

ELEKTRONIKA

dla wszystkich

nr 2/2024 (337) • luty • www.elportal.pl

DIY PLUS
tylko dla prenumeratorów

Rejestrator stanu akumulatorów

PROJEKTY dla elektroników

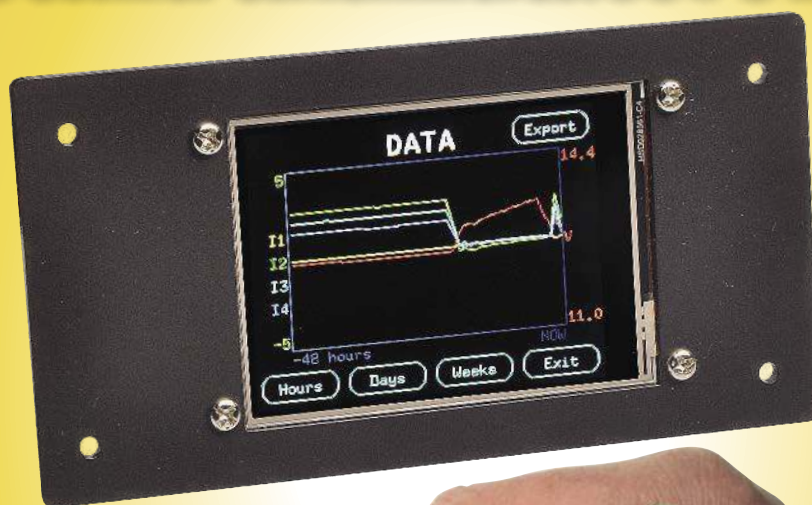
- ▶ Głośnik niskotonowy z dopingiem
- ▶ Pęseta do testów SMD
- ▶ Prosta, liniowa klawiatura muzyczna MIDI

DIY dla wszystkich

- ▶ Tani system automatyki z użyciem uP PIC16F676
- ▶ Robot sterowany gestami i klaśnięciem rąk
- ▶ Elektroniczna maszyna do głosowania w demokratycznych wyborach

TUTORIALE

- ▶ Audio OUT: Budujemy radio tranzystorowe, część 1
- ▶ Chirurgia obwodowa: Rozwiązywanie obwodu: iteracyjne proste obciążenia i symulacja
- ▶ Ekscytacje Maxa: Migające diody LED i śliniacy się inżynierowie
- ▶ Precyzyjny filtr aktywny drugiego rzędu do subwoofera
- ▶ Akumulatory
- ▶ Aneng AN870, multimetr 19999 – test multimetru
- ▶ Elektrochemia
- ▶ Pokój Nauczycielski: Prosty automatyczny wyłącznik światła w łazience



**Pęseta
do testów
SMD**



ISSN 1425-1698 Indeks 33362X
02 >
9 771425 169245
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
przekazniki
radiatory
obudowy
i wiele więcej...

www.piekarz.pl





Najbardziej popularne kity AVT

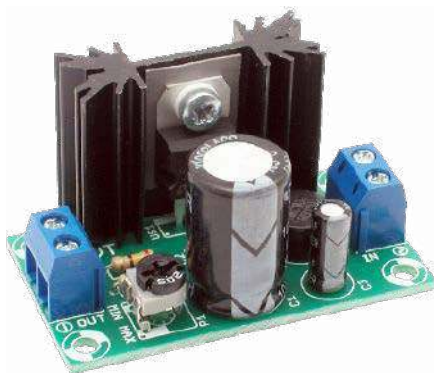
Poznaj listę **TOP 100** na www.elportal.pl/kityavt



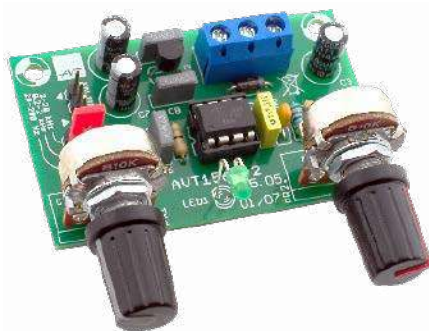
AVT1476 Automatyczny włącznik zmierzchowy
<https://sklep.avt.pl/avt1476.html>



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>



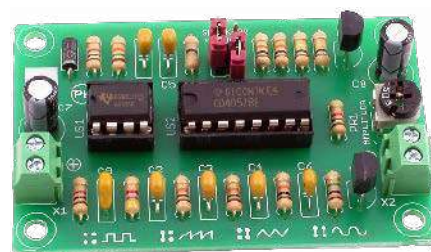
AVT5540 Radio FM z RDS
<https://sklep.avt.pl/avt5540.html>



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



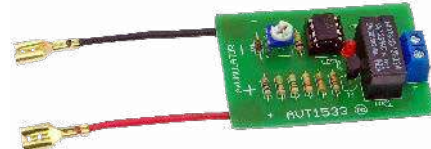
AVT1459 Uniwersalny układ czasowy
<https://sklep.avt.pl/avt1459.html>



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**

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

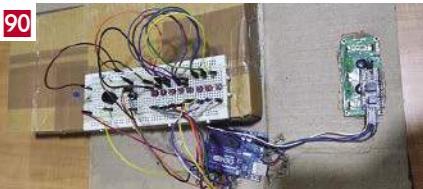


projektów w zbiorze DIY+ na elportal.pl/diy

* 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



Projekty dla elektroników:

- Poza zestawem? W zestawie z podtrzymaniem baterijnym?
- Jak monitorować stan akumulatorów? Rejestrator stanu akumulatorów 8
- Głośnik niskotonowy z dopingiem 17
- Pęseta do testów SMD..... 27
- Prosta, liniowa klawiatura muzyczna MIDI, część 3..... 33

Tutoriale:

- Audio OUT: Budujemy radio tranzystorowe, część 1 37
- Chirurgia obwodowa: Rozwiązanie obwodu:
iteracyjne proste obciążenia i symulacja 41
- Ekscytacje Maxa:
- Migające diody LED i śliniacy się inżynierowie (5) 45
 - Sprytne porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania .. 49
- Precyzyjny filtr aktywny drugiego rzędu do subwoofera 50
- Akumulatory 52
- Aneng AN870, multimetr 19999 – test multimetru 55
- Elektrochemia..... 60
- Edukacja w EdW dla szkół i uczelni: Wykład 15. Obwody mostkowe 64
- Pokój Nauczycielski: Prosty automatyczny wyłącznik światła w łazience ... 76

DIY dla wszystkich:

- Tani system automatyki z użyciem uP PIC16F676 80
- Robot sterowany gestami i klaśnięciem rąk 85
- Elektroniczna maszyna do głosowania w demokratycznych wyborach..... 90

DIY PLUS

- Jednokanałowy bezprzewodowy włącznik/wyłącznik PS3..... 91
- Stereofoniczny procesor audio do telewizora..... 91

Rubryki stałe:

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

A za miesiąc w marcowym EdW



* Programowany hybrydowy zasilacz laboratoryjny z Wi-Fi

W tej konstrukcji unikamy stosowania dużych transformatorów i problemów z wytwarzaniem ciepła, gdyż zastosowano impulsowy przetwornik AC/DC. Dla kompaktowej konstrukcji osiągane parametry są dość imponujące: 0...27 V; 0...5 A do 18 V. Możliwość sterowania za pomocą sieci Wi-Fi lub ekranu dotykowego z pokręteł obrotowym.

* 20 A sterownik szybkości obrotów silnika DC

W wielu zastosowaniach silników DC, np. wózki elektryczne, rowery lub skutery elektryczne, zdalnie sterowane samochody i łodzie potrzebne jest sterowanie szybkości obrotów silnika. Ten projekt idealnie nadaje się do sterowania silników na prąd stały o natężeniu do 20 A i napięciu do 24 V (maksymalnie 30 V). Układ jest prosty w budowie.

* Rejestrator stanu akumulatorów, część 2

Znajomość stanu akumulatorów jest niezbędna do utrzymania ich przez długi czas w dobrej kondycji. W części pierwszej przed miesiącem przedstawiliśmy układ, który na jednej płycie drukowanej łączy wszystkie funkcje rejestratora stanu akumulatorów o prądzie 10 A. Teraz zajmiemy się konstrukcją i procedurami kalibracyjnymi tego układu.

* Trenażer SMD

Nie masz wyboru, musisz się nauczyć montować elementy SMD, które występują obecnie również w wielu projektach dla hobbistów. Żeby ta nauka nie kosztowała Cię zbyt drogo z powodu zmarnowanych płytek PCB i zepsutych komponentów, lepiej skorzystaj z proponowanego rozwiązania – ćwiczy na trenażerze SMD. Do kompletu z trenażerem dobrze jest posłużyć się pęsetą do testów SM, opisaną w lutowym wydaniu EdW.

* Plus kolejna porcja intrygujących projektów DIY.

* Plus wiele artykułów w Twoich ulubionych cyklach Tutoriali.

**W kioskach
od 1 marca**

Refleksyjnych słów parę o mostkach i mostach

Temat wykładu w tym numerze EdW to klasyka elektryki. Trudno byłoby znaleźć elektronika, który nie słyszał o mostku Wheatstone'a lub mostku Graetza. Istnieje jeszcze kilkanaście mniej znanych mostków pomiarowych i układów mostkowych, ale historycznie pierwszy i najbardziej znany jest mostek opisany przez Charlesa Wheatstone'a w roku 1843.

Wynalazcą tego układu, działającego na zasadzie porównania potencjałów dwóch dzielników oporowych o wspólnym zasilaniu, był w rzeczywistości Samuel Hunter Christie, który dokonał tego 10 lat wcześniej, to jest w 1833 roku. Jednak dopiero publikacja Wheatstone'a, w której autor uczciwie wskazał Christie jako wynalazcę, spopularyzowała ten układ, a ponieważ Wheatstone był sławnym uczonym, bardziej znanym niż Christie, to świat jemu przypisał autorstwo tego wynalazku. **Refleksja 1.** W nauce jak w handlu, liczy się marka (brand), czyli znane nazwisko. Aby wynalazek zaistniał trzeba go rozpropagować pod dobrą marką, czyli znanym nazwiskiem.

Historia mostka Graetza pokazuje, że liczy się coś jeszcze. Układ opisany przez niemieckiego fizyka Leo Graetza w 1897 roku był przedstawiony rok wcześniej przez polskiego wynalazcę z Sanoka Karola Pollaka. Nie był to jakiś prowincjonalny amator, tylko świetnie wykształcony Europejczyk, który kierował fabrykami w Anglii, Niemczech i Francji i miał blisko 100 patentów (jest nazywany polskim Edisonem), między innymi wynalazł kondensator elektrolityczny. A jednak mamy mostek Graetza, a nie mostek Pollaka. **Refleksja 2.** Liczy się *country brand*. *Niemiecki fizyk to brzmi lepiej niż polski wynalazca* (zresztą Polski wtedy nie było). Liczy się miękka siła marki kraju.

Popatrzmy jak aktualnie wygląda pozycja Polski w Global Soft Power Index. Zajmujemy 33. miejsce, poniżej naszych realnych możliwości, bo jesteśmy 21. gospodarką świata. Jednak cieszy fakt, że ostatnio w rankingu Soft Power awansowaliśmy z 40. miejsca aż o 7 pozycji. Wiele zyskaliśmy w oczach świata dzięki naszej pomocy Ukraincom. Ta miękka siła kraju przekłada się na poprawę naszej pozycji w handlu zagranicznym, inwestycjach, turystyce itd. **Refleksja 3.** Warto zachować się przyzwoicie.

Wśród wielu typów mostków jest też mostek rezonansowy. Wprowadźmy do Google frazę *resonance bridge* i zobaczymy najpierw spektakularną katastrofę amerykańskiego mostu wiszącego (Tacoma Bridge) w 1940 roku. Początkowo za przyczynę katastrofy uznano zjawisko rezonansu. Później niebicie uodowodniono, że to nieprawda, ale kolaps tego mostu uparcie kojarzony jest z rezonansem. **Refleksja 4.** Katastrofa z nieprawdziwą, ale fajną teorią lepiej się sprzedaje medialnie niż prawdziwa wiedza o mostku rezonansowym.

A jednak ostatecznie imponują nam inne mosty wiszące, np. znane z setek filmów amerykańskich most łączący Manhattan z Long Island, który istnieje już prawie 150 lat i nie zamierza się zawalić.

Brooklyn Bridge. Ten most to dzieło inżynierów Johna Roeblinga i jego syna Washingtona. Powstał w latach 1869–1883. Pomysłodawca i projektant mostu John Roebling zmarł już w pierwszym roku budowy na tęczę po zranieniu stopy na budowie. Kierownictwo budowy przejął jego 32-letni syn Washington, którego niebawem dosięgła choroba dekompresyjna podczas prób kesonów. Został totalnie sparaliżowany i mógł poruszać tylko jednym palcem. Przez jedenaście lat, leżąc na łóżku wystukiwał tym palcem instrukcje dla żony Emily, która w jego imieniu nadzorowała budowę aż do jej zakończenia w 1883 roku. Tytanicznej woli i talentowi tych trojga osób – ojca, syna i jego żony Ameryka zawdzięcza ten wspaniały, przepiękny most. **Refleksja 5.** Warto mieć cel w życiu i walczyć o niego, dopóki choć jednym palcem można poruszać.



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.

Pytanie o Senatora

Już rok temu pocięła mi ślinka, gdy zobaczyłem w EdW 9/2022 projekt kolumn głośnikowych Senator. Są piękne i parametry imponujące. Obiecałem sobie wtedy, że w przyszłości muszę je zrobić. Nadszedł ten czas. Zebrałem już elementy do zwrotnicy i przymierzam się do gromadzenia materiałów na obudowy. Autor artykułu wykorzystał zestaw szafek kuchennych Bunnings, popularnych w Australii. Poszukuję podobnego rozwiązania w sieci IKEA lub z płyt MDF w Castoramie. Są to płyty o grubości 12 mm lub 18 mm, podczas gdy w projekcie autor zastosował płyty 16 mm. Wymiar liniowy, wewnętrzny odstęp między ściankami obudowy wzrośnie o 4 mm (płyta 18 mm) lub zmaleje o 8 mm (płyta 12 mm). Czy zmiana grubości, a więc również objętości powietrza w kolumnie może mieć istotne znaczenie dla jakości brzmienia?

Red. Tak małe zmiany rozmiarów nie powinny mieć znaczenia. Zmiana objętości o 1...2% jest znikoma wobec przybliżeń zastosowanych przy projektowaniu rozmiarów kolumn. Nie były to obliczenia o porównywalnej precyzji.

Nie tylko Prezes

Mam 32 lata, żonę i cudowną dwuletnią córeczkę. Jestem ekonomistą z zawodu i pasjonatem elektroniki. Podaję te personalia, bo dla mojego pokolenia (przynajmniej dla grona moich przyjaciół) jest typowe nie posiadanie telewizora w domu. Tak też wychowujemy nasze dziecko. Trochę ryzykujemy, bo starszy ode mnie kolega, też nie mający telewizora w domu, został wezwany do szkoły z podejrzeniem, że jego syn boryka się z problemami rodziny patologicznej i nie oglądanie telewizji ogranicza jego rozwój umysłowy. Ale wróćmy do tematu. Otóż znany ostatnio przypadek zaskoczenia zanikiem odbioru kanałów TVP dotyczy nie tylko Prezesa Narodu, ale też mojej babci, korzystającej z telewizji naziemnej. Babcia nie miała numeru telefonu do Prezesa TVP, więc zadzwoniła do mnie (na szczęście nie o 3-ej w nocy), bo w rodzinie jestem uznanym guru od elektryki. Wstyd się przyznać, ale nie wiedziałem, o co chodzi. Podszkoliłem się w internecie i wyjaśniłem babci istotę problemu oraz zaproponowałem jej kupno dekodera za 117 zł. Na to babcia, że ona w zasadzie zebrała trochę grosza (chyba te trzynaste, czternaste, itd emerytury) i chętnie kupiłaby nowy telewizor o trochę większym ekranie, bo mężczyźni czytanie pasków. Oczywiście liczy na mój wybór telewizora. Od lat w ogóle nie interesowałem się telewizją, więc od zera zacząłem studiować temat. Jest taka różnorodność technologii, że głowa boli. LED, LCD, QD-mini LED, OLED, Neo QLED i wiele innych cech, funkcji itd. Poszukując informacji w internecie znajduję przede wszystkim opisy na poziomie komercyjnym, który nie wystarczą mi, bo chciałbym też rozumieć skąd się biorą różnice poszczególnych rozwiązań. W zasadzie trochę jestem zażenowany tym, że brakuje mi tej wiedzy. Przecież jestem pasjonatem elektroniki. Stąd prośba do redakcji EdW, żeby włączyć do profilu Waszego fantastycznego miesięcznika również tematykę sprzętu elektronicznego w domu.

Ryszard

Red. Przed wieloma laty była w EdW rubryka MEU (Magazyn Elektroniki Użytkowej). Może warto ją reaktywować w formie cyklu artykułów monotematycznych lub pytań i odpowiedzi, albo leksykonu. Chętnie to zrobimy, ale najpierw chcielibyśmy wysłuchać opinii szerszego grona Czytelników. Prosimy o listy z propozycjami na adres edw@elportal.pl

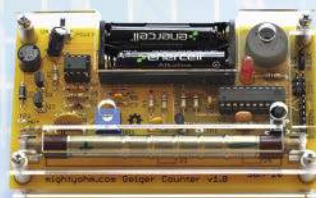
Patronat AVT

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

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

Elektor Bestsellers

SAVE UP TO
26% NOW!



www.elektor.com/sale/deals



eprasa.pl/08b200960c

Poza zestawem? W zestawie z podtrzymaniem bateryjnym? Jak monitorować stan akumulatorów?



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/cbpq1>
Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania

Rejestrator stanu akumulatorów

Znajomość stanu baterii/akumulatorów jest niezbędna do utrzymania ich przez długi czas w dobrej kondycji. System, który może monitorować i rejestrować istotne statystyki akumulatora lub ich baterii, jest bardzo przydatny i może pomóc uniknąć konieczności kosztownych wymian. Może być również używany do rozwiązywania problemów, na przykład, gdy nie wiadomo, który moduł czy panel jest odpowiedzialny za okresowe rozładowywanie baterii.

Energia słoneczna i wiatrowa jest coraz częściej wykorzystywana i coraz tańsza, więc istnieje potrzeba konserwacji akumulatorów związanych z takimi systemami. Możesz także mieć duży akumulator w garażu, przyczepie kempingowej, łodzi lub innym pojeździe, który musisz sprawdzać. Baterie zapasowe na wypadek awarii zasilania sieciowego to kolejny przypadek, w którym może być potrzebny monitor lub rejestrator stanu baterii.

Nasz nowy „Battery Monitor Logger” czyli strażnik stanu akumulatorów, jest wszechstronny i wydajny, będąc w stanie, po wyjęciu z pudełka, obsłużyć ładowarkę i dwa oddzielne obciążenia. Wykorzystuje schemat Micromite LCD BackPack, więc może być przeprogramowany w MMBasic, wariacie języka BASIC Micromite. Ale ponieważ napisaliśmy oprogramowanie z wieloma przydatnymi funkcjami, nie musisz zajmować się programowaniem.

Ostatnio w Silicon Chip-ie opublikowaliśmy miernik pojemności akumulatorów w czerwcu i lipcu 2009 roku (www.siliconchip.com.au/Series/44).

Zawierał on mikroprocesor PIC zdolny do monitorowania napięcia i prądu akumulatora poprzez zewnętrzny boczny prądowy. Mógł on rejestrować dane, a także obliczać takie parametry jak pojemność baterii i szacowany czas jej pracy.

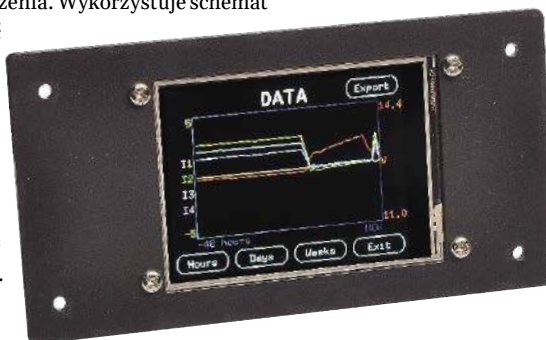
Nowe funkcje

Miernik pojemności akumulatora z 2009 roku wykorzystywał pojedynczy boczny, wskutek czego mógł monitorować tylko całkowity prąd dopływający lub wypływający z podłączonego akumulatora.

Nasza nowa konstrukcja obsługuje do trzech boczników, dzięki czemu może monitorować trzy oddzielne ścieżki prądowe, pomagając rozdzielić dane ładowania lub rozładowania na wiele obciążeń i/lub generatorów czy innych zasilaczy.

Zawiera nawet czwarty wewnętrzny boczny do monitorowania własnego zużycia energii.

Na przykład, możesz mieć panel słoneczny i generator wiatrowy (lub nawet



kilka) i chcesz oddzielnie śledzić energię, którą generują.

Możesz też mieć kilka odbiorników, takich jak lodówka, oświetlenie i czajnik, i chcesz sprawdzić, który z nich zużywa najwięcej energii.

Stary projekt był również ograniczony do zasilania z baterii o maksymalnym napięciu do około 60 V na wejściu (w porównaniu do 100 V w tym aktualnym projekcie) i mógł również przechowywać minimalną ilość danych w mikroprocesorze PIC. Mikroprocesor PIC32, którego użyliśmy w tym rozwiązaniu, ma znacznie większą pamięć, więc może rejestrować więcej danych przez dłuższy czas.

Napięcie i prądy akumulatora są próbkowane w odstępach 10-sekundowych. Dane te są uśredniane co godzinę, co daje do dwóch dni próbek godzinowych. Próbki godzinowe są również uśredniane każdego dnia, aby uzyskać około dwutygodniowe wartości dzienne.

Przepływ zarówno ładunku, jak i energii, jest rejestrowany, aby zapewnić wartości pojemności w Ah (amperogodzinach) i Wh (watogodzinach). Użytkownik określa napięcie akumulatora w pełni naładowanego i rozładowanego, a także pojemność akumulatora, dzięki czemu urządzenie może samo skalibrować się, gdy akumulator jest w pełni naładowany lub rozładowany.

Obliczana jest również nieskomplikowana, liniowa zależność napięcia od stanu naładowania, dająca przybliżone wskazanie stanu akumulatora, gdy dokładniejsze informacje nie są dostępne.

Koncepcja działania

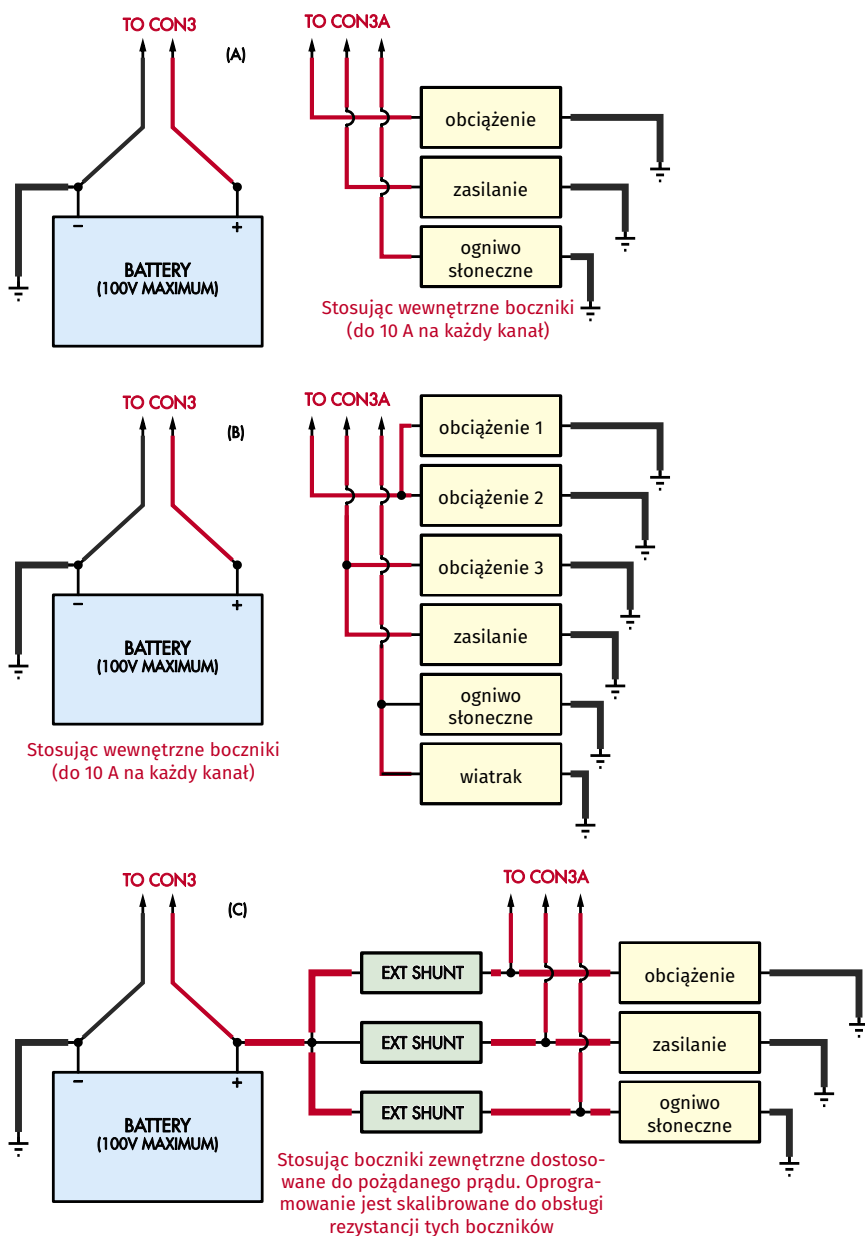
Rysunek 1a przedstawia najprostszy sposób korzystania z rejestratora kondycji akumulatora. Akumulator lub ich zestaw, o maksymalnym napięciu 100 V, podłącza się do dwudrożnego zacisku śrubowego (CON3), podczas gdy dodatnie wyprowadzenia maksymalnie trzech obciążeń lub źródeł ładowania podłącza się do styków trójdrożnego zacisku śrubowego CON3a.

Ujemne końce tych obciążeń/źródeł ładowania łączą się bezpośrednio z ujemnym biegunem akumulatora (masą).

Umożliwia to rejestratorowi MB (skrót od „monitor baterii”) niezależny pomiar i wyświetlanie prądu płynącego do lub z każdego obciążenia lub zasilacza.

Oblicza on również całkowity prąd wejściowy/wyjściowy i wykorzystuje go do śledzenia stanu naładowania akumulatora w amperogodzinach (Ah). Pomnożenie tej wartości przez bieżące napięcie akumulatora daje nominalną liczbę zmagazynowanych watogodzin (Wh) dla bieżącego stanu naładowania.

Jeśli masz do podłączenia więcej niż trzy urządzenia zewnętrzne, mogą one współdzielić zaciski na CON3a, jak pokazano na **rysunku 1b**. Na przykład, jeden zacisk jest współdzielony przez dwa obciążenia (LOAD1 i LOAD2). Pomiar na tym kanale będzie całkowitym



Rysunek 1. Trzy przykłady zastosowania rejestratora/monitora akumulatorów/baterii. Najprostszą konfiguracją, na górze, wykorzystuje wewnętrzne boczniki do monitorowania prądów (do 10 A) do lub z trzech obciążeń/źródeł ładowania. Alternatywnie, jak pokazano na (B), można podłączyć więcej niż trzy obciążenia/źródła ładowania, przy czym niektóre z nich współdzielą boczniki. W przypadku zastosowań o wyższym natężeniu prądu (do setek amperów) można użyć zewnętrznych boczników, jak pokazano na (C)

prądem obciążenia dla tych dwóch urządzeń. Inny zacisk jest współdzielony przez dwa źródła ładowania (SOLAR i WIND) i podobnie ich prądy będą sumowane.

Trzeci zacisk jest współdzielony przez LOAD3 i ładowarkę sieciową. W tym przypadku urządzenie będzie mierzyć przepływ prądu netto do/z akumulatora – tj. zauważy ładowanie akumulatora, jeśli prąd ładowarki przekroczy prąd pobierany przez LOAD3; rozładowywanie, jeśli sytuacja jest odwrotna, i wykryje zerowy prąd, jeśli te oba prądy są równe (tj. prąd LOAD3 jest dostarczany przez ładowarkę).

Jeśli chcesz monitorować prądy powyżej 10 A, możesz użyć tego samego układu, ale z zewnętrznymi bocznikami.

Zwykle mają one niższą rezystancję, a także mogą obsługiwać wyższe rozpraszanie mocy, co pozwala na bezpieczny przepływ większych prądów. Na przykład, można dość łatwo zakupić boczniki 100 A, a nawet 500 A.

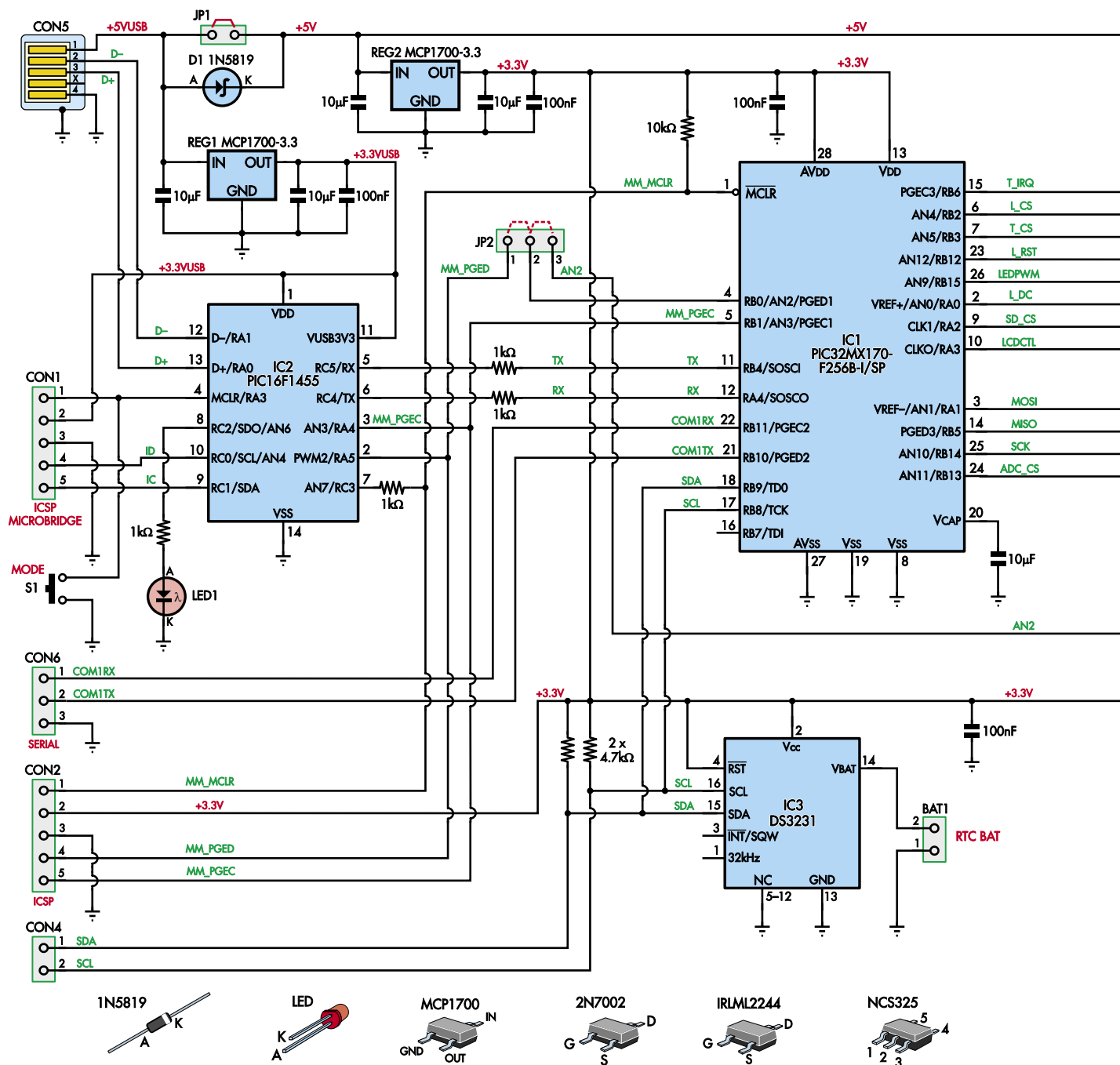
Opis obwodu

Schemat ideowy obwodu rejestratora/monitora stanu akumulatora pokazano na **rysunku 2**. Został on zaprojektowany jako kompletna płytką kompatybilna z Micromite, a nie jako płytką dodatkową do Micromite LCD Backpack.

Pozwala nam to lepiej nadzorować zużycie energii, zmniejszając prąd pobierany z baterii.

Podobnie jak w przypadku każdego urządzenia zasilanego bateryjnie, ważne jest, aby podczas fazy projektowania wziąć pod uwagę zużycie energii przez samo urządzenie.

