operacyjnych jest zaprojektowanych do sterowania stosunkowo niskich impedancji obciążenia, takich jak 600 Ω . Działają one całkiem dobrze jako podstawowe sterowniki słuchawek, z relatywnie niskimi zniekształceniami dla typowych słuchawek, nawet o tak niskiej impedancji jak 16 Ω . Nie są one w stanie dostarczyć ogromnej mocy, ale wystarczającej dla większości słuchawek, aby zapewnić przyzwoitą głośność, przy użyciu jednego taniego układu scalonego.

Wzmacniacz operacyjny może być również używany jako wzmacniacz błędu w sterowaniu ze sprzężeniem zwrotnym. Na przykład do regulacji napędu silnika w celu utrzymania stałej prędkości mimo zmiennego obciążenia.

Wzmacniacz operacyjny może stanowić podstawę różnych oscylatorów do generowania przebiegów o stałej lub zmiennej częstotliwości; głównie sinusoid, ale także fal trójkątnych lub piłokształtnych. Wzmacniacz operacyjny (zwłaszcza typu CMOS) może być używany jako wzmacniacz buforowy o wysokiej impedancji wejściowej lub pierścień ochronny do monitorowania czujników, które nie są w stanie wytrzymać żadnego obciążenia, takich jak wąskopasmowe czujniki tlenu i czujniki pH. Wzmacniacze operacyjne CMOS mogą mieć impedancję wejściową w zakresie teraomów (ponad bilion omów)!

Bufory wzmacniaczy operacyjnych CMOS można również łączyć z przełącznikami analogowymi i kondensatorami o niskim poziomie upływu, tworząc obwody próbkowania i podtrzymywania, często używane do próbkowania napięć w małych oknach czasowych w celu sterowania przetwornika ADC lub podobnego.

Ograniczenia amplitudy sygnału

Przez bardzo długi czas sygnały na wejściach i wyjściach wzmacniacza operacyjnego mogły mieć tylko znacznie mniejszy zakres zmian amplitudy niż zakres zasilania wzmacniacza operacyjnego. Na przykład, jeśli wzmacniacz operacyjny jest zasilany napięciem 12 V, wejścia i wyjścia mogą być ograniczone do zakresu 3...9 V. Lub, przy podwójnym zasilaniu, takim jak ± 15 V, możesz być ograniczony do amplitudy sygnału ± 12 V. Dzieje się tak dlatego, że wewnętrzne obwody wzmacniacza operacyjnego potrzebują pewnego „zapasu” napięcia do działania.

Jednak teraz są dostępne wzmacniacze operacyjne z pojedynczym zasilaniem i wyjściem typu rail-to-rail. Wzmacniacze operacyjne z pojedynczym zasilaniem zazwyczaj umożliwiają zejście wejść i wyjść do szyny ujemnej (np. 0 V). Tak więc wzmacniacz operacyjny z pojedynczym zasilaniem z 12 V może obsługiwać sygnały powiedzmy 0...9 V.

Wyjściowe wzmacniacze operacyjne typu rail-to-rail mają generalnie takie same ograniczenia odnośnie wejścia jak standardowe wzmacniacze operacyjne, ale ich napięcie wyjściowe może przyjmować wartości praktycznie w całym zakresie zasilania. Jest to szczególnie przydatne przy stosowaniu wzmocnienia do sygnałów AC, ponieważ w takim przypadku amplitudy sygnałów wejściowych i tak nigdy nie osiągną napięcia szyn zasilania (przynajmniej nie bez „nasylenia” wzmacniacza operacyjnego).

W dzisiejszych czasach wzmacniacze operacyjne z wejściem/wyjściem typu rail-to-rail (RRIO) są bardzo powszechne. Niektóre z nich mogą nawet pracować przy bardzo niskim napięciu zasilania, np. 1,8 V! Te wzmacniacze operacyjne zasadniczo usuwają ograniczenia, amplitudy sygnałów wejściowych i wyjściowych, które mogą przyjmować w dowolne wartości między szynami zasilającymi. Niektóre z nich obsługują nawet sygnały wejściowe wychodzące poza napięcia szyn zasilających, choć zwykle tylko w jednym kierunku (np. dodatnim) i o ograniczoną liczbę woltów. Należy pamiętać, że wzmacniacze operacyjne RRIO czasami obniżają jakość parametrów na inne sposoby, takie jak wyższy poziom szumów lub zniekształceń, lub po prostu kosztują więcej niż „zwykłe” wzmacniacze operacyjne.

Wiele wzmacniaczy operacyjnych

Gdy wzmacniacze operacyjne stały się tańsze i bardziej wszechstronne, popularne stały się podwójne i poczwórne wzmacniacze operacyjne. Oszczędzają one pieniądze i miejsce; układ scalony z poczwórnym wzmacniaczem operacyjnym często kosztuje mniej niż dwa razy tyle, co pojedynczy i wymaga tylko dwóch ścieżek zasilania do poprowadzenia i jednego kondensatora odsprężającego. Większość podwójnych (8-pinowych) i poczwórnych (14-pinowych) układów scalonych wzmacniaczy operacyjnych ma ten sam układ wyprowadzeń, dzięki czemu można je wymieniać.

Pojedyncze wzmacniacze operacyjne nie są tak wymienne, ponieważ są one zwykle dostarczane w 8-pinowej obudowie. Po uwzględnieniu dwóch szyn zasilających, dwóch wejść i jednego wyjścia, pozostałe trzy piny mogą być używane do regulacji offsetu, zewnętrznych kondensatorów kompensacyjnych lub różnych innych funkcji. Niektóre z nich są wymienne (nawet jeśli nie mają dokładnie tych samych funkcji), ale nie wszystkie. Obecnie pojedyncze wzmacniacze operacyjne są również dostępne w niewielkich 5-pinowych obudowach SMD, które sprawdzają się tam, gdzie przestrzeń jest na wagę złota.

Wnioski

Wzmacniacz operacyjny jest niezwykle elastycznym układem, dostępnym obecnie po bardzo niskich kosztach i w szerokiej gamie różnych wersji, zoptymalizowanych do różnych zadań. Choć możliwe jest przetwarzanie sygnałów analogowych bez użycia wzmacniaczy operacyjnych, generalnie wyniki będą gorsze. Dlatego też większość projektantów układów analogowych stosuje wzmacniacze operacyjne w swoich obwodach. Są one niezbędną cegiełką, bez której większość projektantów nie mogłaby się obejść. Dziękujemy Haroldowi S. Blackowi za ułatwienie nam życia! ■

Roderick Wall i Nicholas Vinen

KickStart

Część 3: Cewki indukcyjne, czym są?

Nasza okazjonalna seria KickStart ma na celu pokazanie czytelnikom, jak używać łatwo dostępnych, niedrogich komponentów i urządzeń do rozwiązywania szerokiej gamy typowych problemów, w możliwie najkrótszym czasie. Każdy z przykładów i projektów można wykonać w ciągu zaledwie kilku godzin, przy użyciu gotowych części. Oprócz krótkiego wyjaśnienia podstawowych zasad i zastosowanej technologii, seria ta zawiera szereg reprezentatywnych rozwiązań i przykładów, a także wystarczającą ilość informacji, aby umożliwić dostosowanie i rozszerzenie ich do własnego użytku.

Co to są cewki indukcyjne?

Są to elementy magazynujące energię w postaci pola magnetycznego i podobnie jak rezystory czy kondensatory zaliczane są do komponentów pasywnych. Kiedy prąd zmienny lub przemienny zostanie przyłożony do cewki indukcyjnej, zareaguje ona chwilowym magazynowaniem energii w polu magnetycznym, zwracając ją, z powrotem, w późniejszym czasie. To tymczasowe magazynowanie energii powoduje spowolnienie zmian prądu, która przez cewkę przepływa.

Kiedy prąd przemienny (np. przebieg sinusoidalny) jest przykładany do zacisków cewki indukcyjnej, efekt (w sensie przeciwstawienia się przepływowi prądu) jest określany jako „reaktancja indukcyjna”. Im szybciej zmienia się prąd, tym większa będzie reaktancja i odwrotnie. Stąd reaktancja indukcyjna jest wprost proporcjonalna do częstotliwości przyłożonego prądu i jest wyrażana wzorem:

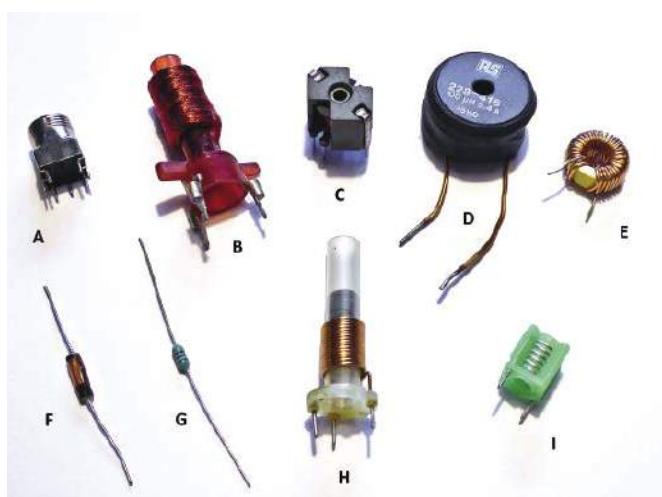
$$Z_L = 2 \cdot \pi \cdot f \cdot L$$

gdzie:

f – częstotliwość w hercach (Hz),

L – indukcyjność w henrach (H).

Podobnie jak rezystancja (która ściśle odnosi się do prądu stałego), reaktancja jest mierzona w omach (Ω) i jest definiowana jako stosunek przyłożonego napięcia przemiennego do przepływającego prądu. Cewka indukcyjna 100 mH będzie wykazywać reaktancję około 63 Ω przy 100 Hz, 630 Ω przy 10 kHz i 6,3 k Ω przy 100 kHz.



Rysunek 3.1. Typowy wygląd małych cewek indukcyjnych

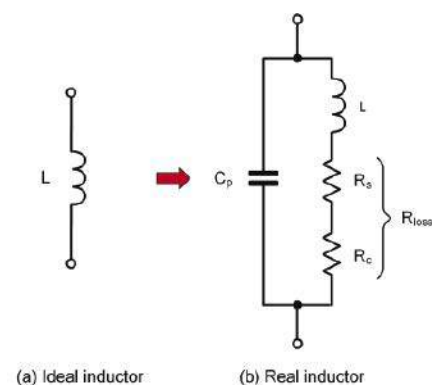
Cewki indukcyjne są używane w zasilaczach do odfiltrowania prądu przemiennego (składowej zmiennej) z prądu stałego, w którym to przypadku są czasami określane jako „dławiki”. W systemach audio i komunikacyjnych, cewki indukcyjne są często stosowane w filtrach i innych obwodach selektywnych częstotliwościowo. Cewki indukcyjne są również używane w aplikacjach do tłumienia szumów i zakłóceń w szerokim zakresie częstotliwości, od sieci (50/60 Hz) po częstotliwości radiowe.

Najbardziej podstawową formą cewki powietrznej jest nic innego jak zwojnica z izolowanego drutu nawinięta na nieprzewodzący i niemagnetyczny karkas. Większe wartości indukcyjności można uzyskać stosując rdzenie wykonane z materiału ferromagnetycznego, takiego jak stal (często laminowana) lub ferryt (materiał ceramiczny o właściwościach magnetycznych).

Cewki indukcyjne są szeroko dostępne w handlu w wielu typowych wartościach indukcyjności, ale można również stosunkowo łatwo wykonać je w domu, chociaż czasami może to być proces tworzenia na chybił trafił. Zanim to wyjaśnimy, spójrz na typowy wybór małych cewek indukcyjnych, które możesz napotkać w elektronice, pokazany na **rysunku 3.1** i opisany w **tabeli 3.1**.

Tabela 3.1. Krótkie zestawienie parametrów cewek z rysunku 3.1

Poz.	L	Opis
A	3,5 mH	Regulowana cewka ferrytowa do montażu PCB
B	2,8 mH	Dławik do zastosowań radiowych niskiej częstotliwości
C	1 mH	Mała cewka ferrytowa z rdzeniem kubkowym do zastosowań ogólnych
D	100 μ H	Duża ferrytowa cewka indukcyjna z rdzeniem kubkowym do zastosowań wysokoprądowych
E	100 μ H	Stała cewka toroidalna do stosowania w filtrach
F	33 μ H	Dławik osiowy zakończony drutem do zastosowań RF
G	3,9 μ H	Mała cewka indukcyjna z wyprowadzeniami osiowymi do zastosowań ogólnych
H	2 μ H	Regulowana cewka indukcyjna z rdzeniem ferrytowym do użytku w obwodach strojonych RF
I	300 nH	Cewka z rdzeniem powietrznym do zastosowań radiowych VHF i UHF



Rysunek 3.2. Schemat „idealnej” i „rzeczywistej” cewki indukcyjnej



Rysunek 3.3. Autorska cewka toroidalna 1,1 mH

Tabela 3.2 Porównanie dwóch cewek indukcyjnych

Parametr	Cewka A	Cewka B
L	100 mH	100 mH
C _p	11 pF	18 pF
R _s	1 Ω	5 Ω
R _c	29 Ω	95 Ω
Q	209	63

Kiedy cewka nie jest cewką?

Niestety, w porównaniu z rezystorami, czy nawet kondensatorami, cewki indukcyjne rzadko zachowują się idealnie, co może być ważne przy ich projektowaniu i stosowaniu. Na **rysunku 3.2** został pokazany symbol „idealnej” cewki indukcyjnej wraz z jej rzeczywistym odpowiednikiem, który uwzględnia straty i pojemność rozproszoną, występujące w „rzeczywistych” komponentach. Podstawowe cechy pokazane na rysunku 3.2(b) obejmują:

- L – indukcyjność określoną przez liczbę zwojów i właściwości magnetyczne materiału rdzenia.
- C_p – pojemność rozproszona (tj. pojemność wypadkowa składająca się z tej między sąsiednimi zwojami, przewodami łączącymi i/lub zaciskami). Pojemność ta może stanowić problem przy wysokich częstotliwościach, gdzie reaktancja pojemnościowa może spaść do stosunkowo niskiej wartości w porównaniu z indukcyjną.
- R_{loss} – całkowita rezystancja strat cewki indukcyjnej. Jest to suma rezystancji strat uzwojenia miedzianego (R_s) i rezystancji strat rdzenia (R_c).
- R_s – rezystancja strat uzwojenia miedzianego, którą można zmierzyć między zaciskami cewki indukcyjnej za pomocą omomierza.
- R_c – rezystancja rdzenia wynika z ciągłego procesu magnesowania i rozmagne-sowywania rdzenia podczas każdego cyklu przyłożonego prądu przemiennego. Należy zauważyć, że w przypadku dużych cewek indukcyjnych R_c może być znacznie większa niż R_s.

Dobroć cewki

Dobroć cewki (Q – Quality), to stosunek jej efektywnej reaktancji przy częstotliwości

roboczej, do jej rezystancji strat, może być zatem traktowana jako wskaźnik tego, jak „dobry” jest to komponent. Współczynnik Q można obliczyć z wzoru:

$$Q = \frac{X_L}{R_{loss}} = \frac{2 \cdot \pi \cdot f \cdot L}{R_S + R_C}$$

Współczynnik Q jest ważny w wielu zastosowaniach i warto zilustrować to przykładem. Załóżmy, że mamy dwie cewki o właściwościach pokazanych w **tabeli 3.2**. Współczynnik Q dla każdego komponentu został obliczony dla częstotliwości roboczej 10 kHz – zwróć uwagę na duże obniżenie dobroci ze względu na stosunkowo niedoskonały materiał rdzenia użyty do produkcji komponentu B.

Indukcyjność

Wartość indukcyjności cewki L, jest określona przez stałą magnetyczną materiału AL i ilość zwojów podniesioną do kwadratu. Stała AL jest z kolei determinowana efektywną przenikalnością (μe) materiału rdzenia, jego wymiarami fizycznymi oraz budową. Stała magnetyczna danego rdzenia jest zwykle podawana w nH/zw² i można ją znaleźć na podstawie opublikowanych danych producentów. Wartość indukcyjności można obliczyć korzystając z zależności:

$$L = n^2 \cdot AL [nH]$$

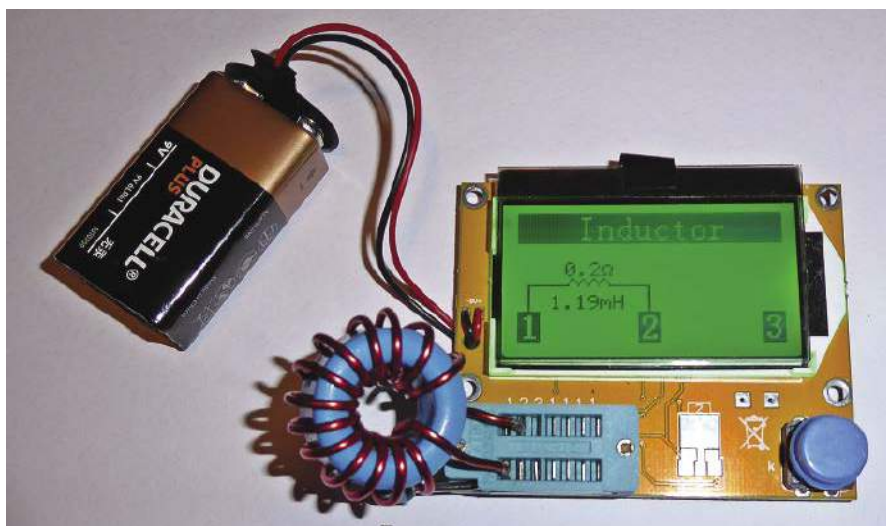
gdzie:

- AL to stała magnetyczna rdzenia podana w nanohenrach (1 nH=10⁻⁹ H) na zwój do kwadratu. Użyjemy tej formuły w następnym rozdziale.

Produkcja cewek toroidalnych

Na szczęście cewki indukcyjne o wartościach w zakresie od około 10 μH do 100 mH bardzo łatwo można wyprodukować przy użyciu standardowych ferrytowych rdzeni toroidalnych. Jednak przy zakupie odpowiedniego rdzenia może być konieczne rozważenie kilku właściwości, w tym:

- materiał rdzenia i jego właściwości magnetyczne,
- fizyczne wymiary rdzenia (zwłaszcza zewnętrzna i wewnętrzna średnica oraz grubość),
- wymagana obciążalność prądowa i maksymalna dopuszczalna rezystancja uzwojenia (te z kolei określają wymaganą grubość drutu nawojowego),
- liczba zwojów wymagana dla uzyskania określonej wartości indukcyjności (należy pamiętać, że wymiary fizyczne rdzenia ograniczają maksymalną liczbę zwojów, które można zastosować, przy danej średnicy drutu).