Zaciski akumulatora i obciążenia/ladowarki znajdują się w prawym dolnym rogu, a dolna połowa prawej strony pokazuje obwody czujników. Inne połączenia zewnętrzne (USB, szeregowo, programowania itp.) są rozmieszczone wzdłuż lewej strony, a obwody Backpack-a zajmują większość



Wieloczynnościowy rejestrator stanu baterii/akumulatorów

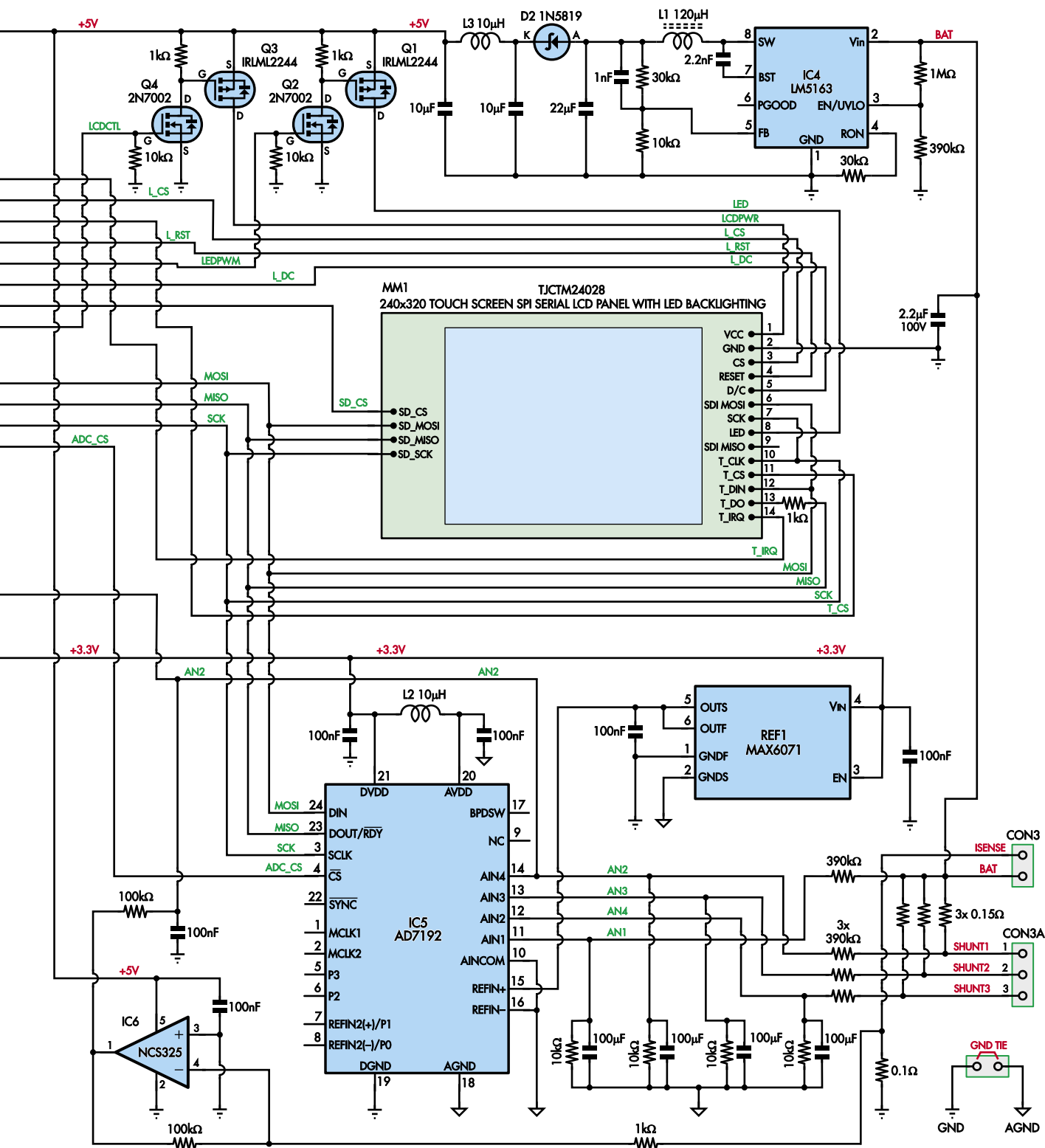
Rysunek 2. Schemat ideowy obwodu zawiera odpowiednik całego modułu Micromite V2 Backpack, precyzyjny wielokanałowy przetwornik ADC i regulator impulsowy zdolny do zasilania urządzenia prądem stałym ze źródła o napięciu od 6 V do 100 V. Urządzenie monitoruje napięcie akumulatora, prąd do/z trzech punktów zewnętrznych oraz własny pobór prądu i rejestruje to wszystko (plus aktualny stan naładowania akumulatora) w wewnętrznej pamięci FLASH mikroprocesora IC1

lewej strony, plus wyświetlacz w środku po prawej stronie. Zasilanie urządzenia znajduje się w górnej części obu stron.

Micromite V2 Backpack (maj 2017; siliconchip.com.au/Article/10652) jest najbardziej zbliżonym do naszego projektu wariantem Backpack'a. To porównanie ma na celu jedynie wyjaśnienie niektórych naszych wyborów projektowych; nie ma znaczenia, jeśli podchodzisz do tego obwodu bez znajomości wcześniejszych projektów.

W tym projekcie zdecydowaliśmy się użyć ekranu dotykowego LCD o przekątnej 2,8 cala (7 cm), zamiast wersji 3,5 cala (9 cm), której używaliśmy ostatnio (np. w V3 Backpack), ponieważ mniejszy wyświetlacz zużywa nieco mniej energii.

V3 Backpack ma również wiele funkcji, które po prostu nie są w tym przypadku potrzebne, stąd nasz wybór V2 Backpack'a jako podstawy dla tego projektu. Główną zaletą w porównaniu



Cechy i parametry:

- Napięcie akumulatora: 6...100 V
- Monitorowanie prądu: do trzech ładowarek lub obciążeń, monitorowanych oddzielnie
- Obsługa prądu: ograniczona tylko przez zastosowane boczniki (10 A z wbudowanymi bocznikami)
- Rozdzielczość prądu: 0,1% (10 mA z wbudowanymi bocznikami)
- Prąd pracy: <1 mA podczas rejestrowania (z wyłączeniem wyświetlacza)
- Interfejs użytkownika: 2,8-calowy kolorowy ekran dotykowy LCD
- Oprogramowanie sprzętowe: napisane w języku BASIC
- Rejestrowanie danych: można przeglądać na wyjściu graficznym lub pobrać jako plik CSV
- Pomiary: aktualny poziom naładowania (Ah) i energia (Wh)
- Stan naładowania: wyświetlany na podstawie napięcia i naładowania

do oryginalnego Micromite BackPack'a jest wbudowany interfejs USB-Serial.

Wykrywanie baterii

Główny obwód wykrywania baterii skoncentrowany jest w obwodach układów scalonych IC5 (AD7192) i REF1 (MAX6071). IC5 to czterokanałowy 24-bitowy przetwornik ADC (analogowo-cyfrowy) z interfejsem szeregowym SPI. Jest on zasilany z wyjścia REG2 o napięciu 3,3 V, a jego szyna analogowa jest filtrowana przy użyciu dławika 10 μ H. Każde z wejść zasilania 3,3 V jest zbocznikowane kondensatorem 100 nF.

IC5 współdzieli magistralę SPI z ekranem dotykowym LCD, a styk 24 IC1 ma funkcję CS, aby wskazać, kiedy IC5 jest adresowany.

IC5 potrzebuje stabilnego napięcia wzorcowego do konwersji napięć na wartości cyfrowe. Pochodzi ono z REF1, układu napięcia wzorcowego MAX6071 2,5 V. Jest to wyjątkowo niskoszumny i precyzyjny układ napięcia odniesienia, zasilany napięciem 3,3 V z REG2, z kondensatorami 100 nF na wejściu i wyjściu. Jego wyjście zasilaja REFIN1+ (styk 15) układu IC5, podczas gdy REFIN1- (styk 16) układu IC5 jest podłączony do masy analogowej.

Każde z czterech wejść analogowych IC5 jest zasilane przez dzielnik 390 k Ω /10 k Ω , zbocznikowany przy mniejszej

rezystancji kondensatorem 100 μ F. Oznacza to, że nominalny odczyt w pełnej skali wynosi 100 V z rozdzielczością około 6 μ V i czasem ustalania około dziesięciu sekund. Używamy przetwornika ADC do wykonywania cyklu konwersji (wszystkich kanałów) mniej więcej raz na dziesięć sekund, co jest tempem potrzebnym do uzyskania maksymalnej rozdzielczości.

Jeden z dzielników jest podłączony bezpośrednio do akumulatora na złączu CON3. Pozostałe trzy monitorują napięcie na końcu obciążenia/ładowarki trzech boczników, które łączą zacisk BAT CON3 z zaciskami CON3A. Mierzając różnicę między napięciami podawanymi do przetwornika ADC, możemy określić przepływ prądu do lub z każdego zacisku.

Płyta drukowana zawiera pola dla rezystorów bocznikowych 15 m Ω , które umożliwiają teoretyczną rozdzielczość poniżej 10 mA. Są to elementy o mocy 3 W, które teoretycznie pozwalają na wykrycie prądu do 14 A. W praktyce jest to mniej i zaciski mają ograniczenie do około 10 A.

Jeśli zamiast tych rezystorów używane są większe boczniki zewnętrzne, wystarczy poprowadzić niskoprądowe przewody pomiarowe z obu ich końców z powrotem do CON3/CON3A. Wartości bocznika można ustawić w oprogramowaniu, aby uwzględnić praktycznie dowolną wartość rezystancji.

Lokalna analogowa sieć uziemienia oddziela napięcia analogowe od cyfrowych sygnałów SPI.

Prąd zasilania

Prąd pobierany przez sam obwód jest niewielki, ale nie bez znaczenia, i należy go uwzględnić, aby uzyskać dokładne pomiary. Ponieważ jest to dość niski prąd, używamy innej techniki do jego monitorowania. Każdy prąd płynący na złączu CON3 z akumulatora do naszego obwodu przepływa przez rezystor bocznikowy 100 m Ω , generując ujemne napięcie w stosunku do masy, proporcjonalne do prądu zasilania.

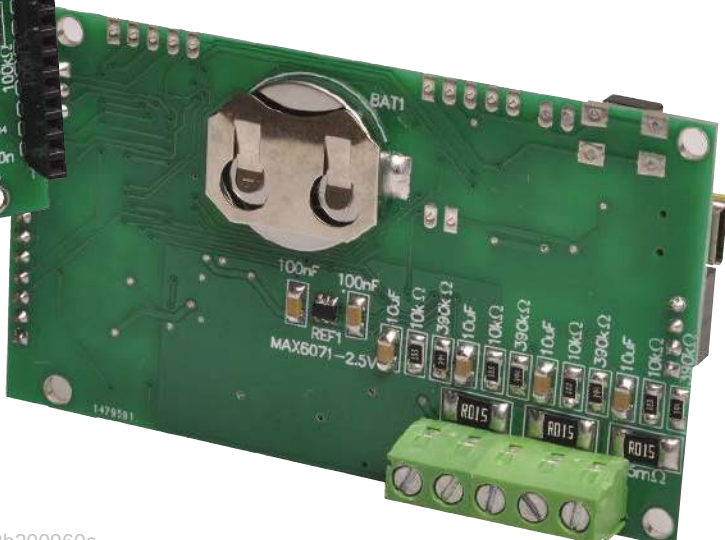
IC6 to jednokanałowy wzmacniacz operacyjny w pięciostykowej obudowie SMD SOT23-5. Jest on podłączony jako wzmacniacz odwracający o wzmocnieniu 100 (100 k Ω /1 k Ω), wysyłający to napięcie na styk 4 układu IC1, gdzie wewnętrzny przetwornik ADC mikroprocesora może je odczytać.

Kondensator 100 nF i rezystor 100 k Ω zapewniają podobne wygładzenie tego sygnału (stała czasowa około dziesięciu sekund), dzięki czemu można go również próbować w podobnych odstępach czasu jak inne kanały.

Podczas pracy rejestratora MB, podświetlenie LED panelu LCD zużywa najwięcej energii, więc używana jest wysoka częstotliwość PWM w celu zapewnienia dokładności tego pomiaru.



Te zdjęcia pokazują wcześniejszy prototyp, w którym brakowało rezystora szeregowego MISO i CON6 (który nie jest używany przez obecną wersję oprogramowania). Niektóre wartości rezystorów i kondensatorów są również nieco inne, ale ogólnie wygląda to dość podobnie do ostatecznej wersji. Podczas budowy zwróć uwagę na wartości pokazane na sitodruku z opisem PCB





Zasilanie

Istnieją dwa możliwe źródła zasilania w tym obwodzie: gniazdo USB CON5 może dostarczać 5 V, podczas gdy złącze baterii na CON3 obsługuje do 100 V z monitorowanej baterii. Na płytce znajduje się kilka komponentów o maksymalnym napięciu 100 V, więc jest to skrajny limit i nie należy go przekraczać.

Układ impulsowego regulatora obniżającego napięcie (buck), IC4 (LM5163), skutecznie obniża napięcie akumulatora do 5 V. Jego zasilanie z akumulatora przez CON3 jest zbocznikowane kondensatorem 2,2 μF i doprowadzone do wejść 2 (VIN) i 1 (GND).

Napięcie powyżej 1,5 V na styku 3 (EN) włącza regulator, co odpowiada napięciu około 5,5 V na CON3 ze względu na dzielnik rezystancyjny 1 M Ω /390 k Ω .

Oprócz akceptowania napięcia do 100 V na wejściu, IC4 ma również wyjątkowo niski prąd spoczynkowy wynoszący zaledwie 10,5 μA bez obciążenia i niewiele więcej przy niewielkich obciążeniach. Jego sprawność zmienia się w zależności od napięcia wejściowego i prądu obciążenia, ale zazwyczaj mieści się w zakresie 75-90%. Więcej szczegółów na temat tego poręcznego małego układu znajduje się w panelu poniżej.

Przełącza on swoje wyjście na styku 8 (SW) naprzemiennie między VIN i GND za pomocą pary wewnętrznych N-kanalowych MOSFET-ów. Górny MOSFET ma napięcie bramki dostarczane z kondensatora 2,2 nF na styku 7 (BOOST).

Impulsy są wygładzane przez cewkę 120 μH i kondensator 22 μF w celu zapewnienia stabilnego napięcia wyjściowego. Napięcie na styku sprzężenia zwrotnego 5 (FB) jest wewnętrznie porównywane z napięciem odniesienia 1,2 V, więc dzielnik 30 k Ω /10 k Ω ustawia napięcie wyjściowe na 4,8 V.

Jest ono ustawione na nieco mniej niż 5 V, więc jeśli dostępne jest alternatywne zasilanie 5 V, przejmie ono zasilanie z akumulatora. Dioda Schottky'ego D2 doprowadza napięcie 4,8 V do filtra π utworzonego z dwóch kolejnych kondensatorów 10 μF i dławika 10 μH .

Kondensator 1 nF bocznikujący rezystor 30 k Ω na górze dzielnika sprzężenia zwrotnego polepsza stabilność obwodu, który zasilają impulsy wyjściowe, zapewniając wystarczające tłumienie tętnień na styku sprzężenia zwrotnego (FB), aby obwód działał poprawnie. Więcej szczegółów na ten temat można znaleźć w naszym panelu.

Szczegóły obwodu mikroprocesora

Opisana wyżej szyna o napięciu około 5 V zasilą następnie sekcję zasilania obwodu Micromite. Układ scalonego stabilizatora MCP1700-3.3 REG2 i powiązane z nim kondensatory filtrujące zapewniają zasilanie 3,3 V dla mikroprocesora IC1. Jest to 32-bitowy mikroprocesor 50 MHz (PIC32MX170F256B), z odpowiednimi kondensatorami odsprężającymi.

IC1 jest zaprogramowany za pomocą oprogramowania układowego MMBasic i uruchamia program BASIC do implementacji funkcji monitorowania i rejestracji.

Podczas gdy niektóre urządzenia Micromite Backpack używały 28-końcówkowej wersji DIP tego układu scalonego, MB wykorzystuje 28-stykowy układ w odudowie SMD (SOIC). Działa on identycznie, ale jest mniejszy, więc możemy zmieścić więcej elementów na płytce drukowanej, a większość innych układów scalonych i tak jest dostępna tylko jako SMD. W tym przypadku końcówki mP są stosunkowo odległe od siebie (raster 1,27 mm/0,05 cala), więc nie ma problemów z przylutowaniem.

Aby oszczędzać energię, mP może za pośrednictwem 14-stykowego złącza LCD włączać i wyłączać zasilanie 5 V ekranu dotykowego.

Opcja zegara DS3231 MEMS

Układ scalony zegara czasu rzeczywistego DS3231 MEMS od około pięciu lat jest najczęściej wybieranym układem do śledzenia czasu. Jego atrakcyjność bez wątpienia zwiększa fakt, że jest on dostępny w postaci łatwego w użyciu modułu, sprzedawanego zwykle jako akcesoria Arduino.

Taki moduł był przedmiotem naszego pierwszego artykułu El Cheapo Modules z października 2016 roku (siliconchip.com.au/Article/102960), który wykorzystaliśmy w kilku projektach, zwykle w połączeniu z Micromite. Moduł zawiera rezystory polaryzujące magistrali I²C, pamięć I²C EEPROM i uchwyt na ogniwo.

Moduł upraszcza połączenie, ponieważ zawiera wszystko, co jest potrzebne do działania układu DS3231, ale czasami jest zbyt duży. Użyliśmy gołego układu scalonego DS3231 (który jest dostarczany w szerokiej 16-stykowej obudowie SOIC SMD) w naszym Micromite Backpack V3 (sierpień 2019; siliconchip.com.au/Article/11764) i zegarze Ol' Timer II (lipiec 2020; siliconchip.com.au/Article/14493).

Aby wesprzeć te projekty, utrzymywaliśmy zapas tych układów scalonych. Pewnego dnia byliśmy zaskoczeni, gdy otrzymaliśmy pakiet małych 8-stykowych elementów SOIC zamiast szerokich 16-stykowych SOIC, których się spodziewaliśmy. Czy zostaliśmy oszukani?

Nie; zamiast oczekiwanego układu otrzymaliśmy wariant DS3231M. Osoby zaznajomione z układem DS3231 wiedzą, że wykorzystuje on tylko osiem wyprowadzeń; dolne końcówki

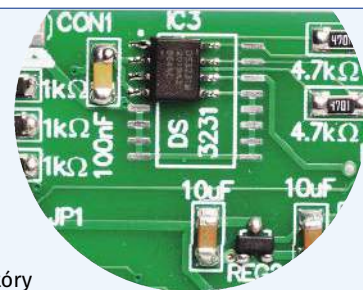
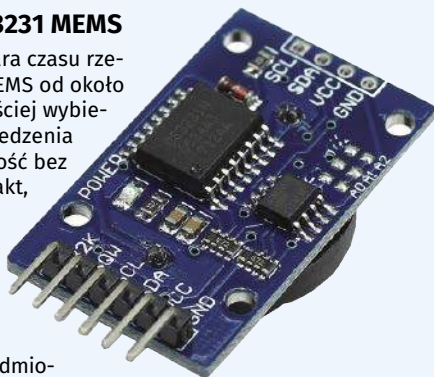
są oznaczone jako NC („niepodłączone”). Powodem dużego opakowania układu nie jest to, że potrzebuje 16 wyprowadzeń, ale dlatego, że zawiera wewnątrz plastikowej obudowy układu scalonego oscylator kwarcowy z kompensacją temperatury, który nie zmieściłby się w 8-stykowym chipie.

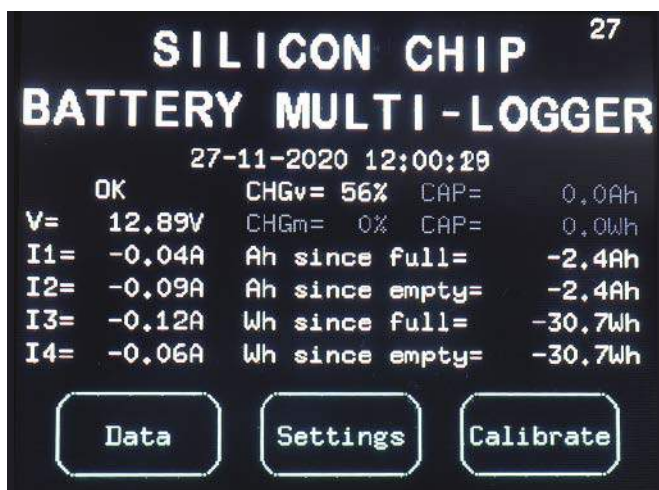
Jednak wraz z postępem technologii MEMS (zobacz nasz artykuł w wydaniu z listopada 2020 r. siliconchip.com.au/Article/14635), oscylator kwarcowy wewnątrz DS3231 został zastąpiony mniejszym układem MEMS – mikroukładem elektromechanicznym.

Biorąc zatem pod uwagę ich niewielkie rozmiary i przyzwoitą wydajność, postanowiliśmy wypróbować je w tym projekcie. Okazało się, że DS3231M działa tak samo jak DS3231. Nominalna dokładność jest nieco gorsza i wynosi ± 5 ppm w porównaniu do $\pm 3,5$ ppm „dużego” chipa, ale w sytuacjach, w których rozmiar ma znaczenie, mniejszy pakiet jest decydujący.

Część MEMS nie wydaje się również cierpieć z powodu starzenia się kwarcu, co oznacza, że w dłuższej perspektywie może być dokładniejsza, chyba, że zostanie to skompensowane we wcześniejszej wersji układu. Pobór prądu z baterii podtrzymania w typowych przypadkach wydaje się być wyższy w przypadku części MEMS, ale w większości przypadków żywotność baterii będzie nadal zbliżona do jej okresu trwałości.

W tym konkretnym projekcie uwzględniliśmy w schemacie PCB obie wersje podzespołu, z podwójnym polem montażowym SMD, który pasuje zarówno do szerokiej 16-stykowej części SOIC, jak i węższej 8-końcówkowej części SOIC MEMS. Nie wiemy, czy DS3231M będzie bardziej popularny niż oryginalny DS3231, ale jesteśmy gotowi na każdą ewentualność.





Ekran 1. Ekran główny zawiera wszystkie krytyczne statystyki dotyczące akumulatora, a także trzy proste opcje menu umożliwiające dostęp do innych funkcji. Wartości pokazane na szaro to obliczenia pojemności, które nie są jeszcze prawidłowe, ponieważ rejestrator nie wykrył pełnego cyklu ładowania i rozładowania; zaświecą się jaśniej, gdy to nastąpi

Wysoki poziom na styku 10 układu IC1 włącza N-kanalowy MOSFET Q4, który w przeciwnym razie nie przewodzi wskutek obniżenia potencjału bramki przez rezystor 10 kΩ. Gdy Q4 przewodzi, ustawia bramkę P-kanalowego MOSFET-a Q3 na niskim poziomie, co pozwala na przepływ napięcia 5 V od źródła Q3 do drenu i do wejścia zasilania panelu LCD.

Podobny układ, sterowany przez wyjście 26 układu IC1, poprzez MOSFET-y Q2 i Q1 przełącza zasilanie podświetlenia LED panelu LCD. Typowo, na wyjście 26 podawany jest sygnał PWM, ustawiając w ten sposób jasność podświetlenia.

W przeciwieństwie do Micromite BackPack V2, który miał ręczną kontrolę jasności przez PWM, pominieliśmy tę opcję sterowania podświetleniem, ponieważ podświetlenie jest niewątpliwie największym odbiornikiem energii w obwodzie.

Dlatego musi być całkowicie wyłączony podczas rejestrowania i monitorowania.

Komunikacja szeregową

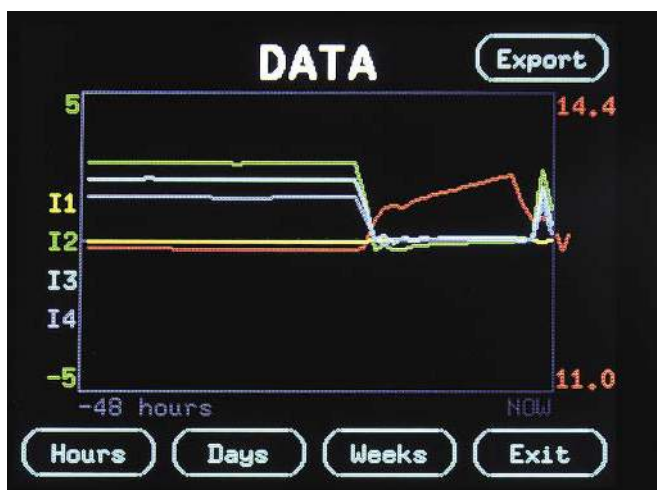
Układ IC1 wysyła dane do wyświetlacza i odbiera sygnały dotykowe z ekranu dotykowego za pomocą magistrali szeregowej SPI na stykach 3, 14 i 25 (MOSI, MISO i SCK). Łączą się one z końcówkami 6 i 12 (MOSI), stykiem 13 (MISO) oraz stykami 7 i 10 (SCK) panelu LCD. MISO oznacza „master in, slave out”, podczas gdy MOSI oznacza „master out, slave in”.

Linia MISO ma szeregowy rezystor 1 kΩ, dzięki czemu może nadal działać, gdy panel LCD jest wyłączony. Sygnały te, plus sygnał wyboru układu z końcówki 9 układu IC1, łączą się również ze złączem karty SD na drugim końcu płytki drukowanej panelu LCD za pośrednictwem czterospilkowej listwy.

Planowaliśmy użyć karty SD do przechowywania danych, ale ograniczenia rozmiaru pamięci FLASH w mP oznaczają, że nie ma wystarczająco dużo miejsca, aby dołączyć (raczej duże) biblioteki potrzebne do tego.

IC2 to 8-bitowy mikroprocesor PIC16F1455 zaprogramowany oprogramowaniem Microbridge. Dzięki temu może on działać jako mostek USB-Serial, a także programować mikroprocesor PIC32.

Przycisk S1 służy do przełączania IC2 między trybami USB-Serial i programowania, przy czym dioda LED1 miga, wskazując, że przekazuje dane szeregowo, lub świeci światłem ciągłym, gdy jest w trybie programowania.



Ekran 2. Ekran danych zapewnia graficzną prezentację zarejestrowanych danych. Można wyświetlać różne przedziały czasowe, a wyświetlacz będzie automatycznie przewijany raz na minutę, aby pokazać bieżące dane. Opcja Weeks (Tygodnie) zapewnia dane z okresu około dwóch tygodni. Dane mogą być również zrzucane jako arkusze CSV przez port szeregowy konsoli za pomocą przycisku Export

Gniazdo mini-USB typu B CON5 służy zarówno do komunikacji USB (D+/D-), jak i opcjonalnie do zasilania 5 V. Dioda Schottky'ego D1 dostarcza napięcie USB 5 V do szyny Micromite 5 V. Zworka JP1 umożliwia w razie potrzeby pominięcie diody D1.

REG1 jest identyczny z REG2 i niezależnie dostarcza 3,3 V do IC2. Sygnały szeregowo TX i RX są mostkowane do i z wirtualnego portu USB-Serial przez IC2. Łączą się one między jego stykami 5 i 6, poprzez rezystory 1 kΩ, ze stykami 11 i 12 konsoli Micromite na IC1.

Styki 2, 3 i 7 układu IC2 mogą być używane do programowania układu IC1 za pośrednictwem interfejsu ICSP; są one podłączone odpowiednio do końcówek 4, 5 i 1 układu IC1. Sygnał PGD przechodzi przez JP2, co pozwala na użycie styku 4 układu IC1 jako wejścia analogowego, gdy nie jest on używany do programowania.

Zarówno IC1, jak i IC2 mają styki programowania szeregowego w obwodzie (ICSP) wyprowadzone na krawędź płytki drukowanej, odpowiednio na CON2 i CON1. Jest to funkcja niespotykana w innych BackPack'ach, ale uwzględniliśmy ją tutaj, ponieważ zastosowane tutaj układy scalone SMD są trudniejsze do zaprogramowania poza obwodem niż układy do montażu przewlekane (DIP).

Zegar czasu rzeczywistego DS3231, IC3, zapewnia dokładne odmierzenie czasu w długich okresach. Jego styki 15 i 16 magistrali szeregowej I²C (SDA i SCL) łączą się z IC1 na stykach 18 i 17, stykach I²C używanych przez oprogramowanie Micromite. Dwa rezystory 4,7 kΩ zapewniają polaryzację wymaganą przez protokół I²C.

Płytką drukowaną jest również wyposażona w pola stykowe SMD SOIC-8, aby umożliwić użycie podobnego układu DS3231M (który wykorzystuje oscylator MEMS zamiast kwarcu). Zobacz oddzielny panel wyjaśniający różnice.

Obsługa oprogramowania

Niektóre z poniższych informacji mogą wydawać się niejasne dla osób nie zaznajomionych z MMBasic, ale informacje te mogą się przydać, jeśli będziesz chciał zmienić kod.

MMBasic z pewnością ułatwia sterowanie panelem LCD (TFT), ponieważ wykonuje inicjalizację rozruchu i ma wbudowane polecenia BASIC do rysowania i pisania na wyświetlaczu. Potrzebuje jednak pewnej pomocy, aby współpracować z naszym obwodem, który rozpoczyna pracę z wyłączonym panelem LCD, a zatem nie jest gotowy do zaakceptowania poleceń inicjalizacyjnych, które są wysyłane automatycznie.

Układ scalony LM5163 regulatora impulsowego redukującego napięcie

Nasze początkowe plany projektowe dla rejestratora baterii zakładały ambitny cel zaprojektowania go tak, aby działał przy zasilaniu z akumulatora o napięciu do 80 V, poprawiając limit 60 V starego miernika pojemności baterii. W tym celu wykorzystano zintegrowany układ scalony LM2574HV pracujący ze stałą częstotliwością 50 kHz, wymagający sporej cewki toroidalnej i kondensatora elektrolitycznego.

Mając nadzieję, że w ciągu ostatniej dekady nastąpił postęp w tej dziedzinie, postanowiliśmy poszukać nowszych części. Znaleźliśmy wiele części zdolnych do pracy z zasilaniem 100 V, co jest imponujące.

Częstotliwość przełączania 1 MHz nie są już rzadkością.

Ta znacznie wyższa częstotliwość przełączania oznacza, że potrzebna jest mniejsza cewka i kondensatory, co pomaga nam zachować kompaktowość naszej płytki.

Wiele znalezionych przez nas części mogło dostarczyć jedynie prądu 100 mA. Chociaż mogło to być wystarczające przy starannym pilnowaniu prądu podświetlenia LCD, chcieliśmy mieć większy zapas. LM5163 był najtańszą częścią zdolną do dostarczenia więcej niż 100 mA (czyli 500 mA) w łatwej do lutowania obudowie SOIC-8, co jest dobrym kompromisem między rozmiarem a łatwością obsługi.

Jak to jest typowe dla nowoczesnych konstrukcji impulsowych regulatorów redukujących napięcie, jest to typ synchroniczny, co oznacza, że ma dwa wewnętrzne przełączniki. Napięcie wejściowe jest przełączane na cewkę indukcyjną przez wewnętrzny MOSFET po stronie wysokiej (high-side). Gdy ten MOSFET jest wyłączony, drugi MOSFET po stronie niskiej (low-side) jest włączany, aby zapewnić ścieżkę dla przepływu prądu z cewki indukcyjnej do obwodu. Eliminuje to potrzebę stosowania zewnętrznej diody pełniącej tę rolę i zwiększa wydajność układu.

LM5163 to konstrukcja COT (stały czas włączenia); czas, w którym MOSFET po stronie wysokiej jest włączony, jest ustawiany przez zewnętrzny rezystor, po czym MOSFET jest wyłączany. Styk sprzężenia zwrotnego monitoruje napięcie wyjściowe, a gdy napięcie wyjściowe spadnie, rozpoczyna się kolejny cykl włączania.

Tak więc cykl pracy jest modulowany w celu utrzymania pożądanego napięcia wyjściowego, ale stały czas włączenia oznacza, że częstotliwość przełączania zmienia się, chociaż można ją przewidzieć.

Kiedy zbudowaliśmy nasz pierwszy prototyp, wszystko działało zgodnie z oczekiwaniami; byliśmy naprawdę pod wrażeniem tego, jak elastyczna i łatwa w użyciu była ta niewielka kość. Ale potem zaczęło coś pisać! Ton zmieniał się wraz z obciążeniem (które mogliśmy łatwo zmieniać poprzez regulację intensywności podświetlenia LCD) i napięciem wyjściowym. Było na tyle złe, szczególnie w okolicach 12 V, że musieliśmy coś z tym zrobić.

Przyczyną były zakłócenia elektryczne, które wpływały na moment włączenia. Układ może włączyć się wcześniej, co powoduje wzrost napięcia wyjściowego. Spowoduje to opóźnienie następnego włączenia, ponieważ sterownik będzie czekał, aż napięcie wyjściowe spadnie poniżej progu.

Impulsy wyjściowe zaczynają grupować się w grupy impulsów i to właśnie te klastry występują na słyszalnych częstotliwościach, powodując wysokie piski, które słyszeliśmy („oscylacja subharmoniczna”) – patrz poniżej.

Jak stwierdziliśmy w przypadku naszego zamiennika Switchmode 78xx (siliconchip.com.au/Article/14533), próba uzyskania

optymalnego działania tego rodzaju podzespołu w szerokim zakresie napięć wejściowych może być trudna. W tym przypadku pomogła dodatkowa pojemność wyjściowa.