Rysunek 3.4. Sprawdzanie cewki indukcyjnej 1,1 mH za pomocą niedrogiego testera elementów

Wykaz elementów, kupuj w sklepie:
(W-wa, ul. Leszczyńska 11, tel. +48222578451,
e-mail: handlowy@avt.pl):

Rezystory: (wszystkie stałe są na 0,25 W, 5%)

R1: 100 kΩ

R2: 47 Ω

R3: 47 kΩ

R4: 10 kΩ

R5: 2,2 kΩ

R6, R7: 1 kΩ

RV1: 200 kΩ (220 kΩ) potencjometr montażowy

Kondensatory:

C1, C2, C5: 100 nF/100 V miniaturowe poliestrowe (patrz tabela 3.3)

C3, C4: 220 μF/16 V elektrolityczne radialne

Układy scalone:

IC1: TL082 8-pin podwójny wzmacniacz operacyjny (lub lepiej TL072)

Pozostałe:

Jednostronna płytka drukowana 74×50 mm, KS3-2021, dostępna w PE PCB Service

P1: listwa 8 szpilek, 2,54 mm

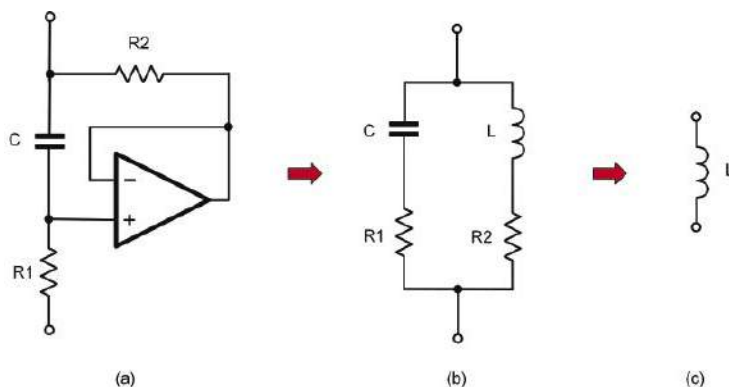
P2: listwa 3 szpilki, 2,54 mm

Zwora dwudrożna do P2

Podstawa DIL8 pod IC1

B1: 9 V PP3 (6F22 lub 6LR61)

Złącze do baterii 9 V



Rysunek 3.5. Elektroniczny ekwiwalent cewki indukcyjnej na bazie żyrtora – zwróć uwagę, że w obu obwodach $R1 \gg R2$, a $R1$ jest uziemiony. Dzięki wysokiemu wzmocnieniu i ujemnemu sprzężeniu zwrotnemu wzmacniacza operacyjnego impedancja efektywna patrząc na zacisk $R2/C$ wynosi: w przybliżeniu $R2 + j\omega L$, gdzie $L = R1 \times R2 \times C$

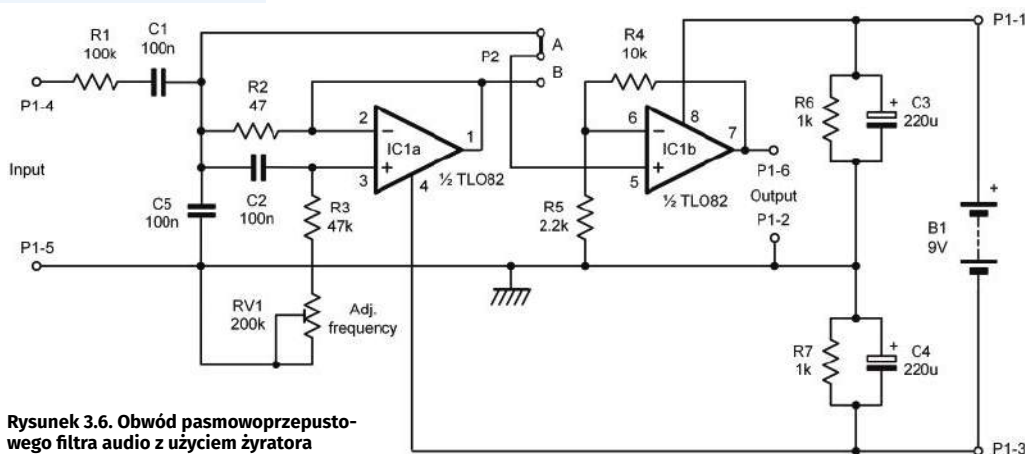
Proces projektowania najlepiej pokazać na prostym przykładzie. Założymy, że potrzebujemy cewki indukcyjnej o określonych parametrach:

- indukcyjność 1,1 mH,
- obciążalność prądowa 5 A,
- częstotliwość robocza 300 kHz,
- maksymalna rezystancja uzwojenia 0,2 Ω,
- do zastosowania w wysokoprądowym zasilaczu impulsowym. Wybraliśmy gotowy rdzeń toroidalny wyprodukowany przez TDK z materiału N30 SIFFERIT. Odnosząc się do danych producenta, rdzeń ma deklarowaną przenikalność (μ) 4300 i wynikającą z tego stałą magnetyczną rdzenia AL równą 5460 (w nH/zw²). Przekształcenie formuły, którą poznaliśmy wcześniej, pozwala nam obliczyć wymaganą liczbę zwojów, w następujący sposób:

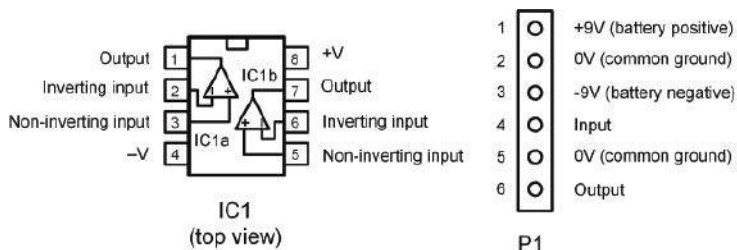
$$n = \sqrt{\frac{L}{AL}}$$

gdzie:

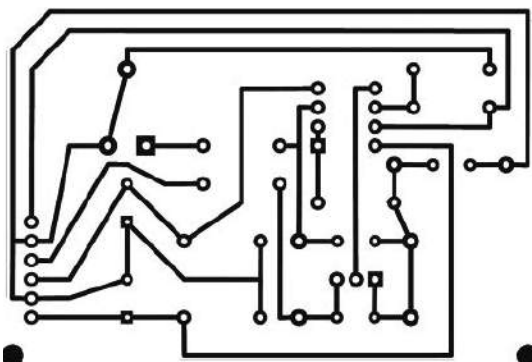
- L to wymagana wartość indukcyjności (w henrach),
- AL to stała rdzenia (5460 w przypadku rdzenia N30).



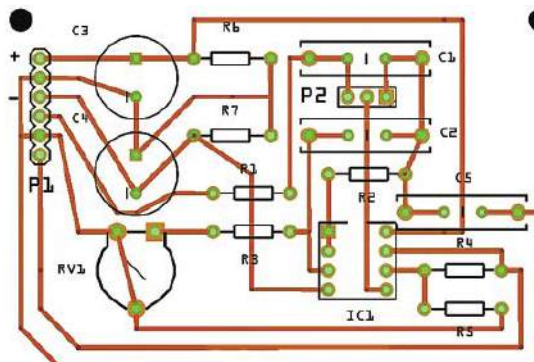
Rysunek 3.6. Obwód pasmowoprzepustowego filtra audio z użyciem żyrtora



Rysunek 3.9. Przyporządkowanie styków złącza na płytce filtru



Rysunek 3.7. Układ ścieżek płytki dla filtra pasmowoprzepustowego audio



Rysunek 3.8. Opisy na stronie elementowej płytki filtru

Temat	Źródło	Uwagi
Cewki	Centrum produktów TDK szczegółowo opisuje szeroką gamę komponentów indukcyjnych: http://bit.ly/pe-jun21-tdk Witryna Coilcraft zawiera przydatne wprowadzenie do cewek indukcyjnych i dławików: http://bit.ly/pe-jun21-cc	Publikacja „TDK Inductor's World” zapewnia zabawne i przystępne wprowadzenie do cewek indukcyjnych. Zobacz http://bit.ly/pe-jun21-tdk2 lub pobierz ze strony internetowej PE z czerwca 2021 r. Witryna Coilcraft zawiera również przydatne uwagi aplikacyjne dotyczące projektowania filtrów L-C.
Indukcyjność i dobroć	Własna książka autora, Electronic Circuits: Fundamentals and Applications (5th Ed, 2020 opublikowana przez Routledge 9780367421984) lub P. Horowitz „Sztuka elektroniki”	Ogólne wprowadzenie do napięcia i prądu przemiennego, wraz z rozdziałami dotyczącymi indukcyjności, współczynnika Q i obwodów rezonansowych.
Tina-TI symulator	Oprogramowanie do symulacji analogowej oparte na Tina-TI SPICE można bezpłatnie pobrać ze strony Texas Instruments pod adresem: www.ti.com/tool/TINA-TI	Pełną wersję 11 Tina można kupić i pobrać ze strony Practical Electronics strona internetowa. Budżetowe wersje oprogramowania są dostępne dla studentów i hobbystów.
Filtry	Filtry Część 8 Electronics Teach-in 4 (dostępna w PE Magazine)	Zawiera ogólne wprowadzenie do zastosowań obwodów analogowych, w tym filtrów pasywnych i aktywnych.
TL082	Arkusz danych TL082 (TL072) firmy Texas Instruments	Szczegóły dotyczące wzmacniacza TL082.

Stąd wymagana liczba zwojów jest określona wzorem:

$$n = \sqrt{\frac{1,1 \cdot 10^{-3}}{5460 \cdot 10^{-9}}} = \sqrt{\frac{1100}{5,46}} = \sqrt{201,5} \approx 14 \text{ zwojów}$$

Rdzeń ma wymiary 35,5×13,6 mm i z łatwością pomieści wymaganą liczbę zwojów. Aby bezpiecznie przewodzić wymagany prąd i zminimalizować straty w uzwojeniu miedzianym, drut musi mieć odpowiedni przekrój i wynikającą z tego średnicę. Wybrana średnica (1,5 mm) emaliowanego drutu miedzianego ma rezystancję mniejszą niż 0,01 Ω/m, a zatem spadek napięcia wynikający z ciągłego prądu stałego o natężeniu 5 A i całkowitej długości uzwojenia wynoszącej 500 mm powinien być mniejszy niż 25 mV. Więcej informacji na temat wyboru odpowiedniego drutu znajduje się w rozdziale *Idąc dalej z cewkami*.

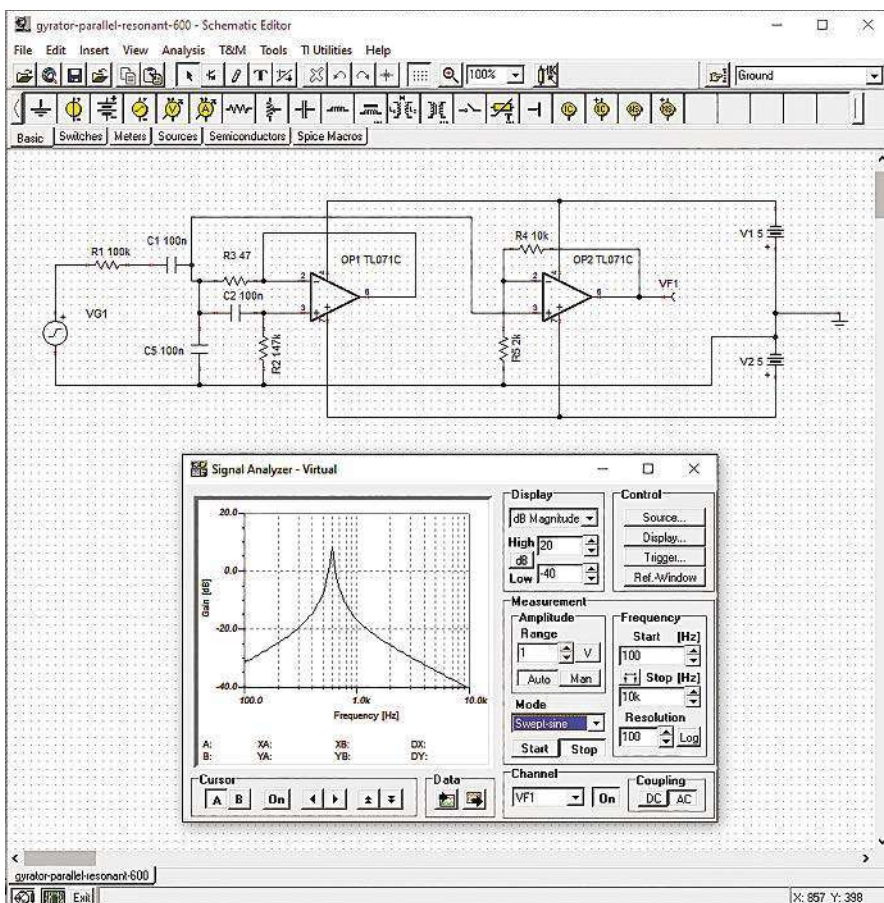
Wartość indukcyjności można sprawdzić za pomocą miernika z funkcją pomiaru indukcyjności lub za pomocą analizatora sieci ewentualnie mostka prądu zmiennego. Alternatywnie można zastosować niedrogi tester komponentów, taki jak ten pokazany na **rysunku 3.4**, aby uzyskać wskazanie wystarczające do potwierdzenia, że osiągnięto wymaganą wartość indukcyjności. Tolerancja do około ±10% jest zwykle akceptowalna dla większości zastosowań, ale zawsze dobrze jest sprawdzić kalibrację przyrządu testowego przy użyciu znanego komponentu, przed przystąpieniem do pomiaru.

W tym momencie kilka ostrzeżeń jak coś pójdzie nie tak. Po pierwsze, przenikalność



Rysunek 3.10. Widok zmontowanej płytki filtra

ferrytu i innych stopów ferromagnetycznych ma tendencję do bardzo gwałtownego spadku w wysokich temperaturach (tj. powyżej około 125°C). Powyżej tego punktu właściwości magnetyczne materiału mogą bardzo szybko zanikać. Po drugie, nie jest niczym niezwykłym, że rdzenie nagrzewają się podczas pracy, zwłaszcza gdy występują duże gęstości strumienia wywołane prądem o wysokiej częstotliwości. Może to dodatkowo zaostrzyć zmniejszenie przenikalności i poważnie ograniczyć wydajność rdzenia. We wszystkich przypadkach warto odwołać się do danych producentów, sprawdzając zarówno charakterystykę przenikalności, jak i zalecany zakres częstotliwości pracy rdzenia.



Rysunek 3.11. Symulacja filtra z żyratorem w programie Tina-TI SPICE

Tabela 3.3 Zalecane wartości C1, C2 i C5 dla różnych zakresów częstotliwości audio

f_{nom}	C	Zakres regulacji
200 Hz	680 nF	180 Hz do 400 Hz
300 Hz	470 nF	220 Hz do 480 Hz
400 Hz	220 nF	350 Hz do 700 Hz
600 Hz	100 nF	500 Hz do 1 kHz
800 Hz	68 nF	600 Hz do 1,2 kHz
1 kHz	47 nF	700 Hz do 1,5 kHz
1,2 kHz	33 nF	900 Hz do 1,8 kHz
1,4 kHz	22 nF	1 kHz do 2 kHz
1,8 kHz	10 nF	1,5 kHz do 3 kHz

Symulacja cewek indukcyjnych

Niestety, czasami może być trudno wyprodukować stosunkowo duże cewki indukcyjne potrzebne do zastosowań audio i niskoczęstotliwościowych. Istnieje jednak potencjalne rozwiązanie w postaci obwodu, który może symulować cewkę indukcyjną za pomocą samych kondensatorów i rezystorów wraz z aktywnym elementem, takim jak tranzystor lub wzmacniacz operacyjny. Układ taki nazywa się „żyratorem” i działa jak lustro, odwracające charakter impedancji obecnej na jego wejściu. W naszym przypadku zamienia impedancję pojemnościową na indukcyjną. Obwód ekwiwalentnej cewki bazującej na żyratorze pokazano na **rysunku 3.5**. Należy zauważyć, że podobnie jak w przypadku większości obwodów żyratorowych, ten układ oczekuje, że będzie zasilany ze źródła o niskiej impedancji względem 0 V.

Filtr pasmowoprzepustowy audio bazujący na żyratorze

Przyjrzyjmy się praktycznemu zastosowaniu żyratora w postaci przestrzajalnego filtra pasmowoprzepustowego audio, w którym żyrator zastępuje zmienną cewkę indukcyjną o wartości kilkuset milihenrów (mH). **Rysunek 3.6** pokazuje kompletny schemat filtra, w którym IC1a działa jako wzmacniacz nieodwracający o wzmocnieniu jednostkowym dla żyratora, podczas gdy IC1b jest nieodwracającym buforem wyjściowym skonfigurowanym na skromne wzmocnienie napięciowe (ok. 5,5 V/V).

Przy podanych wartościach elementów składowych, filtr pokrywa nominalny zakres częstotliwości od 500 Hz do 1,5 kHz i jest przestrzajany za pomocą RV1 (co skutecznie zmienia wartość indukcyjności symulowanej przez żyrator). Podczas normalnej pracy zworka P2 jest zwykle ustawiona w pozycji A, ale bezpośrednie wyjście z IC1b może być użyte poprzez umieszczenie zworki w pozycji B.

Impedancja wejściowa filtra (100 kΩ) jest wyznaczona przez R1, podczas gdy C1 działa jak kondensator sprzęgający, odcinający składową stałą. Strojony obwód filtra składa się z kondensatora C5 wraz z symulowaną indukcyjnością żyratora składającego się z IC1b, C2, R2 i kombinacji szeregowej RV1 i R3. Symetryczne napięcie zasilania dla IC1 uzyskuje się ze standardowej baterii 9 V (6F22 lub 6LR61) i układu dzielnika napięcia zawierającego R6 i R7 wraz z kondensatorami odsprężającymi, odpowiednio C3 i C4.

Zmierzone parametry filtra są następujące:

- przybliżone wzmocnienie napięciowe: 6 dB,
- nominalna częstotliwość środkowa: 600 Hz,
- zakres regulacji: 450 Hz (min.) do 1,09 kHz (maks.),
- częstotliwość środkowa: 567 Hz (w środkowej pozycji RV1),
- szerokość pasma: 60 Hz przy -3 dB (przy 600 Hz),

- impedancja wejściowa: 100 kΩ,
- impedancja wyjściowa: mniejsza niż 100 Ω,
- zasilanie: 9 V DC przy poborze prądu poniżej 9 mA.

Układ ścieżek na płycie drukowanej i warstwę opisową na stronie elementów dla filtra audio bazującego na żyratorze pokazano odpowiednio na **rysunkach 3.7 i 3.8**. Przyporządkowanie pinów dla IC1 i P1 pokazano na **rysunku 3.9**, a gotową płytkę na **rysunku 3.10**. Należy pamiętać, że dla C1 i C2 przewidziano kilka zestawów pól lutowniczych, dzięki czemu można dopasować komponenty o różnych rozmiarach

Od red. EdW: do celów akustycznych lepiej jest zastosować układ TL072, który wyróżnia się niższymi szumami.

Tabela 3.3 zawiera zalecane wartości C1, C2 i C5 dla różnych zakresów częstotliwości audio.

Jeśli obwód wykazuje tendencję do oscylacji bez podłączonego sygnału wejściowego, to można temu zaradzić, podłączając diodę 1N4148 równolegle do C5.

Od red. EdW: zwiększając wartość R2 można zmniejszyć dobroć Q symulowanej cewki ale wtedy trzeba dostosować indukcyjność do wymaganej wartości przy pomocy RV1. Zasada jest prosta: R2 rośnie, RV1 maleje, by iloczyn $R2 \cdot (R3 + RV1) \cdot C2$ pozostał niezmienny.

Diodę można łatwo dodać do spodu płytki i przylutować bezpośrednio do pól lutowniczych kondensatora C5.

Wreszcie, **rysunek 3.11** pokazuje symulację filtra audio bazującego na żyratorach, pokazanego na rysunku 3.5. Ta symulacja została stworzona przy użyciu bezpłatnej wersji popularnego pakietu Tina-TI SPICE (patrz podrozdział *Idąc dalej...*), a czytelnicy mogą pobrać tę aplikację ze strony: www.ti.com/tool/TINA-TI

Zwróć uwagę, jak wirtualny analizator sygnału Tiny wykreślił odpowiedź filtra i pokazuje ostry szczyt przy 600 Hz wraz z szybkim spadkiem po obu stronach.