Na szczęście sekcja arkusza danych LM5163

(przedstawiona na **rysunku 4**) opisuje metody uniknięcia tego problemu. Celem jest zwiększenie tętnienia widzianego przez (FB) – sprzężenie zwrotne, tak, aby regulator miał jasno określony czas włączenia, pomimo obecności szumów.

Wypróbowaliśmy metodę typu 1, która polega na dodaniu rezystancji szeregowej do kondensatora wyjściowego. Dodatkowa rezystancja oznacza, że na napięcie widoczne na styku FB mniejszy wpływ ma kondensator, a większy impulsy z cewki indukcyjnej.

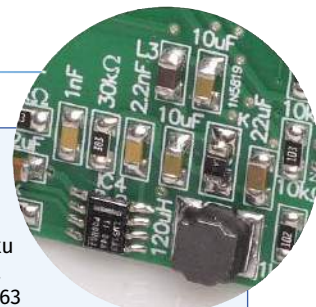
Oznacza to jednak również, że kondensator wyjściowy jest mniej skuteczny w filtrowaniu napięcia wyjściowego i stwierdziliśmy, że w niewielkim stopniu redukuje piski.

Wypróbowaliśmy więc część metody typu 2 (pomijając rezystor szeregowy z typu 1) i po prostu dodaliśmy kondensator „sprężenia zwrotnego” równolegle z górnym rezystorem dzielnika sprzężenia zwrotnego. Oznacza to, że styk FB widzi pełną amplitudę tętnienia napięcia wyjściowego, ponieważ jest on sprzężony dla sygnałów zmiennoprądowych na wyjściu bezpośrednio przez kondensator, a nie przez prosty dzielnik z łańcucha rezystorów.

Skutecznie, czterokrotnie zwiększa to tętnienia widziane przez styk FB za pomocą naszego dzielnika 30 kΩ/10 kΩ, bez pogorszenia filtrowania. Wyeliminowało to piski, więc zachowaliśmy to rozwiązanie w naszym ostatecznym projekcie.

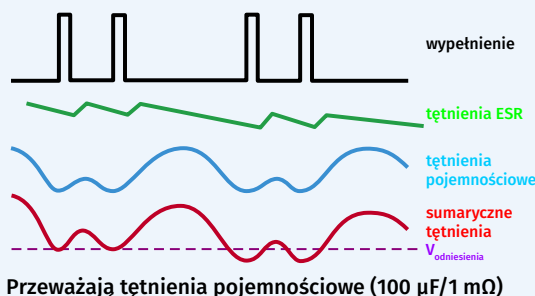
Każde urządzenie przełączające, które zależy od napięcia sprzężenia zwrotnego z dzielnika do przełączania elementów wyjściowych, może skorzystać z kondensatora sprzężenia zwrotnego. Zależy to jednak od częstotliwości pracy, wartości kondensatora i współczynnika dzielnika.

Słowo ostrzeżenia: chociaż kondensator ten może wydawać się lekarstwem na wszystko, jego efektem ubocznym jest spowolnienie reakcji na stany nieustalone, ponieważ zmniejsza on wzmocnienie zamkniętej pętli dla składowych o wyższej częstotliwości.

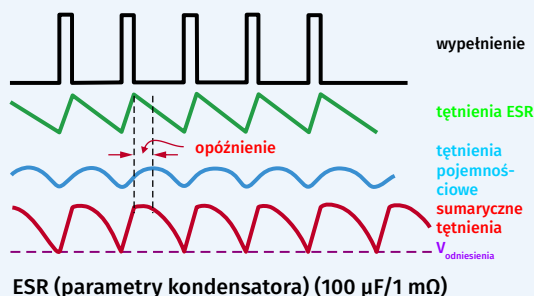


TYPE 1 Lowest Cost	TYPE 2 Reduced Ripple	TYPE 3 Minimum Ripple
$R_{ESR} \geq \frac{20\text{mV} \cdot V_{OUT}}{V_{FB1} \cdot \Delta I_{(nom)}}$ $R_{ESR} \geq \frac{V_{OUT}}{2 \cdot V_{IN} \cdot F_{SW} \cdot C_{OUT}}$	$R_{ESR} \geq \frac{20\text{mV}}{\Delta I_{(nom)}}$ $R_{ESR} \geq \frac{V_{OUT}}{2 \cdot V_{IN} \cdot F_{SW} \cdot C_{OUT}}$ $C_{FF} \geq \frac{1}{2\pi \cdot F_{SW} \cdot (R_{FB1} \parallel R_{FB2})}$	$C_A \geq \frac{10}{F_{SW} \cdot (R_{FB1} \parallel R_{FB2})}$ $R_A C_A \leq (V_{IN(nom)} - V_{OUT}) \cdot t_{ON @ V_{IN(nom)}}$ $C_B \geq \frac{t_{FB-setting}}{3 \cdot R_{FB1}}$

Rysunek 4. Zalecane przez Texas Instruments rozwiązania dla oscylacji subharmonicznych (lub „squegging”) – brak dobrego tłumaczenia, odpowiada „piskom” w radioodbiorniku podczas przestrojania) w LM5163. Wypróbowaliśmy typ 1 i nie zadziałał, ale typ 2 już tak. Wymaga on jedynie dodania kondensatora sprzężenia zwrotnego o niskiej wartości, C_{ff}, w górnej części dzielnika sprzężenia zwrotnego. Typ 3 jest podobny, ale dodaje kolejne łączące i elementy dla lepszej reakcji w stanach przejściowych; to przesada w naszym zastosowaniu



Rysunek 3. Zwykle w kondensatorze niski współczynnik ESR jest uważany za pożądany, ponieważ zapewnia lepsze filtrowanie, ale gdy kondensator zbyt skutecznie filtruje tętnienia, wpływa to na zdolność regulatora do cyklicznego, regularnego generowania impulsów przełączających



Musimy więc dodać procedurę (w podprogramie MM.STARTUP), aby ustawić styk 10 jako wyjście i ustawić go w stan wysoki, a następnie ponownie uruchomić kod inicjalizacji LCD. Za każdym razem, gdy włączamy wyświetlacz po jego wyłączeniu, musimy uruchomić ten kod.

Musimy również monitorować inne linie, które biegną do panelu LCD, ponieważ niektóre z nich są domyślnie w stanie wysokim i dlatego marnują energię. MMBasic nie pozwala na bezpośredni nadzór nad nimi, ponieważ oprogramowanie układowe rezerwuje je do sterowania panelem LCD, więc musimy komendą „POKE” wpisać bezpośrednio kod maszynowy do rejestrów IC1, a następnie uruchomić polecenie, aby ponownie zainicjować sterownik LCD.

Podobnie, wyłączenie procesora wymaga bezpośrednio komendy „POKE” z odpowiednim kodem maszynowym, aby wyłączyć te styki. Nie jest wymagana deinicjalizacja oprogramowania, ponieważ wyświetlacz LCD można po prostu wyłączyć z dowolnego stanu.

Pomimo tej komplikacji, oprogramowanie stosunkowo łatwo wyczuwa dotknięcia panelu LCD, nawet jeśli jest on wyłączony. Jest to konieczne, ponieważ użytkownik potrzebuje sposobu na wybudzenie urządzenia, jeśli jest ono w stanie niskiego poboru mocy.

Nawet gdy wyświetlacz LCD jest wyłączony, styk TIRQ (który jest podłączony do styku 15 układu IC1) ma potencjał obniżający do potencjału masy za każdym razem, gdy panel zostanie dotknięty. Ponieważ oprogramowanie układowe Micromite zapewnia słabą polaryzację tego styku, wystarczy monitorować stan tego wejścia, aby wiedzieć, czy nastąpiło dotknięcie.

Głównym zadaniem programu MMBasic jest odczyt napięcia akumulatora i napięcia na trzech bocznikach w celu określenia napięcia i prądu akumulatora. Rejestruje je w zmiennych, które są przechowywane w pamięci RAM i są regularnie zapisywane w wewnętrznej pamięci FLASH.

Ponieważ zasilanie obwodu działa z monitorowanym akumulatora, wyłączenie go i utrata zawartości pamięci RAM prowadziłyby do poważnej usterki, więc tylko długoterminowe próbki są zapisywane co godzinę w pamięci FLASH. Jeśli urządzenie musi zostać odłączone podczas pracy na akumulatorze, utracona zostanie co najwyżej jedna godzina danych.

Podczas zapisywania w pamięci FLASH dane są uśredniane przez pewien okres, zanim zostaną zarchiwizowane. Oznacza to, że mniej danych musi być przechowywanych, ale duża ilość danych może być przechowywana do celów historycznych.

Można na przykład porównać, ile energii panele słoneczne dostarczają do akumulatora w okresie kilku tygodni. Dane o prądzie

Wykaz elementów:

- 1 dwustronna płytka drukowana o kodzie 11106201 i wymiarach 86x50 mm
- 1 panel dotykowy LCD 2,8-cal z kontrolerem ILI9341
- 1 obudowa UB3 Jiffy (opcjonalna, w zależności od pożądanego montażu)
- 1 wycięty laserowo panel akrylowy pasujący do LCD i obudowy UB3 [SC3456, SC3337, SC5063 lub podobny].
- 2 5-szpilkowe odcinki kątowej listwy kołkowej (CON1, CON2; oba opcjonalne, do programowania IC2 i IC1)
- 1 2-drożny zacisk śrubowy, raster 5/5,08 mm (CON3)
- 1 3-drożny zacisk śrubowy, raster 5/5,08 mm (CON3A)
- 2 2-stykowe złącza kołkowe (CON4 i JP1; oba opcjonalne)
- 1 gniazdo mini-USB do montażu SMD (CON5)
- 1 3-stykowe złącze kołkowe (CON6, port szeregowy; opcjonalnie, brak na fotografii PCB)
- 1 3-stykowe złącze kołkowe (JP2)
- 2 zworki (JP1, JP2)
- 1 uchwyt SMD na baterię pastylkową CR2032 (BAT1) [BAT-HLD-001 – Digi-key, Mouser itp.]
- 1 ogniwo CR2032/CR2025 lub podobne (BAT1)
- 1 cewka indukcyjna 120 µH SMD 6x6 mm (L1) [np. SRN6045TA-121M – Digi-Key, Mouser itp.]
- 2 cewki indukcyjne 10 µH SMD 1206 (L2, L3)
- 1 4-stykowy dotykowy przetwornik przyciskowy SMD lub do montażu przewlekane (S1)
- 1 14-stykowa listwa gniazda żeńskiego (dla wyświetlacza LCD)
- 1 4-stykowa listwa gniazda żeńskiego (dla wyświetlacza LCD)
- 8 śrub z tłem walcowym M3x6
- 4 gwintowane kołki dystansowe M3x12
- 4 podkładki dystansowe M3x1 bez gwintu (np. stos podkładek ID 3 mm)
- 3 boczniki wysokoprądowe [np. Jaycar QP5415, Altronics Q0480 – opcjonalnie, patrz tekst] i wytrzymałe okablowanie pasujące do boczników, baterii i obciążenia (patrz tekst).

Półprzewodniki:

- 1 32-bitowy mikroprocesor PIC32MX170F256B-I/SO programowany za pomocą MMBasic lub z wsadem 11110620A.hex, SOIC-28 (IC1)
- 1 8-bitowy mikroprocesor PIC16F1455-I/SL programowany za pomocą Microbridge, SOIC-14 (IC2)
- 1 układ scalony zegara czasu rzeczywistego DS3231/DS3231M, SMD szeroki SOIC-16 lub wąski SOIC-8 (IC3)
- 1 impulsowy synchroniczny regulator obniżający (buck) LM5163DDAR, SOIC-8 (IC4)
- 1 24-bitowy przetwornik ADC AD7192BRUZ, TSSOP-24 (IC5)
- 1 op-amp NCS325 CMOS, SOT-23-5 (IC6)
- 1 precyzyjny układ napięcia odniesienia MAX6071AAT25+TT 2,5 V, SOT23-6 (REF1)
- 2 regulatory 3,3 V o niskim spadku napięcia MCP1700-3.3, SOT-23 (REG1, REG2)
- 2 P-kanalowe tranzystory MOSFET IRLML2244TRPBF, SOT-23 (Q1, Q3)
- 2 N-kanalowe tranzystory MOSFET N7002, SOT-23 (Q2, Q4)
- 1 dioda LED 3 mm lub SMD 1206 (LED1)
- 2 diody Schottky'ego SS14 (lub równoważne) 40 V 1 A SMD, DO-214AC (D1, D2)

Kondensatory: ceramiczne (wszystkie SMD 1206)

- | | | |
|--------------------------|-------------------------|----------------------------|
| 4 szt. 100 µF 6,3 V X5R | 1 szt. 22 µF 16 V X5R | 7 szt. 10 µF 50 V X7R |
| 1 szt. 2,2 µF 100 V X7R | 10 szt. 100 nF 50 V X7R | 1 szt. 2,2 nF 50 V C0G/NPO |
| 1 szt. 1 nF 50 V C0G/NPO | | |

Rezystory: (wszystkie 1% SMD 1206 1/8W metalizowane, chyba, że zaznaczono inaczej)

- | | | | | |
|--|---------------|---------------|--------------|--------------|
| 1 szt. 1 MΩ | 5 szt. 390 kΩ | 2 szt. 100 kΩ | 2 szt. 30 kΩ | 8 szt. 10 kΩ |
| 2 szt. 4,7 kΩ | 8 szt. 1 kΩ | 1 szt. 0,1 Ω | | |
| 3 szt. 15 mΩ 1% 3 W (SMD 2512; nie jest potrzebny, jeśli używane są zewnętrzne boczniki prądowe) | | | | |

i zużyciu energii są również wykorzystywane do obliczania parametrów, takich jak pojemność baterii i stan naładowania.

Program MMBasic zapewnia również interfejs użytkownika umożliwiający zmianę ustawień oraz tworzenie wykresów i przeglądanie parametrów. Ponadto istnieje opcja zrzucania danych przez port szeregowy, dzięki czemu można je wyeksportować do programu komputerowego w celu sporządzenia wykresów i dalszej analizy.

W przyszłym miesiącu podczas procedury konfiguracji zagłębimy się bardziej w działanie oprogramowania.

W następnym miesiącu

W drugiej i ostatniej części opisu urządzenia, podamy kompletne szczegóły montażu PCB, procedury programowania mikroprocesora, instrukcje konfiguracji i obsługi, informacje o kalibracji wraz z końcową instrukcją budowy. ■

Tim Blythman

Adaptacja do wydania polskiego – Andrzej Nowicki

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

Drodzy Czytelnicy

Podczas prac redakcyjnych nad tym artykułem w samej Redakcji EdW ujawnili się chętni do zbudowania tego układu.

W przypadku szerszego zainteresowania moglibyśmy wytworzyć serię płytek drukowanych i zaprogramowanych mikroprocesorów, które byłyby dostępne w ofercie AVT. Chcielibyśmy wysondować skalę tego zainteresowania. Prosimy, przyslijcie do nas na adres redakcja@elportal.pl odpowiedzi na dwa pytania:

1. czy jesteś zainteresowany zakupem PCB do tego projektu?
2. czy jesteś zainteresowany zakupem zaprogramowanego mikroprocesora do tego projektu?

Twój mail nie będzie traktowany jako zamówienie, jest tylko głosem w sondażu.

Redakcja EdW



Głośnik niskotonowy z dopingiem

Ta kolumna niskotonowa (dla prostoty będziemy dalej nazywać ją subwooferem) wykorzystuje tylko jeden przetwornik 8-calowy (200 mm), ale jej pasmo przenoszenia schodzi aż poniżej 30 Hz i jest w stanie dostarczyć ponad 100 dB SPL! Dzieje się tak pomimo skromnych rozmiarów obudowy, która ma mniej niż 30 cm szerokości, dzięki czemu jest stosunkowo łatwa do ukrycia w pokoju. Jak więc udało się to osiągnąć? Czytaj dalej, aby się dowiedzieć.

Ten subwoofer jest stosunkowo niedrogi i nie jest też wcale taki trudny w budowie, dzięki swojej sprytniej konstrukcji. Jeśli posiadasz już większość prostych narzędzi stolarskich, jego koszt wyniesie około 200 USD (w zależności od miejsca zakupu elementów). Zważywszy na wysoką sprawność kolumny, do jej zasilania można użyć stosunkowo niewielkiego wzmacniacza, aczkolwiek potrzebny będzie aktywny filtr pasmowo-przepustowy, który zostanie opisany w przyszłym miesiącu.

Subwoofer typu „tapped horn” oznacza, że jego jedyny przetwornik niskotonowy jest umieszczony wewnątrz tuby czy raczej długiego tunelu dźwiękowego. Ten typ subwoofera

został rozslawiony przez Thomasa Danleya z Danley Sound Labs. Są one często używane do wzmacniania dźwięku; kilka przykładów można znaleźć na stronie siliconchip.com.au/link/ab9q.

Jeśli chcesz zobaczyć demonstrację subwoofera typu „tapped horn” w ekstremalnie postaci, obejrzyj wideo na stronie <https://youtu.be/Zbf3bzpgml8>.

Określenie „tapped horn”, które w tym przypadku możemy rozumieć jako „kolumnę z dopingiem”, nie pasuje do inżyniera, ponieważ w rzeczywistości nie jest on jakoś dodatkowo wysilony. Zamiast tego, prawdopodobnie bardziej trafne byłoby nazwanie

tego dopasowania zwrotną rurą rezonansową. Odłóżmy jednak semantykę na bok i użyjmy tej potocznej nazwy.

Po przeczytaniu kilku artykułów na temat tego podejścia do tworzenia subwoofera, Autor niniejszego tekstu postanowił sprawdzić, jak to działa. Celem było zaprezentowanie konstrukcji tubowej (lub, jak kto woli, labiryntowej), która pasuje do warunków domowych, umożliwiając Czytelnikom poznanie w przystępny sposób tej koncepcji. Tak więc, jeśli kiedykolwiek zastanawiałeś się nad tego rodzaju subwooferem, oto Twoja szansa na spędzenie weekendu i przekonanie się, jak ta konstrukcja działa!

Ten subwoofer jest więcej niż wystarczający do salonu, gabinetu lub sypialni – został utrzymany w skromnych wymiarach i kosztach. Przedstawiona konstrukcja została uproszczona, aby uniknąć dziwnych kątów cięcia oraz usunięto niepotrzebne zaokrąglenia narożników, aby montaż był jak najprostszy. Skrzynka została nawet zwymiarowana tak, aby można było użyć standardowych arkuszy MDF przy minimum wymaganych cięć.

Przy projektowaniu głośników, projektant musi zonglować kilkoma parametrami, w szczególności: rozmiarem obudowy, poziom ciśnienia dźwięku (SPL – Sound Pressure Level), dopasowaniem niskich i wysokich częstotliwości (pasmo przenoszenia) oraz sprawnością (ile mocy potrzeba do osiągnięcia określonego poziomu dźwięku).

Obudowa tubowa może zwiększyć efektywność kolumny, pasmo przenoszenia i poziom SPL znacznie ponad to, co oferuje obudowa konwencjonalna, szczelna lub wentylowana. Osiąga to poprzez umieszczenie przetwornika wewnątrz toru akustycznego i poprowadzenie drogi propagacji dźwięku dookoła, tak, aby emisja dźwięku z tylnej części przetwornika dodawała się do propagacji z przedniej części przetwornika.

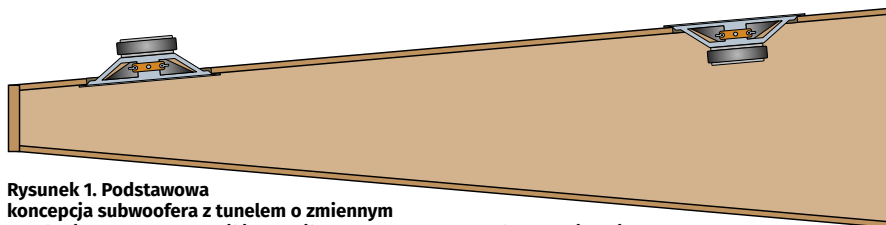
Ale nie ma czegoś takiego jak darmowy obiad, więc płaci się za to złożonością.

Jak pokazano na **rysunku 1**, przednia strona jednego głośnika emituje dźwięk do tunelu blisko jego końca, a tylna strona drugiego głośnika zasila go dźwiękiem blisko jego wyjścia. Jeśli oba przetworniki są zasilane tym samym sygnałem, dostarczają do tuby sygnały w przeciwfazie, ponieważ są skierowane w przeciwnych kierunkach. Daje to symulowaną odpowiedź pokazaną na **rysunku 2**; zwróć uwagę na rozszerzony zakres basów.

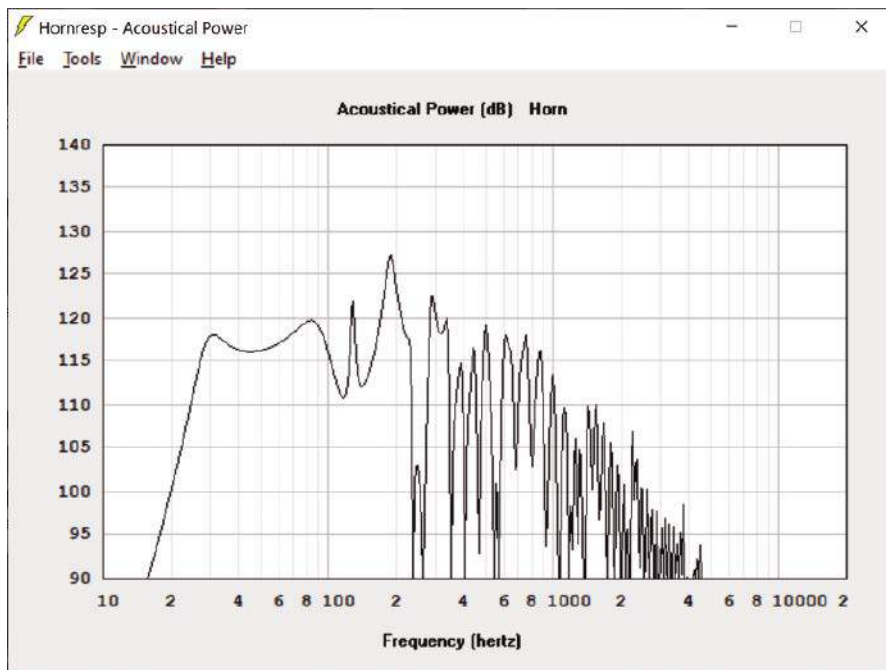
Jednak ten sam przetwornik nie może znajdować się w dwóch różnych miejscach, więc aby przetwornik mógł emitować dźwięk do toru akustycznego zarówno z przedniej, jak i z tylnej strony głośnika, obudowa jest labiryntem – patrz **rysunek 3**. Ta pojedyncza konstrukcja jest nadal bardzo długa i niezbyt wygodna. Obudowę można skrócić na kilka sposobów. Wybraną przez nas konfigurację pokazano na **rysunku 4**.

Idealnie byłoby wykonać ją z rozszerzających się, a nie prostopadłościennych, sekcji, ale cięcie w takim przypadku jest naprawdę kłopotliwe. Dlatego zauważysz, że trochę nagięliśmy konstrukcję i zrobiliśmy sekcje toru dźwiękowego w formie prostopadłościennych. Nasze testy wykazały, że pogorszenie parametrów nie jest na tyle znaczące, aby się nim martwić.

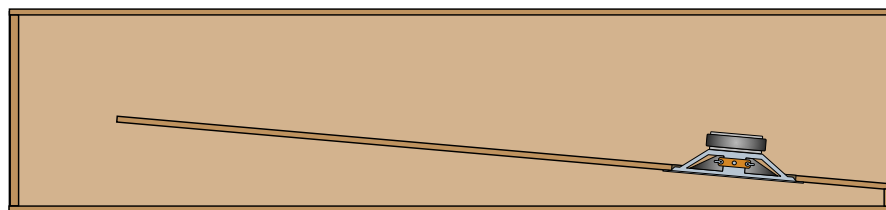
Należy pamiętać, że konwencjonalna, szczelna obudowa ma za zadanie pochłoniąć



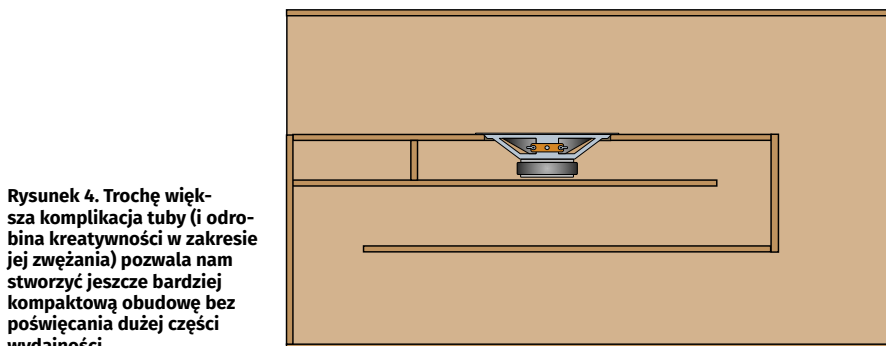
Rysunek 1. Podstawowa koncepcja subwoofera z tunelem o zmiennym przekroju. Dwa przetworniki są zasilane tym samym sygnałem. Ponieważ zamontowano je obrócone względem siebie o 180°, sygnały generowane przez nie w tubie są w przeciwfazie. Jednak przemieszczenie się dźwięku w dół tuby wymaga czasu, więc w pewnym zakresie częstotliwości dźwięk docierający do zewnętrznego przetwornika jest z nim w fazie, co powoduje efektywną interferencję i wzmocnienie dźwięku. UWAGA. Dla uproszczenia pominięto kwestię dźwięku emitowanego po zewnętrznej stronie tunelu!



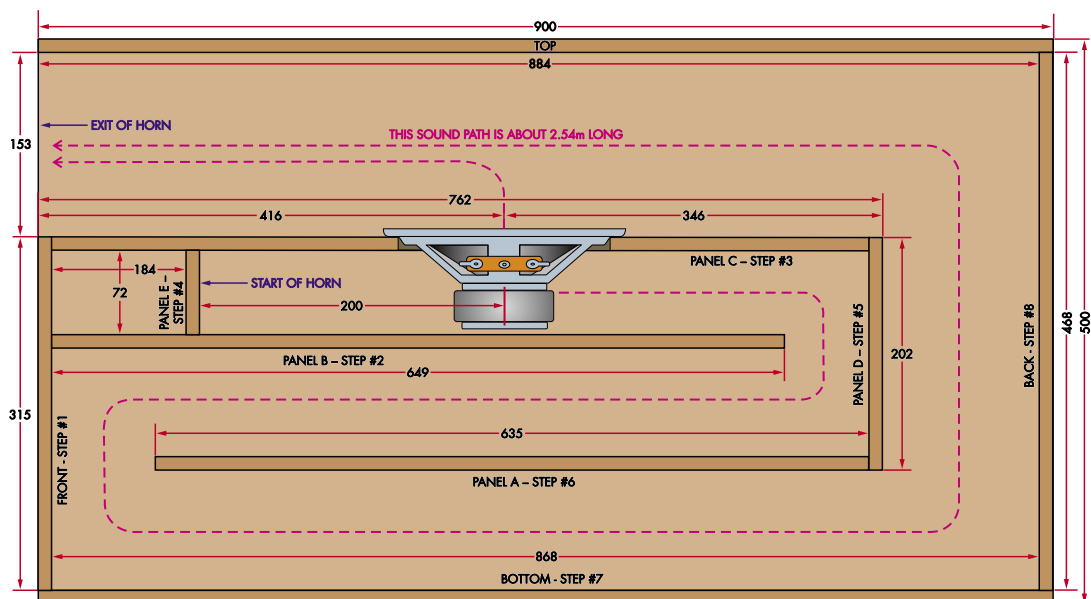
Rysunek 2. Symulowana odpowiedź konstrukcji złożonej wg rysunku 1. Daje to ładne szerokie plateau w zakresie od nieco poniżej 30 Hz do około 100 Hz oraz serię szczytów i minimów przy wyższych częstotliwościach, ponieważ fale dźwiękowe interferują w fazie lub przeciwfazie w zależności od konkretnej częstotliwości. Potrzebujemy więc filtra dolnoprzepustowego, aby wyeliminować sygnały powyżej 100 Hz, aby dźwięk brzmiał dobrze



Rysunek 3. Ta modyfikacja tuby o zmiennym przekroju pokazanej na rysunku 1 jest bardziej praktyczna w budowie, ponieważ jest krótsza i wykorzystuje tylko jeden przetwornik zamiast dwóch, ale osiąga ten sam rezultat



Rysunek 4. Trochę większa komplikacja tuby (i odrobina kreatywności w zakresie jej zwężania) pozwala nam stworzyć jeszcze bardziej kompaktową obudowę bez poświęcania dużej części wydajności



Rysunek 5. Ten schemat pokazuje, w jakiej kolejności sugerujemy mocować panele wewnętrzne do boku i pokazuje w postaci przerywanych linii dwie ścieżki akustyczne. Pokazuje również większość ważnych wymiarów, dzięki czemu możesz sprawdzić, czy budujesz subwoofer prawidłowo, ale ponieważ jest mało prawdopodobne, aby przyciąć panele dokładnie do odpowiednich rozmiarów, nie oczekuj idealnego dopasowania. Należy również pamiętać, że na rysunku 5 górny i dolny panel znajdują się powyżej i poniżej panelu bocznego, a nie nad i pod nim

emisję dźwięku z tylnej strony głośnika. Żonglując długością i ukształtowaniem toru dźwiękowego od tylnej części przetwornika do wylotu, uzyskujemy konstruktywną interferencję dźwięku w określonym paśmie. Zwiększa to wydajność i pozwala nam przesunąć przenoszenie niskich częstotliwości w dół.

Oczywiście wiąże się to z kompromisami. Tuba o zmiennym przekroju działa tylko w ograniczonym paśmie, po czym sygnał wyjściowy staje się serią szczytów i zaników. Dlatego musimy ustawić częstotliwość zwrotnicy od góry wystarczająco nisko, aby odciąć wszystkie niepożądane wyższe częstotliwości. Ponadto, poniżej odcięcia niskich częstotliwości od dołu, wychylenie membrany staje się niesterowalne, podobnie jak w przypadku wentylowanej obudowy.

Rozwiązaniem jest zasilanie subwoofera z aktywnej zwrotnicy pasmowej, która odcina zarówno wysokie częstotliwości, jak i zapewnia filtr subsoniczny do usuwania niepożądanych niskich częstotliwości.

Każdy profesjonalny system dźwiękowy zawiera filtr poddźwiękowy dla swoich subwooferów. Chroni to przetworniki przed nadmiernym wzbudzeniem, bezużytecznym wychyleniem membrany i zapobiega marnowaniu mocy przez wzmacniacze, które w przeciwnym przypadku zasilatyby głośniki sygnałami, których nie są one w stanie wygenerować w postaci dźwięku.

Ten artykuł przedstawia tylko subwoofer. Powinien on być zasilany sygnałem, który przeszedł przez filtr subsoniczny 20 Hz (górnoprzepustowy) 24 dB/oktawę i filtr dolnoprzepustowy -24 dB/oktawę z punktem odcięcia -3 dB przy 80 Hz.

W przyszłym miesiącu przedstawimy aktywną zwrotnicę, która to zapewni. Niemniej

jednak, prawdopodobnie można naszą kolumnę niskotonową wysterować z wyjścia subwoofera w wielu systemach kina domowego, zauważając jednak, że te rzadko zawierają filtr subsoniczny.

Konstrukcja

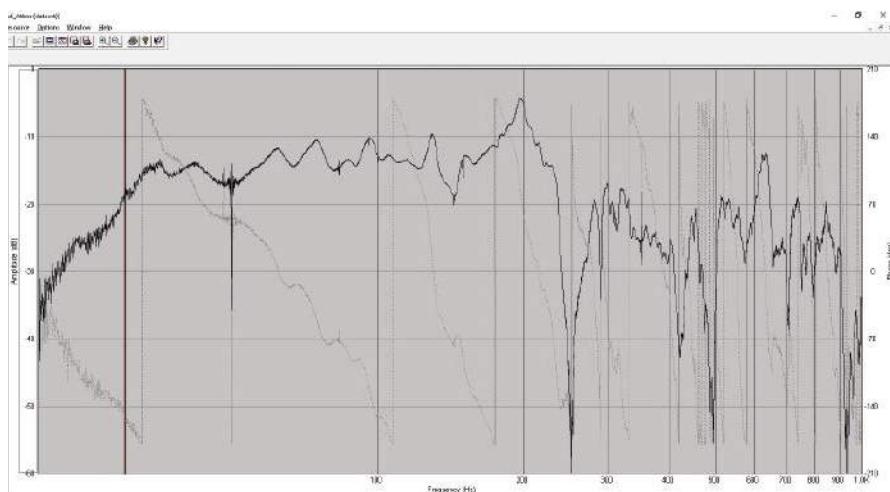
Subwoofer został zaprojektowany przy użyciu programu o nazwie „Hornresp”, napisanego przez Davida McBeana. Jest on dostępny za darmo na stronie www.hornresp.net i wspierany na kilku forach DIY Audio. Można powiedzieć, że ten program nie jest jakoś super łatwy w użyciu, ale pozwala nam modelować różne długości i wymiary sekcji labiryntu.