Idąc dalej z cewkami

Ta sekcja zawiera szczegółowe informacje na temat różnych źródeł, które pomogą Ci zlokalizować części składowe oraz dalsze informacje, które pozwolą Ci nawijać cewki indukcyjne i wykorzystywać je we własnych aplikacjach. Zawiera również łącza do przydatnych stron internetowych, które pomogą Ci szybko zapoznać się z podstawową wiedzą na kluczowe omawiane tematy. ■

Mike Tooley

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, czerwiec 2021 (www.epemag3.com)

REKLAMA

ELMAX 1988

Certyfikat Underwriters Laboratories

№ BA1-0 480348 TYPE 1

Zakład produkcyjny:

05-660 Warszawa
ul. M. Rópelewskiej 17
tel. 22 781 63 95
22 781 05 40
fax. 22 781 63 95 w.23
www.elmax.waw.pl
elmax@elmax.waw.pl

OBWODY DRUKOWANE

Produkcja, Projektowanie, Montaż

Płytki jednostronne	Serie dowolne	Dokumentacja technologiczna	Montaż elektroniczny
Płytki dwustronne	Prototypy	Dokumentacja konstrukcyjna	Ilości modelowe produkcyjne
Płytki na podłożu aluminium	Maksymalny wymiar płytek 1w 630 mm		
Aktywny kalkulator prototypów na stronie internetowej	Pokrycie Sn lub SnPb line na życzenie	Płyty czosnek FR4	Krótkie terminy
	Maski, opisy montażowe w różnych kolorach	Trawione szablony SMD	Wykonania super ekspresowe

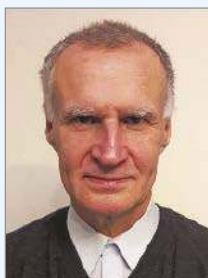
Uczmy się na cudzych błędach

Celem tej rubryki jest kształtowanie u Czytelników EdW umiejętności krytycznego czytania schematów i opisów projektów autorskich. Wszyscy jesteśmy omylni. Konstruktorzy projektów elektronicznych też. W projektach publikowanych w Internecie, ale też w artykułach drukowanych zdarzają się błędy różnej wagi, w tym też takie, które sprawiają, że układ nie może działać prawidłowo. Uczmy się wykrywać te błędy na przykładach projektów sprawdzonych w naszym redakcyjnym Pokoju Nauczycielskim.

Pamiętajmy! Nie oceniamy Autorów, tylko uczymy się na cudzych błędach.

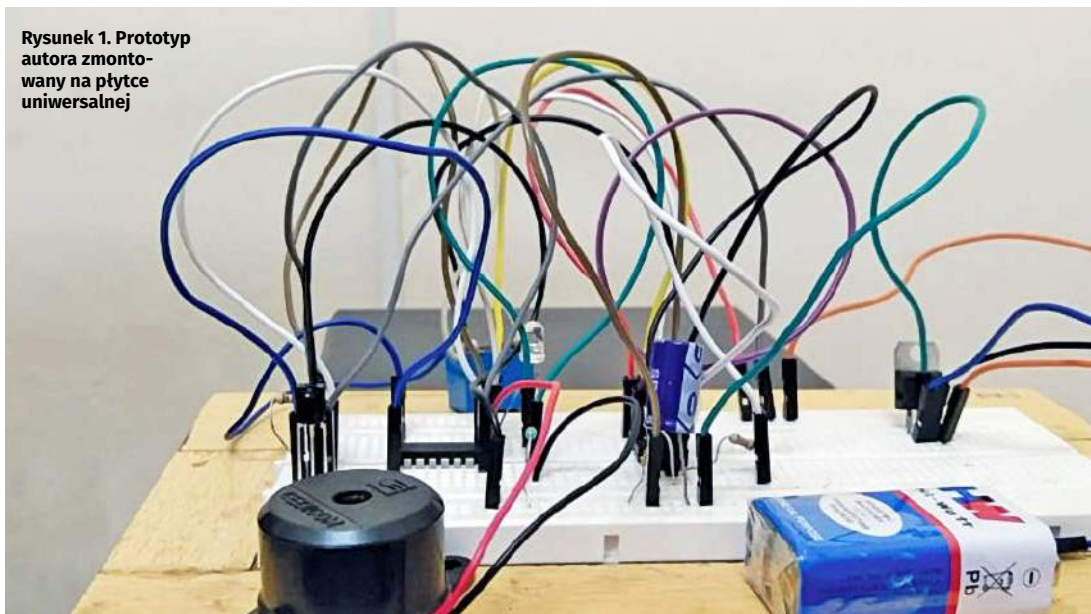
Zapraszamy Czytelników do współpracy z naszym Pokojem Nauczycielskim. Jeśli natrafiłicie w Internecie lub źródłach drukowanych na opis projektu z poważnymi Waszym zdaniem błędami, to przysyłajcie takie opisy do naszej redakcji (redakcja@elportal.pl w tytule wiadomości: Pokój Nauczycielski) wraz z Waszymi uwagami.

Projekt sprawdza i poprawia Karol Świerc



Mgr inż. elektronik – absolwent Wydziału Automatyki i Informatyki Politechniki Śląskiej z 1980 roku. Przez 25 lat prowadził serwis RTV. Mówi o sobie: „z elektroniką łączy mnie związek „z rozsądku”, moją pierwszą miłością była matematyka i fizyka”. Autor wielu artykułów publikowanych w EdW.

Rysunek 1. Prototyp autora zmontowany na płycie uniwersalnej

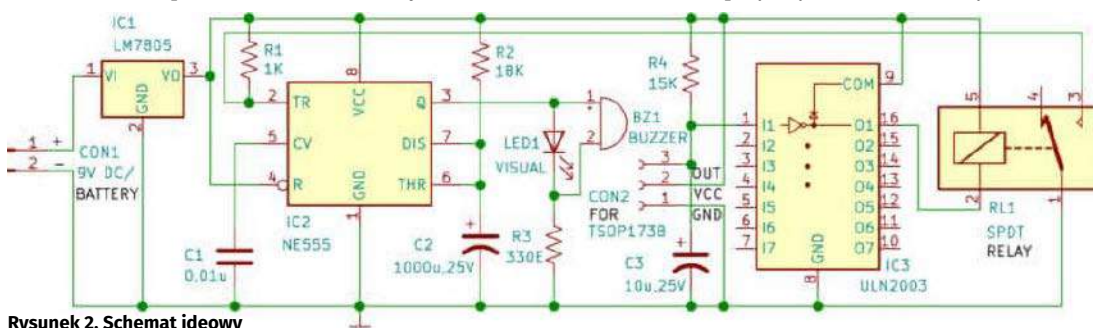


Zdalnie sterowany sygnalizator na czas pandemii

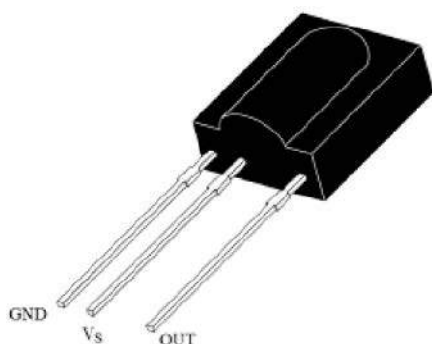
Zaprezentowany projekt wychodził na przeciw zaleceniom Światowej Organizacji Zdrowia WHO na czas pandemii koronawirusa. Wśród zaleceń wymieniono staranne mycie rąk, czyli takie, gdy przez co najmniej 20 sekund wykonujemy dokładne namydlenie i dopiero potem spłukujemy pod bieżącą wodą. Dwadzieścia sekund może wydawać się czasem niedługim, ale trzeba o tym na bieżąco przypominać. Czas ten można odmierzyć zegarkiem na ręce lub stoperem w smartfonie, ale jest to niewygodne i niepraktyczne. Prezentujemy układ, który pozwoli kontrolować poprawność mycia rąk w sposób automatyczny.

Zaproponowany układ zasygnalizuje, iż upłynął zadany przedział czasowy. Sygnalizacja jest dwójakiego rodzaju, diodą LED oraz (dla tych którzy preferują sygnał akustyczny) „brzęczykiem”. Prototyp wykonany przez autora widzimy na **rysunku 1**. Schemat ideowy tego „alarmu” pokazuje **rysunek 2**.

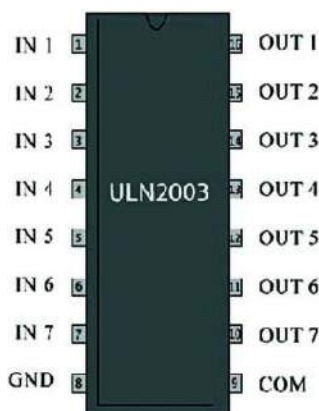
W układzie zastosowano baterijkę 9 V, stabilizator 5 V typu LM7805 (IC1), timer NE555 (IC2), driver przełącznika ULN2003 (IC3), przełącznik SPDT z cewką 5 V (RL1) oraz kilka elementów pasywnych. Czas timera wyznacza stała



Rysunek 2. Schemat ideowy



Rysunek 3. Rozmieszczenie wyprowadzeń odbornika TSOP1738



Rysunek 4. Wyprowadzenia układu scalonego ULN2003

Wykaz elementów, kupuj w sklepie.avt.pl
(W-wa, ul. Leszczyńska 11, tel. +48222578451,
e-mail: handlowy@avt.pl):

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

R1: 1 kΩ

R2: 18 kΩ

R3: 330 Ω

R4: 15 kΩ

Kondensatory:

C1: 0,01 μF ceramiczny

C2: 1000 μF/25 V elektrolityczny

C3: 10 μF/25 V elektrolityczny

Półprzewodniki:

IC1: LM7805 – stabilizator 5 V

IC2: NE555 – timer

IC3: ULN2003 – driver przełącznika

Inne:

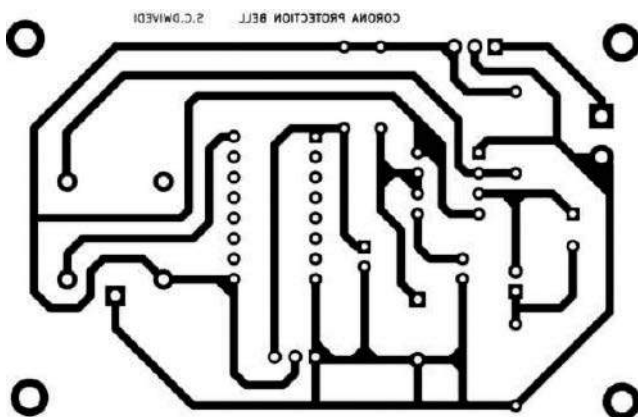
CON1: złącze 2-pinowe

BZ1: buzzer

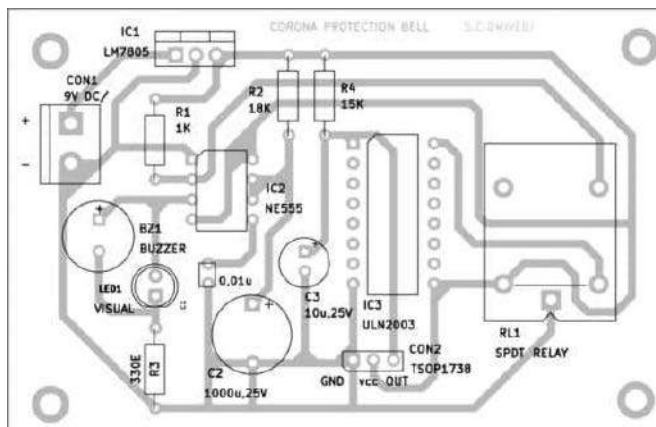
RL1: przełącznik SPDT 5 V

TSOP1738 odbiornik pilota podczerwieni

Bateria 9 V



Rysunek 5. Schemat płytki PCB



Rysunek 6. Schemat ułożenia elementów na PCB

czasowa R2-C2 i powinien wynosić ok. 20 sekund. W timerze skonfigurowanym jako przerzutnik monostabilny, obowiązuje relacja $T=1,1RC$. W naszym przypadku $R=R2=18\text{ k}\Omega$ i $C=C2=1000\text{ }\mu\text{F}$. Więc $T=1,1 \times 18\text{ k}\Omega \times 1000\text{ }\mu\text{F} = \text{ok. } 20\text{ sek.}$

W układzie zastosowano popularny odbiornik sygnałów z pilotów RTV – TSOP1738, pracujący w zakresie podczerwieni, wyposażony w wyjście z aktywnym stanem niskim. Jeśli odbiornik zostanie oświetlony (w podczerwieni) przez naciśnięcie dowolnego przycisku pilota, stan niski zostanie podany na wejście I1 układu ULN2003 – drivera przełącznika. Za pośrednictwem przełącznika zostanie wyzwolony przerzutnik zbudowany na układzie scalonym NE555. W rezultacie wyjście Q układu IC2 przyjmie stan wysoki na czas ok. 20 sek, co zaświeci diodę LED1 i włączy sygnał dźwiękowy buzzera.

Seria odbiorników TSOP17XX pracuje na różnych częstotliwościach nośnych (co jest zawarte w oznaczeniu XX układu scalonego). Elementy te współpracują z pilotami urządzeń powszechnego użytku, takimi jak telewizory, odtwarzacze DVD itp. TSOP1738 przeznaczony do sygnału o częstotliwości 38 kHz jest jednym z popularniejszych. Układ ten zawiera diodę PIN i przedwzmacniacz oraz filtr dla bezbłędnej detekcji sygnału PCM. TSOP17XX zamknięty jest w obudowie o trzech wyprowadzeniach – ich rozmieszczenie i funkcje pokazuje rysunek 3. Trzecim układem scalonym jest driver przełącznika ULN2003 (IC3). Układ wyprowadzeń tego IC widzimy na rysunku 4. Tutaj zastosowano tylko jeden tor, który steruje bezpośrednio 5-woltową cewką przełącznika RL1.

Działanie urządzenia wymaga jeszcze pilota zdalnego sterowania, który powinien być umocowany lub umieszczony w pobliżu umywalki. Chcąc umyć ręce należy najpierw nacisnąć dowolny przycisk pilota,

REKLAMA

KEY PRODUCENT AUTOMATYKI GRZEWCZEJ
11-200 Bartoszyce ul. Bohaterów Warszawy 67 pwkey@onet.pl
tel. (89)7635050 fax (89)7635051

TANIE REGULATORY

DO KOTŁÓW WĘGLOWYCH I NA DREWNO

z wbudowanym termostatem pokojowym
zapewniającym komfort i oszczędność



REGULATORY DO KOTŁÓW Z PODAJNIKIEM

REGULATORY POGODOWE

- Prosta obsługa, bogate możliwości programowania
- Możliwość dopasowania do każdego kotła i rodzaju paliwa
- Wysoka jakość
- Gwarancja 24 miesiące

www.pwkey.pl

a świecenie diody LED i dźwięk buzzera wskaże jak długo powinien się myć ręce zanim opłuczesz je wodą. Na **rysunku 5** pokazano projekt jednostronnej płytki PCB dla tego układu. **Rysunek 6** pokazuje ułożenie elementów na tej płycie.

Po zmontowaniu układu należy go zamknąć w odpowiednio przygotowanej obudowie. Z przodu należy umieścić odbiornik TSOP1738, diodę LED i złącze CON1. Wewnątrz, oprócz PCB należy zamontować

baterijkę 9 V, a buzzer również wewnątrz lub na tylnej ścianie obudowy. Przed odbiornikiem TSOP1738 trzeba wykonać okienko ułożone tak, aby strumień podczerwonego światła diody nadajnika je oświetlał. ■

Rakesh Jain

Ten artykuł był wcześniej opublikowany na łamach „EFY”, styczeń 2023 (efymag.com)

Uwagi i poprawki

Od Redakcji EdW:

Obstrzeżenia zostały już dawno zniesione. Jednak pomijając sens tego układu, można zadać pytanie: na ile sposobów można wykonać 20-sekundowy wyłącznik monostabilny? Pewnie setki lub tysiące, ale bardziej fikuśnego i skomplikowanego rozwiązania niż to zaprezentowane przez autora, trudno byłoby wymyślić.

W roli monoflopa pracuje „samotny” 555, a wyzwolić można go wiele prościej. Autor wymyślił wyzwolenie z pilota (choć w tym zastosowaniu nie będzie to zbyt wygodne). W takim razie użycie odbiornika TSOP jest raczej obligatoryjne. Wystarczyłoby między wyjście odbiornika a Trigger timera 555 dać zwykłą diodę (a i bez niej też układ powinien działać). W takim razie, po co komplikować dodatkowym układem ULN2003 i przełącznikiem? Skoro już mamy przełącznik, to prosi się, żeby zastosować go przynajmniej do sterowania zewnętrznym obciążeniem. Buzzer i diodę LED autor zasilą z wyjścia Q timera 555. Wyjście to ma niewielką obciążalność, szczególnie przy zasilaniu 5-voltowym.

Użycie timera 555 w roli przerzutnika monostabilnego jest wręcz katalogowe. Odmierzany czas ustala stała czasowa R2-C2 przemnożona przez ln3 (logarytm naturalny z 3), czyli tak jak autor pisze – współczynnik ok. 1,1.

Wejście Trigger jest aktywne w stanie niskim, co w tym przypadku wprowadza komplikacje. Układ ULN2003 to 7 tranzystorów Darlingtona (zauważmy na rysunku 4, że ten układ scalony nie ma zasilania). W układzie używany jest tylko jeden tranzystor, ale co gorsza, stopień ten wykazuje negację sygnału. TSOP1738 ma wyjście aktywne stanem niskim, a więc żeby dostosować się do Trigger timera, potrzebna jest jeszcze jedna negacja. Realizuje ją przełącznik – styki NC, ale to oznacza, że przez cewkę przełącznika cały czas

plynie prąd, a jedynie na krótką chwilę wyzwolenia przełącznik „puszcza”. Obciążalność wyjść ULN2003 wynosi 0,5 A, i powiedzmy, że zastosowano taki przełącznik, że to wystarcza. Jest to rozwiązanie (łagodnie mówiąc) niedobre. Szczególnie zważywszy na fakt, iż układ jest zasilany z baterijki. A jeśli tak, to i stabilizator 7805 też nie ma tu wielkiego sensu. ULN2003 i NE555 mogą pracować przy 9 V, a czas odmierzany przez timer jest praktycznie niezależny od zasilania (jeśli nawet by miało znaczenie dokładne odmierzenie 20 sekund). W razie potrzeby prąd cewki przełącznika można by ograniczyć niewielkim rezystorem.

Powyższe uwagi nie są stricte błędami, ale należy je rozpatrywać jako „usterki”. Można by jeszcze pomyśleć, że błędem jest brak diody zwrotnej równoległej do cewki przełącznika. To akurat jest w porządku. Taka dioda mieści się w układzie scalonym ULN2003. Natomiast zastosowanie tu „całego ULN” może sugerować jedynie, iż „mam taki w szufladzie i chcę koniecznie go jakoś wykorzystać”. Wątpliwości może też budzić równoległe połączenie buzzera i diody LED ze wspólnym rezystorem ograniczającym prąd. Jaki buzzer będzie tu dobrze pracował i jakiego koloru jest dioda LED (napięcie na niej)?