Jeśli wypróbujesz ten program, zalecamy skorzystanie z „Kreatora głośników” w menu „Narzędzia”. Pozwala to na zmianę długości i przekroju każdej sekcji tuby, przy

jednoczesnej obserwacji reakcji sekcji na zasilanie dźwiękiem.

Prezentowana przez nas tuba spełnia następujące wymagania:

- Punkt odcięcia z dołu pasma -3 dB poniżej 30 Hz,
- Tętnienia pasma przepustowego nie większe niż 4 dB; w realnym świecie pomieszczenia też mają różnego rodzaju rezonanse,
- Zastosowanie łatwo dostępnego, niedrogiego przetwornika,
- Materiał na obudowę, arkusze MDF, które mogą być transportowane w małym samochodzie; powiedzmy VW Golf,
- Tylko małe arkusze materiału wymagane do wykonania obudowy, najlepiej z minimalnymi liniami cięcia,
- Obudowa, którą można ukryć pod biurkiem lub za kanapą.



Rysunek 6. Zmierzona odpowiedź prototypowego subwoofera bez filtra pasmowoprzepustowego. Ciemna linia to amplituda, a jaśniejsza przerywana szara linia to faza. Zgadza się to całkiem dobrze z symulacją, chociaż odpowiedź faktycznie rozciąga się do ponad 200 Hz, zanim zaczną pojawiać się poważne maksima i minima



Subwoofer został przetestowany w warsztacie Autora, jak pokazano powyżej, oraz w przestronnej sali kościelnej pokazanej obok



Jeśli chodzi o przetwornik, Altronics C3088 (niestety niedostępny, podany w spisie części zamiennik jest oferowany w Polsce) zapewnia dobrą równowagę między rozmiarem, mocą i ceną, zapewniając jednocześnie całkiem przyzwoite wychylenie membrany w porównaniu do swoich odpowiedników. Wychylenie membrany w subwooferach jest naprawdę ważne i często jest pomijane. Przy danym SPL, im niższa częstotliwość, tym gwałtowniej rośnie skok membrany.

Uwzględnienie tego faktu jest niezbędne przy projektowaniu subwoofera. Ostatecznie, niezbędny jest przetwornik z „dobrym Xmax”. Inaczej mówiąc, zależy nam na tym, aby głośnik był w stanie poruszyć/przetłoczyć jak największą objętość powietrza.

C3088 ma 4,5 mm skoku cewki drgającej, a w naszych testach otrzymaliśmy ponad 5 mm efektywnego Xmax, co jest całkiem dobrym wynikiem.

Składając tubę w sposób pokazany na rysunku 4, aby osiągnąć powyższe, potrzebujemy

następujących parametrów – patrz zaznaczone wymiary na rysunku 5:

- 200 mm od początku tuby do tylnej części głośnika,
- 2,54 m od tyłu głośnika do wylotu tuby,
- ~420 mm (416 mm w rzeczywistości) od przodu głośnika do wylotu tuby.

To definiuje naszą ogólną obudowę jako posiadającą następujące wymiary:

- Szerokość wewnętrzna (oś „z” na rysunku 5): 250 mm, zewnętrzna 282 mm,
- Głębokość wewnętrzna: 868 mm, zewnętrzna 900 mm,
- Wysokość zewnętrzna: 500 mm.

Wydajność

Wynikową kolumnę niskotonową z tubą o zmiennym przekroju pokazano na rysunku 5. Pomiar wydajności subwooferów jest znacznie trudniejszy niż głośników pełnozakresowych ze względu na odbicia i rezonanse w pomieszczeniu. Pomiary pokazane na **rysunku 6** Autor wykonał z odległości jednego metra, ale

nie w rogu pomieszczenia. Umieszczenie subwoofera w rogu pomieszczenia, z około 20 cm odstępem między subwooferem a ścianami, zapewni lepszą wydajność (tj. więcej basu!).

Poziom dźwięku jest pokazany czarną linią (oś w dB po lewej stronie), podczas gdy jaśniejsza linia to faza (oś w stopniach po prawej stronie). Zwróć uwagę na szczyt odpowiedzi przy 200 Hz. Naprawdę potrzebujesz zwrotnicy, która zapewni minimum 18 dB tłumienia w tym punkcie, w przeciwnym razie będziesz w stanie usłyszeć rezonanse toru dźwiękowego.

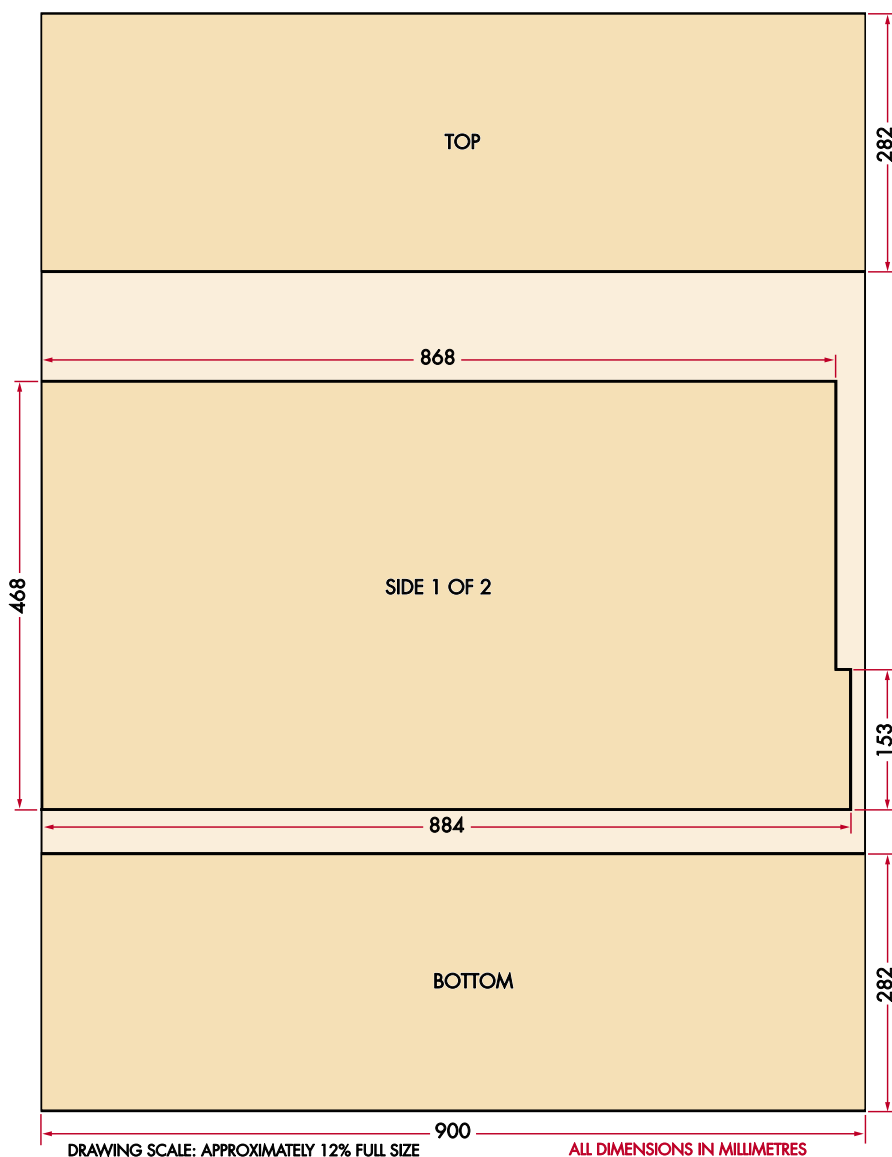
Pasma przenoszenia jest nieco równomierniejsze niż oczekiwaliśmy, ale wykazuje przewidywane tętnienia powyżej 100 Hz, szczyt przy 200 Hz i głęboki spadek przy około 250 Hz. Nie ma wątpliwości, że ten subwoofer potrzebuje stromej zwrotnicy.

Dalsze testy zostały przeprowadzone przez Autora w jego warsztacie a raczej przerobionej szopie, o powierzchni 60 m². Subwoofer generował w nich bardzo solidny bas i potrafił nawet zagrzecotać blaszaną puszką (patrz wyżej). Zintegrował się bardzo dobrze z kilkoma małymi głośnikami monitorowymi wykorzystującymi pięciocalowe przetworniki nisko-średniotonowe Vifa. Autor ustawił punkt odcięcia zwrotnicy między subwooferem a głośnikiem głównym na 80 Hz i nie zastosował żadnego tłumienia ani do subwoofera, ani do głośnika średniotonowego.

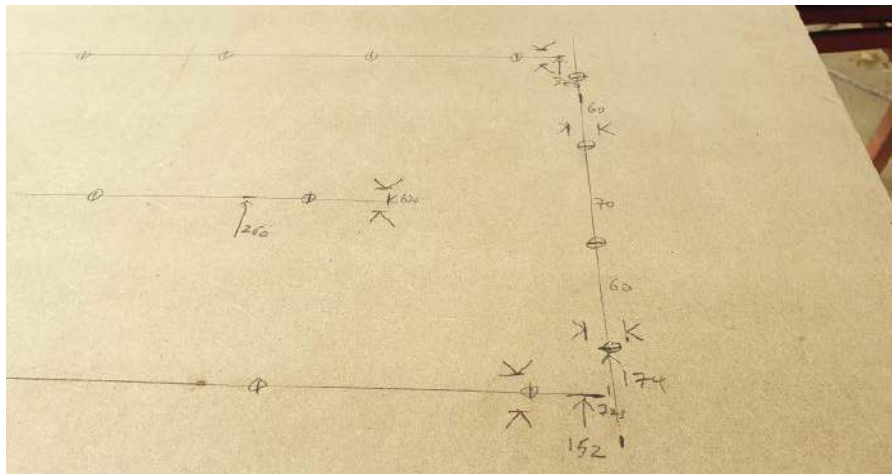
Następną próbą było przetestowanie subwoofera. Po pomalowaniu Autor zabrał kolumnę do dość dużej sali kościelnej i zintegrował z jakimś starym, ale niezwykle wydajnym 10-calowym przetwornikiem nisko-średniotonowym. Miał on z odległości 1 m

Wykaz elementów:

- 1 głośnik niskotonowy Altronics C3088 8-calowy 70W [niedostępny] [lub Wagner SB20PFC30-8] [jedytny dystrybutor w Polsce]
- 3 płyty MDF 1200×900 mm o grubości 16 mm
- 134 wkręty do drewna 8-10G o długości 50 mm z tłem stożkowym (pudełko 250 sztuk)
- 8 wkrętów 16 mm 8G (do montażu głośnika)
- 1 samoprzylepna taśma piankowa o szerokości 10 mm i długości 1 m (można ją wyciąć z szerszego paska).
- 1 para zacisków głośnikowych (użyliśmy złącza Speakon, ale można użyć dowolnego typu)
- 1 kabel głośnikowy o długości 1 m (dwużyłowy, 17AWG) [Altronics W1936].
- 1 butelka kleju typu Wikol, co najmniej 200 ml
- 1 pojemnik szpachlówki budowlanej, co najmniej 200 ml
- 1 puszkę farby podkładowej odpowiedniej do drewna
- 1 arkusz papieru ściernego o ziarnie 80 (kup więcej niż jeden, aby mieć zapas)
- 1 arkusz papieru ściernego o ziarnie 120 (kup więcej niż jeden, aby mieć zapas)
- 1 arkusz papieru ściernego o ziarnie 240 (kup więcej niż jeden, aby mieć zapas)
- 1 litr teksturowanej farby DuraTex (lub jakaś podobna farba do drewna) [www.cannonsound.com].
- 1 tubka akrylowego wypełniacza szczelin (na wypadek nieoczekiwanych szczelin w połączeniach)
- 4 noży (my użyliśmy czterech okrągłych noży Surface Gard 38 mm Side Glide z Bunnings)



Rysunek 7...9. Oto panele, które należy wyciąć z trzech arkuszy o wymiarach 1200×900 mm. Możliwe, że będziesz w stanie wyciąć je wszystkie z jednego arkusza o wymiarach 1200×2400 mm, jeśli masz sposób na jego transport (lub dostarczenie), chociaż nie sprawdziliśmy tego. Praca z mniejszymi arkuszami jest też nieco łatwiejsza. Jeszcze lepszym rozwiązaniem jest zlecenie sklepowi lub warsztatowi stolarskiemu wykonania wstępnych cięć, co pozwoli uzyskać trzy paski o szerokości 292 mm, trzy paski o szerokości 250 mm i dwa paski o szerokości 468 mm. Następnie wystarczy wykonać kilka dodatkowych cięć, aby uzyskać wszystkie potrzebne elementy



skuteczność znacznie powyżej 90 dB przy 1 W. Zwrotnica była ustawiona na 80 Hz, ale subwoofer był sporo podkręcony, aby dopasować poziom niskich i średnich tonów.

W tej sali o powierzchni 110 m² (pokazanej w prawym górnym rogu), która ma 10 metrów wysokości, prezentowana w tekście kolumna zaprezentowała się dobrze. Choć nie dałoby się z jej użyciem zorganizować dyskoteki, poradziła sobie z muzyką pop i bluesem na „entuzjastycznym”, choć nie „ekstremalnym” poziomie. Autor ma jednak, według własnej deklaracji, dość wysoką tolerancję na hałas.

Będąc w kościele, Autor wypróbował brzmienie kilku bardzo obciążonych pasmem niskotonowym chorałów gregoriańskich i był pod wrażeniem, że mógł wręcz poczuć bas.

Budowa

Zobacz listę części, aby zorientować się, jakie materiały musisz kupić. Potrzebne będą również następujące narzędzia:

- Prosta ręczna piła tarczowa; nie potrzebujesz wymyślnej piły stołowej. Alternatywnie, można poprosić pracownika lokalnego sklepu z narzędziami o wykonanie długich cięć i użycie piły ręcznej do pozostałych, krótszych cięć. [Od Red. EdW: Jak w innych artykułach o budowie kolumn, gorąco zalecamy skorzystanie z profesjonalnego warsztatu stolarskiego, choćby w supermarkiecie budowlanym z płytami MDF – zaoszczędzi nam to dużo czasu, nerwów, a nawet pieniędzy, w przypadku zniszczenia zakupionych płyt cięciem we własnym zakresie.](#)
- Wiertarka ręczna z wiertłami o średnicy 3 mm i 4 mm, stożkową końcówką do pogłębiania i końcówką do wkrętów z łbem krzyżakowym. Przydatna byłaby przystawka lub uchwyt umożliwiający dokładne pionowe wiercenie w płytach.
- Długa metalowa linijka lub prosta krawędź typu łąty budowlanej (mało zniszczonej).
- Zaciski typu G lub zaciski do stołu warsztatowego, do przytrzymywania płyty MDF podczas cięcia i wiercenia.
- Frezarka z frezem o promieniu 12 mm do wykańczania krawędzi.
- Rolka ścierna o średnicy 10 mm i długości 100 mm, z trzpieniem do zacisku wiertarki

Aby dowiedzieć się, gdzie będą leżeć panele i gdzie wywiercić otwory, należy tymczasowo ułożyć wycięte panele, jak pokazano na rysunku 5, a następnie użyć ołówka lub innego markera, aby odrysować ich kontury. Następnie możesz użyć linijki, aby narysować linie wzdłuż środka położenia każdego panelu, a miejsca wiercenia otworów będą znajdować się wzdłuż tych linii. Możesz zobaczyć na zdjęciach, jak to zrobił Autor (choć nie zaznaczył krawędzi paneli, tylko ich środki, ponieważ ma doświadczenie, jak to zrobić)

- Szpachelka i skrobak do mieszania i nakładania wypełniacza na otwory śrub.
- Wanna z wodą i ściereczka do naczyń, aby wyczyścić rozlany klej i nadmiar wyciśnięty z połączeń.

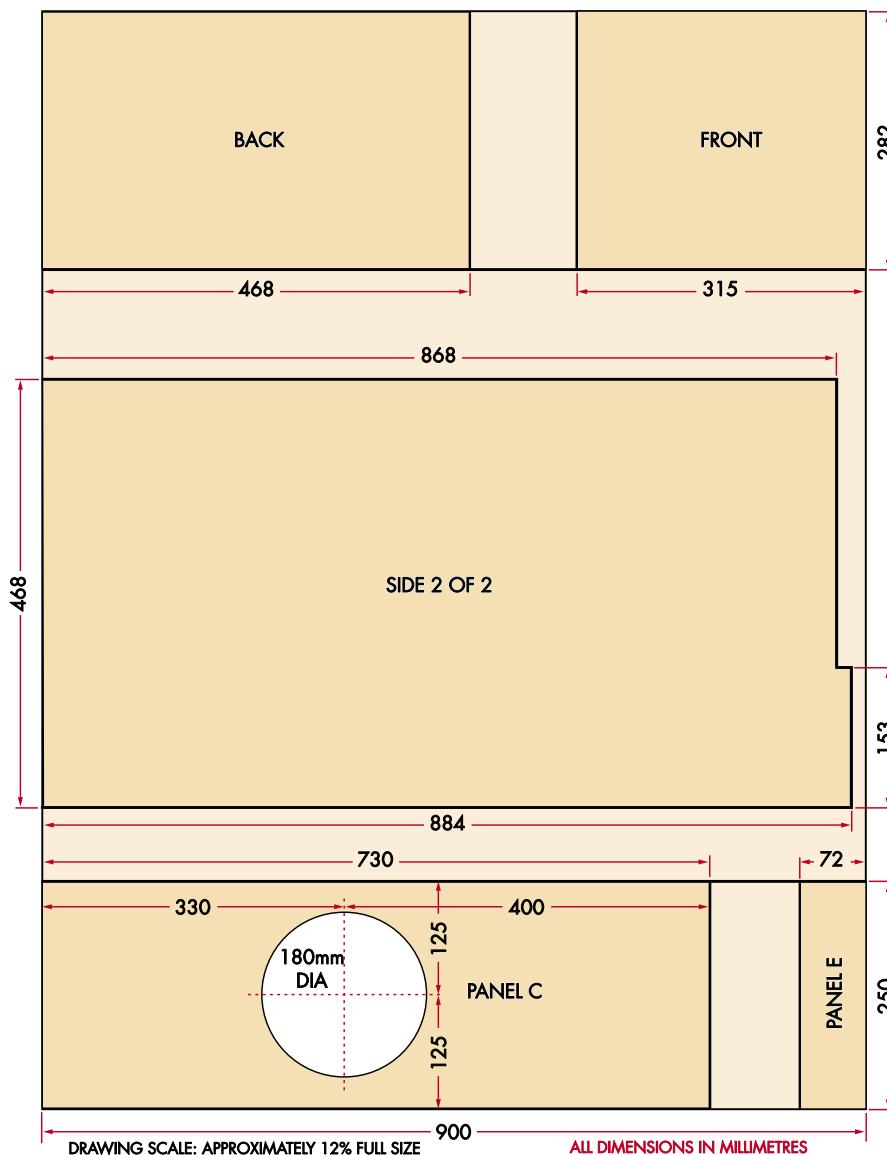
Cięcie arkuszy

Rozplanowaliśmy panele na trzech arkuszach płyt, które można przewozić w VW Golfie lub większym samochodzie, zgodnie z wcześniejszymi wymaganiami. Przed rozpoczęciem cięcia przejrzyj rysunki (rysunki 7...9). Większość potrzebnych elementów ma szerokość 250 mm lub 282 mm.

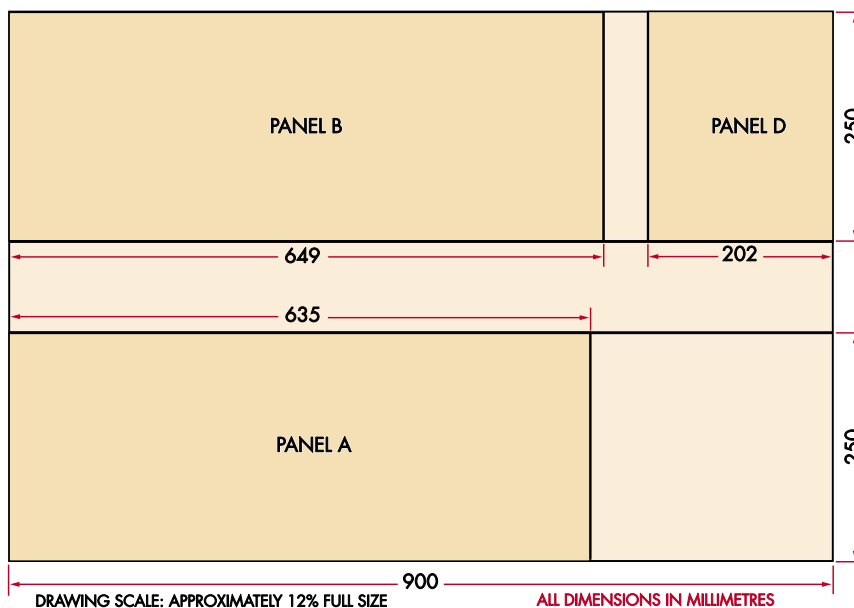
Po wykonaniu głównych cięć na pasy, można wyciąć boki ze ścinków, a także szereg długości z paneli o szerokości 250 mm i 282 mm.

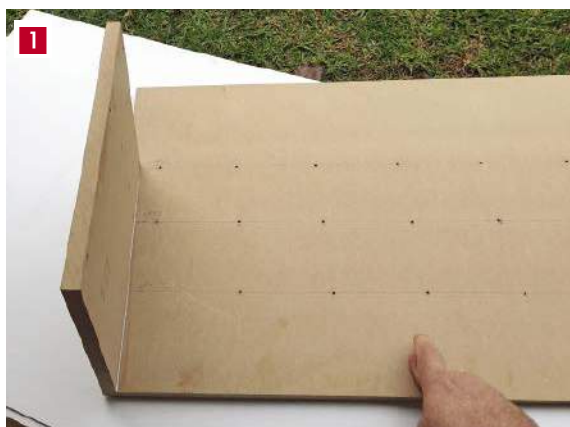
Zmierz dokładnie i sprawdź, czy panele boczne nie są zbyt wysokie lub głębokie, ponieważ błąd w tym wymiarze spowoduje zwis górnych lub tylnych paneli. Kilka wskazówek:

- Sprawdź, czy wszystkie wymiary części mieszczą się w tolerancji ± 2 mm, choć w przypadku „mebli do salonu” będziesz musiał się bardziej postarać.
- Zaznacz lokalizację otworów (patrz zdjęcie poniżej). Użyj ołówka, aby zaznaczyć panele od wewnątrz. Nie bój się mierzyć i zaznaczać wyraźnie, długimi kreskami, ponieważ po zmontowaniu pudła kolumny będą one ukryte.
- Nie spiesz się i sprawdź, czy wszystkie oznaczenia znajdują się we właściwym miejscu. Po zmontowaniu tego głośnika będzie to naprawdę uciążliwe, jeśli będziesz musiał coś przesunąć!
- W panelach bocznych znajduje się wiele otworów na śruby. Upewnij się, że są one oznaczone z dokładnością do około 2 mm. Pomiary te są niezbędne do wkręcenia śrub w panele (przegrody) wewnętrzne.
- Panel C ma wycięcie na głośnik, które należy wykonać po wycięciu panelu, ale



Jeśli używany jest głośnik niskotonowy Wagner, średnica otworu powinna wynosić 187 mm





Podczas klejenia paneli należy użyć zestawu zacisków i/lub obciążników na czas wiązania. Klej do drewna Titebond (krajowy Wikol) jest całkiem dobry do tego typu zadań. Należy pamiętać, że panele te są również łączone za pomocą śrub, a nie tylko kleju

przed złożeniem obudowy. Użyj przykładnicy z podziałką (zwijana taśma jest mało dokładna), aby zaznaczyć otwór ołówkiem, a następnie użyj wyrzynarki, aby go wyciąć. Jeśli masz szczęście i posiadasz

odpowiednią pilę otwornicę, to jeszcze lepiej. Jeśli jej nie masz, możesz użyć małej piły ręcznej, choć wymaga to trochę wytrwałości!

Teraz wykonaj wszelkie oznaczenia montażowe, które Twoim zdaniem pomogą Ci wyrównać panele. Odnies się do zdjęć; umieszczenie znaków „V” pomoże ci w ustawieniu paneli we właściwym położeniu.

Otwory na śruby określają linie, wzdłuż których zostaną przymocowane środki paneli o grubości 16 mm, więc krawędzie tych paneli będą znajdować się 8 mm po obu stronach linii wyznaczających rzędę otworów. Po zaznaczeniu otworów na śruby, zaznacz środkową linię lokalizacji paneli i dodaj $V_s = 8$ mm po obu stronach tej linii, abyś mógł podczas montażu zobaczyć, jak dobrze panel jest wyśrodkowany wzdłuż linii otworów na śruby.

Gdy wszystko się zgadza, wywierć otwory o średnicy 4 mm pod wkręty. **Wierć od wewnątrz.** Tam, gdzie wiertło wychodzi z płyty MDF, pojawią się odpryski, ale zajmijmy się tym w następnym kroku.

Następnie pogłęb wszystkie otwory od zewnątrz, aby panele miały schludny i uporządkowany wygląd. Nawierć otwory na tyle

głęboko, aby łby śrub były zagłębione w panelach na ok. 2 mm (jak pokazano poniżej na przedostatniej fotografii).

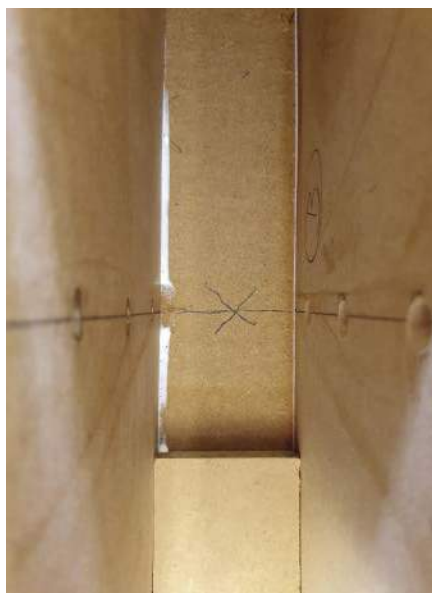
Montaż

Na rysunku 5 i numerowanych fotografiach przedstawiono kolejność, w jakiej należy zmontować panele. Prześledźmy te kroki po kolei.

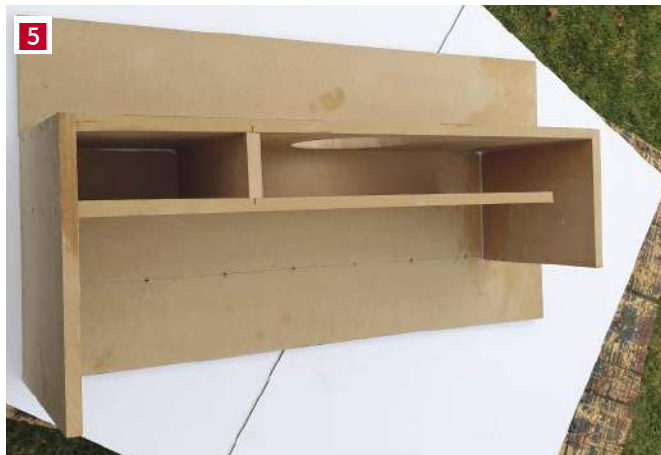
Krok 1 to przymocowanie panelu przedniego, który znajduje się w wycięciu w rogu panelu bocznego (fotografia oznaczona „1”). Jeśli to konieczne, spiłuj wycięcie w panelu bocznym, tak, aby panel przedni znajdował się w rogu. Poświęć trochę czasu, aby zrobić to dobrze, ponieważ wszystkie kolejne panele są do niego dopasowane.

Sprawdź, czy znaczniki na wewnętrznej stronie paneli przedniego i bocznego są dobrze wyrównane, a następnie nałóż niewielką ilość kleju na połączenie.

Wiertło 3 mm służy do wstępnego wiercenia otworów w bokach paneli MDF, w które zostaną następnie wkręcone śruby. Jest to ważne, aby zapobiec pękaniu paneli. Gdy wszystko jest już wyrównane, wywierć wstępnie jeden otwór (3 mm) na minimalną głębokość 50 mm w boku panelu, a następnie włóż śrubę.



Jest to zbliżenie panelu z kroku 3, pokazujące odstęp między panelami B i C





Po nałożeniu kleju i zaciśnięciu połączeń należy pozostawić je zaciśnięte przez co najmniej godzinę, a następnie pozostawić do utwardzenia na około jeden dzień

Skorzystaj teraz z okazji, aby popchnąć panel tak, aby był prosty i dobrze wyrównany. Zrób to przed wywierceniem pozostałych otworów. Upewnij się, że jesteś zadowolony, ponieważ wszystko, co nastąpi później, zależy od tego panelu! Pracę ułatwi Ci solidny metalowy przymiar o kącie prostym.

Po upewnieniu się, że wszystko jest w porządku, wywierć pozostałe otwory, a następnie skreć oba panele ze sobą.

Kroki 2 i 3 to wewnętrzne przegrody B i C (fotografie oznaczone „2” i „3”). Nałóż klej wzdłuż dolnej i przedniej krawędzi paneli 2 i 3, ale rób to pojedynczo. Dosuń panel do wyrównania i użyj oznaczeń, które wykonałeś, aby wszystko wyrównać. Oznaczenia w kształcie litery V pomogą ustawić każdy panel prostopadle do wywierconych otworów.

Wciskając panel na miejsce, wywierć i przykręć dolny otwór w panelu przednim (od kroku 1). Należy pamiętać, że rozpoczęcie od wkręcenia śruby w dolny otwór pozwoli na dociśnięcie paneli B i C do panelu przedniego przy minimalnym błędzie.

Kontynuuj wstępne wiercenie i wkręcanie śrub we wszystkie pozostałe otwory. Wyczyść klej, który wyciekł z połączeń.

Krok 4 (fotografia oznaczona „4”) polega na zamontowaniu wewnętrznej przegrody E. Wciśnij ją między panele B i C. Będzie ciasno. Spróbuj nałożyć tam trochę kleju, ale zakładając, że masz dobre dopasowanie, nie powinno to być krytyczne. Jeśli tu i ówdzie pojawi się szczelina, nałóż na nią warstwę wypełniacza akrylowego. Wywierć otwory i przykręć śruby z obu paneli B i C oraz przez panel boczny.

Krok 5 (fotografia oznaczona „5”) polega na zamontowaniu wewnętrznej przegrody D. Wyrównaj panel D z panelem C. Ponownie użyj znaków V na panelu, aby ustawić koniec panelu A panelu D we właściwym miejscu. Sztuczka polega na uzyskaniu dobrego wyrównania w rogu paneli C i D. Pomocna może być przykładnica kątowna.

Ponownie zacznij od śruby na dole połączenia paneli C i D. Po umieszczeniu jej na miejscu, nawierć i wkręć wszystkie wkręty, sprawdzając wyrównanie.

W kroku 6 (fotografia oznaczona „6”) zamontuj panel A podobnie do panelu D.

Kroki 7 i 8 (fotografie oznaczone „7” i „8”) polegają na zamontowaniu górnego i tylnego panelu. Zacznij od górnego panelu, upewniając się, że na styku

panelu przedniego i górnego znajduje się czysta krawędź. Wyczyść ją i przykręć wzdłuż panelu przedniego i bocznego. Wstępnie nawierć i przykręć wszystkie śruby do tego panelu. Następnie przykręć tylny panel tylko dwoma śrubami – nie przyklejaj go jeszcze.

Sprawdź dopasowanie dolnego panelu, starając się uzyskać dobre wyrównanie z panelem tylnym i przednią krawędzią panelu bocznego. Poruszaj nim, aby uzyskać jak najlepsze dopasowanie. Jeśli występuje niewielka niewspółosiowość, znacznie łatwiej jest ją teraz poprawić. Pamiętaj, że przed malowaniem będziesz szpachlować i szlifować – więc drobne niedokładności znikną.

Jeśli konieczne będzie nieznaczne przesunięcie tylnego panelu, usuń dwie śruby i wywierć nowe otwory, aby zamocować panel w wybranym miejscu. Po upewnieniu się, że wszystko jest w porządku, wywierć, przyklej i przykręć wkręty przez pozostałe otwory w tylnym panelu. **Nie wierć, nie przyklejaj ani nie przykręcaj jeszcze dolnego panelu!**

Krok 9 polega na zamontowaniu drugiego panelu bocznego. Do przyklejenia panelu bocznego Autor użył wypełniacza akrylowego zamiast kleju typu Wikol, ale nie jest to konieczne, zwłaszcza jeśli panele zostały





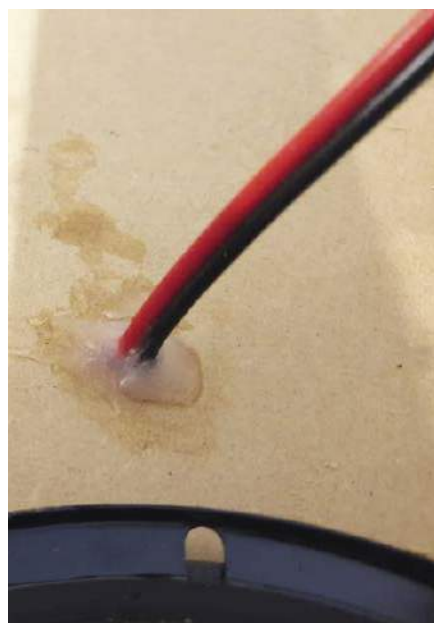
dokładnie przycięte. Po nałożeniu kleju wsuń panel boczny na miejsce, a następnie upuść go na wewnętrzne przegrody. Prawie gotową obudowę – przed montażem głośnika i przykręceniem ostatniego panelu – przedstawia fotografia oznaczona „9”.

Wepchnij panel boczny na miejsce tak, aby pasował wzdłuż górnego panelu, a następnie wywierć otwory i wkręć wkręty wzdłuż tej krawędzi. Następnie dociśnij przednią i tylną krawędź panelu bocznego, aby uzyskać dobre wyrównanie z panelem przednim i tylnym, i ponownie wywierć otwory i wkręć wkręty.



Frezarka znacznie ułatwia wykończenie krawędzi, ale można to również zrobić papierem ściernym. Wszelkie szczeliny i pęknięcia można wypełnić za pomocą mieszanki kleju do drewna i trocin lub wypełniacza do drewna.

Gotowy subwoofer został pokryty podkładem, a następnie pomalowany na czarno. Można też po prostu nałożyć lakier lub wosk do drewnianych podtóg, w zależności od tego, jak ma wyglądać



Po zakończeniu należy uszczelnić przewody głośnikowe

Teraz wywierć i przykręć śruby przez wszystkie otwory w panelu bocznym. Jeśli pomiary były prawidłowe, wszystkie śruby wejdą w wewnętrzne przegrody. Jeśli wiertło wypadnie przez otwory i ominie wewnętrzną przegrodę, należy wiercić pod kątem, który pozwoli śrubie złapać wewnętrzną przegrodę (nie powinno to mieć miejsca!).

Montaż przetwornika

Głośnik należy zamontować przed zainstalowaniem dolnego panelu. Gdy wszystko jest na swoim miejscu, poruszaj głośnikiem C3088,

aby upewnić się, że jest dobrze osadzony w wyciętym otworze. Jeśli otwór jest nieco niewymiarowy, głośnik nie będzie dobrze osadzony. Jeśli tak jest, nadszedł czas, aby to naprawić! Stolarze mogą kręcić na to nosem, ale możesz użyć pilnika tarnika, aby nieznacznie powiększyć otwór, biorąc pod uwagę, że jest on ukryty w środku.

Następnie przyklej taśmę piankową wokół krawędzi otworu głośnika. Zapewni to dobre uszczelnienie między głośnikiem a panelem C. Potem zainstaluj przewód głośnikowy, jak pokazano na kolejnym zdjęciu, upewniając się,



że ma wystarczającą długość, aby przeciągnąć go przez otwór pod zaciski i przylutować do zacisków. Upewnij się, że możesz zidentyfikować przewód „+” do odpowiedniego gniazda zacisków głośnikowych, ponieważ musi on być podłączony do końcówki „+” głośnika.

Poprowadź przewód głośnikowy do złącza głośnikowego. Użyliśmy złącza Speakon, ponieważ wiele naszych głośników korzysta z takich złączy, chociaż możesz preferować użycie w subwooferze gniazda bananowych i/lub zacisków.

Wywierciliśmy otwór 25 mm na tylnym panelu, aby zamontować używane przez nas złącza. Nie pokazaliśmy lokalizacji ani rozmiaru tego otworu na schematach cięcia, ponieważ jego rozmiar i kształt będą zależeć od złącza i może znajdować się on w dowolnym miejscu na tylnym panelu. Prawdopodobnie najlepiej będzie wyglądał, jeśli będzie znajdował się gdzieś wzdłuż pionowej linii środkowej.

Teraz umieść głośnik w otworze i zamontuj go za pomocą śrub 16 mm 8G, których łąby nie przechodzą przez otwory w koszu głośnika (tj. z wystarczająco dużymi łbami lub podkładkami, jeśli to konieczne). Wywierć otwory na głębokość 2 mm, a następnie wkręć osiem

śrub. Stopniowo dokręcaj śruby po przeciwnych stronach głośnika, aż wszystkie będą dokręcone. Nie dokręcaj ich zbyt mocno, ponieważ taśma piankowa zapewni dobre uszczelnienie.

Po zamontowaniu głośnika przymocuj dolny panel. Nie przyklejaj go; po prostu przykręć go za pomocą dużej liczby wkrętów. Umożliwi to późniejszy dostęp do głośnika, jeśli zajdzie potrzeba jego wymiany.

Wykończenie obudowy

Wszystkie zewnętrzne krawędzie Autor wyfrezował frezem o promieniu 12 mm. Jeśli nie masz frezarki, nie ma to znaczenia. Użyj papieru ściernego o ziarnie 80, a następnie 120, aby zaokrąglić krawędzie, aż będą wyglądać gładko.

Następnie użyj szpachlówki budowlanej do wypełnienia wszystkich otworów po wkrętach z łbem stożkowym. Po wyschnięciu przeszlifuj te obszary, a następnie nałóż drugą warstwę szpachlówki, aby uzyskać naprawdę gładkie obszary. Nie wypełniaj jednak otworów w dolnym panelu! Musisz być w stanie go odkręcić.

Po stwierdzeniu, że obudowa jest wystarczająco gładka, a wszystkie otwory na śruby

– z wyjątkiem tych w dolnym panelu – są teraz zlicowane z płytą MDF, pokryj obudowę lakierem; my użyliśmy lakieru DuraTex. Najpierw nałożyliśmy cienką warstwę, a po godzinie drugą, grubszą, używając wałka o średnicy 10 mm.

DuraTex to teksturowana farba sprzedawana do stosowania na profesjonalnych obudowach głośnikowych. Jest ona wytrzymała i tworzy teksturę, dzięki czemu dobrze toleruje i ukrywa nierówności. Pomaga również ukryć wszelkie niedoskonałości naszej pracy.

Na koniec przykręć nóżki i subwoofer jest gotowy. Zgodnie z wcześniejszą obietnicą, w przyszłym miesiącu zostanie opisana aktywna zwrotnica, która idealnie nadaje się do użytku z tym subwooferem (lub dowolnym dwudrożnym lub trójdrożnym systemem głośnikowym). ■

Phil Prosser

Adaptacja do wydania
polskiego – Andrzej Nowicki

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



Elektrochemia

Rozwiązanie znajdziesz na www.elportal.pl/quizy

1. Luigi Galvani eksperymentował z żabimi nogami ponieważ:

- ☐ badał wpływ elektryczności statycznej na płazy;
- ☐ badał działanie nerwów i mięśni;
- ☐ szukał metody leczenia drgawek.

2. Za przyczynę zaobserwowanego zjawiska Galvani uważał:

- ☐ zwierzęcy magnetyzm;
- ☐ zwierzęcą elektryczność;
- ☐ naturalną elektryczność.

3. Galwanoskop żabi to prymitywne narzędzie do:

- ☐ wykrywania elektryczności;
- ☐ pomiaru wielkości ładunku elektrycznego;
- ☐ wykrywania impulsów nerwowych.

4. Alessandro Volta chcąc obalić teorię Galvaniego stworzył:

- ☐ woltomierz;
- ☐ ulepszony galwanoskop;
- ☐ ogniwo elektrochemiczne.

5. Zjawisko Seebecka polega na:

- ☐ powstawaniu różnicy napięć między dwoma metalami, jeśli te się stykają;
- ☐ przepływie jonów w roztworze elektrolitu między elektrodami wykonanymi z różnych metali, jeśli te są ze sobą połączone;
- ☐ wydzielaniu się soli lub gazów na powierzchniach elektrod ogniwa elektrochemicznego.

6. Szereg napięciowy przedstawia:

- ☐ szereg wartości napięć różnych typów ogni;
- ☐ szereg wartości napięć dla różnych substancji używanych do budowy ogni;
- ☐ szereg wartości napięć dla różnych elektrolitów, gdy użyjemy ich w stosie Volty.

7. Za wartość zerową (odniesienia) uznaje się napięcie dla:

- ☐ tlenu;
- ☐ siarki;
- ☐ wodoru.

8. Leclanché stworzył ogniwo:

- ☐ niklowo-żelazowe;
- ☐ cynkowo-żelazowe;
- ☐ cynkowo-węglowe.

9. Salmiak to zwyczajowa nazwa:

- ☐ chlorku amonu;
- ☐ azotanu amonu;
- ☐ siarczanu amonu.

10. Depolaryzator służy do:

- ☐ usuwania nadmiaru ładunków z ogniwa;
- ☐ wychwytywania gazów z powierzchni elektrody w ogniwie;
- ☐ ochrony przed gromadzeniem się ładunków statycznych.



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

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



Pęseta do testów SMD

To sprytne małe urządzenie składa się zaledwie z 11 komponentów. Mimo to może mierzyć wartości wielu rezystorów SMD i kondensatorów, a także pokazywać orientację diod i LED oraz mierzyć ich napięcia przewodzenia. Jest szybkie i łatwe w użyciu, zasilane z wbudowanego ogniwa pastylkowego, z ekranem OLED o wysokim kontraście do wyświetlania odczytów.

Praca z częściami SMD czasami jest trudna. Odczytywanie oznaczeń komponentów może być męczące dla oczu, jeśli podzespół w ogóle jest oznaczony! Elementy takie jak kondensatory SMD są całkowicie anonimowe i po wyjęciu z opakowania prawie niemożliwe do odróżnienia. Pęseta testowa SMD ułatwia to zadanie, informując użytkownika o komponencie poprzez zwykłe jego zmierzenie.

W niektórych przypadkach pęsety te mogą również mierzyć właściwości komponentu po jego przylutowaniu do płytki (choć, w zależności od konfiguracji obwodu, czasami odczyty nie będą dokładne).

W miarę upływu czasu coraz mniej części elektronicznych jest dostępnych w wariantach do montażu przewlekane i coraz częściej producenci wypuszczają produkty głównie lub w całości jako SMD. Są one mniejsze i tańsze niż części przewlekane, mogą być montowane po obu stronach płytki (często z wewnętrznymi ścieżkami biegnącymi pod spodem), a także są mniej wrażliwe na wstrząsy i wibracje.

Oczywiście, chociaż mniejsze części mogą być korzystne, stwarzają również problemy podczas pracy z nimi. Niektóre narzędzia, takie jak pęseta i lupa, są niezbędne.

Gdy tylko będziesz miał okazję wypróbować naszą pęsetę testową SMD, myślimy, że dodasz ją do swojego niezbędnika SMD!

Pęseta

Pęseta

Wartości części SMD są bardzo trudne do odczytania za pomocą multimetru. Wielokrotnie zdarzało nam się wciskać sondy multimetru w wyprowadzenia części SMD, próbując uzyskać odczyt, tylko po to, aby wypłynęły, odleciały i nigdy więcej nie zostały znalezione. Pęseta zapewnia znacznie bardziej naturalny sposób na wykonanie tej czynności, a ponieważ nie trzeba wywierać dużego nacisku, istnieje mniejsze prawdopodobieństwo, że część wystrzeli w kosmos.

Co więcej, ponieważ pęseta jest wygodnym sposobem na podnoszenie i przenoszenie takich części, jeśli dołączymy narzędzie pomiarowe do pęsety, może Ci ono powiedzieć, jaką częścią się zajmujesz, gdy jesteś w trakcie umieszczania jej na płytce.

Cechy i specyfikacje:

- Identyfikuje i mierzy rezystory, kondensatory, diody i diody LED
- Kompaktowy wyświetlacz OLED
- Zasilanie z pojedynczego ogniwa litowego, około pięciu lat pracy w trybie czuwania
- Automatyczne włączanie i wyłączanie zasilania
- Wyświetla napięcie własne ogniwa, gdy żaden komponent nie jest podłączony
- Może w pewnych okolicznościach mierzyć komponenty w obwodzie
- Może wykonać tysiące pomiarów przed wyczerpaniem ogniwa
- Pomiar rezystancji: 10 Ω do 1 MΩ
- Pomiar diod: polaryzacja i napięcie przewodzenia do około 3 V
- Pomiar pojemności: 1 nF do 10 μF

napięcia zasilania przez rezystor 10 kΩ, dzięki czemu mikroprocesor działa normalnie tak długo, jak podawane jest zasilanie.

CON1 to złącze programowania szeregowego w obwodzie (ICSP), którego szpilki łączą się odpowiednio ze stykami 4, 1, 8, 7 i 6 układu IC1. W razie potrzeby można go użyć do zaprogramowania układu IC1 na płytce. Nie jest to konieczne w przypadku zakupu wstępnie zaprogramowanego układu PIC.

Wykrywanie podzespołów

Oznaczenia IOTOP i IOBOT na schemacie oznaczają normalne stany IO tych styków. W stanie bezczynności styk 2 jest ustawiany w stanie wysokim, a styk 3 w stanie niskim. Odpowiada to oznaczeniom CON+ i CON-.

W każdym cyklu pomiarowym układ IC1 porównuje wewnętrzne napięcie referencyjne 1,024 V z napięciem szyn zasilających i na tej podstawie oblicza napięcie ogniwa. Może to być później wykorzystane do obliczenia napięcia przewodzenia diody. Jeśli żaden element nie zostanie wykryty, wyświetlane jest napięcie ogniwa.

Następnym testem jest sprawdzenie obecności kondensatora. Styk 2 jest w stanie

niskim i pobierana jest seria próbek napięcia na styku 5, aż napięcie na styku 5 spadnie poniżej połowy napięcia ogniwa lub pobranych zostanie 255 próbek.

Jeśli układ IC1 nie widzi spadku napięcia przypominającego rozładowanie kondensatora, zgłasza, że nie zidentyfikuje kondensatora. Może się to również zdarzyć, jeśli pojemność jest zbyt niska (co powoduje, że napięcie spada szybciej niż IC1 może wykonać pomiary) lub zbyt wysoka (co powoduje, że napięcie nie zmienia się wystarczająco w okresie próbkowania).

Pojemność jest obliczana na podstawie spadku napięcia w funkcji czasu, chociaż w celu uniknięcia kosztownej obliczeniowo funkcji logarytmicznej stosowane jest przybliżenie; nasz kod mieści się w niewielkiej ilości bajtów wypełniających dostępną przestrzeń programową.

Dokładność przybliżenia jest jako taka tylko przy wartościach zbliżonych do górnej granicy zakresu pomiaru. Biorąc jednak pod uwagę, że wiele kondensatorów jest dostarczanych w klasie dokładności 20%, jest to wystarczające do większości celów i będzie odpowiednie

do rozróżnienia komponentów, chyba że są one bardzo zbliżone pod względem wartości.

Test pojemności jest wykonywany jako pierwszy, ponieważ oznacza to, że czas od ostatniej próbki może być wykorzystany do upewnienia się, że kondensator jest jak najbliżej pełnego naładowania.

Należy pamiętać, że nie należy podłączać naładowanego kondensatora do miernika pęsety (lub innego podobnego miernika). Jeśli jest naładowany do więcej niż kilku woltów, gdy jest podłączony, lub polaryzacja jest odwrócona, może łatwo uszkodzić mikroprocesor IC1. Nawet jeśli tak się nie stanie, prawdopodobnie pomiar nie będzie prawidłowy.

Jeśli kondensator nie zostanie wykryty, stan bezczynności jest przywracany na 200 μs (aby umożliwić ustabilizowanie się napięcia). Następnie mikroprocesor wykonuje pomiar napięcia na styku 5, odwraca polaryzację na kolejne 200 μs, wykonuje kolejny pomiar, a następnie odwraca polaryzację. Algorytm uśrednia 16 próbek dla każdej polaryzacji, aby poprawić dokładność.

Co drugi wynik surowego pomiaru ADC jest modyfikowany w celu uwzględnienia faktu, że został wykonany z odwróconą polaryzacją. Jeśli dwa pomiary napięcia są zbliżone, przyjmuje się, że część jest rezystorem, a jego wartość jest podawana zgodnie ze wzorem na dzielnik napięcia (patrz rysunek 2).

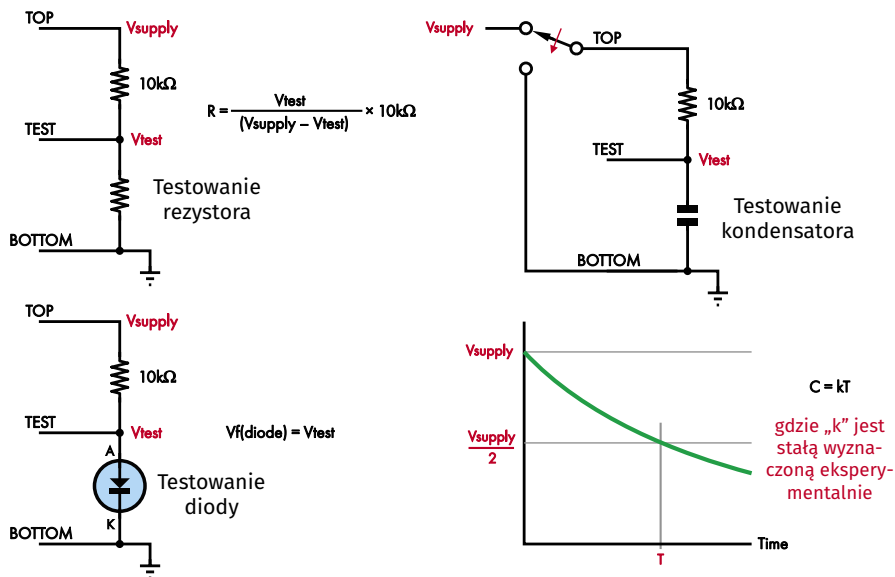
Jeśli jedna wartość jest bliska pełnego napięcia szyny, a druga nie, to część jest prawdopodobnie pewnego rodzaju diodą, więc są zgłaszane napięcie przewodzenia i jego kierunek.

Może to obejmować diody LED, diody krzemowe i diody Schottky'ego. Część LED fototranzystorów i optoizolatorów również powinna wskazywać odczyt diody. Dwukolorowe diody LED i inne układy diod mogą nie zostać wykryte, ponieważ będą przewodzić w obu kierunkach i nie będą wykazywać otwartego obwodu w odwrotnym kierunku.

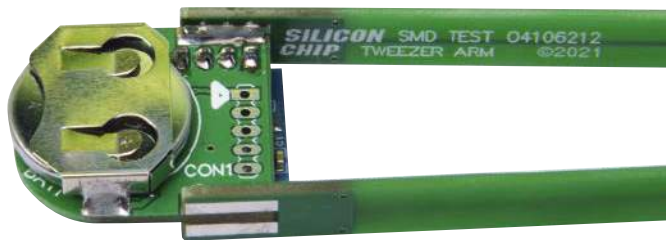
Jeśli jesteś zręczny, prawdopodobnie możesz zidentyfikować tranzystory bipolarne, podłączając pęsetę do ich przypuszczalnych wyprowadzeń bazy i emitera i identyfikując polaryzację złącza; powinno zostać wykryte jak dioda.

Diody LED podłączone anodami do CON+ i katodami do CON- będą polaryzowane w kierunku przewodzenia przez prąd spoczynkowy i zasilane prądem o natężeniu kilkuset mikroadmperów, co powinno wystarczyć do ich przyćmionego świecenia i wskazania, że działają.

Prąd testowy jest dość niski ze względu na rezystor 10 kΩ, nie więcej niż około



Rysunek 2. Pokazuje różne sposoby pomiaru wartości komponentów za pomocą pęsety. Rezystancja jest mierzona przy użyciu dobrze znanego wzoru na dzielnik rezystancji, podczas gdy test diody mierzy napięcie na elemencie w obu kierunkach. Pomiar pojemności opiera się na zmianie napięcia w przedziale czasowym po rozładowaniu przez znaną rezystancję



Z tyłu pęsety nie widać zbyt wiele, ale zwróć uwagę, że jedno ramię, listwa kołkowa OLED (CON2) i uchwyt ogniwa (BAT1) znajdują się dość blisko siebie. Przed zamontowaniem ogniwa pastylkowego należy dwukrotnie sprawdzić, czy nie ma zwarcia

300 μ A. Dlatego wskazane napięcie przewodzenia może być nieco niższe niż można by się spodziewać (np. czytając kartę katalogową). Na przykład diody krzemowe wykazują napięcie przewodzenia około 0,5...0,6 V.

Po określeniu typu i wartości części (lub napięcia ogniwa) są one wyświetlane po prostu jako liczba z odpowiednimi jednostkami i mnożnikami. Aby odróżnić napięcie ogniwa od napięcia przewodzenia diody, wyświetlany jest symbol diody z polaryzacją dopasowaną do usytuowania diody w stosunku do sond pęsety.

Po pięciu sekundach niewykrycia żadnej części, ekran OLED przechodzi w tryb niskiego poboru mocy, styk 5 zostaje włączony jako źródło przerwania, a mikroprocesor przechodzi w tryb uśpienia. Mikroprocesor można wybudzić, po prostu zwierając sondy pęsety.

Pokazaliśmy więc, jak taki prosty obwód może wykonywać różne testy w celu wykrycia i pomiaru szeregu komponentów. Rysunek 2 pokazuje działanie tych algorytmów w nieco bardziej szczegółowy sposób.

Gdy ekran OLED jest aktywny, pobór prądu wynosi około 4 mA. Spada on do 5 μ A, gdy mikroprocesor jest w stanie uśpienia, a wyświetlacz OLED jest wyłączony. Tak więc żywotność ogniwa będzie zależeć głównie od czasu, w którym pęseta jest faktycznie używana. Typowe ogniwo pastylkowe CR2032 ma pojemność 220 mAh. Daje to żywotność około pięciu lat, co jest dobrym wynikiem, biorąc pod uwagę, że typowy „okres przechowywania” ogniwa pastylkowego wynosi 10 lat.

Budowa

Jeśli jeszcze nie zacząłeś pracować z częściami SMD, zaczniesz teraz, ponieważ oczywiście zaprojektowaliśmy pęsetę testową SMD z komponentami SMD. Użyj górnego i dolnego schematu montażowego PCB



Aby pokazać szczegóły konstrukcyjne, pozostawiliśmy naszą pęsetę nagą, ale warto przykryć główną płytkę drukowaną krótkim kawałkiem szerokiej folii termokurczliwej. Posłuży to również do przytrzymania na miejscu ogniwa pastylkowego

pokazanego na **rysunku 3** jako przewodnika podczas budowy. Główna część pęsety SMD jest zamontowana na płytce PCB o kodzie 04106211 i wymiarach 28×26 mm.

Zalecamy użycie topnika lutowniczego (najlepiej pasty, chociaż strzykawka z płynnym topnikiem jest lepsza niż nic), lutownicy z regulowaną temperaturą grotu, miedzianej plecionki do usuwania nadmiaru lutu i lupy. Sugerujemy również użycie zwykłej pęsety.

Ponieważ topnik może generować dym po podgrzaniu, należy pracować w miejscu z dobrą wentylacją. Sprawdź również, czy twój topnik ma zalecany rozpuszczalnik czyszczący; niewielkie ilości alkoholu izopropylowego są dobrym, uniwersalnym środkiem czyszczącym, natomiast wymieniany niekiedy spirytus metylowy (inaczej drzewny) z oczywistych powodów uważamy za nieakceptowalny. Należy również uważać ze stosowaniem spirytusu denaturowanego, ponieważ często zawiera dodatki acetonu lub rozpuszczalnika nitro,

które mogą uszkodzić sitodruk PCB, a nawet niektóre podzespoły.

Zacznij od przymocowania płytki PCB do powierzchni roboczej stroną montażu komponentów skierowaną do góry. Jeśli nie masz imadła lub uchwytu do PCB, użyj masy Blu-Tack jak gumy do żucia, aby przykleić ją do biurka.

Nałóż topnik na pola lutownicze komponentów SMD, a następnie przytrzymaj IC1 na miejscu. Jeśli wszystkie wyprowadzenia znajdują się wewnątrz pól stykowych, to wszystko jest w porządku. IC1 powinien mieć małą kropkę oznaczającą styk 1; upewnij się, że znajduje się on na końcu najbliższego miejsca dla kondensatora 100 nF, jak zaznaczono na PCB.

Wyczyść grot lutownicy i nałóż niewielką ilość świeżego lutu. Następnie dotknij grotiem jednego z narożnych styków układu IC1. Powinno to spowodować przepływ lutu na miejsce zetknięcia. Jeśli część wygląda na przylegającą do płytki drukowanej, a jej



Rysunek 3. Pomimo obecności tylko kilku komponentów, wykorzystaliśmy obie strony płytki drukowanej. Jedną z zalet komponentów SMD, w porównaniu z częściami przewlekany, jest to, że znacznie łatwiej jest mieć części po obu stronach bez obawy o to, gdzie idą przewody. Zwróć uwagę na orientację układu IC1; po jego zamontowaniu reszta montażu jest dość prosta

Rysunek 4. Na płytce PCB ramion nie ma zamontowanych żadnych komponentów; są to w zasadzie tylko elastyczne przewody, które są przylutowane do głównej płytki drukowanej i zaciskają testowany element na drugich końcach

wyprowadzenia nadal znajdują się w obrębie wszystkich pól, przylutuj pozostałe wyprowadzenia, dotykając ich grotami lutownicy.

W razie potrzeby można dodać więcej lutu na grot, a także więcej topnika. Jedynym problemem związanym z użyciem zbyt dużej ilości topnika jest to, że będzie on generował więcej dymu, a jego czyszczenie zajmie nieco więcej czasu. Czyli nie zawsze więcej oznacza lepiej.

Jeśli okaże się, że zmostkowałeś jakieś styki, najłatwiej jest przylutować pozostałe wyprowadzenia przed naprawieniem tego, ponieważ pomoże to utrzymać układ scalony we właściwym miejscu. Następnie nałóż więcej topnika, docisnij miedzianą plecionkę do zmostkowanych styków za pomocą lutownicy i delikatnie odsuń plecionkę, gdy pochłonie nadmiar lutu.

Sprawdź styki za pomocą lupy przed przystąpieniem do dalszych czynności i w razie potrzeby powtórz powyższe kroki. Konieczne może być usunięcie pozostałości topnika, jeśli utrudnia on widoczność między stykami.

Pozostałe części można przylutować podobnie, z tą różnicą, że żadna nie jest spolaryzowana, a wszystkie mają znacznie większe wyprowadzenia i pola stykowe.

W następnej kolejności umieść jedyny kondensator; prawdopodobnie będzie to jedyna część bez oznaczeń. Przylutuj jedno wyprowadzenie, sprawdź poprawność ułożenia na stykach lutowniczych PCB, a następnie przylutuj drugie wyprowadzenie. W razie potrzeby wyretuszuj pierwszy styk.

Następnie zamontuj rezystory; oba mają taką samą wartość. Nie są one spolaryzowane, ale dobrą praktyką jest dopasowanie oznaczeń do optu na płytce drukowanej, aby pomóc w rozwiązywaniu problemów.

Odwróć płytkę drukowaną, aby zamontować uchwyt ogniwa. Podobna technika lutowania będzie działać dla uchwytu ogniwa, z tą różnicą, że jest on nieco większy, więc będzie

potrzebował więcej ciepła. Podkręć lutownicę, jeśli jest regulowana.

Umieść uchwyt ogniwa, upewniając się, że otwór jest skierowany w stronę zakrzywionego końca płytki drukowanej. Jeśli wygląda na to, że nie można włożyć lub wyjąć ogniwa, to prawdopodobnie jest to niewłaściwa strona. Nałóż trochę topnika i przylutuj jedno wyprowadzenie. Sprawdź, czy wszystko jest prawidłowo wyrównane, a następnie przylutuj drugie wyprowadzenie. W razie potrzeby możesz wyretuszować pierwszy styk.

Na tym kończą się części montowane powierzchniowo i jest to dobry moment na wyczyszczenie pozostałości topnika. Ponieważ wiele środków do czyszczenia topnika to łatwopalne rozpuszczalniki, po tym kroku należy pozwolić płytce PCB dokładnie wyschnąć.

Jeśli masz pusty mikroprocesor, teraz jest dobry moment, aby go zaprogramować. Należy to zrobić przed zainstalowaniem modułu OLED, ponieważ może on po podłączeniu zakłócać programowanie.

Programowanie IC1

Możesz pominąć tę sekcję, jeśli masz wstępnie zaprogramowany mikroprocesor, co będzie miało miejsce, jeśli kupiłeś go w sklepie internetowym Silicon Chip-a.

W przeciwnym razie do zaprogramowania tego układu potrzebny będzie programator PICkit 3 lub PICkit 4 oraz oprogramowanie MPLAB X IDE (zintegrowane środowisko programowania), które można pobrać bezpłatnie ze strony internetowej Microchip (zwykle w pakiecie z MPLAB X IDE).

Można również użyć programatora Snap, jeśli zmodyfikujesz go zgodnie z instrukcjami na stronie 69 wydania Silicon Chip-a z czerwca 2021 r. (patrz siliconchip.com.au/Article/148890). Jest to konieczne, ponieważ programator Snap nie może dostarczać

zasilania w inny sposób (lub można wymyślić inny sposób tymczasowego zasilania mikroukładu podczas programowania).

Chociaż możliwe jest przylutowanie złącza programowania do płytki drukowanej pęsety, ponieważ będzie on używany tylko raz i będzie przeszkadzał, wolimy użyć delikatnej siły, aby przytrzymać złącze na miejscu podczas programowania.

Wybierz PIC12F1572 jako mikroukład docelowy w IPE, a następnie otwórz plik 0410621A.hex. Następnie wystarczy nacisnąć przycisk „Programm”, aby rozpocząć proces (tuż przed tym należy zacząć naciskać złącze, aby przytrzymać szpilki złącza na płytce drukowanej).

Jeśli pojawi się komunikat „Programming/Verify complete”, programowanie zakończyło się pomyślnie. W przeciwnym razie spróbuj ponownie.

Odłącz programator przed przejściem do następnego kroku.

Zakończenie

Jeśli chcesz dodać metalowe końcówki do ramion pęsety (wykonanych z płytek drukowanych oznaczonych kodem 04106212 o wymiarach 100×8 mm), łatwiej jest to zrobić przed zamontowaniem ich na pęsecie. Wytnij kawałki mosiężnego paska mniej więcej na wymiar. Kawałki można precyzyjnie przyciąć do odpowiedniej długości po zmontowaniu pęsety.

Przylutuj po jednym pasku do końca każdego ramienia, pozostawiając nadmiar około 5...10 mm. Należy pamiętać, że po zakończeniu montażu paski powinny znajdować się po wewnętrznej stronie ramion (szczegóły można znaleźć na naszych zdjęciach).

Postaraj się wprowadzić trochę lutu w otwory w płytce drukowanej, ponieważ zwiększy to wytrzymałość mechaniczną. Miedziane podkładki do montażu powierzchniowego są zasadniczo przyklejone do płytki drukowanej, więc nie trzeba wiele, aby je oderwać.

Jeśli nie masz mosiężnego paska, opłaca się dodać kilka małych kropelek lutu do końcówek pęsety. Zapewni to większy obszar styku, a także pewną odporność na zużycie końcówek.

Umieść ramiona na płytce PCB pęsety na polach CON+ i CON- i z grubsza wyrównaj ich pozycje. Ich końce bez wywierania nacisku powinny być oddalone od siebie o około 10...15 mm; zapewnia to rozsądną siłę roboczą i zasięg. Odstęp ten oznacza również, że pęseta może być używana do testowania

Wykaz elementów:

- 1 dwustronna płytkę drukowaną o kodzie 04106211 i wymiarach 28×26 mm (główna płytkę drukowaną)
- 2 dwustronne płytki drukowane o kodzie 04106212 i wymiarach 100×8 mm (ramiona pęsety)
- 1 8-bitowy mikroprocesor PIC12F1572-I/SN lub PIC12F1572-E/SN zaprogramowany wsadem 0410621A.hex, SOIC-8 (IC1)
- 1 moduł OLED 0,49" 64×32 (Silicon Chip Online Shop nr kat. SC5602)
- 1 uchwyt ogniwa pastylkowego, do montażu powierzchniowego (BAT1) [Digi-key BAT-HLD-001-ND, Mouser 712-BAT-HLD-001 lub podobny]
- 1 litowa bateria pastylkowa CR2032 lub CR2025
- 1 5-szpilekowa jednorzędowa kątowna listwa kołkowa (CON1; opcjonalnie, potrzebna tylko do programowania IC1)
- 1 kondensator ceramiczny 100 nF SMD 1206 50 V X7R, 1206 [Altronics R9935]
- 2 rezystory 10 kΩ 1%, SMD 1206 [Altronics R8188]
- 2 krótkie 15×2 mm kawałki cienkiej (np. 1 mm) blachy mosiężnej na końcówki pęsety (opcjonalnie)
- 1 przezroczysta rurka termokurczliwa o długości 40 mm i średnicy 30 mm (opcjonalnie; patrz tekst)
- 2 odcinki 100 mm rurki termokurczliwej o średnicy 10 mm (opcjonalnie; patrz tekst)



Można nabyć gotowe pęsety z przewodami przeznaczonymi do podłączenia do innych urządzeń, takich jak multimetr. Jeśli wolisz takie, możesz odciąć wtyczki bananowe i przylutować je do naszej płytki głównej zamiast naszych ramion przymocowanych do PCB. W takim przypadku należy upewnić się, że przewód dodatni jest podłączony do pola CON+ na płytce drukowanej, a czarny przewód do CON-

części przewlekanych, takich jak rezystory osiowe, diody i kondensatory.