Kolejną kwestią jest „interfejs” między wyjściem odbiornika TSOP i wejściem ULN2003, gdzie mamy praktycznie bazę tranzystora Darlingtona. TSOP wypuszcza krótkie impulsy, gdyż powiela kod otrzymany z pilota. To nie byłoby dobre dla sterowania cewką przełącznika. Obwód R4-C3 w połączeniu z charakterystyką wyjścia odbiornika TSOP stanowi tłumik o długiej stałej czasowej w górę (150 milisekund, to wystarczy) i krótkiej w dół. To powinno sprawę załatwić. Po co tak komplikować coś, co można wykonać o wiele prościej? Pytanie to pozostanie chyba tajemnicą autora. Można tylko powiedzieć – rozwiązanie oryginalne i tylko dlatego ciekawe.

REKLAMA

Sięgnij po archiwalne wydania ELEKTRONIKI dla WSZYSTKICH

Prenumeratorzy mają bezpłatny dostęp do e-wydawnictw archiwalnych EdW starszych niż 24 miesiące.



Przesyłka GRATIS

Zamów wygodnie na www.UlubionyKiosk.pl

ep@epi.pl 7037036930



Rysunek 1. Silnik prognostyczny mojej wybranki

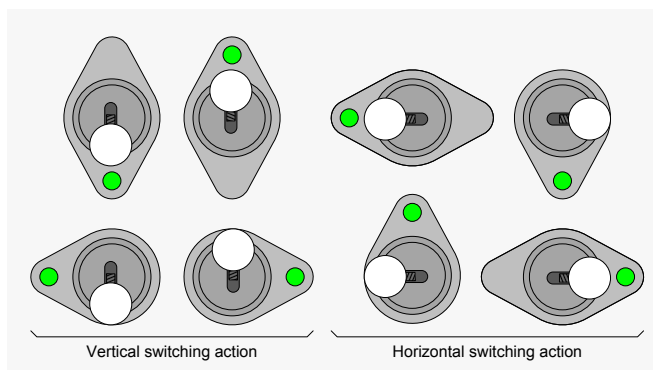
Migające diody LED i śliniący się inżynierowie

Mówiłem to już wcześniej i powtórzę to jeszcze raz: „Pokaż mi migającą diodę LED, a pokażę ci śliniącego się mężczyznę”. Prosta prawda jest taka, że nie można mieć zbyt wielu diod LED w swoich projektach hobbyistycznych lub systemach elektronicznych.

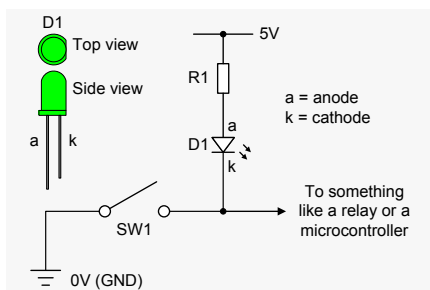
Weźmy na przykład mój Inamorata Prognostication Engine (silnik prognostyczny Inamorata – **rysunek 1**), którego zadaniem jest przewidywanie nastroju mojej żony (Gina the Gorgeous), gdy wychodzę z biura i zmierzam do domu (ironicznym zrządzeniem losu, jeśli Gina kiedykolwiek odkryje przeznaczenie silnika, nie będę potrzebował go do przewidywania jej nastroju w danym momencie). Oprócz trójkolorowych diod LED oświetlających lampy próżniowe na górze i oświetlających wzmacniacz w górnej szafce, główny panel sterowania jest ozdobiony małymi światełkami: po dwa na każdy przełącznik dwustanowy i przełącznik przyciskowy oraz 16 na każdy potencjometr. W tej chwili używam tych diod LED do uzyskania efektu tęczy. Teraz, gdy wiem, że wszystko działa zgodnie z planem, jestem gotowy do rozpoczęcia programowania ich na poważnie, co skłoniło mnie do refleksji nad tym, jak możemy ulepszyć nasze projekty przy użyciu tylko jednej jednokolorowej diody LED, a następnie przechodząc do większych rzeczy.

Pojedyncza jednokolorowa dioda LED

Samo powiedzenie „pojedyncza jednokolorowa dioda LED” wywołuje łezkę w oku, ale staram się być odważny i założyć się, że przy odrobinie pomyślnie możemy uczynić nawet to interesującym. Zaczniemy od rozważenia jednobiegunowego przełącznika SPST (*single-pole, single-throw*), którego używamy do sterowania czymś. Należy pamiętać, że mamy do czynienia z bardzo zwykłą, standardową diodą LED w dopasowanej obudowie, która różni się od przełącznika. Oznacza to, że nie rozważamy komponentu ze zintegrowaną diodą LED wyglądającą na przykład jak aureola otaczająca środek przełącznika.



Rysunek 2. Alternatywne orientacje przełączników i lokalizacje diod LED



Rysunek 3. Przełącznik bezpośrednio sterujący diodą LED

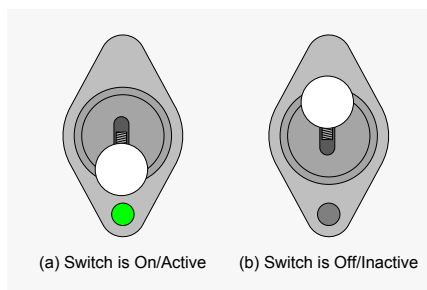
Pierwszą rzeczą do rozważenia jest to, czy przełączanie jest zorientowane pionowo czy poziomo (obie opcje są używane w Inamorata Prognostication Engine), a także, czy nasza dioda LED ma być zamontowana powyżej, poniżej, po lewej czy po prawej stronie przełącznika (jak na **rysunku 2**).

Istnieją dwa powody, dla których prezentuję diodę LED jako zieloną na tych zdjęciach. Po pierwsze, ludzkie oko jest najbardziej wrażliwe na kolor żółtozielony. Drugim jest to, że oryginalna dioda LED, którą wyciągnąłem z mojego skarbca części zamiennych, kiedy zacząłem pisać tę kolumnę, była zielona (takie są nieprzewidywalne kaprysy losu). Niezależnie od tego, którą kombinację przełącznik-orientacja/lokalizacja diody LED wybierzemy, zacznijmy od rozważenia przełącznika sterującego diodą LED bezpośrednio, bez pomocy mikrokontrolera (jak na **rysunku 3**). Zauważ, że powodem, dla którego pokazuję tutaj zasilanie 5 V, jest to, że planuję użyć Arduino Uno w późniejszych eksperymentach. Należy również pamiętać, że dodatnia końcówka diody LED, anoda (a), jest dłuższym przewodem, podczas gdy jej ujemna końcówka, katoda (k), jest oznaczona płaską stroną na plastikowej obudowie.

Aby przypomnieć sobie, jak działają te elementy i upewnić się, że wszystko jest zrozumiałe – dwa główne parametry diody LED to spadek napięcia przewodzenia (V_r) i maksymalny prąd przewodzenia (I_r). Załóżmy, że V_r i I_r dla naszej diody LED wynoszą odpowiednio 2 V i 20 mA (0,02 A). Korzystając z prawa Ohma, wiemy, że $V = IR$, co oznacza, że $R = V/I$. Ponieważ nasze zasilanie wynosi 5 V, a nasza dioda obniża napięcie o 2 V, oznacza to, że nasz rezystor $R_1 = (5 - 2) / 0,02 = 3 / 0,02 = 150 \Omega$.

Dla celów niniejszej dyskusji przyjmijmy, że nasz przełącznik jest zorientowany tak, aby miał pionowe działanie przełączające, a dioda LED jest zamontowana poniżej przełącznika. Przyjmijmy również, że jeśli dźwignia przełącznika jest skierowana w stronę diody LED, to przełącznik jest w stanie włączonym. Z kolei jeśli dźwignia jest skierowana od diody LED, przełącznik jest w stanie wyłączonym (jak na **rysunku 4**).

To prowadzi nas do interesującej kwestii, że w Wielkiej Brytanii zwykle naciska się przełącznik w dół, aby go aktywować i w górę, aby



Rysunek 4. Stany aktywne/nieaktywne naszego przełącznika

go dezaktywować. W przypadku przełączników kołyskowych naciska się dolną część przełącznika, aby go aktywować, a górną, aby go dezaktywować. (Wyjątkiem od tej reguły jest lotnictwo, gdzie przełączenie przełącznika w górę powoduje jego włączenie). Dla porównania, w Stanach Zjednoczonych sytuacja wygląda zazwyczaj odwrotnie, a w innych częściach świata nie obowiązuje żadna reguła. W przypadku własnych projektów możesz robić, co chcesz, ale spójność jest prawdopodobnie ważniejsza niż orientacja i zawsze dobrze jest uczynić swoje systemy „przyjaznymi lokalnie”.

Dodawanie Arduino

Będziemy musieli uzgodnić jakąś formę terminologii, która pomoże nam śledzić rzeczy w miarę postępów. Zacznijmy od poniższego, aby opisać poprzedni przykład przełącznika bezpośrednio sterującego diodą LED:

Przełącznik → Włączony; Dioda LED → Włączona

Przełącznik → Wyłączony; Dioda LED → Wyłączona

Chociaż jest to funkcjonalne, nie jest inspirujące. Czy jest jakiś sposób, aby dodać trochę czadu? Cóż, tak, a najprostszym sposobem na to jest wprowadzenie do akcji mikrokontrolera. Na potrzeby tej dyskusji użyjemy Arduino Uno. Co więcej, musimy planować z wyprzedzeniem, ponieważ wraz ze wzrostem zaawansowania naszych efektów, drganie styków przełączników może stać się problemem.

Dla przypomnienia, gdy włączamy lub wyłączamy przełącznik, jego styki mogą w rzeczywistości załączać i rozłączać wiele razy, zanim ustabilizują się w nowym stanie. Dotyczy to przełączników, przełączników kołyskowych, przełączników przyciskowych itp. Nie stanowi to problemu na przykład w przypadku włącznika światła w domu, ponieważ wszelkie migotanie światła, które może wystąpić, dzieje się zbyt szybko, aby nasze oczy i mózgi mogły to dostrzec. Inaczej jest jednak w przypadku mikroprocesorów, które mogą próbować swoje dane wejściowe miliony razy na sekundę, ponieważ mogą one postrzegać to, co uważamy za pojedyncze przejście przełącznika, jako składające się z wielu zdarzeń.

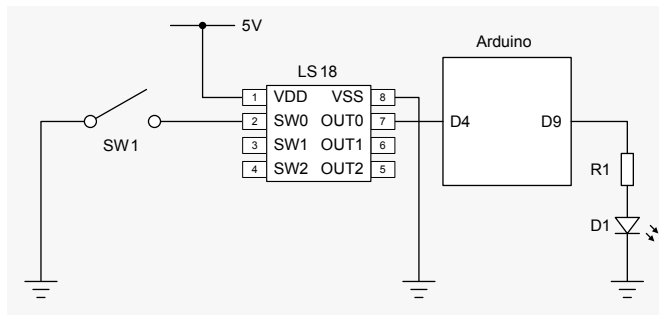
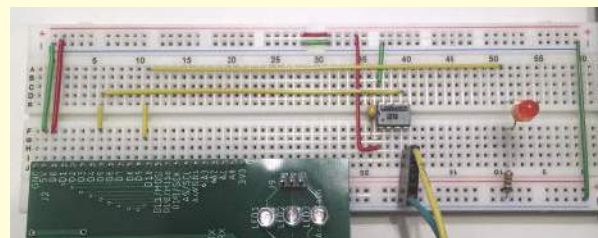
Słowo do mądrych

To tylko mała rada ode mnie dla ciebie. Nie ma znaczenia, jak dobre zdanie masz o sobie lub jak prosty jest układ, który budujesz, twoje życie będzie o wiele mniej frustrujące, jeśli wykonasz następujące czynności:

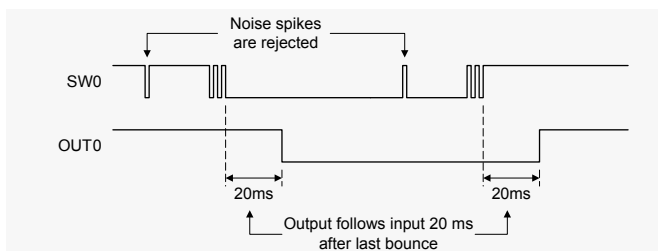
Narysuj schemat układu na kartce papieru.

Porównaj fizyczny układ z jego schematem i zaznacz przewody jeden po drugim.

Kiedy podłączałem płytkę drukowaną do pierwszego programu, zapomniałem dołączyć jeden mały przewód potężniejszy. Gdybym postąpił zgodnie z powyższą radą, oszczędziłbym sobie wielu kłopotów. A tak spędziłem mnóstwo czasu na tropieniu tego małego szkraba. Przeoczenie tego przewodu spowodowało wiele frustracji. Tak, wiem, wolałbyś tego nie robić, ale zastosowanie się do moich rad pomoże ci uniknąć zbędnych frustracji.



Rysunek 5. Używanie LS18 do usuwania drgań przełącznika i Arduino do sterowania naszą diodą LED



Rysunek 6. Działanie układu odblokowującego przełącznik LS18

Istnieje wiele różnych technik sprzętowych i/lub programowych, których możemy użyć do eliminacji drgania styków przełączników. W tym przypadku korzystam z układu LS18 w obudowie DIL firmy LogiSwitch.net (**rysunek 5**). Jest to układ 3-kanalowy, co oznacza, że może obsługiwać trzy przełączniki. Wejścia układu usuwają zakłócenia, a wyjścia z LS18 podążają za odpowiednimi wejściami z opóźnieniem 20 ms po ustaniu drgań styków (co pokazano na **rysunku 6**).

Adnotacje „D4„ i „D9„ na rysunku 5 oznaczają odpowiednio wyprowadzenia 4 i 9 – cyfrowe wejścia/wyjścia (I/O) Arduino. Zauważ, że nie używamy rezystorów podciągających na wejściu SW0 układu LS18 ani na wejściu D4 układu Arduino, ponieważ układ LS18 zajmuje się tym za nas (uwielbiam ten układ i używam go – oraz jego kuzyna LS118 – we wszystkich moich projektach).

Dla ułatwienia udostępnię każdy kod, który przygotuję dla Arduino w ramach tej serii artykułów, do pobrania, przeglądania i rozważania. Aby zapewnić punkt odniesienia, napisałem szkic (program) Arduino, który replikuje efekty przełącznika bezpośrednio sterującego diodą LED (<https://bit.ly/2tfzZo>).

```
#define SWITCH_ON LOW
#define SWITCH_OFF HIGH
#define LED_ON HIGH
#define LED_OFF LOW
int PinSwitch = 4;
int PinLed = 9;
int OldSwitchState;
int NewSwitchState;
void setup ()
{
  pinMode(PinSwitch, INPUT);
  pinMode(PinLed, OUTPUT);
  OldSwitchState = digitalRead(PinSwitch);
  if (OldSwitchState == SWITCH_ON)
  {
    digitalWrite(PinLed, LED_ON);
  }
  inny
  {
    digitalWrite(PinLed, LED_OFF);
  }
}
void loop ()
{
  NewSwitchState = digitalRead(PinSwitch);
  if (NewSwitchState != OldSwitchState)
  {
    if (NewSwitchState == SWITCH_ON)
    {
      digitalWrite(PinLed, LED_ON);
    }
    inny
    {
      digitalWrite(PinLed, LED_OFF);
    }
    OldSwitchState = NewSwitchState;
  }
}
```

Obawiam się, że gdy spojrzysz na mój program, możesz pomyśleć, że sposób, w jaki to zrobiłem, jest „przesadny”. Na przykład, ostatecznie użyłem około 50 linii kodu, ale mogłem napisać ten pierwszy szkic o wiele bardziej zwięźle, jak poniżej:

```
int PinSwitch = 4;
int PinLed = 9;
void setup() {
  pinMode(PinSwitch, INPUT);
  pinMode(PinLed, OUTPUT);
}
void loop() {
  if (digitalRead(PinSwitch) == LOW)
    digitalWrite(PinLed, HIGH);
  else digitalWrite(PinLed, LOW);
}
```

Właściwie moglibyśmy usunąć instrukcję `if()` i zredukować treść naszej funkcji `loop()` do jednej linii, gdybyśmy czuli się szczególnie odważni. Istnieje kilka powodów, dla których napisałem pierwszy szkic w taki sposób, jak to zrobiłem, a najważniejszym z nich jest przygotowanie sceny na to, co ma nadejść, na przykład specjalna sekwencja uruchamiania po pierwszym podłączeniu zasilania do systemu.

Do tej pory pewnie myślisz, że jestem geniuszem kodowania – moje dobre maniere nie pozwalają mi zaprzeczyć. Po drodze nauczyłem się kilku sztuczek, a jeśli jesteś nowy w tym wszystkim, to na końcu artykułu zebrałem kilka małych wskazówek, podpowiedzi i praktycznych zasad, które mogą ułatwić ci życie.

Pretensjonalny? Ja?

Możesz nazwać mnie pretensjonalnym, jeśli chcesz, ale czuję, że po prostu włączając i wyłączając naszą diodę LED brakuje trochę refleksji „po co”? Chodzi o to, że teraz, gdy dodaliśmy mikrokontroler, możemy robić wszelkiego rodzaju przebiegłe rzeczy. Na przykład, gdy przełącznik się zamyka, możemy błysnąć diodą LED trzy razy, po czym zrobić pauzę i włączyć ją na stałe (<https://bit.ly/2FbmWqY>):

Przełącznik → Wł.; Dioda LED → Błysk-Błysk-Pauza-Wł.

Przełącznik → Wyłączony; Dioda LED → Wyłączona

Jeśli spojrzysz na mój kod, zobaczysz, że używam wywołań funkcji `delay()`. Robię to, aby kod był jak najłatwiejszy do zrozumienia. W praktyce jednak użycie tej funkcji nie byłoby dobrym pomysłem w tej sytuacji. Dzieje się tak dlatego, że `delay()` jest funkcją blokującą, co oznacza, że uniemożliwia Arduino robienie czegokolwiek innego (poza obsługą przerwań). Na szczęście istnieją lepsze sposoby robienia rzeczy, ale są one nieco bardziej skomplikowane, więc zostawimy je jako temat na inny dzień.

Wracając do naszej dyskusji, inną możliwością jest to, że – zamiast włączać i wyłączać – wygaszamy diodę LED odpowiednio w górę i w dół:

Przełącznik → Wł.; Dioda LED → Zanik-Wł.

Przełącznik → Wyłączony; Dioda LED → Wygaszona

Jak więc wykonać to wygaszanie? Możliwe jest przyciemnienie diody LED poprzez obniżenie jej napięcia zasilania, ale zazwyczaj nie działa to tak dobrze, jak można by się spodziewać, zwłaszcza dlatego, że wszystko pójdzie w gruzy, gdy napięcie zasilania zbliży się do wartości spadku napięcia przewodzenia diody LED. Alternatywą jest bardzo szybkie włączanie i wyłączanie diody LED. Jeśli na przykład dioda LED jest włączona tylko przez 50% czasu, będzie wydawać się tylko w połowie tak jasna, jak gdyby była włączona przez 100% czasu. Na szczęście mikrokontrolery mogą włączać i wyłączać rzeczy bardzo szybko, co pozwala nam zaimplementować coś, co nazywamy modulacją szerokości impulsu (PWM). W przypadku Arduino mamy funkcję o nazwie `analogWrite()`, która pozwala nam określić numer pinu i parametr PWM, gdzie ten parametr jest 8-bitową liczbą całkowitą bez znaku, która może reprezentować wartości od 0 (całkowicie wyłączony) do 255 (całkowicie włączony) – **rysunek 7**.