Odkryliśmy, że dopasowanie ramion równo z krawędzią płytki PCB ułatwiło lutowanie i utrzymało ramię CON+ z dala od połączenia CON2 i ekranu OLED. Wygląda to również schludniej; zobacz nasze zdjęcia.

Gdy będziesz zadowolony z pozycji ramion, nałóż dużą ilość lutu po obu stronach połączeń, aby zabezpieczyć je na miejscu. Wypróbuj działanie, sprężystość i wyrównanie ramion i dostosuj je w razie potrzeby.

Możesz także przyciąć i dopasować końcówki, jeśli tego wymagają. Ściśnięcie ramion razem i przeciągnięcie drobnym pilnikiem po końcówkach wyrówna je, jeśli są nieco różnej długości.

Aby końcówki ramion były równoległe, umieść drobny papier ścierny lub płaski pilnik między końcówkami i doszlifuj je, aż będą zadowolające. Pomoże to również dodać trochę

szorstkości do końcówek, aby pomóc im uchwycić komponenty i uniknąć możliwości ich lotu w nieznane!

Ekran OLED

Moduł OLED jest ostatnim elementem do zamontowania. Złącze dostarczone z modulem ma przekładkę o odpowiedniej grubości, aby zamontować OLED równoległe do głównej płytki drukowanej, chociaż szpilki prawdopodobnie wymagają przycięcia.

Zacznij od przylutowania listwy kołkowej do PCB w stykach CON2, najlepiej z dłuższymi szpilkami skierowanymi do góry. Ułatwi to ich późniejsze przycięcie. Sprawdź, czy nie ma mostków między wyprowadzeniami CON2, ramieniem CON+ i uchwytem ogniwa.

Przylutuj jeden zestyk ekranu OLED do górnej części listwy kołkowej i sprawdź, czy wszystko wygląda prawidłowo, a ekran nie dotyka niczego pod spodem; w razie

potrzeby wyreguluj jego położenie. Przylutuj pozostałe szpilki, a następnie przytnij ich nadmiar od góry, uważając, aby nie uszkodzić ekranu OLED. Następnie usuń folię ochronną z wyświetlacza.

Używanie

Włóż ogniwo litowe biegunem ujemnym w kierunku do płytki drukowanej. OLED powinien ożyć i pokazać odczyt nieco ponad 3 V dla świeżego ogniwa. Ściśnięcie ramion powinno pokazać rezystancję kilku omów.

Jeśli wyświetlacz w ogóle się nie uruchamia, sprawdź połączenia OLED. Jeśli nie ma pomiaru rezystancji, może to oznaczać problem z obwodem testowym; sprawdź rezystory, IC1 i ramiona pęsety.

Po przejściu pęsety w tryb uśpienia, do wybudzenia wykorzystuje ona cyfrowe czujniki o niskim poborze mocy. Dlatego mogą się wybudzić, jeśli są podłączone do niektórych, ale nie wszystkich części. Odwrotnie podłączone diody i rezystory o wysokiej wartości mogą nie wybudzić pęsety, ale prawie wszystkie kondensatory (po rozładowaniu) wydają się to robić.

W takim przypadku należy po prostu zerwać końcówki pęsety, a następnie zbadać element. Po wykryciu elementu, żaden następny element nie zostanie wykryty przez pięć sekund.

Uwaga

Podobnie jak w przypadku każdego projektu wykorzystującego ogniwo pastylkowe, pęseta powinna być trzymana z dala od dzieci, które mogą połączyć ogniwo. Pęseta ma również dość spiczaste końcówki, co jest kolejnym powodem, dla którego należy trzymać ją poza zasięgiem ciekawskich palców.

Na główną płytkę drukowaną można nałożyć kawałek szerokiej, przezroczystej rurki termokurczliwej, aby ją zaizolować i zabezpieczyć. Może to również posłużyć do zabezpieczenia ogniwa pastylkowego; nie powinno ono wymagać zbyt częstej wymiany, a termokurczliwą rurkę można wymienić po takim czasie.

Można również zamontować cieńszą folię termokurczliwą na ramionach. Zapewni to lepszą izolację, a także poprawi uchwyt powierzchni ramion pęsety. ■

Tim Blythman

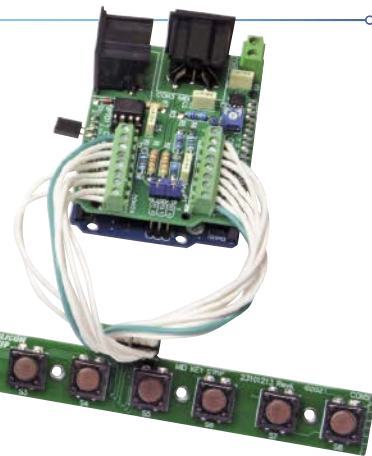
Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. www.silicon-chip.com.au

REKLAMA

www.facebook.com/ElportalPL



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



Prosta, liniowa klawiatura muzyczna MIDI, część 3

Ta klawiatura MIDI jest następcą naszej 64-klawiszowej matrycy MIDI, którą zaprezentowaliśmy w numerach 12-2023 i 1-2024 EdW. Jest podobnie modyfikowalna i oferuje sposób na łatwe tworzenie muzyki, chociaż można ją wykorzystać w wielu innych zastosowaniach.

Podczas gdy panele MIDI typu matryca są popularne jako kompaktowy sposób sterowania i łączenia się ze sprzętem MIDI, liniowy układ klawiatury, taki jak w fortepianie, jest bardziej „standardowy” i dla wielu osób dość intuicyjny.

Jest to modułowy dodatek do sprzętu MIDI, który opisaliśmy w EdW w grudniu i styczniu 2023/24, na podstawie oryginalnego tekstu w Silicon Chip-ie (siliconchip.com.au/Series/363).

Podobnie jak opisana tam matryca MIDI, niniejsza prościutka klawiatura nie musi być używany wyłącznie do celów sterowania MIDI lub muzycznych.

Matryca MIDI została zaprojektowana do użytku z płytą Arduino Leonardo, ponieważ Leonardo może z łatwością zapewnić natywny interfejs USB MIDI dzięki wszechstronnym bibliotekom Arduino MIDI.

Zademonstrowaliśmy również kilka różnych szkiców programów, które można uruchomić na Leonardo, aby uzyskać różne funkcje, a także pokazaliśmy kilka sposobów łączenia się z oprogramowaniem zarówno na komputerze PC, jak i smartfonie z systemem Android.

W tym samym tekście zaprezentowaliśmy nakładkę Arduino, która umożliwia podłączenie sprzętu do szerokiej gamy urządzeń MIDI za pomocą standardowych złączy DIN.

Opisana obecnie klawiatura ma zastąpić matrycę 64 przycisków, jako część większej konstrukcji, tak jak zostało to przedstawione we wcześniejszych częściach tej serii. Zapoznaj się z tymi artykułami, zwłaszcza z pierwszą

częścią, aby zrozumieć, w jaki sposób można używać matrycy (a teraz klawiatury). Do przekształcenia przedstawionej tu klawiatury w minimalny koder MIDI potrzebna jest co najmniej płytka Arduino Leonardo i kilka przewodów połączeniowych.

Matryca

Oryginalna matryca jest w zasadzie tylko tablicą przycisków, którą moduł Leonardo może skanować, aby odbierać dane wejściowe użytkownika. W naszym oprogramowaniu MIDI każde naciśnięcie klawisza jest konwertowane na nutę.

Każdy rząd lub kolumna matrycy jest podłączona do cyfrowego wejścia modułu Leonardo. Korzystając z tradycyjnej techniki skanowania każdego rzędu po kolei, można wykryć poszczególne naciśnięcia przycisków.

W naszej wersji oprogramowania wiersze są podłączone do końcówek skonfigurowanych jako wejścia z wysokorezystancyjną polaryzacją do wysokiego poziomu logicznego. Początkowo wszystkie końcówki kolumn są również ustawione na tryb wejścia o wysokiej impedancji.

Każda kolumna jest kolejno konfigurowana jako wyjście i ustawiana w stan niski. Jeśli dowolny przycisk podłączony do tej kolumny zostanie naciśnięty, odpowiadający mu styk wiersza zostanie również ustawiony przez styki przycisku w stanie niskim. Skanując kolejno kolumny, możemy wykrywać naciśnięcia poszczególnych przycisków.

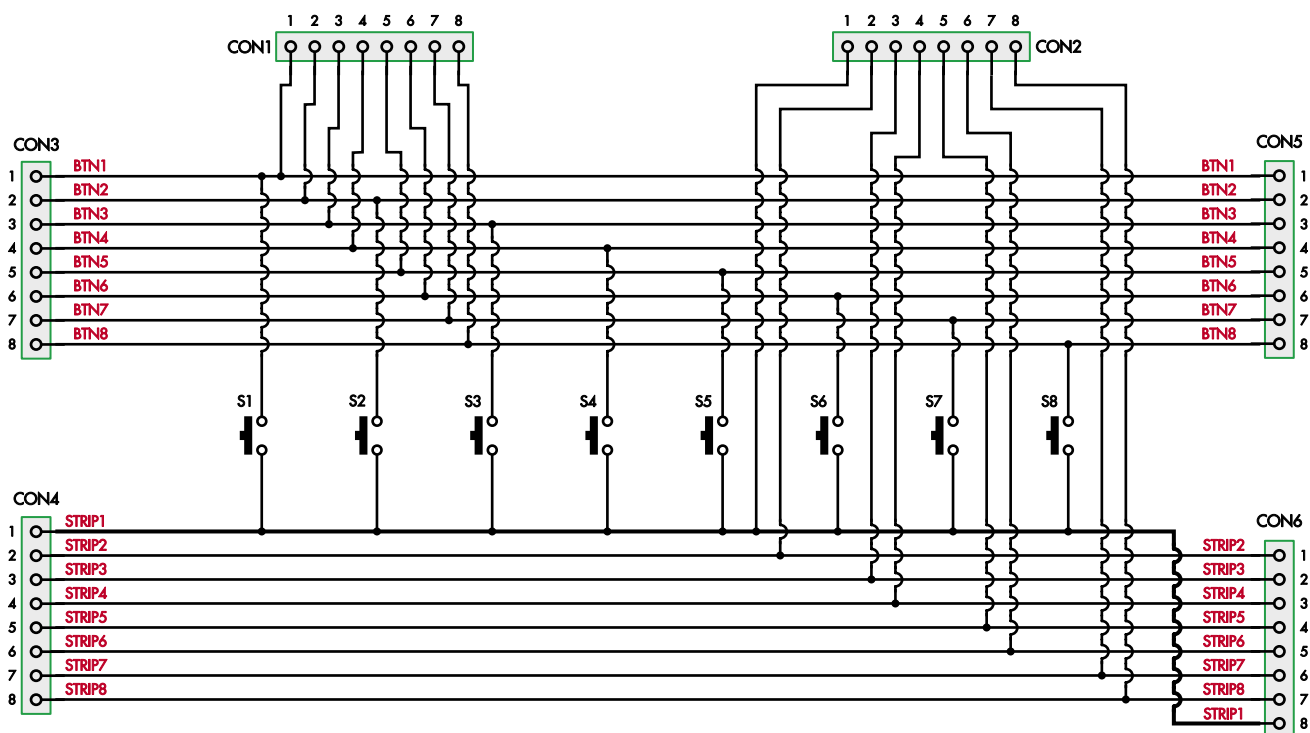
Chociaż ten system jest prosty, nie może wykryć wielu jednoczesnych naciśnień klawiszy; w tym celu każdy przycisk musiałby być wyposażony w diodę, aby zapobiec niejednoznaczności zwarcia linii i kolumn przenoszonymi się przez matrycę (komputerowi gracze znają to jako problem „trzech klawiszy”, generujący wirtualne „naciśnięcie” czwartego). Nasza matryca pomija te diody na rzecz prostoty i kompaktowości, a ta liniowa klawiatura również nie próbuje rozwiązać tego problemu.

Nowa klawiatura

Rozważaliśmy klawiaturę liniową dla naszego oryginalnego projektu, ale nie mogliśmy znaleźć sposobu, aby była ona zarówno kompaktowa, jak i funkcjonalna.



Nasz prototyp wykorzystuje trzy z tych płytek PCB, ponieważ klawiatura wykonana z pełnego zestawu ośmiu płytek PCB miałaby ponad metr szerokości. Zachowaliśmy pola CON1 i CON2 na niektórych płytkach, aby zademonstrować i przetestować różne opcje połączeń. W praktyce potrzebny jest tylko jeden zestaw; należy pamiętać, że podłączenie do CON3 i CON4 skrajnej lewej listwy jest równoważne



Rysunek 1. Jest to schemat ideowy prostego obwodu pojedynczej płytki drukowanej z ośmioma przyciskami. Przesunięcie między CON4 i CON6 umożliwia łatwą rozbudowę do ośmiu płytek drukowanych i 64 przycisków. Choć płytka ma tylko 8 przycisków, to układ połączeń jest bardziej skomplikowany i trudniejszy do prześledzenia, niż dla matrycy 8×8

Teraz opracowaliśmy modułową konstrukcję, dzięki czemu można zbudować użyteczną klawiaturę, która nadal jest kompaktowa, lub można ją rozszerzyć aż do 64 klawiszy, co skutkuje urządzeniem o długości ponad metra! Ale nadal potrzebujesz tylko 16 przewodów, aby podłączyć ją do Arduino.

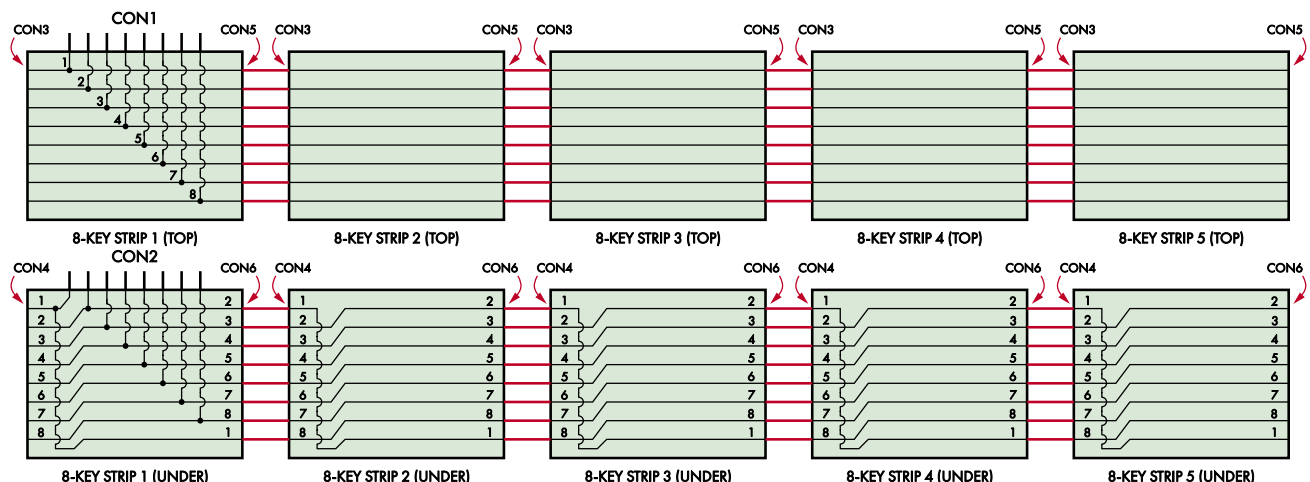
Podstawową jednostką klawiatury jest pojedyncza płytka drukowana z ośmioma przyciskami. Każdy klawisz jest podłączony do tego samego styku w rzędzie, co pozostałe, a także do jednego z ośmiu styków w kolumnie. Pojedynczy moduł klawiatury jest identyczny z jednym rzędem matrycy.

Rysunek 1 przedstawia schemat ideowy obwodu. CON1 jest podłączony do kolumn, a każdy zacisk na CON1 jest podłączony do jednej strony każdego z przełączników dotykowych, S1-S8. Pozycja 1 CON2 jest podłączona do drugiej strony przełączników S1-S8.

Na każdym końcu płytki drukowanej modułu klawiatury znajdują się złącza CON3-CON6, które mogą być używane do łańcuchowego łączenia kolejnych płytek drukowanych w celu rozszerzenia klawiatury. Są to ośmiostykowe dwurzędowe złącza szpilkowe przystosowane przez nas do montażu powierzchniowego, o rastrze 2,54 mm.

Złącza CON3 i CON5 (na górnej stronie płytki drukowanej) są podłączone w tej samej kolejności i równolegle do złącza CON1. W ten sposób sygnały kolumn mogą przechodzić między płytkami drukowanymi, łącząc sąsiednie złącza CON3 i CON5. Są one podłączone jako magistrala równoległa.

Podobnie, z tyłu płytki PCB, CON4 na jednej płytce PCB łączy się z CON6 na następnej. CON4 jest okablowany tak samo jak CON2, ale sprytnym rozwiązaniem jest sposób okablowania CON6. Styk 1 CON6 jest podłączony do styku 2 na CON4 i tak dalej, wszystkie przesunięte o jedną pozycję.



Rysunek 2. Pokazuje, jak połączyć wiele płytek drukowanych z 8 przyciskami na każdej w większą klawiaturę, aby moduł Arduino mógł stwierdzić, który klawisz został naciśnięty. Każda płytka PCB wzdłuż łańcucha przesuwają pozycję, w której połączenie jest ostatecznie wykonywane w CON2, umożliwiając wykrycie do 64 klawiszy. Pokazanych jest 5 płytek z połączeniami na górze i na spodzie

A co z czarnymi klawiszami?

Być może zauważysz w tym miejscu, że pianina i fortepiany mają dwa rzędy klawiszy, białe i czarne, i miałbyś rację. Ponadto na każdą oktawę przypada siedem białych klawiszy, a nie osiem. Zachowaliśmy liniowy układ ośmiu klawiszy, aby uczynić go prostym i możliwym do zastosowania w szerokim zakresie aplikacji.

Planujemy zaprojektować w późniejszym terminie 12-klawiszową płytkę drukowaną, w której klawisze są rozmieszczone naprzemiennie i pogrupowane jak w pianinie. W międzyczasie, jeśli chcesz używać tej płytki jak prawdziwego pianina, możesz zbudować ją w dwóch rzędach, z górnym rzędem przesuniętym poziomo o 6 mm (Od Red. EdW: czy raczej 10 mm, co stanowi połowę odległości pomiędzy przyciskami na PCB?) od dolnego rzędu i z przerwami w klawiszach u góry, aby uzyskać odpowiednią konfigurację. Oba rzędy można połączyć szeregowo (zakładając, że zawierają łącznie nie więcej niż 8 płytek drukowanych).

Oprogramowanie można stosunkowo łatwo zmodyfikować, aby przekształcić dwa rzędy klawiszy w prawidłową sekwencję, tak aby mogły działać jak klawiatura pianina. Pozostałoby jednak ograniczenie wykrywania tylko jednego naciśnięcia klawisza na raz. Nasza planowana płytkę PCB przyszłej klawiatury fortepianowej usunie to ograniczenie.

Zalóżmy, że podłączyliśmy zestaw ośmiu takich modułów, numerując je od 1 do 8 od lewej do prawej, z CON3 i CON4 podłączonymi odpowiednio do CON5 i CON6. Łącząc się z CON1 i CON2 na pierwszym module, otrzymalibyśmy odpowiednik pełnej matrycy 8x8, tyle tylko, że z przyciskami w jednym rzędzie.

Rysunek 2 pokazuje, jak „rzędy” są mapowane z powrotem do CON2 na pierwszej płytce drukowanej. CON1, CON3 i CON5 są po prostu połączone równolegle i nie są modyfikowane przez ten system.

Inne konfiguracje

Jeśli przyjrzyj się uważnie płytce drukowanej, zobaczysz, że mała wypustka na górze, z której wystają CON1 i CON2, może być usunięta. Pozwala to na odłamanie/odcięcie tych zakładek we wszystkich modułach z wyjątkiem jednego.

W rzeczywistości, ponieważ CON3 i CON4 są okablowane identycznie jak CON1 i CON2, można nawet usunąć zakładkę ze wszystkich płytek i zamiast tego po prostu wykorzystać połączenia matrycy z CON3 i CON4 na końcu ostatniej lewej płytki.

Jeśli nie masz nic przeciwko przeprogramowaniu styków w oprogramowaniu (lub zmianie sposobu ich podłączenia z powrotem do płytki Leonardo), połączenia CON1 i CON2 nie muszą być wykonane na pierwszej płytce. Można nawet pobrać te połączenia ze środka tablicy.

Zaprojektowaliśmy płytkę drukowaną tak, aby używała dużych

12 mm przełączników dotykowych, ponieważ są one znacznie przyjemniejsze w dotyku dzięki większej powierzchni pod palec. Może się okazać, że niektóre mniejsze przełączniki można dopasować, wyginając ich przewody, chociaż nie próbowaliśmy tego.

Ponieważ na tej płytce drukowanej jest mniej miejsca na prowadzenie przewodów niż w matrycy, nie ma możliwości zamontowania podświetlanych przycisków, którą posiada matryca 8x8.

Sprzęt

Podobnie jak matryca, prezentowana przez nas klawiatura ma dość podstawową konstrukcję, dzięki czemu można ją dostosować do własnych wymagań. Przyciski są umieszczone w odstępach 20 mm, a na każdej płytce drukowanej znajdują się cztery otwory montażowe pod śruby M3.

Nominalnie otwory montażowe będą rozmieszczone w odstępach 40 mm, choć zależy to od dokładności montażu sąsiednich płytek. Płytki PCB mają szerokość 20 mm, nie licząc zakładek dla CON1 i CON2; 28 mm z zakładką na miejscu.

Zdecydowanie zalecamy montaż klawiatury na listwie wsporczej, aby płytki drukowane nie ugięły się podczas naciskania klawiszy. Połączenia dla CON3–CON6 nie zapewnią dużej wytrzymałości mechanicznej, ponieważ są one lutowane powierzchniowo i są tylko powierzchniowo połączone z płytką drukowaną.

Budowa

Większość Czytelników będzie chciała zbudować klawiaturę z wieloma płytkami drukowanymi ułożonymi w linii, więc opiszemy, co jest potrzebne, aby to osiągnąć. Klawiatura jest montowana na płytce PCB o kodzie 23101213 i wymiarach 158x28 mm. Użyj schematu montażowego PCB, pokazanego na **rysunku 3**, jako przewodnika montażu komponentów.

Zaplanuj i ułóż moduły przed rozpoczęciem budowy. Aby zachować kompaktowość, połączenia między płytkami są nieco ciasne i łatwiej będzie je połączyć przed zamontowaniem innych podzespołów.

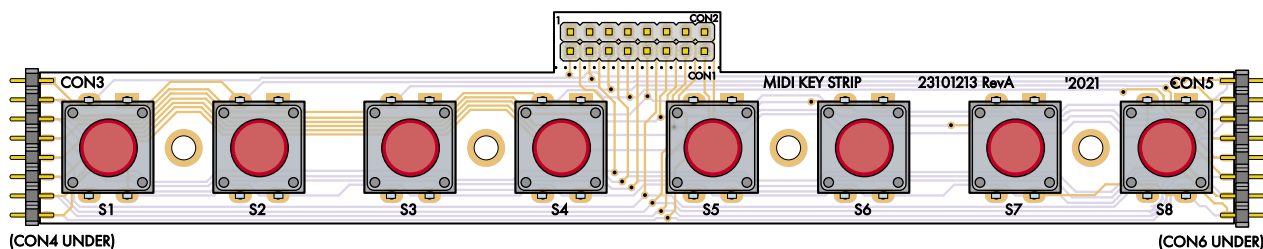
Jeśli chcesz zastosować inny układ połączeń, prawie każda metoda okablowania, odpowiednio CON3 i CON5 oraz CON4 i CON6 będzie działać. Możesz nawet użyć gniazd żeńskich na jednym końcu i listew kołkowych na drugim, aby umożliwić rozłączanie płytek, ale zwiększy to nienaturalnie odstęp między przyciskami różnych płytek.

Aby rozpocząć montaż płytki drukowanej, należy odłamać wszystkie występy złączy CON1/CON2, które nie są potrzebne. Zrób to, nacinając krawędź wzdłuż linii ostrym nożem, aby przeciąć miedziane ścieżki, a następnie ostrożnie wygnij płytkę drukowaną szczypcami, aby wykonać czyste odłamanie.

Możesz spiliować i wyczyścić szorstką krawędź. Oprócz naszych zwykłych ostrzeżeń dotyczących unikania wdychania pyłu PCB (np. poprzez pracę na zewnątrz i noszenie masek), należy uważać, aby nie spiliować ścieżek, które biegną blisko krawędzi PCB, szczególnie z tyłu.

Każda płytkę PCB ma 158 mm długości, co oznacza, że jest 2 mm wolnego miejsca na łącznik, jeśli odstęp między przyciskami mają być równe. Użyliśmy przyciętych dwurzędowych listew kołkowych. Plastikowe oprawki szpilek mają szerokość blisko 2 mm, zapewniając niezbędne odstępy.

Zacznij od przycięcia kołków, które mają być użyte jako łączniki. Jest to kłopotliwe, ale konieczne, ponieważ między sąsiednimi korpusami przełączników na sąsiednich płytkach drukowanych nie ma więcej niż 8 mm, a typowe kołki mają około 11 mm wysokości.



Rysunek 3. Podczas montażu nie ma zbyt wiele okazji do popełnienia błędów, chociaż zalecamy najpierw zamontować mostki pomiędzy PCB, ponieważ przełączniki dotykowe utrudniają dostęp do nich podczas ich lutowania. Przyciski powinny zatrzasknąć się na miejscu, więc ich lutowanie jest łatwe

Wykaz elementów:

8 modułów klawiatury

7 męskich listew kołkowych 2x8, przyciętych na wysokość 1,6 mm z każdej strony (CON3-CON6)

1 złącze 2x8 styków (męskie lub żeńskie, aby pasowało do potężniejszego Leonardo, CON1 i CON2) sprzęt montażowy w zależności od zastosowania (gwintowane kołki dystansowe M3, śruby itp.)

Moduł klawiatury

1 dwustronna płytka drukowana klawiatury o kodzie 23101213 i wymiarach 158x28 mm

8 przełączników dotykowych 12 mm [np. Diptronics DTS-21N-V lub Jaycar SP0608/SP0609, Altronics S1135 + S1138].

Liczbę cięć można zmniejszyć o połowę, przesuwając szpilki w plastikowej listwie. Umieść płytkę drukowaną na twardej, płaskiej powierzchni i umieść listwę 2x8 szpilek w otworach złącza CON1/CON2. Mocno docisnij plastik płaską krawędzią, która zmieści się między szpilkami. Idealnie nadaje się do tego stalowa linijka.

Spowoduje to przesunięcie kołków tak, że z oporaki będzie wystawać tylko 1,6 mm (grubość płytki drukowanej) każdej szpilki. Teraz odwróć listwę 2x8 kołków i użyj głębokości płytki drukowanej jako przyrządu do cięcia szpilek z drugiej strony na długość 1,6 mm.

Końcówki szpilek mogą odlecieć przy cięciu z dużą prędkością, dlatego należy nosić okulary ochronne i celować w szpilki podczas cięcia tak, aby odleciały od Ciebie. Zobacz sąsiednie zdjęcia, które pokazują, jak powinna wyglądać listwa kołkowa po przycięciu, a następnie po przymocowaniu do PCB.

Lutowanie tych złączy jest trochę trudne, ponieważ nie są one dobrze dopasowane. Traktuj je jak części montowane powierzchniowo, nakładając przed lutowaniem pastę topnikową na pola stykowe. Zalecamy zabezpieczenie listew podczas lutowania taśmą wysokotemperaturową (np. Kapton), aby się nie poruszały.

Przymocuj końce na miejscu i sprawdź, czy kołki nie dotykają ścieżek przełącznika dotykowego. Możesz nawet przetestować przełączniki, aby potwierdzić odstępy.

Przylutuj pozostałe kołki i szczerze użyj topnika. Pomoże on lutowi uformować czyste połączenia, które znajdują się tam, gdzie trzeba. Odwróć płytkę i polutuj kołki z tyłu PCB.

Usuń nadmiar topnika za pomocą rozpuszczalnika, np. IPA, i przetestuj odsłonięte ścieżki CON3-CON6 pod kątem ciągłości między oboma końcami paska. Jak widać na rysunku 2, CON3 jest podłączony bezpośrednio do CON5. Ale CON4 będzie przesunięty względem CON6 (chyba, że masz pełny zestaw ośmiu płytek drukowanych), więc sprawdź, czy każda ścieżka na CON4 jest podłączona do jednej i tylko jednej ścieżki na CON6.

Najlepiej zrobić to teraz, ponieważ przerabianie tych połączeń z przylutowanymi przyciskami dotykowymi może być dość kłopotliwe.

Następnie zamontuj przełączniki. Powinny one zatrzasknąć się na swoim miejscu; przed lutowaniem sprawdź, czy leżą równo.

Na koniec należy przylutować złącza CON1 i CON2. Użyliśmy żeńskich gniazd, aby dopasować kable, które przygotowaliśmy dla matrycy 8x8, ale możesz użyć dowolnego rozwiązania, nawet lutując przewody bezpośrednio do PCB.

Podłączanie

Przetestowaliśmy nasze urządzenie ze szkicem MIDI_ENCODER. Jeśli jeszcze tego nie zrobiłeś, zalecamy przeczytanie wcześniejszych części tej serii artykułów, ponieważ opisują one oprogramowanie bardziej szczegółowo.

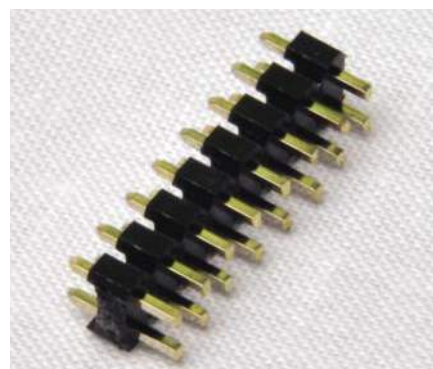
Ponieważ klawiatura jest w rzeczywistości odpowiednikiem matrycy wyposażonej w niepodświetlane przyciski, można zastosować przy posługiwaniu się klawiaturą liniową wiele pomysłów związanych z matrycą.

Podobnie jak w przypadku matrycy, podłącz przewód CON1 klawiatury do CON2 nakładki MIDI (lub odpowiednich styków Leonardo), a CON2 klawiatury do CON1 na nakładce MIDI, łącząc styk 1 ze stykiem 1.

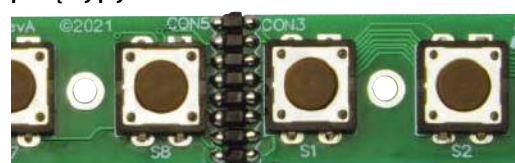
Sprawdź, czy wszystkie przyciski działają zgodnie z oczekiwaniami, korzystając z powiadomień o klawiszach wyświetlanych w emulatorze monitora szeregowego Arduino. Jeśli okaże się, że niektóre klawisze na płycie drukowanej nie działają (ale nie wszystkie), sprawdź połączenia kolumn pod kątem ciągłości; są to CON3 do CON5 z przodu płytki drukowanej.

Jeśli żaden z przycisków na płycie nie działa, może to być problem z połączeniami CON4 do CON6 z tyłu płytki.

Dzięki zestawowi płytek drukowanych klawiatury podłączonych do naszej nakładki MIDI, mamy liniowy układ przycisków, na których można grać jak na pianinie. Należy jednak pamiętać, że domyślnie, w przeciwieństwie do pianina, nie można używać wielu przycisków jednocześnie



Przycięte kołki mostka (pokazane powyżej) mają około 7 mm wysokości, dzięki czemu zmieszczą się między końcowymi przełącznikami na sąsiednich płytkach drukowanych (pokazane poniżej). Plastikowa obejma ma 2 mm wysokości, dzięki czemu uzyskano równomierne odstępy pomiędzy płytkami



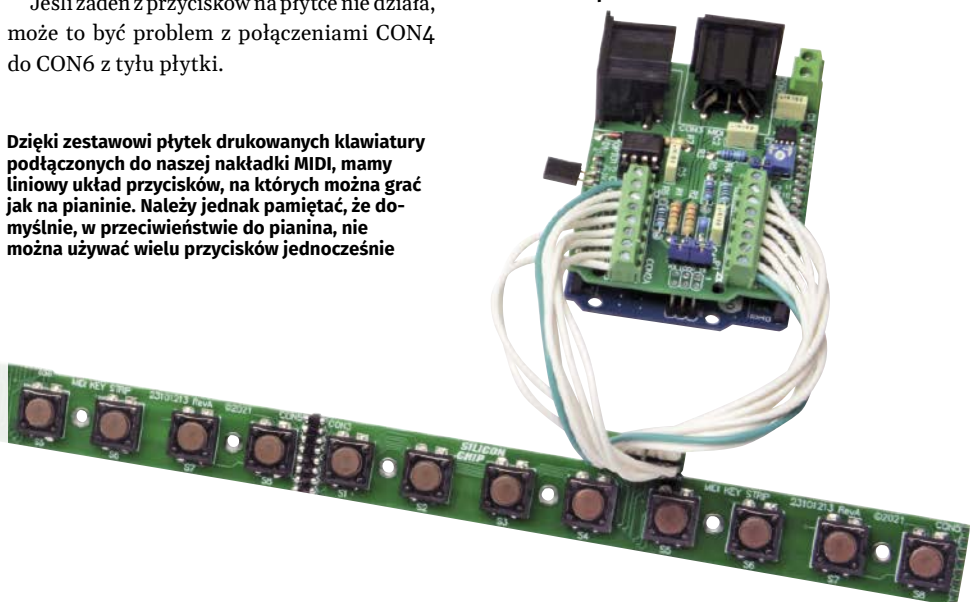
Podsumowanie

Podobnie jak matryca 8x8, klawiatura liniowa została zaprojektowana do współpracy ze sprzętem i oprogramowaniem MIDI opracowanym przez Redakcję Silicon Chip-a. Uważamy jednak, że Czytelnicy znajdą inne zastosowania, zwłaszcza w przypadkach, gdy do mikroprocesora musi być podłączonych wiele przycisków. ■

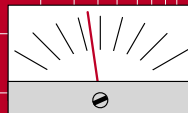
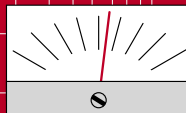
Tim Blythman

Adaptacja do wydania polskiego – Andrzej Nowicki

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



AUDIO OUT



Budujemy radio tranzystorowe, część 1

Prawie 50 lat temu próbowałem zbudować radio. Byłem młody, niedoświadczony... i odbiornik nie zadziałał. Nie byłem jedynym, który próbował. Tak naprawdę to musi być jeden z najpopularniejszych projektów radiowych, jakie kiedykolwiek wydrukowano, ale dziwne jest to, że nie ukazał się w żadnym zwykłym magazynie radiowym ani elektronicznym – nie był nawet skierowany do dorosłych.

Pojawił się w niezwykle popularnej i wpływowej serii książek dla dzieci zatytułowanej „Biedronka”. (Jeśli masz trochę wolnego czasu, możesz spędzić bardzo przyjemną (nieelektroniczną) godzinę oglądając dokument prezentujący tę serię: <http://bit.ly/pe-mar21-lb>).

48 lat, żeby to zadziałało

Przez te wszystkie lata nie udało mi się ukończyć projektu na wczesnym etapie, co mnie dręczyło, ale niniejszy projekt to naprawił! Tak naprawdę, to nie tylko zbudujemy kompletną, oryginalną wersję *Ladybird Radio*, ale mam także do zaseruowania kilka ulepszeń. Najpierw jednak trochę tła.

Specyficzna kultura poważnego amatorstwa w Wielkiej Brytanii sprawia, że twórczość artystyczna i technologiczna kraju przewyższa jego wielkość ekonomiczną. Może było tak dlatego, że była to kolebka rewolucji

przemysłowej, ale podejrzewam, że głównym czynnikiem była mnogość dostępnych magazynów i książek poświęconych elektronice. Wśród nich znalazła się seria ilustrowanych książek dla dzieci „Biedronka” (*Ladybird*), a część mojej kolekcji pokazano na **rysunku 1**.

Życie zaczyna się i kończy

Podejrzewam, że dla wielu z nas kariera inżynierska lub naukowa tak naprawdę zaczyna się w wieku około dziesięciu lat. Dla mnie było to rok 1972, rok, w którym ukazała się książka wielkiego George’a Dobbsa „*Ladybird Book Making a Transistor Radio*” – pokazana na **rysunku 2**. Pomogła mi ona nawet nauczyć się czytać, ponieważ nic innego w szkole nie przyciągało mojego zainteresowania.

Książka ta na jakiś czas znalazła się w pierwszej dziesiątce literatury dziecięcej i wywarła wpływ na całe pokolenie młodych ludzi. Niestety, nigdy nie spotkałem samego George’a Dobbsa. Miałem nadzieję przeprowadzić z nim wywiad na temat jego książki podczas Złotu Welsh Radio 2018 w Newport w październiku 2018 r., ale był chory i odwołał wykład. Zmarł wkrótce potem, w marcu 2019 r., w wieku 75 lat (jednak magazyn „*Practical Wireless*” przeprowadził z nim wywiad w numerze z czerwca 2009 r.).

Stare sklepy

Na szczęście odkryłem, że jeden z moich dostawców komponentów, John Birkett z Lincoln, znał George’a jako przyjaciela rodziny od ponad 50 lat. Sklep Johna powstał w 1962 roku i jest jednym z ostatnich istniejących

sklepów stacjonarnych, obok kilku innych, takich jak Cricklewood Electronics i JPR. Nadal ma trochę książek George’a Dobbsa i dostarcza wszystkie stare komponenty dla tych, którzy chcą zbudować radio Biedronka oraz zabytkowy sprzęt elektroniczny. John ma już dziewięćdziesiątkę, więc dziś sklep prowadzi głównie jego córka Judy, którą ochrzcił, nawiasem mówiąc, Wielebny Dobbs. Wywiad z obojgiem znajduje się w wydaniu „*Practical Wireless*” ze stycznia 2020 roku.

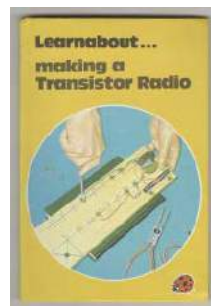
Dzisiaj znalazłem kopie

Nakład „Budujemy radio tranzystorowe” już dawno się wyczerpał, ale możesz mieć szczęście na targach butów lub w sklepie charytatywnym i zapłacić za tę książeczkę 50 pensów. Oczywiście użytkownicy serwisu eBay mogli dotrzeć tam przed Tobą. W Internecie czyste, wczesne wydanie może kosztować do 20 funtów, ale średni koszt wynosi około 5...8 funtów, mimo że na odwrocie często widnieje informacja „15 pensów”! Trochę poszukiwań w Google prawdopodobnie umożliwi ci znalezienie wersji online.

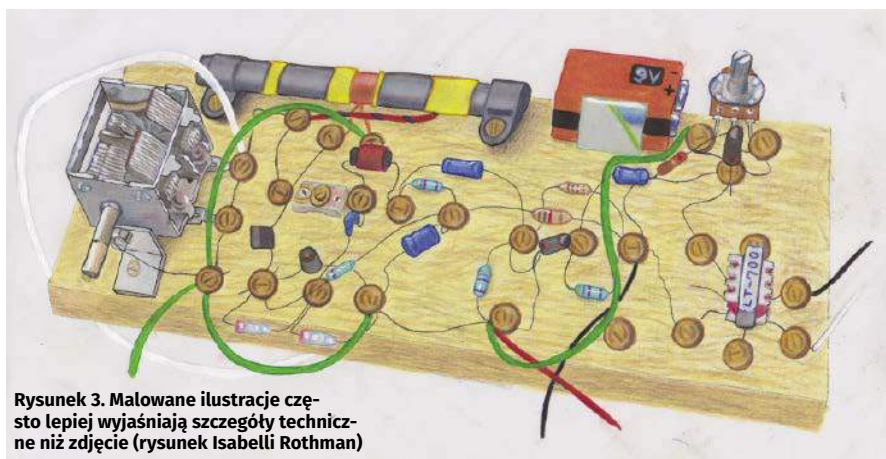
Od red. EdW: Kopię elektroniczną książki „Budujemy radio tranzystorowe”, można znaleźć na stronie <https://>



Rysunek 1. Mała próbka słynnych książek z serii „Biedronka”



Rysunek 2. „Budowanie radia tranzystorowego” – wywarła wpływ na wielu starszych czytelników PE?



Rysunek 3. Malowane ilustracje często lepiej wyjaśniają szczegóły techniczne niż zdjęcie (rysunek Isabelli Rothman)

archive.org/details/MakingATransistorRadio-LadybirdBook.

Spójrz mamó! – Żadnego komputera!

W serii „Biedronka” było wiele znakomitych książek technicznych dla początkujących. Książki były „Internetem” mojego dzieciństwa. Wprowadziły mój „plastyczny” umysł we wszystkie formy technologii.

Migrowałem z książek „Biedronka” w kierunku kolorowych książek w miękkiej oprawie wydawnictwa Hamlyn Publishing, ponieważ miały one raczej charakter wizualny niż algebracyjny, jak większość książek technicznych. Ulubionym tematem była Electronics, napisana przez Rolanda Worcestera (pseudonim FG Rayera, który napisał wiele artykułów do magazynów elektronicznych z lat 70. XX wieku). Miała fantastyczne kolorowe schematy obwodów, a ponieważ nikt nie kolekcjonuje tych książek, chodzą one po cenach śmieciowych (daję je studentom). Następnie zacząłem lutować w wieku 11 lat, usiadłem na podłodze i wycierałem żelazko Antex o biały dywan mojej mamy. (Błąd, który doradzono mi tylko raz!).

Sztuka ilustracji elektroniki

Podobnie jak w książkach o botanice i anatomii, w seriach „Biedronka” i „Hamlyn”



Rysunek 4. Technika łączenia śrubowo/podkładkowego pomysłu George'a Dobbsa

koncepty i konstrukcje zostały wyjaśnione ilustracjami, a nie zdjęciami. Są wyraźniejsze niż zdjęcia, ponieważ można podkreślić główne punkty i pominąć zbędne informacje. Jedną z moich sąsiadek była ilustratorką serii „Biedronka” p.t. „Piotruś i Jane”, niestety to nie ona robiła ilustracje książeczek „radiowych”, tu akurat był to Bernard Robinson – zobacz: <http://bit.ly/pe-mar21-lb1>

Bernard był moim ulubionym ilustratorem „Biedronki”, a szczególnie interesował się muzyką i technologią. Nie możemy tutaj użyć ilustracji z „Księgi Biedronki”, ponieważ wydawnictwo Pingwin jest właścicielem praw autorskich, ale moja córka Isabella narysowała podobny przykład na **rysunku 3**. Oby długo kontynuowała tradycję!

Prawdziwe „płytki stykowe”

Nawet w „naddźwiękowych latach siedemdziesiątych” wydawcy uważali lutowanie za zbyt niebezpieczne dla dzieci, więc George wymyślił unikalny system konstrukcyjny wykorzystujący mosiężne śruby i podkładki. Wiązało się to z zaciśnięciem przewodów na kawałku drewna w celu wykonania połączeń, w wyniku czego powstała prawdziwa płytka prototypowa, jak pokazano na **rysunku 4**. Był jeden problem z tą techniką: trzeba było być wystarczająco silnym fizycznie, aby wykonać stolarkę. Przy pierwszych próbach nic nie osiągnąłem, bo użyłem zbyt twardego drewna – aby dziecko mogło wkręcić wkręty w drewno, musi to być drewno iglaste, np. sosna. Wiercenie wgłębień pod nakrętki doniczkowe za pomocą ręcznego świda również było trudne; Udało mi się przejść całość już przy pierwszej próbie. Gdybym tylko miał wtedy wkrętarke/wiertarkę Makita! Otwory prowadzące o średnicy 2 mm są znacznie lepsze niż sztydło czy przebijak.

Następnie poprosiłem mamę, aby kupiła mi książki o stolarce i metalu z serii „Ladybird” z księgarni w Manchesterze,

w której pracowała. Ja z kolei korzystałem z tych książek, aby uczyć tych samych umiejętności moje dzieci.

Dziwne jest to, że po całym dniu stosowania techniki zakręcania przyzwyczałem się do niej. Używam jej teraz w obszarach prototypowania bez lutowania, takich jak łączenie dużych komponentów z nowoczesnymi płytkami prototypowymi i pasywnymi zwrotnicami głośników. Śruby wymagają jednak okresowego dokręcania. Ponadto połączenie drutów o różnych średnicach pod tą samą podkładką, może spowodować poluzowanie najcieńszego drutu, na przykład końcówki od tranzystora.

Drewno przewodzące

Kiedy już dotrzemy do etapu budowy, uważajmy na lekko zielonkawą kawałki drewna używane do ogrodzeń i tarasów. Są one tanalizowane, aby zapobiec gniciu. Jest to proces impregnacji drewna związkami miedzi. Do 2003 r. w USA stosowano nawet arsen. Oznacza to, że drewno nie nadaje się do „płytek stykowych”, ponieważ pod wpływem wilgoci słabo przewodzi prąd elektryczny. Rezystancja może spaść do około 250 kΩ pomiędzy sondami, jak pokazano na **rysunku 5**. Przy dużej powierzchni techniki „płytki stykowa/śrubowa” „Ladybird” jest ona znacznie niższa. Na szczęście obwody tranzystorów germanowych mają na ogół niską impedancję, więc większość (nieimpregnowanego) drewna jest w porządku. Użyłem sklejkę o grubości 9 mm.

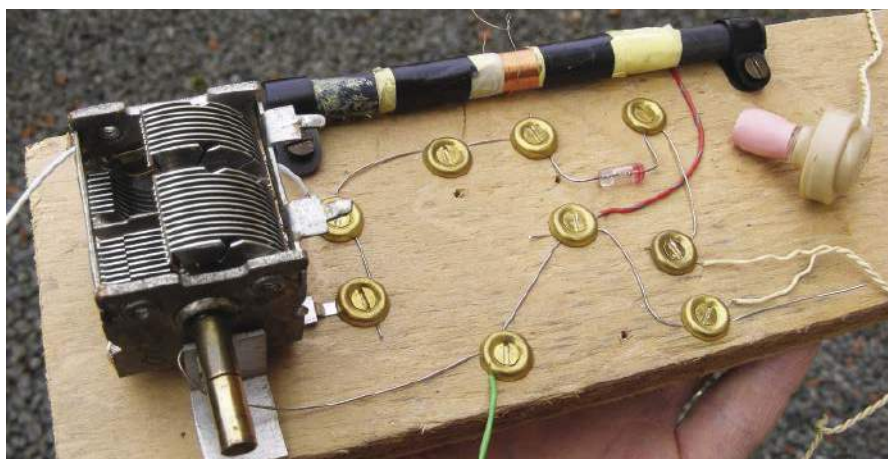
Podaję, że dzisiaj nowa technika konstrukcyjna dla dzieci wykorzystywałaby łączówki śrubowe. Brakowałoby tu jednak wyjątkowej przejrzystości wizualnej metody „Biedronka”.

Zrób to... w końcu

Z przykrością muszę stwierdzić, że moje pierwsze radio „Biedronka” nigdy nie zadziało. Jednak to mnie nie zniechęciło, bo po prostu wiedziałem, że elektronika jest tym, co chcę robić. Udało mi się uruchomić drewniany multiwibrator z drugiej książki Dobbsa z serii „Biedronka”, pt. „Dowiedz się o... prostej



Rysunek 5. Tanalizowane (impregnowane) drewno może być przewodzące – nie używaj go w tym projekcie



Rysunek 6. Najprostsze radio – odbiornik kryształkowy i pierwszy projekt w książce z rysunku 2. Zwróć uwagę na szklaną diodę OA70, pośrodku, po prawej stronie

elektronice”. Nadal używam tego obwodu, aby przedstawić moim uczniom elektronikę.

Wiem, że wszyscy mamy „projekty szkolne w naszych szafach”; ale ten, spóźniony o 48 lat, był moim najgorszym, więc zdecydowałem się jeszcze raz podejść do problemu radia, żeby zobaczyć, co poszło nie tak. Czy to było „złe” drewno – czy ja? Zacząłem od pierwszego obwodu w książce, radia kryształkowego.

Radio kryształkowe

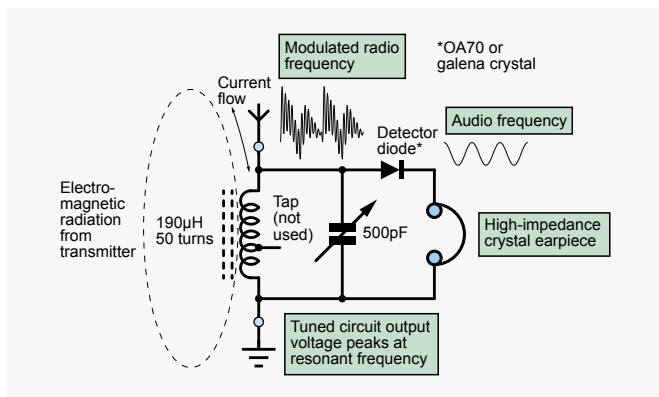
Najprostszą konstrukcją radiową jest odbiornik kryształkowy pokazany na **rysunku 6**. Jest on całkowicie pasywny i nie wymaga baterii, wykorzystując samą energię fal radiowych. Oczywiście odbiorniki kryształkowe odbierają tylko „staromodne” transmisje z modulacją amplitudy (AM), takie jak BBC Radio 4 na częstotliwości 198 kHz ([w Polsce jest to Program 1, na częstotliwości 225 kHz](#)). Mieszkam w pobliżu nadajnika fal średnich nad rzeką Ithon w Llandrindod Wells w środkowej Walii, który zapewnia silny sygnał za pomocą anteny zewnętrznej (muszę wypróbować ramkę antenową nawijaną na parasolkę jak w japońskim radiu kryształkowym Kasa: www.kasradio.com/en/index.html). Niestety, jedyne, co udało

mi się uzyskać w warsztacie, to szum zasilacza impulsowego, przez które nasz świat RF stał się tak zanieczyszczony. Nawet na zewnątrz słyszałem tylko Radio 5 Live. Podejrzewam, że zestaw kryształków stał się obecnie bezużytecznym obiektem historycznym. Podejrzewam również, że wrażliwość ludzkiego słuchu spadła z powodu hałasu ulicznego i korzystania z telefonów. Przy mocy wyjściowej sygnału odbiornika rzędu mikrowatów większość dzisiejszych ludzi po prostu go nie słyszy.

Obwód pokazany na **rysunku 7** nie może być prostszy. Wszystko, czego potrzeba, to równoległy obwód strojony i detektor do demodulacji kształtu fali. Osiąga się to poprzez prostowanie sygnału; wynikowy poziom prądu stałego odzwierciedla modulację lub obwiednię.

Aby uzyskać maksymalną moc do zasilania detektora kryształkowego, konieczna jest długa, wysoka antena i dobre uziemienie. Jeśli którekolwiek z nich będzie niewystarczające, zestaw w ogóle nie będzie działał.

Pamiętaj, że sam pręt ferrytowy nie ma szans zebrać wystarczającej ilości sygnału dla detektora – naprawdę potrzebujesz długiego, dobrze izolowanego przewodnika i dobrego uziemienia, aby zamknąć obwód.



Rysunek 7. Schemat odbiornika kryształkowego „Biedronka”

Jako dziecko wykorzystałem metalową ramę łóżka jako antenę. Tym razem zdecydowałem się na wykorzystanie suszarki do prania wykonanej z drutu i miedzianej rury wodnej jako uziemienia, jak pokazano na **rysunku 8**. Kolejnym udoskonaleniem wartym wypróbowania jest wykonanie zawieszenia anteny z izolatorów szklanych (**rysunek 9**).

Słuchawki

Należy używać słuchawek ze zrównoważoną armaturą o impedancji 4 kΩ, które charakteryzują się doskonałą czułością. Dziś bardzo trudno je dostać. Można także zastosować słuchawkę piezoelektryczną (**rysunek 10**), co jest kolejną najlepszą rzeczą i łatwiejszą do kupienia. To brzydki element, wyglądający jak aparat słuchowy z lat pięćdziesiątych. Jednakże jego przetwornik piezoelektryczny ma wysoką czułość i wysoką impedancję. Do tego stopnia, że usłyszałem kliknięcie, dotykając jednego przewodu na pierścionku na palcu (drugi trzymałem w drugiej dłoni). Irytujące było to, że nie mogłem utrzymać tego przedmiotu w uchu bez nałożenia na niego kawałka różowego silikonu. Ponadto odkryłem, że jestem całkowicie głuchy na bardzo niski poziom dźwięku w prawym uchu.

Najwyraźniej istnieje bardziej nowoczesny typ, który wykorzystuje piezoelektryczne dyski ceramiczne. Przetworniki te działają również jako kondensator wyglądający dla demodulatora, mając pojemność około 15 nF. Wypróbowałem także parę studyjnych słuchawek z ruchomą cewką, Beyer Dynamic DT100 o impedancji 2 kΩ i okazały się zaskakująco nieczułe i mało przydatne.

Wzięty do niewoli

Radia kryształkowe mają fascynującą historię jako podstawa odbiorników jenieckich. Większość z nich opierała się na kryształach detekcyjnych i fragmentach skradzionych ze starych telefonów. Jedynym wymaganym „prawdziwym” elementem elektronicznym była

Rysunek 8. Do anteny radia kryształkowego potrzebne jest dobre uziemienie. Idealne są miedziane rury wodociągowe prowadzone w ziemi. (Uwaga, obecnie w wodociągach często stosuje się niebieskie rury plastikowe). Włóż stalowy pręt wzmacniający do miedzianej rury uziemiającej, jeśli zamierzasz wbić ją w twardą glebę. Jako dziecko używałem rury od kaloryfera





Rysunek 9. Dobra antena to podstawa. Pozioma pętla wokół szyny do wieszania obrazów w wiktoriańskim domu wystarczy, aby odbiornik kryształkowy działał, ale aby uzyskać najlepsze rezultaty, należy zawiesić ją na zewnątrz tak wysoko i długą, jak to możliwe, za pomocą szklanych izolatorów. Znalazłem ten izolator na poboczu nieczynnej linii kolejowej!

słuchawka telefonu, a ich wersje jenieckie były możliwe do wykonania, jeśli można było znaleźć magnesy i cienki drut miedziany. Kondensator strojeniowy może być parą blaszanych puszek wsuwanych i wysuwanych. Detektorem był często kryształ galeny, zasadniczo ruda ołowiu lub siarczeka ołowiu, naturalny półprzewodnik. Minerale ten, pokazany na **rysunku 11**, można znaleźć w ziemi w moim rodzinnym hrabstwie Derbyshire; moi przyjaciele i ja próbowaliśmy zrobić z niego detektory.

George Dobbs ma w swojej książce rozdział poświęcony odbiornikom jenieckim i zaleca zażywanie kawałka koksu; Nie powiodło mi się, gdy spróbowałem koksu ze sterty na szkolnym placu zabaw (to było jeszcze przed ogrzewaniem na gaz ziemny i każda szkoła miała mnóstwo paliwa koksowniczego).

Te stare detektory były urządzeniami typu „punktowego”, prekursorami tranzystora. Punkt styku był krytyczny i konieczne było wiele szturchania, aby uzyskać prostowniczą granicę kryształu. Dobrze sprawdzi się ostry ołówek grafitowy lub zaostrożony drut sprężysty; utleniony drut w ogóle nie działa. Po znalezieniu odpowiedniego miejsca należy je dokładnie i mocno przytrzymać. Drut można zwinąć w sprężynę, jak pokazano na **rysunku 12**.



Rysunek 11. Kryształ Galeny (po prawej) w połączeniu z kwarcem – marzenie elektronicznego hipisa



Rysunek 10. Stuchawki piezoelektryczne – wyglądają okropnie i brzmią jeszcze gorzej, ale reagują już na kilka mV

Co ciekawe, odkryłem, że kryształ galeny pokazany na rysunku 11 miał bardzo niską rezystancję, około 4Ω , gdy był sondowany na nowych błyszczących powierzchniach. Będąc naturalnym minerałem, jego właściwości elektryczne są bardzo zmienne. Każda konkretna próbka będzie zawierać różne zanieczyszczenia. Powstałe działanie domieszkujące może zmienić półprzewodnik z typu P na N, a czasami po prostu na przewodnik. Istnieje możliwość zakupu wybranych kryształów galeny „radiowej”. Są one zazwyczaj montowane w mtseczce ze stopu Wooda. Wszystkie kryształy mogą wykazywać działanie diodowe tylko w określonych pozycjach styku, takich jak obszary utlenione lub granice kryształów. Działanie diodowe jest często bardzo słabe. Mierząc spadek napięcia w kierunku przewodzenia za pomocą funkcji testu diody multimetru, miałem 200 mV w jedną stronę i 300 mV w drugą. Zapewnia to pewien stopień prostowania, a wysoki upływ wsteczny zapobiega pełnemu naładowaniu pojemnościowej, kryształowej słuchawki, eliminując potrzebę stosowania rezystora obciążającego.

Prawdziwe diody

Aby obejść ten problem, opracowano germanową diodę małosygnałową. Są to jedyne nadal produkowane elementy półprzewodnikowe z kontaktem punktowym. Były także jednymi z pierwszych elementów półprzewodnikowych produkowanych masowo, niezbędnych podczas II wojny światowej. Zwykle są one zamknięte w szkło, dzięki czemu kryształ i zaostrożony drut są wyraźnie widoczne, jak pokazano na **rysunku 13**. Ten

typ był wczesnym urządzeniem STC. Podobnie wyglądające OA70 zastosowano w obwodzie pokazanym na rysunkach 6 i 7. Stare diody Siemens OA85 nadal są dostępne, ale są pokryte czarną farbą, aby zatrzymać światło zwiększające prąd upływu. W projekcie „Ladybird” wykorzystano OA81, który był bardzo popularny w brytyjskich radiach lat 60. XX wieku. Wszystkie te duże, stare diody mają obudowę SO-15.

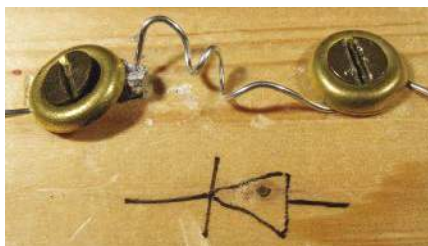
Dostępnych jest kilka fizycznie mniejszych typów, takich jak OA91, 1N34A i CG92, które dobrze wykonają to zadanie. OA91 to OA81 w obudowie DO-7. Jeśli nie możesz uzyskać diody germanowej, typ Schottky’ego, taki jak 1N60, ZC5800 lub BAT42, może działać prawie równie dobrze w radiach kryształkowych. W końcu można by argumentować, że detektor galenowy, z jego złączem typu metal-półprzewodnik jest formą diody Schottky’ego. Normalne diody krzemowe, takie jak 1N4001 lub 1N4148, nie będą działać ze względu na wysokie napięcie przewodzenia i ostre kolano, chociaż 0,5 V polaryzacji przewodzenia z akumulatora sprawi, że będą działać. Strona radiowa Kasa sugeruje użycie diody LED i wystawienie jej na działanie światła w celu wygenerowania własnego napięcia polaryzacji.

W następnym miesiącu

W kolejnej części dojdziemy do sedna książki i zbudujemy prawdziwe radio tranzystorowe. ■

Jake Rothman

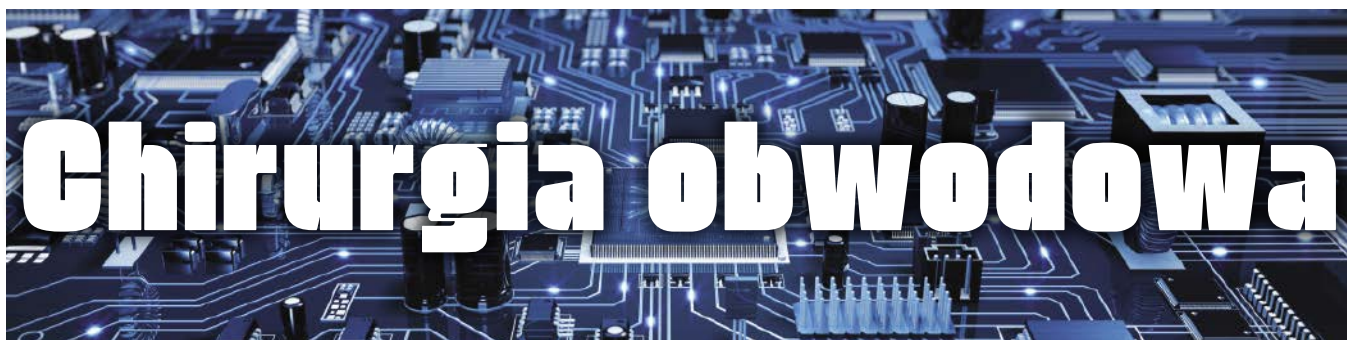
Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, marzec 2021 (www.epemag3.com)



Rysunek 12. Detektor wykorzystujący kryształ galeny, wzorowany na projekcie książki z rysunku 2. Zwróć uwagę na sprężysty, spiczasty drut, czasami nazywany „kocim wąsem”



Rysunek 13. Stara germanowa dioda ostrzowa. Zwróć uwagę na płytkę germanu i sprężynowy drut ostrzowy



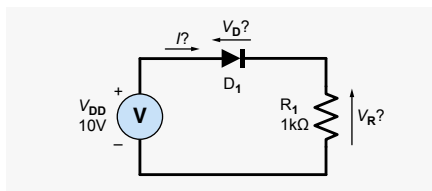
Chirurgia obwodowa

Rozwiązywanie obwodu: iteracyjne proste obciążenia i symulacja

Prosta obciążenia, to graficzny sposób rozwiązywania problemu zachowania się obwodu bez konieczności przeprowadzania jego szczegółowej analizy matematycznej. W tym artykule przyjrzymy się prostym obciążeniom w kontekście różnych sposobów rozwiązywania obwodów nieliniowych – w szczególności obwodu dioda-rezystor pokazanego na rysunku 1.

Jest to obwód prosty – dioda spolaryzowana w kierunku przewodzenia, czyli dioda, przez którą prąd płynie w kierunku wskazanym przez kształt strzałki symbolu diody. Zakładając, że znamy podstawową teorię obwodów, zależności napięcie-prąd diody i rezystora oraz wartości parametrów charakterystycznych diody, możemy założyć, że łatwo będzie znaleźć równanie prądu w obwodzie przedstawionym na rysunku 1. Gdyby zamiast diody i rezystora były dwa rezystory, byłoby to łatwe dzięki podstawowej teorii obwodów; okazuje się jednak, że „rozwiązanie” obwodu diodowego nie jest takie proste.

Zanim przejdziemy do szczegółów obwodu z diodą, przejrzymy potrzebną nam teorię obwodów i rozwiążemy obwód z dwoma rezystorami jako punkt odniesienia. Dwie z kluczowych zasad analizy obwodów znane są jako prawa Kirchhoffa (pochodzące z 1854 r.). Jedno z praw stanowi, że „suma prądów wpływających do węzła w obwodzie wynosi zero (lub całkowity prąd wejściowy = całkowity prąd wyjściowy)”. Oznacza to również, że prąd płynący przez zestaw elementów połączonych szeregowo jest taki sam w każdym elemencie (prądy diody i rezystora są równe w obwodzie na rysunku 1). Drugie prawo stanowi, że „suma napięć wokół dowolnej pętli w obwodzie wynosi zero”.



Rysunek 1. Układ dioda-rezystor

Prawa obwodowe i modele elementów

Prawa Kirchhoffa dotyczą zarówno napięcia, jak i prądu. Aby w pełni przeanalizować obwód, musimy znać zależności między prądami i napięciami w obwodzie. Zależności te opisują równania charakterystyczne poszczególnych składników. W naszym przykładzie charakterystyczne równanie rezystora to $U=IR$, co wielu czytelników rozpozna jako prawo Ohma. Na przykład dla rysunku 1 możemy napisać: $U_R=IR_1$.

Prawo Ohma, a co za tym idzie także charakterystyczne równanie rezystora, zostało po raz pierwszy uzyskane eksperymentalnie przez Georga Ohma w latach dwudziestych XIX wieku. Można je również wyprowadzić z podstawowej fizyki przewodników, chociaż historycznie stało się to znacznie później. W przypadku prawdziwego rezystora prawo Ohma jest tylko przybliżeniem, rezystancja może w rzeczywistości zmieniać się wraz z przyłożonym napięciem i prawdopodobnie będzie się zmieniać wraz z temperaturą, w tym również spowodowaną samonagrzewaniem rezystora. Ponadto rezystory generują szum elektryczny, a prawo Ohma dotyczy tylko średniego prądu. Matematyczny sposób, w jaki wybieramy reprezentację komponentu, nazywany jest „modelem” tego komponentu.

W większości przypadków w przypadku elementów rezystancyjnych używanych w obwodach możemy po prostu użyć zależności $U=IR$ w analizie obwodów. Jeśli zdecydujemy się uwzględnić w naszych obliczeniach inne czynniki (i zrobimy to poprawnie), otrzymamy dokładniejsze wyniki kosztem większej trudności i złożoności obliczeniowej. W przypadku obliczeń ręcznych zbyt duża złożoność może sprawić,

że proces będzie podatny na błędy lub nawet trudny do wykonania, a w przypadku analizy komputerowej (symulacja obwodów) czasy wykonywania wydłużają się dla bardziej złożonych modeli. Jak zawsze w inżynierii, istnieją kompromisy – w tym przypadku pomiędzy dokładnością a złożonością.

W równaniu charakterystycznym rezystora $U=IR$ (rezystancja) jest ogólnie określane jako „parametr” (lub parametr modelu). Wszystkie równania charakterystyki składowej będą miały jeden lub więcej parametrów. Możemy analizować obwód algebraicznie, używając symboli parametrów (np. $R_1, R_2, \dots$ dla rezystorów), ale jeśli potrzebujemy znać rzeczywiste napięcia i prądy, musimy podać określone wartości parametrów dla każdego elementu.