W przypadku Arduino, każdy pełny cykl pokazany na rysunku 7 zajmuje tylko około 1/500 sekundy. Korzystając z tej techniki, dioda LED jest włączana i wyłączana tak szybko, że migotanie nie jest wyczuwalne dla ludzkiego oka. Arduino Uno ma sześć pinów, które

obsługują PWM poprzez funkcję `analogWrite()`: 3, 5, 6, 9, 10 i 11 (na szczęście zdecydowaliśmy się użyć pinu 9 do zasilania naszej diody LED – to prawie tak, jakbyśmy mieli plan).

Zacznijmy od „niezgrabnego” podejścia, w którym „zapalamy”, zwiększając jasność w trzech dużych krokach, zaczynając od 0% i przechodząc do 33%, następnie 66%, a następnie 100%. Zmniejszamy jasność w ten sam niezgrabny sposób (<https://bit.ly/2u7GACa>). Gdy już zobaczymy to w akcji, zmodyfikujemy nasz program, aby wygaszanie odbywało się znacznie płynniej (<https://bit.ly/2MKPqfp>).

Poczuj moc!

Inną rzeczą, którą możemy rozważyć, jest możliwość zaimplementowania sekwencji rozruchowej. Oznacza to, że po pierwszym podłączeniu zasilania do systemu chcemy zrobić coś specjalnego – być może coś takiego jak poniżej (<https://bit.ly/2QGLGBQ>):

Włączenie zasilania (przełącznik jest włączony): LED → ((flash-flash-flash-pause) * 3) → Wł.

Uruchomienie (przełącznik wyłączony): Dioda LED → wydłużony błysk → Wyl.

Wyobraźmy sobie teraz system z wieloma przełącznikami i diodami LED, na przykład Prognostication Engine. W takim przypadku możemy sprawić, że nasza sekwencja startowa będzie uczta dla oczu. Co więcej, jako że ta wspaniała piękność jest wyposażona w trójkolorowe diody LED, możemy zmieniać cały schemat kolorów w zależności od oceny aktualnego stanu rzeczy na świecie (jasne, wesołe kolory, gdy wszystko idzie dobrze; bardziej złowieszcze kolory, gdy sytuacja staje się złowieszcza – <https://en.wikipedia.org/wiki/DEFCON>).

Oddychanie, drzemka i chrapanie

Kolejną rzeczą do rozważenia jest to, co się stanie, jeśli nikt nie aktywuje przełącznika przez określony czas – powiedzmy 10 minut – który możemy nazwać „Okresem nieaktywności”. Dobrze byłoby to w jakiś sposób zaznaczyć.

Istnieją dwa podstawowe scenariusze, kiedy ten okres upływa: albo przełącznik i dioda LED są obecnie włączone, albo są obecnie wyłączone. Zacznijmy od przypadku, w którym są włączone. Jedną z rzeczy, które moglibyśmy zrobić, byłoby wejście w tryb, w którym dajemy naprawdę krótki okresowy błysk:

Nieaktywny limit czasu (przełącznik jest włączony): LED → Flash-Pause-Flash-Pause...

Limit czasu nieaktywności (przełącznik jest wyłączony): Dioda LED → pozostaje wyłączona

Inną alternatywą jest to, że system wchodzi w tryb „oddychania” (słyszałem również, że nazywa się to „drzemaniem” i „chrapaniem”, ponieważ system jest w pewnym sensie uśpiony), w którym dioda LED stopniowo zanika w górę i w dół... i w górę i w dół:

Czas nieaktywności (przełącznik jest włączony): LED → Tryb oddychania

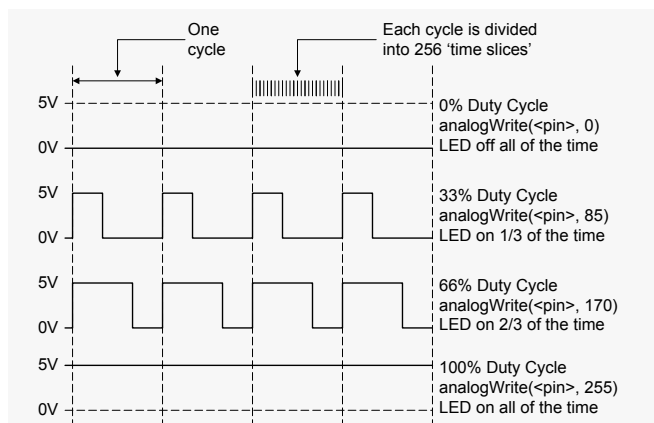
Limit czasu nieaktywności (przełącznik jest wyłączony): Dioda LED → pozostaje wyłączona

Wreszcie, jeśli używamy trybu oddychania, gdy przełącznik jest włączony, może moglibyśmy użyć przerywanego błysku, gdy przełącznik jest wyłączony:

REKLAMA

AVTEDU625
Pipek dręczyciel
sklep.avt.pl





Rysunek 7. Użycie PWM do sterowania jasnością diody LED

Czas nieaktywności (przełącznik jest włączony): LED → Tryb oddychania

Nieaktywny limit czasu (przełącznik jest wyłączony): LED → Flash-Pause-Flash-Pause...

Gdy pojawi się coś, co ponownie wybudzi maszynę, być może mogliśmy uruchomić sekwencję rozruchową, o której mówiliśmy wcześniej.

Istnieje kilka powodów, dla których nie pokazałem żadnych przykładów kodu dla tych przypadków drzemki i chrapania. Pierwszym z nich jest to, że obecnie mamy tylko jeden przełącznik, więc jedynym sposobem na wybudzenie systemu byłoby przełączenie naszego przełącznika w przeciwny stan. Jeśli na przykład przełącznik był włączony, gdy weszliśmy w tryb drzemki, jedynym sposobem na wybudzenie systemu byłoby wyłączenie przełącznika. Tak więc wdrożenie tego trybu miałooby o wiele więcej sensu, gdyby nasz system zawierał kilka przełączników.

Innym problemem jest użycie funkcji `delay()`, która utrudnia nam wykrycie, kiedy jakaś akcja ma miejsce w świecie zewnętrznym i zareagowanie na nią w odpowiednim czasie. Omówimy techniki obejścia tego problemu w kolejnym odcinku tego cyklu.

Podeksytowanie roślinie!

Na wypadek, gdybyś był zainteresowany, przygotowałem film pokazujący wszystkie efekty, które omówiliśmy do tej pory w akcji (<https://bit.ly/2ZUvL1W>) – chciałbym usłyszeć, co o tym myślisz.

Pamiętaj też, że wszystkie zaprezentowane powyżej efekty to tylko pierwsze pomysły, które przysły mi do głowy podczas pisania tej kolumny. Byłoby wspaniale, gdybyś mógł przeprowadzić kilka własnych eksperymentów, a następnie wysłać mi e-mail, jeśli wymyślisz jakieś sprytne efekty, którymi chciałbyś się podzielić.

Naprawdę ekscytujące jest to, że do tej pory rozważaliśmy tylko pojedynczy przełącznik z pojedynczą jednokolorową diodą LED. A co z przełącznikami przyciskowymi (zatrząskowymi i chwilowymi)? Co z dwukolorowymi i trójkolorowymi diodami LED? A co z dwoma diodami LED powiązanymi z każdym przełącznikiem lub przyciskiem?

Nie wiem jak wy, ale ja mam zawroty głowy, biorąc pod uwagę różnorodność potencjałów, perspektyw i możliwości. Nie mogę się doczekać, aby zobaczyć, co wymyślimy w przyszłym miesiącu. W międzyczasie czekam na wasze komentarze, pytania i sugestie. Do następnego razu, miłego oglądania!



Komentarze lub pytania?
Napisz do Maxa na adres:
max@CliveMaxfield.com

Sprytne porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania



Zamierzam udostępnić do pobrania szkice (programy) Arduino, które będą towarzyszyć moim kilku następnym felietonom w tym cyklu, więc pomyślałem, że dobrym pomysłem może być uzasadnienie stylu kodowania, którego używam.

Pierwszą kwestią, na którą należy zwrócić uwagę, jest to, że używasz zintegrowanego środowiska programistycznego Arduino (IDE) do tworzenia szkiców przy użyciu języków programowania C/C++. C++ to zasadniczo C z rozszerzeniami (jak C na sterydach). C++ jest klasyfikowany jako język obiektowy (OPP), ale tak naprawdę nie musimy się martwić o to, co to oznacza w tym momencie. To, co uważamy za „język Arduino”, to zestaw wstępnie zdefiniowanych funkcji C/C++, które ktoś dla nas napisał i które można wywołać z naszego kodu.

Funkcja `main()`

Jedną z rzeczy, która denerwuje niektórych ludzi, jeśli znają już C/C++, jest to, że programy napisane w tych językach zawsze mają funkcję najwyższego poziomu o nazwie `main()`. W przypadku Arduino istnieją jednak dwie funkcje, które są domyślnie umieszczane w szkicu, `setup()` i `loop()`, ale nie ma funkcji `main()` (można również tworzyć własne funkcje, ale szkic musi zawierać funkcje `setup()` i `loop()`).

IDE dodaje również komentarz wewnątrz funkcji `setup()`, mówiący: „Umieść tutaj swój kod konfiguracyjny, aby uruchomić go raz”. Podobnie, wewnątrz funkcji `loop()` znajduje się komentarz mówiący: „Umieść tutaj swój główny kod, aby uruchamiać go wielokrotnie” (użycie słowa „główny” jest niefortunne w tym kontekście).

Za kulisami, po kliknięciu ikony „Kompiluj”, IDE generuje plik tymczasowy, który zawiera cały kod wraz z funkcją `main()` wyglądającą jak poniżej:

```
main ()
{
    setup();
    while (1) loop();
}
```

Powoduje to wywołanie funkcji `setup()` jeden raz, a następnie wywołanie funkcji `loop()` w kółko i w kółko.

Litery i spacje

C/C++ to języki rozróżniające wielkość liter – na przykład zmienne o nazwach takich jak `fred`, `Fred` i `FRED` są traktowane przez kompilator jako zupełnie różne byty. C/C++ nie dba o to, czy używasz białych znaków, czy nie (gdzie białe znaki obejmują spacje, tabulatory i puste linie). Głównym powodem używania białych znaków jest zwiększenie czytelności kodu – są one usuwane podczas kompilacji kodu, więc nie sprawiają, że ostateczny program załadowany do Arduino będzie większy.

A propos tabulatorów... nie używaj ich. Tak, oczywiście, że łatwiej jest nacisnąć klawisz <tab> jeden raz niż wielokrotnie klawisz <space>. Problem polega na tym, że nie wiesz, ile spacji inne systemy skojarzą z tabulatorem, co może na przykład sprawić, że twoje wydruki będą wyglądać dziwnie. Nie zapominajmy też o artykule z 2017 roku na stronie StackOverflow.blog: *Developers Who Use Spaces Make More Money Than Those Who Use Tabs* (<https://bit.ly/37iBoeg>).

Komentarze

Wszystko w linii następującej po dwóch znakach ukośnika (//) jest postrzegane jako komentarz. Komentarze wielowierszowe mogą być rozpoczęte ukośnikiem i gwiazdką (/*) oraz zakończone gwiazdką i ukośnikiem (*/). Komentarze naprawdę pomagają tobie i innym w zrozumieniu i utrzymaniu twoich programów, więc używaj ich, ale nie dodawaj komentarzy, gdy nie ma takiej potrzeby (instrukcja `a =`

`a + 1;` po której następuje komentarz, który mówi: „// Dodaj 1 do a” nie jest dobrym pomysłem). Kieruj się zdrowym rozsądkiem.

Masz styl?

Pamiętasz piosenkę *Style*, śpiewaną przez Franka Sinatrę, Deana Martina i Binga Crosby’ego (<https://bit.ly/39qLpGG>)? Jest tam kilka wersów, które brzmią: *You’ve either have or you haven’t got style* (Jeśli go masz, wyróżniasz się na miłą). Jest to z pewnością prawdą w przypadku pisania kodu.

Styl kodowania w dużej mierze zależy od osobistych preferencji. Jeśli pracujesz dla dużej firmy, oczekuje się, że będziesz podążać za ich wewnętrznym stylem. Jeśli robisz to dla siebie, możesz rozwinąć swój własny sposób robienia rzeczy. Najważniejszą rzeczą jest bycie konsekwentnym, abyś ty (i inni) mógł łatwiej zrozumieć i utrzymać swój kod w przyszłości.

Z zawodu jestem inżynierem projektującym sprzęt. Jestem samoukiem, jeśli chodzi o oprogramowanie, więc naprawdę nie powinieneś słuchać niczego, co mówię o kodowaniu. Z drugiej strony, miałem styczność z dużą ilością kodu i widziałem wiele rzeczy, które mi się podobają (wraz z wieloma, które mi się nie podobają). Mój własny styl ewoluował przez lata, gdy miałem styczność z różnymi technikami. Poniżej znajduje się podzbiór technik kodowania, które działają dla mnie:

`#define`

Zawsze używam wielkich liter alfanumerycznych i podkreśleń w nazwach stałych, na przykład:

```
#define START_COUNT 0
#define END_COUNT 100
```

Uzasadnienie: Pozwala mi to łatwo zidentyfikować stałe globalne w treści programu.

Nie używaj „magicznych liczb”

Przyjrzyjmy się poniższemu stwierdzeniu:

```
for (i = 0; i < 10; i++)
{
    // Rób rzeczy tutaj
}
```

Zarówno 0, jak i 10 byłyby uważane za „magiczne liczby”, ponieważ pojawiły się znikąd, jak za dotknięciem czarodziejskiej różdżki. Ogólnie rzecz biorąc, dobrze jest użyć 0 (lub 1), ponieważ wszyscy wiemy, że liczenie zwykle zaczyna się lub kończy na 0 (lub 1), ale najlepiej jest użyć stałej wartości dla dowolnej innej liczby – na przykład (zakładając, że `MAX_COUNT` został wcześniej zadeklarowany za pomocą instrukcji `#define`):

```
for (i = 0; i < MAX_COUNT; i++)
{
    // Rób rzeczy tutaj
}
```

Następnym razem

W następnym odcinku tego cyklu przyjrzymy się nazwom zmiennych, nazwom funkcji i wywołaniom funkcji. W międzyczasie, jeśli masz jakieś wskazówki i sztuczki dotyczące kodowania, którymi chciałbyś się podzielić, chętnie o nich usłyszę. ■

Clive „Max” Maxfield

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, marzec 2020 (www.epemag3.com)

Pierwszy taki telefon, wielkości palca z ekranem dotykowym E-INK

Wykonać własny telefon? To brzmi niewiarygodnie. Telefony komórkowe stają się coraz bardziej skomplikowane i wyposażane w coraz większą ilość funkcji dodatkowych. Nie zawsze jest to zaletą i wielu osobom utrudnia korzystanie z podstawowych funkcji telefonu. Drugą istotną cechą smartfonów jest ich wyświetlacz. Wyświetlacze LCD z tylnym podświetlaniem lub OLED-owe nie są zdrowe dla oczu, gdy dłużej się w nie wpatrujemy. Można się dziwić jak tak zaawansowane urządzenia są konstruowane i zastanawiać czy jest tu pole do popisu dla konstruktora własnego smartfona. Bieżący projekt wychodzi na przeciw takim pomysłom.

Pokazujemy jak można wykonać smartfon wg własnego projektu z wykorzystaniem wyświetlacza typu E-ink. To wyświetlacz dotykowy niewielkiego rozmiaru. Interfejs użytkownika realizuje system operacyjny

bazujący na Linuxie i łączy VNC HDMI. Oprogramowanie to realizuje wszystkie podstawowe funkcje telefonu. Prototyp telefonu wykonanego przez autora pokazuje zdjęcie na **rysunku 1**. Na **rysunku 2** widzimy

wyświetlacz podczas, gdy ktoś dzwoni na nasz telefon. **Rysunek 3** pozwala ocenić wielkość telefonu. Można go przymocować do palca.

Wyświetlacz typu E-ink jest bardzo korzystny z wielu względów. Nie męczy oczu i pobiera bardzo mało energii z baterii. Ekran tego typu jest alternatywą bliską druku na papierze. Treść wyświetlacza jest czytelna nawet po odłączeniu zasilania, tak długo dopóki zechcemy go „odświeżyć”. To nie wszystkie z zalet jakie oferują „atramentowe” wyświetlacze. W bieżącym projekcie przystępujemy do projektu telefonu, ale i konstrukcja komputera tego typu jest możliwa. Na dodatek, koszt takiego urządzenia własnej konstrukcji okazuje się niski. Zaczynamy zatem projekt od skompletowania potrzebnych podzespołów. Zebrano je poniżej.

Projekt telefonu

Projekt własny „zrób to sam” jest atrakcyjny, gdy cechy jego będą dostosowane do własnych potrzeb. Musimy wykonać własną obudowę tak, aby zmieścić w niej potrzebne podzespoły. Wielkość obudowy zdeterminowana jest przede wszystkim wymiarami wyświetlacza. Na **rysunku 4** pokazano jak wykonać obudowę. W szczególności jak wydrążyć w niej otwory dla modułu GSM oraz dla mikrofonu. Trzeba też przewidzieć miejsce dla



Rysunek 1. Prototyp wykonany przez autora



Rysunek 2. Ktoś dzwoni na nasz telefon

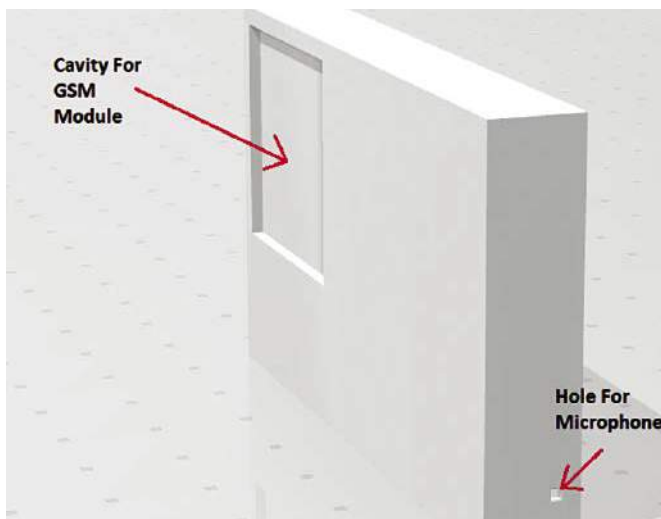


Rysunek 3. Rozmiar telefonu porównywalny z wielkością palca

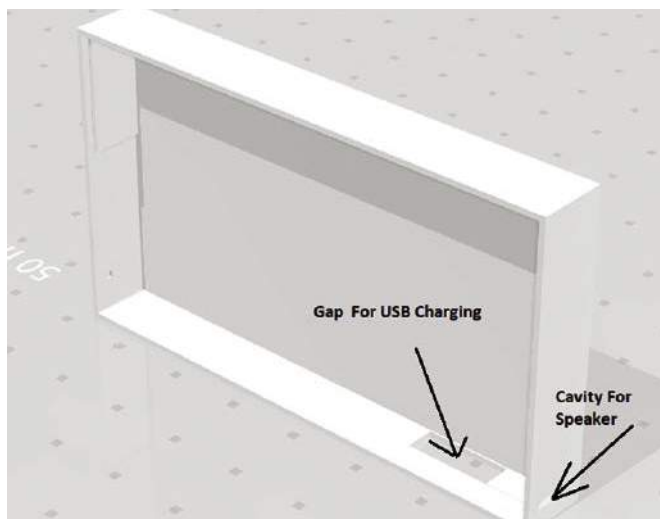
Wykaz elementów, kupuj w sklepie.avt.pl
(W-wa, ul. Leszczyńska 11, tel. +48222578451,
e-mail: handlowy@avt.pl):

Raspberry Pi Zero – 1 szt. – mikrokontroler
SIM800L – 1 szt. – moduł GSM
Głośniczek – 1 szt. – 8 Ω/0,5 W (miniaturowy)
Mikrofon – 1 szt. – mikrofon miniaturowy
Bateria i ładowarka – 1 szt. – akumulator 3,3 V i ładowarka
Ekran dotykowy E-ink 5,4 cm – 1 szt. – wyświetlacz typu „atramentowego”
Obudowa – 1 szt. – obudowa PLA z drukarki 3D

**Kod źródłowy
tego projektu jest dostępny
do pobrania ze strony
<https://tiny.pl/cgsww>**



Rysunek 4. Projekt obudowy



Rysunek 5. Otwory w obudowie dla głośniczka i złącza USB

głośniczka, a także na złącze mikro-USB przez które będziemy ładować baterię. Te szczegóły pokazano na **rysunku 5**. Gdy już zaprojektujemy wymiary obudowy, można ją wykonać w technologii druku 3D.

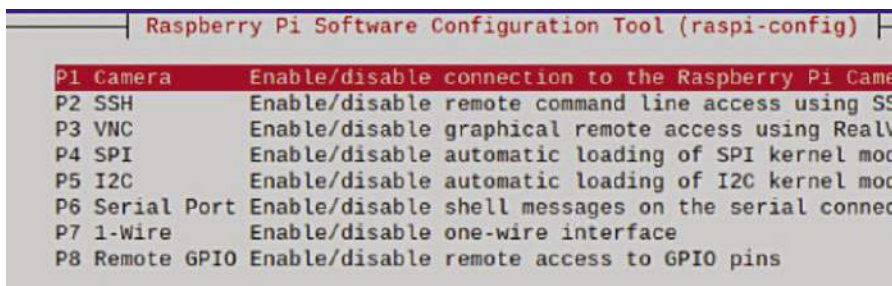
Przystępując do właściwej części projektu, przygotuj płytkę Raspberry Pi z najnowszym systemem operacyjnym i uruchom łącznie szeregowy SPI oraz I²C. Aby to wykonać, należy wpisać kolejno komendy pod Linuxem, co pokazano na **rysunku 6**.

Konfigurację należy rozpocząć wpisując komendę:

```
sudo raspi-config
```

Następnie należy zainstalować drivery dla wyświetlacza dotykowego E-ink. Zadanie sprowadza się do załadowania odpowiednich modułów oprogramowania w Pythonie. Należy wpisać komendy zgodnie z poniższą listą:

```
wget http://www.airspayce.com/mikem/
bcm2835/bcm2835-1.68.tar.gz
tar zxvf bcm2835-1.68.tar.gz
cd bcm2835-1.68/
sudo ./configure && sudo make && sudo
make
check && sudo make install
sudo apt-get install wiringpi
#For Pi 4, you need to update it:
wget https://project-downloads.drogon.
net/wiringpi-latest.deb
sudo dpkg -i wiringpi-latest.deb
gpio -v
#You will get 2.52 information if you
```



Rysunek 6. Programowa konfiguracja systemu

```
install it correctly
sudo apt-get update
sudo apt-get install python3-pip
sudo apt-get install python3-pil
sudo apt-get install python3-numpy
sudo pip3 install RPi.GPIO
sudo pip3 install spidev
cd ~
git clone https://github.com/
waveshare/
Touch_e-Paper_HAT
```

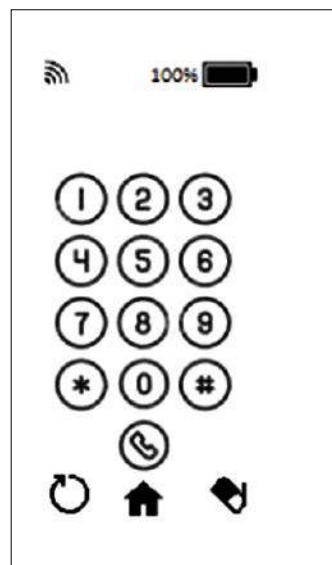
Utworzenie interfejsu użytkownika

Możemy utworzyć własny UI (User Interface), zgodnie z własnymi upodobaniami. Można zaprojektować wygląd ikon oraz nich rozmieszczenie. Efekt uzyskany w prototypie widoczny jest na **rysunku 7**.

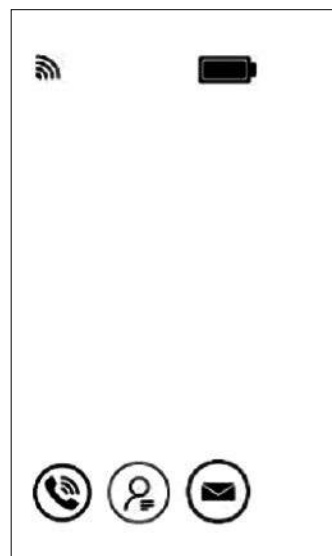
Interfejs użytkownika składa się z kilku podstawowych stron odpowiadających podstawowym funkcjom telefonu. Funkcje te można zebrać wg poniższej listy:

- Ekran główny
- Ekran wywołania numeru telefonu
- Ekran kontaktów
- Wiadomości
- „Call” przychodzący
- Ekran nawiązania rozmowy

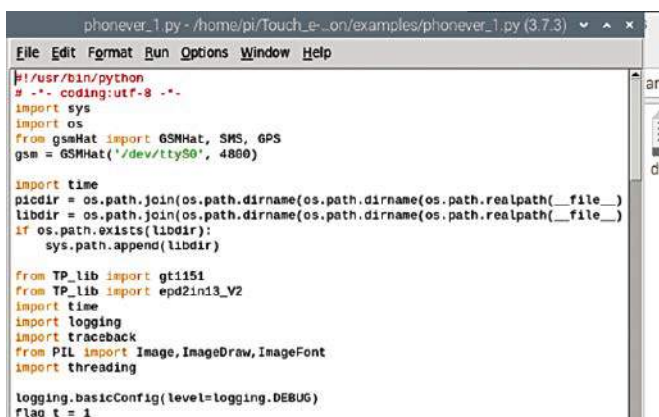
Dla każdego ekranu w interfejsie użytkownika trzeba wybrać ikony, przeskalować ich rozmiar i rozmieścić tak, aby mieściły się na niewielkim ekranie (przekątna 5,4 cm w prototypie).



Rysunek 7. Interfejs użytkownika – strona wybierania numeru



Rysunek 8. Wygląd głównej strony interfejsu użytkownika



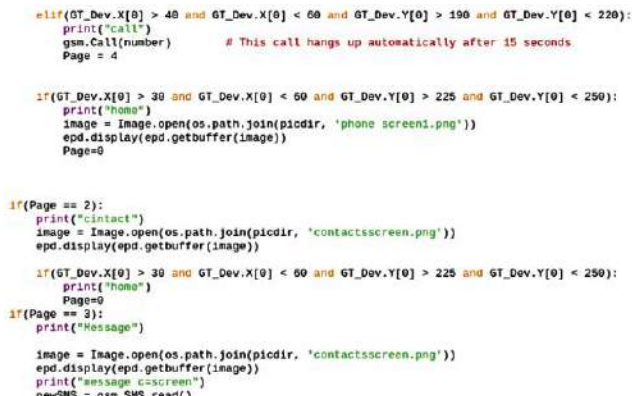
Rysunek 9. Kod programu w Pythonie



Rysunek 10. Kod dla strony głównej



Rysunek 11. Fragment kodu funkcji wybierania numeru



Rysunek 12. Kod programu funkcji – wiadomości

W ekranie głównym użyteczna będzie też obecność ikon z dodatkowymi informacjami jak: poziom naładowania baterii czy siła zasięgu sygnału. Na pozostałych stronach powinny być czytelne ikony na wzór tego, do czego jesteśmy przyzwyczajeni w tradycyjnym smartfonie. Na środku ekranu warto pozostawić puste pole dla takich informacji jak godzina bieżącego czasu, data itp.

Na stronie wywołania numeru użytkownika trzeba zamieścić ikony cyfr i przydzielić im funkcje zgodne z tym co przedstawiają. Tak samo stworzymy UI dla pozostałych stron. Informacje te trzeba zapisać w folderze pic biblioteki Pythona.

Kodowanie systemu operacyjnego telefonu

System operacyjny telefonu składa się z modułów w programie Pythona. Tworząc ten kod należy ściągnąć konieczne moduły i biblioteki. Wśród nich oprogramowanie które pozwoli na komunikację Raspberry z wyświetlaczem E-ink, a także z modułem GSM SIM800L. Następnie należy utworzyć stosowne ścieżki programowe do „obrazków” ikon wchodzących w skład interfejsu użytkownika. Całość oprogramowania składa się z pętli programowych i rozgałęzień tworzonych przez różne warunki pracy tych pętli.

Na pierwszej „domowej” stronie sprawdź zgodność położenia ikon z działaniem ekranu dotykowego. Tu sprawdź działanie ikony nawiązania połączenia, kontaktów i wiadomości. Sprawdź też funkcję rozmowy przychodzącej i poprawność dzwonka twojego telefonu. Interfejs użytkownika zawiera wiele rozgałęzień warunkowych które prowadzą do kolejnych stron. Na stronie nr 2 (wybierania numeru) muszą pokazać się przyciski numeryczne. Wypełnienie ciągu cyfr musi zakończyć się przesłaniem komendy do modułu GSM. Moduł ten realizuje właściwe połączenie, zaś interfejs użytkownika powinien się przełączyć na kolejną stronę. Na stronie „wydzwaniania” musi się też znaleźć ikonka zawieszenia połączenia. Jej naciśnięcie powinno skutkować odpowiednią komendą dla modułu SIM800L i powrót do „domowej” strony interfejsu użytkownika. W analogiczny sposób sprawdź działanie telefonu przypisane funkcjom wszystkich ikon. Kod programu w Pythonie, który funkcje te realizuje pokazano na rysunkach 9...12.

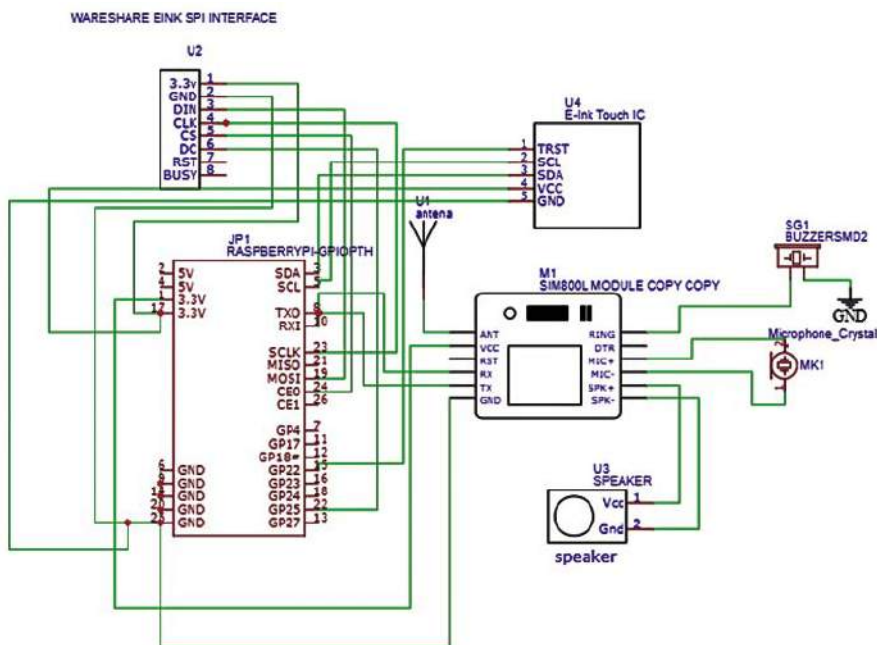
Połączenia podzespołów telefonu

Po wykonaniu softwareowej części projektu twojego telefonu, należy wykonać

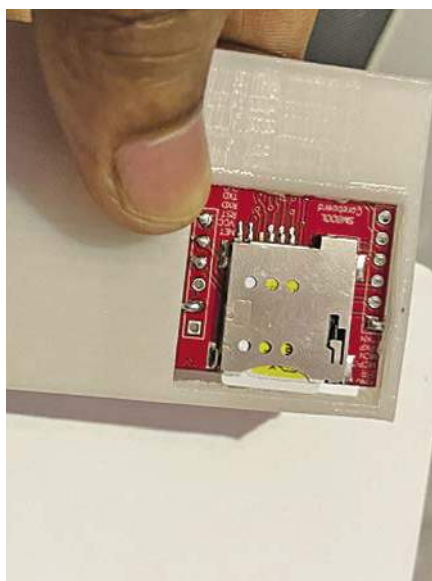
fizyczne połączenia podzespołów zgodnie ze schematem ideowym na **rysunku 13**. Wpierw umieść we właściwych miejscach obudowy płytkę Raspberry Pi, moduł SIM800L a także głośniczek i mikrofon. Między płytką Raspberry i wyświetlaczem powinna się zmieścić bateria i złącze pozwalające na podłączenie ładowarki. Wyświetlacz E-ink montowany jest jako kanapka na Raspberry. Złącze między tymi komponentami widoczne jest na **rysunku 14**. Raspberry zawiera piny męskie, zaś na wyświetlaczu jest gniazdo z tak samo rozmieszczonymi pinami żeńskimi. Zdjęcie na **rysunku 15** pokazuje moduł GSM we właściwym miejscu obudowy. Na **rysunku 16** jest złożony telefon z ekranem wypełniającym przednią część obudowy.

Testy końcowe

Twój „E-ink telefon” jest już gotowy do prezentacji swoich możliwości. Wystarczy umieścić kartę SIM w przewidzianym dla niej slocie na module GSM i uruchomić kod programu systemu operacyjnego. Naciśnij ikonkę „dzwonienie”, powinien pokazać się ekran wyboru numeru. Po wyborze lub wpisaniu numeru, zakończ przyciskiem ikonki „Call”. Sprawdź także działanie funkcji wybierania numeru z kontaktów i uzupełniania ich.



Rysunek 13. Schemat ideowy telefonu



Rysunek 15. Moduł GSM ulokowany we właściwym miejscu obudowy



Rysunek 16. Montaż ekranu wypetnia przednią część obudowy



Rysunek 14. Złącze na Raspberry Pi dla podłączenia ekranu E-ink



Rysunek 17. Końcowy wygląd twojego „E-ink telefonu”

Spróbuj teraz zadzwonić z innego telefonu na swój „E-ink”. Powinien odezwać się dźwięk buzzera i powitalny ekran na wyświetlaczu. Naciśnięcie ikonki przyjęcia rozmowy powinno uruchomić nawiązanie połączenia i rozmowę. W tej fazie testu warto też sprawdzić funkcję zawieszenia rozmowy.

Ciesz się działaniem własnego telefonu. Ostateczny jego wygląd (w prototypie autora) widoczny jest na **rysunku 17**.

Uwaga

Zaprezentowany tu projekt jest pierwszą wersją, którą autor ma zamiar rozwijać

i wyposażać w dodatkowe funkcje. Efekt tych działań możesz śledzić na stronie: www.electronicsforu.com. ■

Ashwini Kumar Sinha

Ten artykuł był wcześniej opublikowany na łamach „EFY”, maj 2022 (efymag.com)

REKLAMA

www.facebook.com/ElportalPL

Zdalnie sterowany samochodzik zabawka

Zabawka której działanie wyjaśnia bieżący projekt, to samochodzik wyposażony w dwa silniczki oraz zespół dwóch przekładni. Jedna napędza koła tylne, druga steruje skrętem kół przednich. Pilot umożliwiający zdalne bezprzewodowe sterowanie mieści się poręcznej obudowie z intuicyjnie rozmieszczonymi przyciskami.

Przykład takiej zabawki pokazuje zdjęcie na **rysunku 1**. Komunikacja między pilotem-nadajnikiem i odbiornikiem w samochodziku odbywa się w paśmie 27 MHz. Opis działania układu rozpoczniemy od nadajnika, którego schemat pokazano na **rysunku 2**.

Nadajnik wykonano na specjalizowanym układzie scalonym YX4116, który zadowala się niskim napięciem zasilania. Stanowi je 3-woltowa bateria, 2×1,5 V AA. Układ scalony IC1 będący sercem nadajnika jest w istocie koderem umożliwiającym wysłanie czterech rozkazów. Chip ten mieści się w ośmio-nóżkowej obudowie, choć aż 4 wyprowadzenia przewidziano dla obsługi 4-rech przycisków. To rozkazy ruchu wprzód, w tył oraz prawo i lewo. Reszta elektroniki w nadajniku to generator fali nośnej na częstotliwości 27 MHz. Wykonano go na dwóch tranzystorach oraz niewielkiej liczbie pasywnych elementów.

Przjrzyjmy się schematowi z **rysunku 2**. To nie prototyp jak zwykle z rubryki DIY (Zrób To Sam), lecz fabryczna płytka o oznaczeniu QF-1694-5T.

Generator fali nośnej pracuje na tranzystorze T1. To wzmacniacz w konfiguracji wspólnego emitera z obciążeniem indukcyjnym (ceweczką 2,2 uH) i stabilizowany kwarcem 27 MHz (w tym modelu, konkretnie 27,145 MHz). Wstępna polaryzacja bazy odbywa się za pośrednictwem rezystora R2 i pochodzi z ósmego wyprowadzenia układu scalonego. Tak utworzony generator oscyluje na częstotliwości własnej wyznaczonej częstotliwością rezonatora kwarcowego. Wyjście

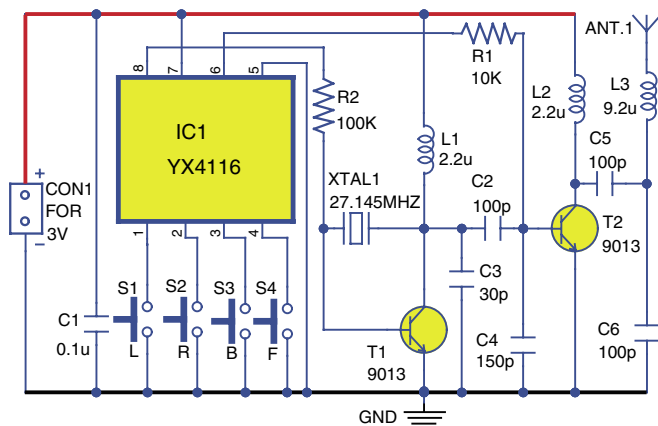


Rysunek 1. Przykład samochodu-zabawki zdalnie sterowanej w paśmie 27 MHz

kodera jednego z czterech rozkazów wyprowadzone jest na nogę 6 układu scalonego. Sygnał ten jest miksowany z częstotliwością nośną oscylatora na bazie tranzystora T2, który jest drugim stopniem wzmacniacza kierującego sygnał do anteny nadawczej. Obciążeniem tranzystora T2 jest także indukcyjność 2,2 uH. Sprzężenie do anteny jest pojemnościowe przez C5, zaś dopasowanie impedancji anteny realizują pojemność C6 i indukcyjność L3 9,2 uH. Działanie nadajnika opiera się na kluczowaniu fali nośnej 27 MHz. W stanie aktywnym wyjścia (pin-6), nadajnik generuje czystą sinusoidę o ciasno zdeterminowanej częstotliwości (dzięki stabilizacji rezonatorem kwarcowym). Kodowane dane zawarte są w przerwach kluczowanej fali nośnej i zdeterminowane są czasowo.

Schemat odbiornika pokazano na **rysunku 3**. Tu układ jest nieco bardziej skomplikowany. Płytkę odbiornika oznaczono symbolem QF-1688-R3, a jej sercem jest także specjalizowany chip o oznaczeniu TXM8D423C. Cały układ ułożono w samochodziku zabawce i zasilany jest on napięciem 4,5 V (trzy baterijki 1,5 V AA). Układ scalony IC1 zamknięto w obudowie 16-to nóżkowej, gdyż oprócz części odbiorczej zawiera on driver dwóch silniczków. Pracują one w konfiguracji H-bridge, co w łatwy sposób pozwala na zmianę kierunku obrotu silników.

„Front End” odbiornika fali o częstotliwości radiowej jest w istocie odbiornikiem regeneracyjnym stosowanym od dziesięcioleci. Stopień ten wykonano na tranzystorze T1 pracującym w konfiguracji wspólnej bazy. To wąskopasmowy wzmacniacz selektywny nastrojony na częstotliwość nadajnika 27 MHz. Obwód rezonansowy złożony jest z L1 i C2, przy czym cewka L1 jest strojona, co umożliwia precyzyjne dostrojenie odbiornika do częstotliwości na jakiej pracuje nadajnik (dla C2=22 pF, L1 powinna mieć indukcyjność ok. 1,5 uH). Jednak informacja (którą chce przekazać nadajnik) nie jest zawarta w fali nośnej, lecz w jej modulacji lub kluczowaniu. Zatem częstotliwość nośna jest następnie odfiltrowana filtrem RC i do układu trafia składowa niskiej częstotliwości którą trzeba zdekodować. Zdekodowana informacja trafia do układów wykonawczych i driverów silników. Układ steruje dwoma silnikami. Napędowym, oznaczonym na schemacie jako



Rysunek 2. Schemat ideowy nadajnika

Back/Front oraz sterującym kierunkiem jazdy prawo/lewo (Right/Left). Sprzężenie obu silniczków z układem scalonym jest bardzo proste. Są one wpięte w przekątną mostka H-bridge, co umożliwia płynną regulację obrotów w obu kierunkach.

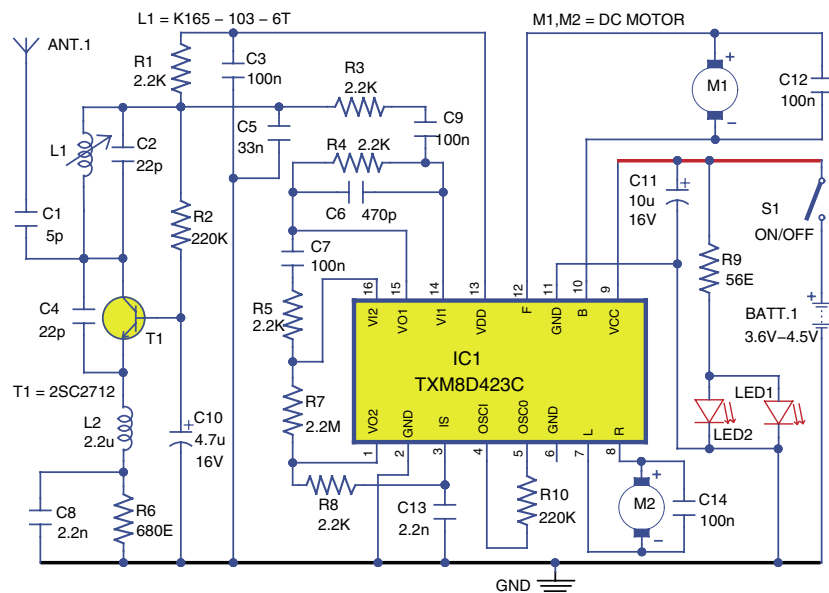
Na rysunku 4 pokazano zestaw podzespołów: nadajnik, odbiornik oraz dwa

silniczki. Zestaw taki jest dostępny w wielu sklepach z zabawkami i przewidziany jest dla własnych konstrukcji lub jako zespół wymienny w zabawkach wyposażonych w zdalne sterowanie. Pełna zabawa elektronika-hobbysty nakazywałaby montaż nadajnika i odbiornika na płytce uniwersalnej. Skoro jednak dostępne są takie podzespoły

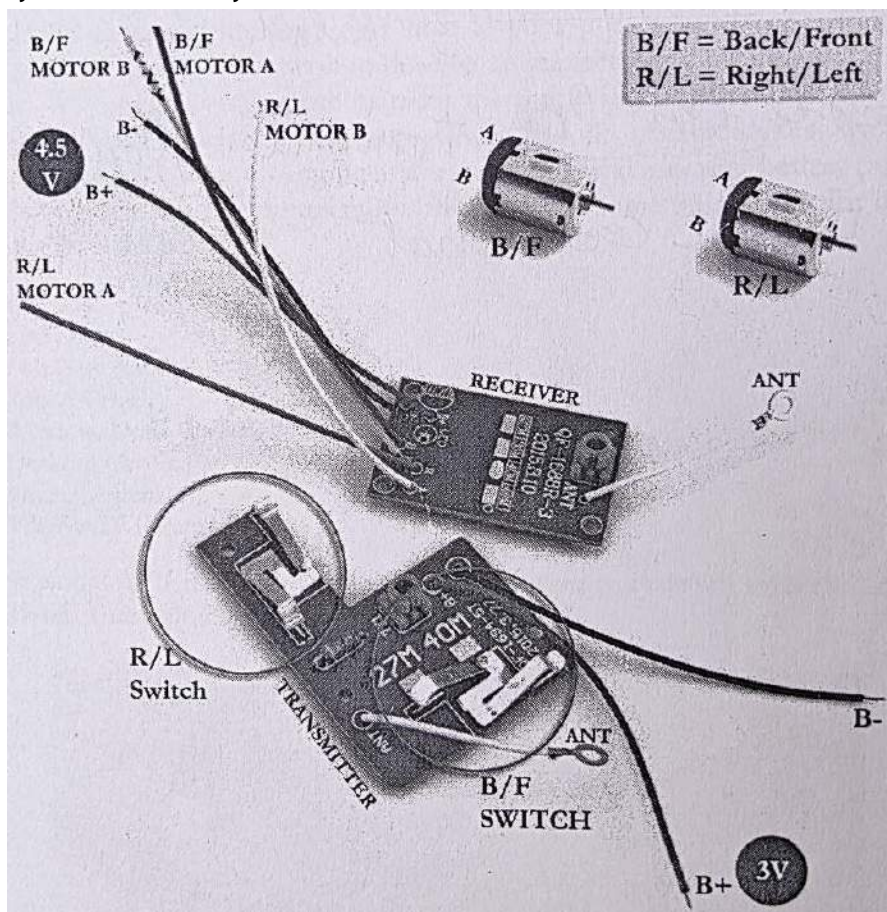
pracujące w paśmie 27 MHz i są to elementy niedrogie, trudno oprzeć się pokusie aby je wykorzystać. Zatem, podłącz tylko 3-woltową baterijkę do nadajnika i zamknij całość w odpowiednio przygotowanej obudowie. Także, do prac mechanicznych ograniczą się czynności w obrębie odbiornika. Możesz skonstruować własną zabawkę podłączając silniki do kół i/lub przekładni. Jakość tak wykonanej zabawki będzie w większym stopniu zdeterminowana konstrukcją mechaniczną. Na rysunku 4 opisane są wszystkie wyprowadzenia płytki nadajnika, odbiornika i wykonawczych silniczków. Wystarczy podzespoły te połączyć zgodnie z tym opisem i odbiornik zasilić baterijką 4,5 V. ■

T.K. Hareendran

Ten artykuł był wcześniej opublikowany na łamach „EFY”, październik 2021 (efymag.com)



Rysunek 3. Schemat ideowy odbiornika



Rysunek 4. Moduły zdalnego sterowania z opisem wyprowadzeń

Wykaz elementów, kupuj w sklepie.avt.pl
(W-wa, ul. Leszczynowa 11, tel. +48222578451,
e-mail: handlowy@avt.pl):

Elementy w nadajniku

Półprzewodniki:

IC1 – YX4116 – układ radiowego nadajnika
T1, T2 – 9013 – tranzystor npn

Rezystory: (wszystkie 0,25 W, ±5%)

R1 – 10 kΩ

R2 – 100 kΩ

Kondensatory:

C1 – 0,1 μF ceramiczny

C2, C5, C6 – 100 pF ceramiczny

C3 – 30 pF ceramiczny

C4 – 150 pF ceramiczny

Pozostałe:

ANT1 – antena 10 cm

CON1 – złącze 2-pinowe

S1, S2, S3, S4 – przyciski dotykowe

L1, L2 – ceweczka 2,2 μH

L3 – cewka 9,2 μH

XTAL1 – 27,145 MHz – rezonator kwarcowy

Bateria 3 V

Elementy w odbiorniku

Półprzewodniki:

IC1 – TXM8D423C – odbiornik radiowego zdalnego sterowania

T1 – 2SC2712 – tranzystor npn

LED1, LED2 – dioda LED 3 mm

Rezystory: (wszystkie 0,25 W, ±5%)

R1, R3, R4, R5, R8 – 2,2 kΩ

R2, R10 – 220 kΩ

R6 – 680 Ω

R7 – 2,2 MΩ

R9 – 56 Ω

Kondensatory:

C1 – 5 pF – ceramiczny

C2, C4 – 22 pF – ceramiczny

C3, C7, C9, C12, C14 – 100 nF – ceramiczny

C5 – 33 nF – ceramiczny

C6 – 470 pF – ceramiczny

C8, C13 – 2,2 nF – ceramiczny

C10 – 4,7 μF/16 V – elektrolityczny

C11 – 10 μF/16 V – elektrolityczny

Pozostałe:

ANT1 – antena 10 cm

CON1 – złącze 2-pinowe

L1 – K165-103-6T – cewka strojona

L2 – cewka 2,2 μH

M1, M2 – silnik 3 V–4,5 VDC

S1 – przełącznik on/off

BATT1 – bateria 3 V lub 4,5 V

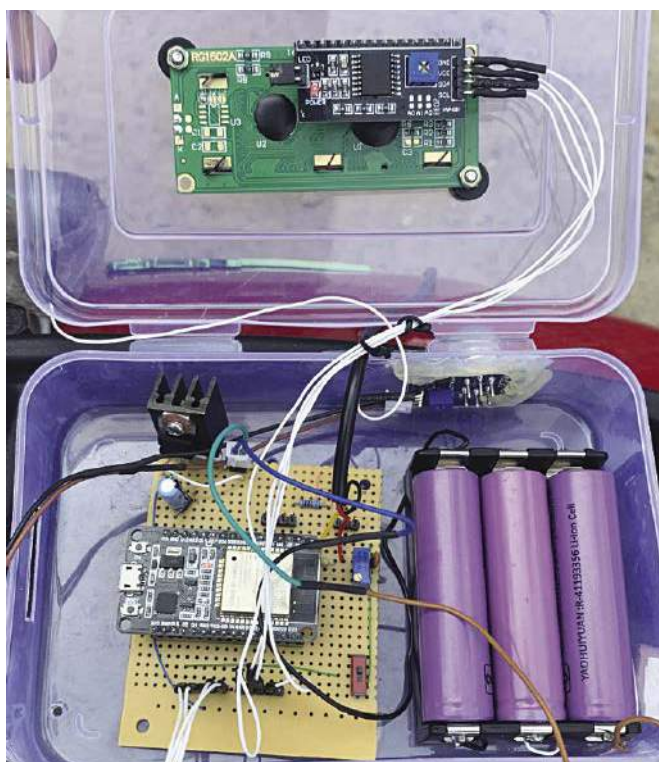
System nadzoru temperatury baterii w pojazdach elektrycznych

Powodem zajęcia się tym tematem są coraz częstsze przypadki pożaru spowodowanych błędami technicznymi baterii w elektrycznych rowerach jak również innych pojazdach zasilanych z dużych baterii. Jak się okazuje, powodem tak poważnych zdarzeń jest brak poprawnego nadzoru pracy akumulatora i jego temperatury. Bieżący projekt wychodzi na przeciw potrzebom w tym zakresie. Monitorowana jest na bieżąco temperatura jak również oznaki dymu mogące świadczyć o możliwości powstania tak niebezpiecznego zdarzenia jak zapalenie się baterii jak również bardziej rozległy pożar.

Niebezpieczeństwo takie występuje we wszystkich pojazdach elektrycznych, gdyż ich baterie gromadzą na tyle dużo energii, że pozwolą na znaczny wzrost temperatury. Nasz monitor ma zaalarmować użytkownika jak również natychmiast odłączyć obciążenie od akumulatora. Prototyp wykonany przez autora pokazuje zdjęcie na **rysunku 1**. Na **rysunku 2** pokazano schemat blokowy systemu mającego zabezpieczyć przed tak niebezpiecznym zdarzeniem.

W pojazdach elektrycznych powszechnie znalazły zastosowanie akumulatory litowo-jonowe. Zakres dopuszczalnej temperatury ich pracy i przechowywania specyfikowany jest w przedziale -20°C do $+60^{\circ}\text{C}$. Aczkolwiek dbałość o żywotność baterii nakazuje aby był to przedział $+10$ do $+40^{\circ}\text{C}$.

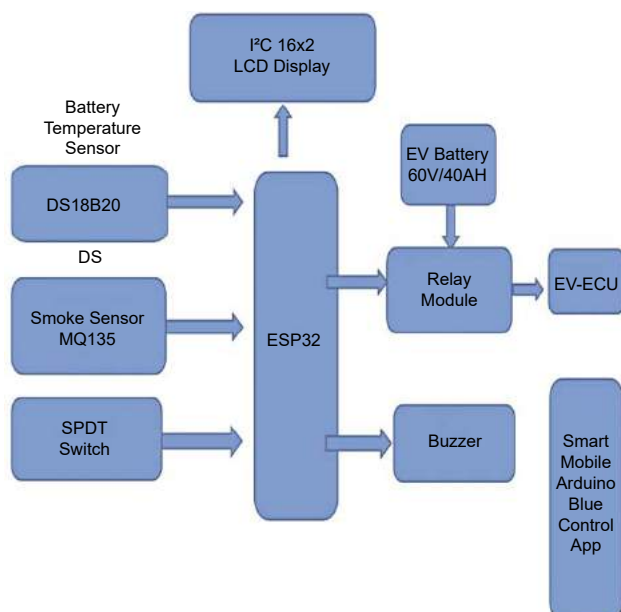
Zaproponowany tu system zasilany jest z dodatkowej baterii 12 V i wykonany jest na bazie takich podzespołów jak: płytka z mikrokontrolerem ESP32 (U1), czujnik temperatury DS18B20 (U2), detektor dymu MQ135 (U3), wyświetlacz LCD 2×16 znaków (U4) oraz regulator napięcia LM7805 (U5) i kilka prostych niedrogich elementów. Jako wejście dla systemu ESP32 należy traktować także dwupozycyjny przełącznik (typu SPDT) oraz potencjometr użyteczny podczas konfiguracji całego urządzenia. Podzespołami widzianymi jako wyjście mikrokontrolera jest wyświetlacz obsługiwany magistralą I²C, buzzer alarmujący i przełącznik którego zadaniem jest odłączenie zasilania w sytuacji zagrożenia. Jako wyjście dla ESP32 należy widzieć także aplikację Blue Control na którą zostaną przesłane dane do odczytu przez użytkownika. Celem zastosowania przełącznika SPDT S1 jest wybór trybu pracy



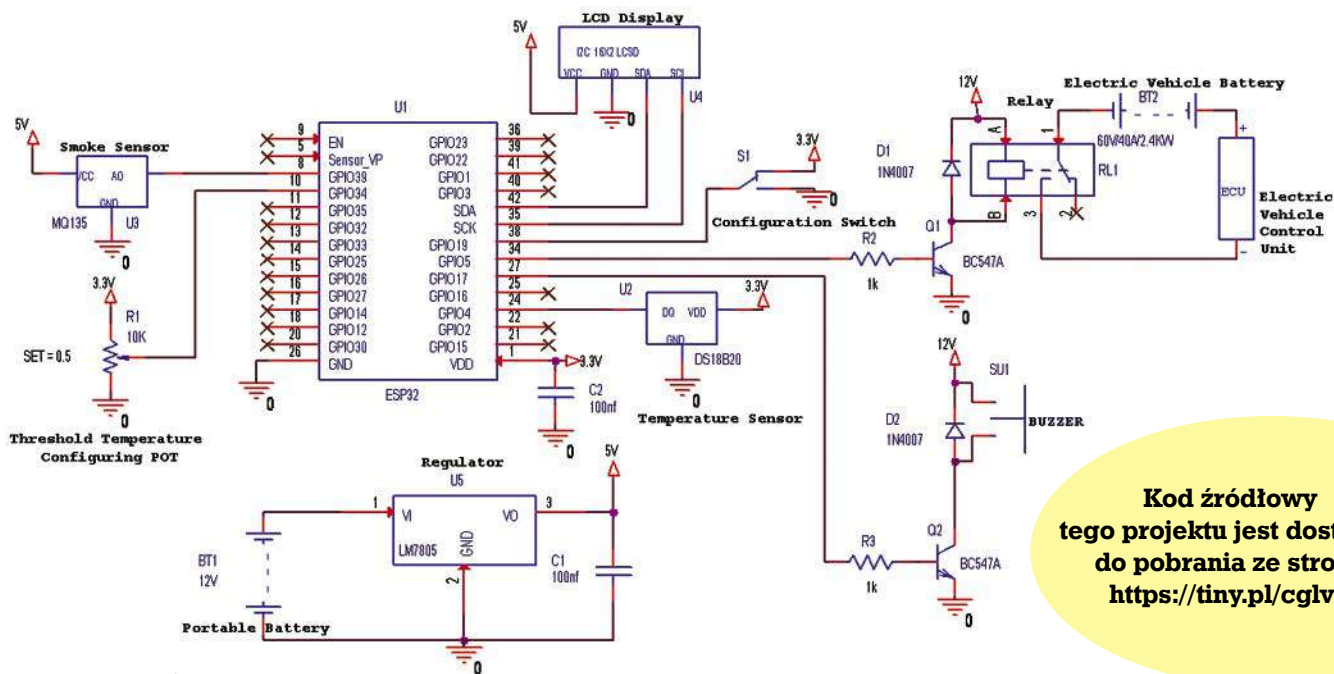
Rysunek 1. Prototyp wykonany przez autora

Wykaz elementów, kupuj w sklep.avt.pl (W-wa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl):

ESP32 (U1) – 1 szt. – płytka mikrokontrolera ESP32
 DS18B20 (U2) – 1 szt. – czujnik temperatury
 MQ135 (U3) – 1 szt. – detektor dymu
 RG1602A (U4) – wyświetlacz LCD z interfejsem I²C – 2×16 znaków
 LM7805 (U5) – 1 szt. – stabilizator 5 V
 10 kΩ (R1) – 1 szt. – potencjometr
 1 kΩ (R2, R3) – 2 szt. – rezystor 0,25 W
 1N4009 (D1, D2) – 2 szt. – dioda prostownicza
 100 nF (C1, C2) – 2 szt. – kondensator ceramiczny
 SU1 – 1 szt. – Buzzer
 GU-SH112D (RL1) – 1 szt. – przełącznik o dużej obciążalności styków
 BC547A (Q1, Q2) – 2 szt. – tranzystor npn
 Przełącznik on/off (S1) – 1 szt. – przełącznik SPDT
 Bateria 12 V (BT1) – 1 szt. – bateria litowo-jonowa 60 V/40 A/2,4 kW (BT2) – 1 szt. – bateria w pojeździe elektrycznym
 Smartfon – 1 szt. – urządzenie mobilne z systemem Android
 Aplikacja Arduino Blue Control – aplikacja wymagana podczas instalacji
 IDE Arduino Software – wymagane podczas instalacji



Rysunek 2. Schemat blokowy systemu nadzoru temperatury akumulatora



Kod źródłowy
tego projektu jest dostępny
do pobrania ze strony
<https://tiny.pl/cglvb>

Rysunek 3. Schemat ideowy systemu

```
tms
#include <OneWire.h>
#include <DallasTemperature.h>
#include <LiquidCrystal_I2C.h>
#include "BluetoothSerial.h"

#if !defined(CONFIG_BT_ENABLED) || !defined(CONFIG_BLUEDROID_ENABLED)
#error Bluetooth is not enabled! Please run 'make menuconfig' to and enable it
#endif

BluetoothSerial SerialBT;

String message = "";
char incomingChar;
String temperatureString = "";
String humidityString = "";

#define ADC_VREF_mV 4755.0
#define ADC_RESOLUTION 4096.0
#define MQ135 39
#define TEMP_REF 34

const int oneWireBus = 4;

LiquidCrystal_I2C lcd(0x27, 16, 2);
const int SW = 19;
const int BUZZAR = 17;
const int EV_ECU = 5;
const float threshold2=400.00;
```

Rysunek 4. Zrzut fragmentu kodu źródłowego programu

mikrokontrolera. Oprócz normalnego trybu użytkowania systemu jest bowiem też niezbędny tryb jego konfiguracji.

Oprogramowanie

W celu zaprogramowania płytki mikrokontrolera ESP konieczne jest Arduino IDE. IDE można zainstalować z poziomu menadżera płytki ESP. Należy ściągnąć najnowszą wersję Arduino IDE. Następnie trzeba zainstalować biblioteki które zapewnią obsługę interfejsu czujników. W bieżącym systemie wykorzystano termometr Dallasa obsługiwany linią 1-wire oraz wyświetlacz ciekłokrystaliczny

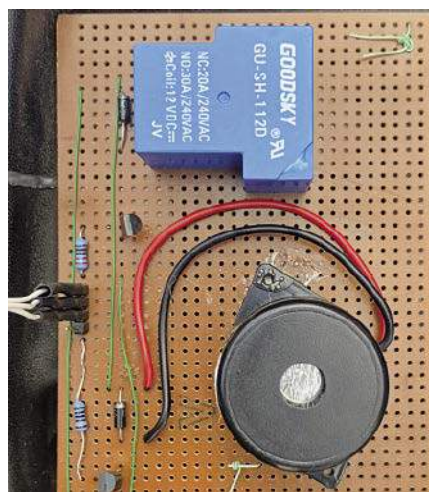
obsługiwany magistralą dwuprzewodową. Zatem z biblioteki menadżera należy wybrać odpowiednie podprogramy. Należy też załadować kod programu który obsługuje nasz projekt. Kod ten jest dostępny pod adresem: <https://tiny.pl/cglvb>.

Poprawność wykonania procesu programowania wymaga ustawienia właściwego portu i typu płytki w programie Arduino IDE. Płytke ESP32 należy skomunikować z komputerem i wymusić tryb pozwalający na uploading kodu programu.

Konstrukcja i testowanie pracy urządzenia

Wpierw należy wykonać czynności programowe konieczne dla przepisania kodu źródłowego tms.ino do ESP32. Następnie należy zmontować układ zgodnie ze schematem na rysunku 3.

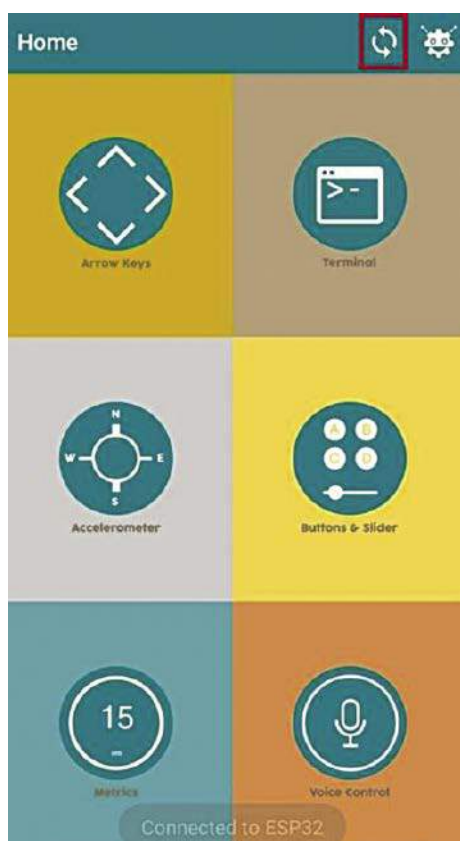
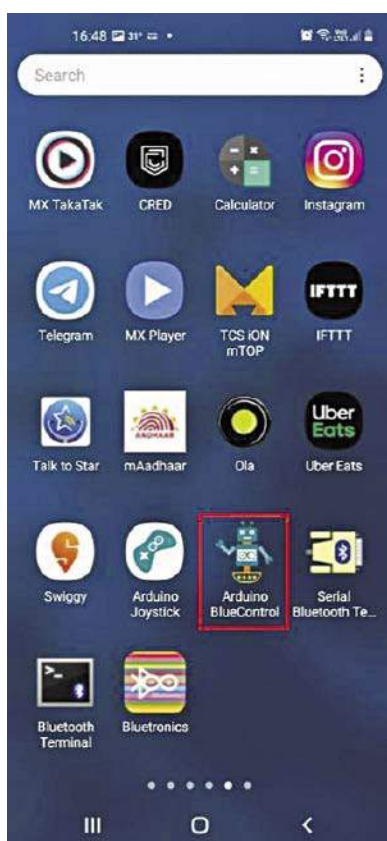
Podczas uruchamiania systemu w miejsce akumulatora i stabilizatora 5-cio woltowego można posłużyć się adapterem 5VDC lub baterią o nominalnym napięciu 5 V. Na rysunku 5 pokazano jak w prototypie wykonanym przez autora wykonano płytke z buzzerem i przekaźnikiem. Na rysunku 6



Rysunek 5. Podłączenie przekaźnika w prototypie



Rysunek 6. Testowanie urządzenia w pojeździe elektrycznym



Rysunek 7. Połączenie Bluetooth z urządzeniem mobilnym

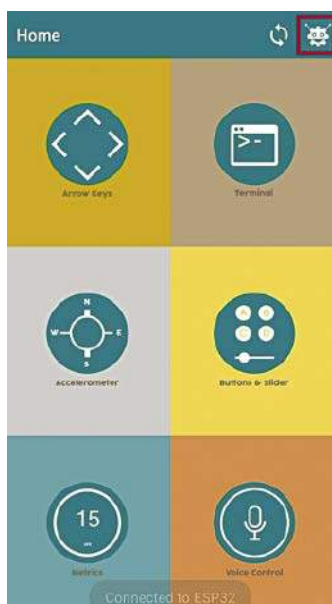
jest zdjęcie jak autor umieścił całość układu w pojeździe elektrycznym. Zdjęcia z prac testowych projektu pokazano na **rysunkach 6 do 9**. Informacja o temperaturze i ew. wykryciu dymu przez smoke detector wyświetlana jest także w aplikacji na smartfonie. Telefon należy skomunikować po łączu Bluetooth. Dla pracy systemu jest więc potrzebna też aplikacja Blue Control. Końcowe

prace sprowadzają się do zasilenia systemu i skomunikowaniu aplikacji na telefonie. Wtedy na telefonie powinieneś odczytać temperaturę którą mierzy czujnik przymocowany do baterii. Jeśli zostanie przekroczona progiowa wartość temperatury (którą wcześniej ustawiłeś) lub zostanie uaktywniony czujnik dymu, powinien zostać uruchomiony alert zarówno w aplikacji jak i za pośrednictwem

buzzera. Ponadto, zastosowany przełącznik powinien spowodować natychmiastowe odcięcie obciążenia od baterii w twoim pojeździe elektrycznym. ■

E.Venkatesan and M.Dinesh

Ten artykuł był wcześniej opublikowany na łamach „EFY”, styczeń 2023 (efymag.com)



Rysunek 8. Ustawienie alarmów

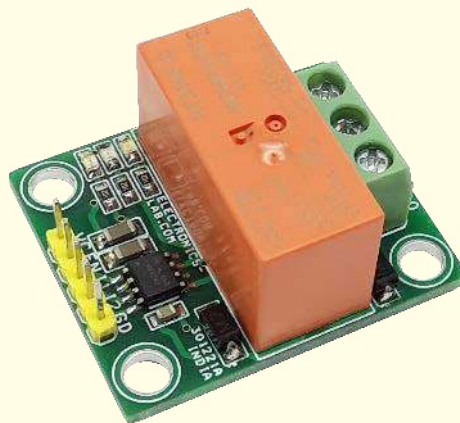


Rysunek 9. Wyświetlanie informacji w aplikacji telefonu

Przedstawiamy początkowe fragmenty dwóch projektów ze zbioru kilkudziesięciu projektów dostępnych wyłącznie dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl. W rubryce **DIY PLUS** zamieszczamy aktualnie najciekawsze projekty publikowane w Internecie w formacie open source. Prenumeratorów EdW zapraszamy do zapoznania się na www.elportal.pl z niezwykle inspirującymi zasobami rubryki **DIY PLUS**.

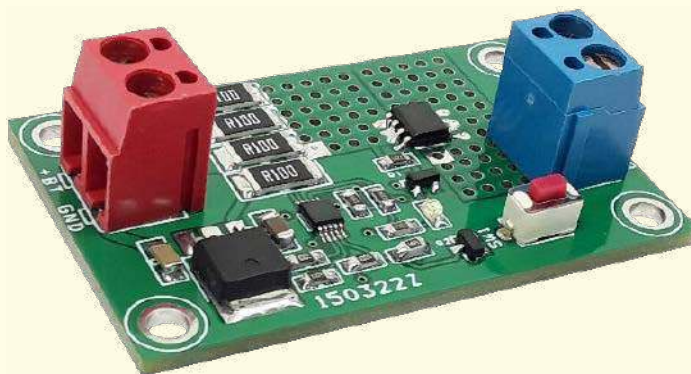
Inteligentny dwucewkowy sterownik przekaźnika blokującego – moduł przekaźnika bistabilnego

Ta inteligentna płyta Dual Coil latching Relay może kontrolować moc ON/OFF urządzenia poprzez zastosowanie krótkiego impulsu napięcia do wejścia 1 i wejścia 2. Ten projekt jest przydatny w aplikacjach o niskiej mocy, ponieważ cewka nie jest zasilana przez cały czas i wymaga tylko krótkiego impulsu napięcia. Cewka przekaźnika pozostaje w tej pozycji nawet po odłączeniu zasilania. Projekt zbudowany jest z wykorzystaniem układów FAN3240 firmy ONSEMI. Układ ten zawiera podwójne wysokoprądowe sterowniki przekaźników zaprojektowane do napędzania dwucewkowych spolaryzowanych przekaźników zatrzaśkowych, które łączą i odłączają zasilanie w inteligentnych licznikach elektronicznych oraz w aplikacjach z falownikami słonecznymi. Układ scalony zawiera automatyczne wyłączenie termiczne, a blok filtrów/timerów ma zapobiegać niezamierzonemu przełączeniu z hałaśliwych sygnałów wejściowych, zapewniając kwalifikację impulsu wejściowego (tQUAL) i ograniczenie maksymalnej szerokości impulsu wyjściowego (tMAX). Przekaźniki polaryzacyjne, dwustabilne, zatrzaśkowe są wykorzystywane w wielu rodzajach urządzeń elektronicznych i różnorodnych aplikacjach.



Zabezpieczenie nadprądowe/ nadrozładowcze do akumulatorów kwasowo-ołowiowych

Opisywany projekt zabezpiecza i monitoruje akumulator kwasowo-ołowiowy przed zbyt niskim napięciem akumulatora i stanami nadprądowymi. Układ składa się ze wzmacniacza MAX4373 current-sense z wewnętrznymi podwójnymi komparatorami oraz P-kanalowego MOSFETu w szeregu z akumulatorem i jego obciążeniem. Układ pracuje jako normalnie zamknięty przełącznik, który może zostać otwarty, gdy wzmacniacz prądowy i komparatory wykryją wysoki prąd obciążenia LUB niskie napięcie baterii. W przypadku wystąpienia nadmiernego prądu, komparator blokuje wyjście i może być zresetowany za pomocą wbudowanego przełącznika SW1.



Niektóre projekty aktualnie dostępne tylko dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl:

- | | | |
|---|---|--|
| 1. Półprzewodnikowy przekaźnik mocy DC z prądowym sprzężeniem zwrotnym | 4. Automatyczny system ogrodniczy z NodeMCU i Blynk, ArduFarmBot 2 | 14. Wyświetlacz EKG z użyciem Arduino |
| 1. Wyłącznik nadprądowy – przekaźnik wyłączający nadprądowy | 5. TinyML – Rozpoznawanie ruchu przy pomocy Raspberry Pi pico | 15. Łatwy do zbudowania robot kroczący |
| 1. Uniwersalny konwerter napięcia AC – wyjście 18 V DC z wejścia 85...265 V AC | 6. Wzmacniacz piezoelektryczny do gitary i skrzypiec | 16. Sonarowy theremin MIDI |
| 1. Moduł procesora echa głosu – urządzenie opóźniające do efektów dźwiękowych, echo, reverb | 7. Wysokowydajny i niezawodny sterownik bipolarnego silnika krokowego | 17. Zamek elektroniczny na kod |
| 2. Najlepszy sposób na próbkowanie dźwięku za pomocą ESP32 | 8. Sterownik silnika prądu stałego z wykorzystaniem przekaźnika i mosfetu – interfejs Arduino | 18. Prosty tester tranzystorów |
| 1. Choinka z Arduino i pikselowymi diodami | 9. Przedwzmacniacz do mikrofonu MEMS | 19. Zegar binarny z użyciem Microbit |
| 2. RPi – stacja pogodowa IoT | 10. Super prosty czuły wykrywacz metali | 20. Przetwornik częstotliwości na napięcie (tachometr) – przetwornik częstotliwości na napięcie z czujnikiem magnetycznym o zmiennej reluktancji |
| 3. Niskobudżetowy monitor jakości powietrza IoT oparty o RaspberryPi 4 | 11. Stymulator czaszkowy Arduino (Bio-BrainTuner) | 21. Izolowany obwód wykrywania napięcia 250 V AC z pojedynczym wyjściem (wejście 250 V prądu przemiennego, wyjście 5 V) |
| | 12. Generator sygnałów AD9833 | |
| | 13. Obserwacja charakterystyk tranzystora | |

Miesięcznik „Elektronika dla Wszystkich” (12 numerów w roku) jest wydawany we współpracy z kilkoma redakcjami zagranicznymi



Wydawnictwo:
AVT-Korporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Adres redakcji:
03-197 Warszawa, ul. Leszczyńska 11
e-mail: edw@elportal.pl, www.elportal.pl

Redaktor merytoryczny:
Paweł Sujko

Dział Reklamy:
Katarzyna Gugala
katarzyna.gugala@elportal.pl, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobanski
jakub.sobanski@elportal.pl

Copyright AVT-Korporacja Sp. z o.o., Warszawa, ul. Leszczyńska 11. Projekty publikowane w „Elektronice dla Wszystkich” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki dla Wszystkich”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice dla Wszystkich” jest dozwolone wyłącznie po uzyskaniu pisemnej zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice dla Wszystkich”.

DTP, okładka, redakcja strony internetowej www.elportal.pl:
MAD Sp. z o.o.

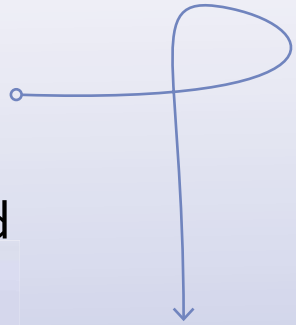
Prenumerata:
W Wydawnictwie AVT, e-mail: prenumerata@avt.pl
tel. 22 257 84 22, (godz. 10:00–14:00)

W RUCH S.A., e-mail: prenumerata@ruch.com.pl
tel. 801 800 803, 22 717 59 59, www.prenumerata.ruch.com.pl

Subscribe to Elektor's newsletter and get the chance to

WIN

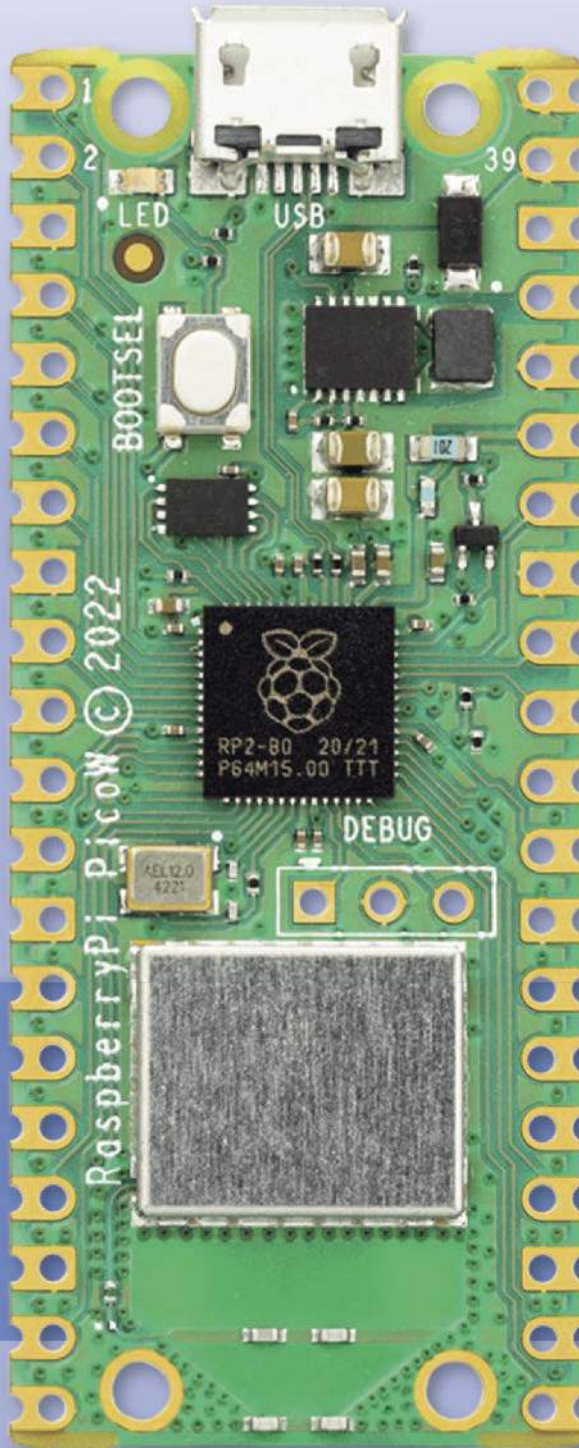
a Raspberry Pi Pico W board



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