Równanie diody

Aby uwzględnić elementy półprzewodnikowe, takie jak dioda na rysunku 1, w analizie obwodu potrzebny jest model; czyli równania charakterystyczne do analizy ogólnej oraz wartości parametrów rzeczywistych do konkretnych obliczeń. Możemy wykorzystać wiedzę z zakresu fizyki półprzewodników, aby wyprowadzić równanie charakterystyczne dla diody – podobnie jak prawa Kirchhoffa i Ohma, jest to dobrze ugruntowana teoria. Dla wyidealizowanego przypadku otrzymujemy, co następuje, łącząc prąd płynący przez diodę (I_D) z napięciem na niej (U_D):

$$I_D = I_S \left(\exp\left(\frac{U_D}{U_T}\right) - 1 \right)$$

Nazywa się to „równaniem Shockleya dla diody idealnej”. U_T nazywa się napięciem termicznym. Wynosi ono około 26 mV w temperaturze pokojowej i jest wyrażane przez kT/e , gdzie T jest temperaturą

złącza (w stopniach Kelvina), a e i k to stałe fizyczne, odpowiednio ładunek elektronu i stała Boltzmanna. Kluczowym parametrem dla poszczególnych diod jest I_s – prąd nasycenia wstecznego – który może zmieniać się o rzędy wielkości w zależności od elementu i temperatury, np. 10^{-14} do 10^{-10} A, dla diod krzemowych (domyślnie w LTspice jest to 10^{-14} A).

Wykładniczy charakter zależności napięcie-prąd oznacza, że przy niskich napięciach przewodzenia płynie bardzo mały prąd, ale szybko rośnie wraz z przyłożonym napięciem, tak że dioda „włącza się” i przewodzi przy napięciu około 0,6 do 0,7 V. W przypadku przewodzenia na poziomach „włączonych” lub w ich pobliżu składnik wykładniczy jest znacznie większy niż 1 i możemy uprościć równanie diody do:

$$I_D = I_s \exp\left(\frac{U_D}{U_T}\right)$$

Jeśli znamy I_D i chcemy znaleźć U_D , możemy zmienić to równanie, biorąc logarytmy naturalne ($\ln$ – odwrotność funkcji wykładniczej), aby otrzymać:

$$U_D = U_T \ln\left(\frac{I_D}{I_s}\right)$$

Obwód z dwoma rezystorami

Jak wspomniano powyżej, zanim zajmujemy się obwodem dioda-rezystor, przeanalizujemy obwód z dwoma rezystorami (**rysunek 2**). Następnie spróbujemy zastosować tę samą procedurę dla obwodu z diodą. Aby przeanalizować obwód z rysunku 2, możemy postępować w następujący sposób:

Prąd w R_1 jest określony przez: $I = U_{R1}/R_1$

Zatem $U_{R1} = IR_1$

Prąd w R_2 (w funkcji U_{R1} i U_{DD}) jest określony wzorem:

$$I = \frac{U_{DD} - U_{R1}}{R_2}$$

Wyeliminuj U_{R1}

$$I = \frac{U_{DD} - IR_1}{R_2}$$

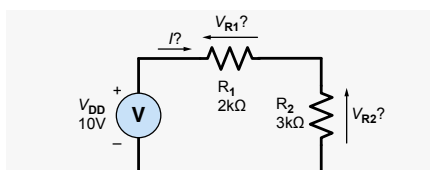
Rozwiąż dla I

$$I = \frac{U_{DD}}{R_1 + R_2}$$

Używając wartości schematowych: $I = 10 \text{ V} / (2 \text{ k}\Omega + 3 \text{ k}\Omega) = 2 \text{ mA}$. Używając $U = IR$, otrzymujemy również $U_{R1} = 4 \text{ V}$ i $U_{R2} = 6 \text{ V}$.

Obwód dioda-rezystor

Spróbujmy teraz wykonać tę samą procedurę dla obwodu rezystor-dioda na rysunku 1. Prąd



Rysunek 2. Układ z dwoma rezystorami

w diodzie wyraża się wzorem: $I = I_s \exp(U_D/U_T)$, więc $U_D = U_T \ln(I/I_s)$. Zatem prąd w R (w kategoriach U_D i U_{DD}) wynosi:

$$I = \frac{U_{DD} - U_D}{R}$$

Wyeliminuj U_D

$$I = \frac{U_{DD} - U_T \ln\left(\frac{I}{I_s}\right)}{R}$$

Następnym krokiem byłoby rozwiązanie równania dla I – przekształcenie równania tak, aby I było podmiotem, ale niestety nie możemy tego zrobić łatwo. Obwody diodowe wytwarzają równania przestępne, które na ogół są trudne lub niemożliwe do rozwiązania. Równanie transcendentalne to takie, które obejmuje funkcję transcendentalną, która obejmuje funkcje wykładnicze/logarytmy, które tu mamy. Funkcja transcendentalna to taka, której nie można wyrazić w kategoriach operacji algebraicznych – nazwa „transcendentalna” odnosi się do idei, że „przekracza” lub wykracza poza algebrę. W rzeczywistości rozwiązanie równania obwodu dioda-rezystor nie jest niemożliwe, ale wymaga bardzo zaawansowanej matematyki, w szczególności funkcji W Lamberta, z którą większość ludzi prawdopodobnie nigdy nie spotka się ani nie użyje (jeśli chcesz, zobacz stronę Wikipedii dotyczącą funkcji W. Lamberta aby wiedzieć więcej).

Praktyczna zasada

Jeśli nie możemy rozwiązać obwodu z rysunku 1, manipulując algebraicznie równaniami obwodu, co możemy zamiast tego zrobić, aby znaleźć prąd diody? Prostą odpowiedzią, powszechnie stosowaną w szybkich obliczeniach projektowanych obwodów, jest założenie, że napięcie na diodzie (spadek napięcia w kierunku przewodzenia) wynosi 0,7 V. Powszechnie wiadomo, że w przypadku typowych prądów roboczych w większości obwodów napięcie na standardowej diodzie krzemowej przewodzącej będzie prawdopodobnie mieścić się w zakresie od 0,6 do 0,8 V; dlatego często przyjmuje się, że 0,7 V jest wartością praktyczną. W przypadku innych typów diod, które mogą mieć różne napięcia (np. diod LED), spadek napięcia w kierunku przewodzenia jest często podawany w arkuszu danych, więc możemy zastosować to samo podejście. W obwodzie rezystor-dioda pokazanym na rysunku 1, jeśli napięcie zasilania nie jest zbyt bliskie oczekiwanemu napięciu diody, oczekiwany zakres napięć przewodzenia nie ma większego wpływu na prąd rezystora, więc możemy uzyskać rozsądne oszacowanie po prostu zakładając, że $U_D = 0,7 \text{ V}$. Jeśli to zrobimy, otrzymamy $U_{R1} = 9,3 \text{ V}$ i $I = U_{R1}/R_1 = 9,3/1000 = 9,3 \text{ mA}$ dla obwodu z rysunku 1.

Praktyczna zasada 0,7 V daje nam przybliżone rozwiązanie obwodu, ale może nie być wystarczająco dokładna. Co możemy zrobić, jeśli chcemy znaleźć dokładniejszą wartość? Wybór wartości 0,7 V, to w zasadzie świadome przypuszczenie, co dzieje się w obwodzie. Po takim przypuszczeniu możemy użyć równania diody, aby znaleźć odpowiedni prąd, a co za tym idzie, napięcie rezystora (U_{R1}) przy przepływającym przez niego prądzie. Następnie możemy dodać napięcia rezystora i diody, które powinny być równe napięciu zasilania (jeśli nie, prawo napięciowe Kirchhoffa nie jest spełnione).

Gra w zgadywanie

Jeśli nasze sugerowane napięcie zasilania jest błędne, możemy dokonać nowego, miejmy nadzieję lepszego, przypuszczenia – na przykład, jeśli obliczone przez nas sugerowane napięcie zasilania jest zbyt wysokie, możemy nieco zmniejszyć założone napięcie diody (np. oszacować 0,68 V zamiast 0,7 V) i zobacz, czy to poprawi sytuację. Możemy zgadywać i udoskonalać wartość napięcia diody, przeliczając prąd i implikowane U_{DD} i zgadując ponownie, coraz bardziej zbliżając się do prawidłowej wartości. Formalnie ten proces zgadywania nazywany jest iteracyjnym rozwiązywaniem obwodu. Oto przykład dla rysunku 1, w którym użyjemy $I_s = 1,0 \times 10^{-14} \text{ A}$ i $U_T = 25,85 \text{ mV}$ (dla temperatury 27°C):

Zgadnij: $U_D = 0,7 \text{ V}$, więc $I = I_s \exp(U_D/U_T) = 5,70 \text{ mA}$ (równanie diody)

Zatem: $U_{DD} = U_R + U_D = 6,400 \text{ V}$ – za nisko – więc zwiększ U_D , przypuszczając, do 0,725 V

Przy: $U_D = 0,725 \text{ V}$, $I_D = 4,986 \text{ mA}$ (z równania diody)

Zatem: $U_R + U_D = 15,711 \text{ V}$ – za wysokie – więc zmniejsz U_D , przypuszczając, do 0,7125 V

Przy: $U_D = 0,7125 \text{ V}$, $I_D = 9,242 \text{ mA}$ (z równania diody)

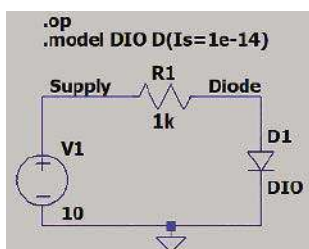
Zatem: $U_R + U_D = 9,955 \text{ V}$ po prostu za niskie – więc zwiększ U_{DD} do 0,7126 V

Przy: $U_D = 0,7126 \text{ V}$, $I_D = 9,278 \text{ mA}$ i $U_R + U_D = 9,990 \text{ V}$.

Ten końcowy wynik oznacza, że zasilanie wynosi 9,990 V, a nie 10 V, co stanowi błąd 0,1%. Moglibyśmy w dalszym ciągu dodawać więcej znaczących cyfr, ale musimy przerwać w pewnym momencie, gdy wynik jest „wystarczająco zbliżony” – sposób, w jaki decydujemy, który jest wystarczająco bliski, zależy od tego, czego potrzebujemy w zakresie dokładności.

Metody numeryczne

Proces zgadywania i ponownego obliczania, którego właśnie użyliśmy, to uproszczona, ręczna wersja tego, co robią symulatory, takie



Rysunek 3. Schemat LTspice dla obwodu na rysunku 1

jak LTspice, podczas rozwiązywania obwodu. Podejście do znalezienia następnego przybliżenia jest bardziej wyrafinowane niż podejście „wymyśl liczbę” sugerowane powyżej. Stosowane techniki matematyczne nazywane są „metodami numerycznymi” lub „analizą numeryczną” i obejmują raczej przybliżenie numeryczne niż manipulację symboliczną. Mają one długą historię. Na przykład metoda Newtona-Raphsona, którą można zastosować w symulatorach obwodów, wywodzi się z prac Izaaka Newtona i Josepha Raphsona z XVII wieku. Inne metody numeryczne są jeszcze starsze i sięgają najwcześniejszych zapisów matematycznych na babilońskich tabliczkach glinianych.

Podobnie jak w przypadku naszej ręcznej iteracji, symulatory obwodów muszą zdecydować, kiedy zatrzymać dane obliczenia – kiedy są one „wystarczająco blisko” – co w terminologii symulatorów określa się jako zbieżność. Symulatory SPICE mierzą zbieżność na kilka sposobów: różnicę między iteracjami i stopień, w jakim obecne prawo Kirchhoffa zostało spełnione. Jeśli różnica między iteracjami jest wystarczająco mała, oznacza to, że rozwiązanie zostało osiągnięte i nowe kroki nie będą znacząco przybliżone. W powyższej ręcznej iteracji sprawdziliśmy domniemane napięcie zasilania – skutecznie sprawdzając, czy spełniono prawo Kirchhoffa dotyczące napięcia – sprawdzenie bieżącego prawa przez SPICE jest podobne – jeśli wartości liczbowe prądu w obwodzie są w wystarczającym stopniu zgodne z prawem,

wówczas rozwiązanie prawdopodobnie będzie dobre. SPICE wykorzystuje szereg wartości numerycznych, aby zdecydować, kiedy nastąpiła zbieżność, na przykład abszol., tolerancja błędu bezwzględnego prądu i rełtol., tolerancja błędu względnego. Wartości te można zmienić za pomocą opcji użytkownika, na przykład w LTspice poprzez zakładkę SPICE w Panelu sterowania.

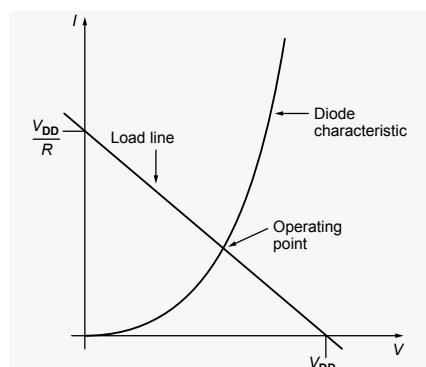
Przykład symulacji

Rysunek 3 przedstawia obwód z rysunku 1 narysowany jako schemat w LTspice i zawiera instrukcję modelu diody, aby ustawić parametr I_s na 1×10^{-14} A. Nie jest to ściśle wymagane, ponieważ jest takie samo jak domyślne, ale ilustruje jak można ustawić ten kluczowy parametr. Przeprowadzamy analizę punktu pracy w LTspice (dyrektywa .op) w celu uzyskania prądu i napięcia. LTspice symuluje przy domyślnej temperaturze 27°C, więc powinniśmy uzyskać wyniki podobne do naszych ręcznych obliczeń. Wyniki pokazane poniżej są zbliżone do naszych obliczeń.

```
--- Punkt pracy ---
V (dioda):      0,712763      napięcie
V (zasilanie):  10           napięcie
I(D1):          0,00928724    prąd_elementu
I(R1):          -0,00928724    prąd_elementu
I(V1):          -0,00928724    prąd_elementu
```

Proste obciążenia

Alternatywnym podejściem do obliczeń iteracyjnych, które nie wykorzystuje symulatora, jest wykreślenie wykresu prostej obciążenia. Jest to wykres charakterystyk diody i rezystora na tych samych osiach – punkt przecięcia obu wykresów jest rozwiązaniem obwodu, zwanym także punktem pracy. Termin „prosta obciążenia” wywodzi się z pomysłu, że mamy aktywny element (diodę lub tranzystor) sterujący sygnałem (wyjściem) do obciążenia (często rezystora). Jak widzieliśmy, obwodów takich jak ten na rysunku 1 nie można łatwo rozwiązać za pomocą algebry, linia obciążenia jest graficznym podejściem



Rysunek 4. Prosta obciążenia dla układu rezystora-dioda

do rozwiązywania obwodu, które nie wymaga pełnego rozwiązania algebraicznego.

Dla wykresu linii obciążenia obwodu dioda-rezystor (rysunek 1) na tym samym wykresie (patrz rysunek 4) nanosimy charakterystykę diody:

$$I_D = I_s \exp \frac{U_D}{U_T}$$

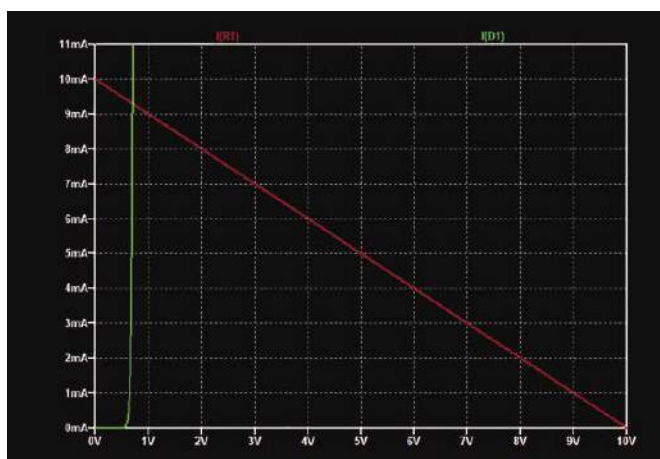
I prąd rezystora:

$$I = \frac{U_{DD} - U_D}{R}$$

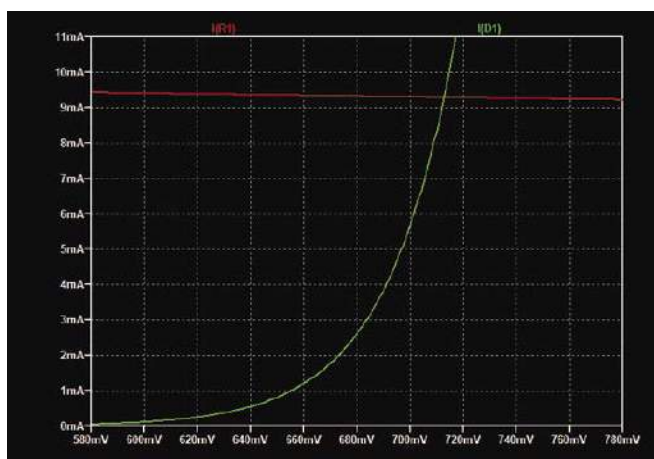
Wykres charakterystyki diody jest po prostu wykresem zależności prądu od napięcia diody. Wykres prądu rezystora – prosta obciążenia – jest linią prostą, ponieważ rezystor ma liniową zależność prąd-napięcie. Wykres ten jest odwrotnością samodzielnie wykreślanej zależności prądu od napięcia rezystora ponieważ jako oś używamy napięcia diody – zwiększenie U_D zmniejsza napięcie rezystora (razem dają U_{DD}), a tym samym zmniejsza prąd rezystora.

Prosta obciążenia jest łatwa do narysowania, ponieważ przecina osie w punktach $I=0$ i $U_D=0$, co można łatwo znaleźć z powyższego równania jako odpowiednio $U_D=U_{DD}$ i $I=U_{DD}/R$ (patrz rysunek 4).

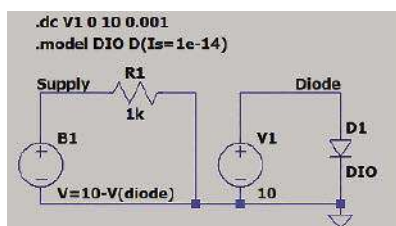
Ręczne wykreślenie prostej obciążenia na papierze milimetrowym, specjalnie w celu rozwiązania obwodu takiego jak na rysunku 1, jest prawdopodobnie zbyt dużym wysiłkiem, aby



Rysunek 5. Rzeczywista prosta obciążenia dla obwodu z rysunku 1



Rysunek 6. Prosta obciążenia z rysunku 5, powiększona, aby pokazać kształt charakterystyki diody



Rysunek 7. Obwód LTspice używany do wykreślenia prostej obciążenia, wykresy (rysunki 5 i 6)

było opłacalne, ale wykresy prostej obciążenia, takie jak z rysunku 4, są przydatne do wizualizacji zachowania się obwodów. Widzimy, jak zmiana wartości rezystora wpływa na punkt pracy i jak różne punkty pracy zmieniają wpływ małych zmian napięcia diody na prąd diody. Wykresy prostych obciążenia diody są często rysowane z pewną artystyczną swobodą, gdy są wykorzystywane do ogólnego omówienia zachowania obwodu – tak aby kształt krzywej diody można było zobaczyć razem z prostą pełnego obciążenia (np. rysunek 4). Aby uzyskać wykres wyglądający jak na rysunku 4, wymagane byłoby zasilanie około 0,8 V i rezystor około 80 Ω – niezbyt typowe wartości.

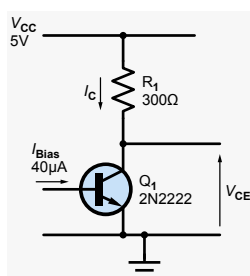
Proste obciążenia w LTspice

Rysunek 5 przedstawia rzeczywisty wykres prostej obciążenia dla obwodu z rysunku 1, uzyskany za pomocą symulatora LTspice. Nie jest zwyczajem wyznaczanie linii obciążenia za pomocą SPICE, ponieważ, jak widzieliśmy, symulator może bezpośrednio określić punkt pracy obwodu. Jednakże rysunek 5 ładnie ilustruje fakt, że dioda ma prawie stałe napięcie około 0,7 V w szerokim zakresie prądów roboczych, gdy jest wykreślone na osi napięcia obejmującej cały zakres zasilania (krzywa diody jest prawie linią pionową).

Rysunek 6 przedstawia te same wyniki wykreślone dla małego zakresu napięcia diody bli-

skiego 0,7 V. Tutaj wyraźniej widać kształt krzywej charakterystyki diody. Należy zauważyć, że prąd rezystora jest prawie stały w tym zakresie (linia pozioma). To ilustruje, że przy typowych napięciach zasilania kilkukrotnie przekraczających napięcie przewodzenia diody, prąd rezystora nie zmienia się zbytnio dla różnych napięć diody. To prowadzi nas z powrotem do punktu wyjścia naszej dyskusji na temat tego obwodu, gdzie po prostu założyliśmy, że spadek diody wynosi 0,7 V i w rezultacie otrzymaliśmy prąd o natężeniu 9,3 mA. Jeśli powiększymy bardziej i użyjemy kursora, okaże się, że punkt przecięcia krzywych jest ściśle zgodny z wynikami symulacji punktu pracy podanymi powyżej.

Wykres prostej obciążenia uwzględnia prądy komponentów niezależnie w całym zakresie zasilania, tak więc nie możemy bezpośrednio użyć obwodu z rysunku 3 do jego wykreślenia. Do uzyskania wykresów na **rysunkach 5 i 6** wykorzystano schemat LTspice z **rysunku 7**. Zastosowano liniową analizę przemiatania DC w celu stopniowania napięcia ze źródła napięcia U1 od 0 do 10 V w krokach co 0,001 V. Napięcie to przykłada się bezpośrednio na diodę, dzięki czemu możemy uzyskać prąd charakterystyczny. Małe kroki są przydatne, ponieważ zamierzamy przybliżyć stromą krzywą diody. Należy zauważyć, że symulując samą diodę w zakresie od 0 do 10 V, uzyskujemy niemożliwie duże prądy diody – symulator nie jest ograniczony maksymalnym rozpraszaniem mocy



Rysunek 9. Podstawowy układ ze wspólnym emitere tranzystora bipolarnego

ani możliwościami zasilania! W rzeczywistym obwodzie maksymalny prąd diody jest ograniczony przez rezystor.

Napięcie rezystora uzyskuje się poprzez odjęcie aktualnie

przyłożonego napięcia diody ($U(\text{dioda})$) od napięcia zasilania 10 V (w miarę przemiatania U1) przy użyciu behawioralnego źródła napięcia (B1). Źródła behawioralne pozwalają nam pisać równania na wyjściu źródła pod względem innych parametrów obwodu.

Proste obciążenia tranzystora

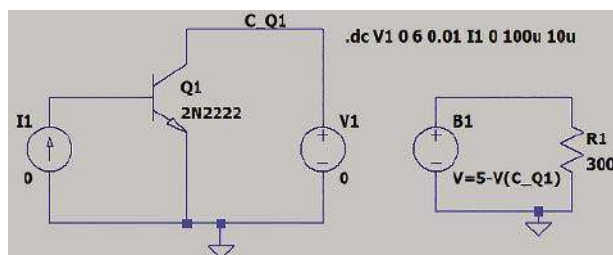
Aby nieco rozszerzyć temat prostych obciążenia, na **rysunku 8** przedstawiono prostą obciążenia dla podstawowego obwodu wspólnego emitiera z **rysunku 9**. Podobnie jak w przypadku obwodu diodowego, oddzielnie symulujemy charakterystykę tranzystora i rezystora obciążenia, aby uzyskać wykres prostej obciążenia – zastosowany obwód pokazano na **rysunku 10**. Zakładamy, że tranzystor został spolaryzowany prądem bazy 40 μA , ale nie określiliśmy, jak tego dokonano. Podobnie jak prosta obciążenia diody-rezystora, prosta obciążenia w tym obwodzie przecina osie przy napięciu zasilania (5 V) i prądzie, który płynąłby przez rezystor, gdyby był on podłączony do zasilania (5/300=16,7 mA).

Punktem pracy obwodu jest miejsce, w którym linia obciążenia przecina charakterystykę wyjściową tranzystora ($I_C=f(U_{CE})$) odpowiadająca polaryzacji bazy. Rysunek 8 pokazuje, że w tym przykładzie napięcie wyjściowe przy braku sygnału będzie wynosić około 2,6 V. Jeżeli tranzystor wzmacniał sygnał, wówczas prąd bazowy zmieniałby się odpowiednio. Prosta obciążenia pozwala w łatwy sposób sprawdzić, jak będzie się zmieniać napięcie wyjściowe – punkt pracy przesuwają się wzdłuż prostej obciążenia wraz ze zmianą prądu bazy. Na przykład, jeśli sygnał wejściowy spowodowałby zmianę prądu bazy od 30 do 50 μA , wówczas napięcie wyjściowe zmieniałoby się od około 2,1 do 3,2 V. ■

Ian Bell



Rysunek 8. Prosta obciążenia dla obwodu z tranzystorem bipolarnym z rysunku 9. Wykres przedstawia zmianę prądu kolektora (I_C) w funkcji napięcia kolektor-emitery (U_{CE}) przy różnych prądach bazy (I_B). Punkt pracy obwodu się przesuwa wzdłuż prostej obciążenia gdy zmienia się prąd bazy i wskazuje też sdpczynkowy punkt pracy (tj. bez sygnału)



Rysunek 10. Układ LTspice użyty do uzyskania prostej obciążenia z rysunku 8



Migające diody LED i śliniący się inżynierowie (5)

W poprzednim odcinku (EdW 01/2024) rozważaliśmy tradycyjną trójkolorową diodę LED zawierającą trzy diody LED RGB ze wspólną katodą. Na potrzeby tej dyskusji będziemy je nazywać „podręcznymi diodami LED”, aby przypomnieć, że znajdują się one wewnątrz pojedynczego elementu. Kiedy byłem młodszy, myślałem, że te elementy są ekstra fantastyczne. Szczerze mówiąc, z biegiem lat nieco się do nich przyzwyczaiłem, do tego stopnia, że zacząłem być nimi nieco rozzaczarowany.

Po pierwsze, każda taka dioda LED wymaga trzech cyfrowych pinów wejścia/wyjścia (I/O) w mikrokontrolerze (MCU) do jej obsługi. Po drugie, jeśli chcesz uzyskać dostęp do więcej niż ośmiu podstawowych kolorów zapewnianych przez włączanie i wyłączanie trzech podręcznych diod LED – czerwony, zielony, niebieski, żółty (czerwony + zielony), cyjan (zielony + niebieski), gorący różowy (czerwony + niebieski), biały (wszystkie włączone) i czarny (wszystkie wyłączone) – wtedy te piny muszą obsługiwać modulację szerokości impulsu (PWM), jak opisano w części 1 tego cyklu (EdW 10/2023). I po trzecie, naprawdę nie byłem pod wrażeniem efektu końcowego.

Fajne NeoPixela

Przedstawię ci inny rodzaj trójkolorowej diody LED, która nigdy nie przestaje zachwycać. To WS2812, znana również jako „NeoPixel” (termin pierwotnie ukojony przez ludzi z Adafruit). Ta mała piękność ma około 5×5 mm kwadratowych i 2 mm grubości (rysunek 1). Oprócz trzech bardzo jasnych diod LED RGB, WS2812 zawiera również niewielki układ kontrolera WS2811, który zawiera trzy 8-bitowe generatory PWM – po jednym dla każdej z diod LED – i obsługuje prosty protokół komunikacji szeregowej.

Każdy NeoPixel ma cztery piny (niektóre są dostarczane w obudowach 6-pinowych, ale tylko cztery z nich pełnią jakąkolwiek funkcję): 0 V, 5 V, Data-In i Data-Out, gdzie Data-Out z jednego NeoPixela może sterować Data-In innego. Umożliwia to łączenie łańcuchów długich łańcuchów NeoPixeli i sterowanie całym łańcuchem za pomocą pojedynczego cyfrowego wyjścia MCU.

NeoPixela są dostępne w różnych opcjach obudów, w tym w postaci surowych chipów (<https://bit.ly/2SOBe8b>), pojedynczych płytek Flora (<https://bit.ly/3ck1ZZS>) i tradycyjnych obudów przelotowych 8 mm (<https://bit.ly/3bkJbs5>). W przypadku modułów Flora, lubię kupować je w arkuszach po 20 sztuk (<https://bit.ly/2LdFRVq>).

Te małe światełka można również kupić w postaci pasków, pierścieni, biżuterii i wyświetlaczy (wystarczy wejść na stronę Adafruit.com, wyszukać „NeoPixel” i przejść przez 28 wspaniałych stron produktów NeoPixel). Do celów projektu, który dalej omówię, używam 5-metrowego paska zawierającego 30 NeoPixeli na metr (<https://bit.ly/2SNHTzL>).



Rysunek 1. WS2812 aka „NeoPixel”

Istnieją paski z 60 i 144 NeoPixelami na metr, ale ja podzielę mój pasek na 144 pojedyncze segmenty. Z doświadczenia wiem, że cięcie pasków 30 na metr pozostawia mi większe miedziane pady do lutowania niż cięcie pasków 60 lub 144 na metr. W rzeczywistości, dla tego konkretnego projektu wolalbym użyć modułów Flora, jak omówiono powyżej, ale paski dają mi 30 NeoPixeli za 17 USD, podczas gdy arkusze Flora dają tylko 20 NeoPixeli za 35 USD. Ponieważ do tego projektu potrzebuję 144 NeoPixeli, cięcie pasków zapewnia znacznie tańszą opcję.

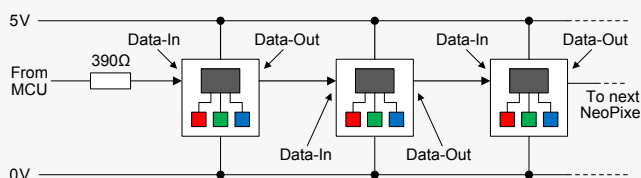
Najdrobniejsze szczegóły

Abyśmy mieli o czym rozmawiać, założmy, że chcemy sterować paskiem NeoPixeli (rysunek 2). Jest kilka punktów do odnotowania na froncie fizycznej implementacji. Teoretycznie, każdy NeoPixel może pobierać nawet 60 mA, jeśli wszystkie trzy diody LED są w pełni włączone (ich prąd znamionowy wynosi 20 mA każda). W praktyce, kiedykolwiek mierzyłem je w pełni włączone było tylko 45 mA, więc jest to wartość, którą się kieruję.

Praca z zaledwie kilkoma NeoPixelami to jedno, ale jeśli sterujesz dużą ich liczbą, musisz zacząć zwracać uwagę na to, ile prądu potrzebujesz. W przypadku projektu, który omówimy za chwilę, mój najgorszy pobór prądu wyniosłby 50 mA dla Arduino i (144×45 mA) dla moich NeoPixeli, gdybym miał je wszystkie w pełniysterować, aby uzyskać wspaniałą biel. Daje nam to łącznie 6530 mA, co oznacza, że będę potrzebował zasilacza, który zapewni co najmniej 7 A, aby dać mi trochę zapasu. W praktyce, oczywiście, prawdopodobnie będę oświetlał tylko podzbiór NeoPixeli i zazwyczaj będę je oświetlał tylko jedną lub dwiema diodami LED, ale zawsze najlepiej jest projektować pod kątem najgorszego scenariusza.

Osobiście zazwyczaj używam tego samego zasilacza do zasilania mojego MCU i moich NeoPixeli, ale różni ludzie robią to na różne sposoby. Jeśli na przykład zdecydujesz się użyć kabla USB do zasilania MCU i oddzielnego zasilacza do zasilania NeoPixeli, to bardzo ważne jest, aby upewnić się, że sygnały 0 V (GND) MCU i NeoPixeli są połączone razem.

Dobrym pomysłem jest również dodanie dużego kondensatora elektrolitycznego (1000 µF, 6,3 V lub więcej) pomiędzy zaciskami



Rysunek 2. Sterowanie łańcuchem NeoPixeli

0 V (GND) i 5 V zasilacza. Jeśli używasz surowych diod WS2812 lub NeoPixels w obudowach przelotowych, ważne jest, aby dodać kondensator ceramiczny 100 nF między zaciskami 0 V i 5 V każdego elementu. Jest to kolejny dobry powód do korzystania z Floras, pierścieni lub pasków Adafruit, ponieważ mają one już te kondensatory na miejscu.

Ponadto ważne jest, aby dołączyć rezystor szeregowy jak najbliżej pierwszego elementu w łańcuchu NeoPixel (pierścienie NeoPixel firmy Adafruit mają już zainstalowany taki rezystor). Kiedy po raz pierwszy zacząłem używać pasków NeoPixel, od czasu do czasu pierwszy element w łańcuchu umierał. Kiedy użyłem oscyloskopu, aby spojrzeć na sygnał Data-In pochodzący z mojego Arduino Uno, zobaczyłem, że częstotliwość zboczyła tak szybko, że w stanach niestabilnych pojawiały się piki ponad 5 V poniżej 0 V. Po kilku próbach ustaliłem, że rezystor o wartości 390 Ω dobrze tłumi zakłócenia i od tamtej pory nie miałem żadnych problemów. To powiedziawszy, powinienem zauważyć, że w instrukcji do Neopixeli (<https://bit.ly/2SR14eu>), którą warto przeczytać i „przetrawić”, ludzie z Adafruit zalecają użycie rezystora 470 Ω . Widziałem też inne osoby oferujące różne sugestie, ale dawno temu kupiłem kilkadziesiąt rezystorów 390 Ω właśnie w tym celu, więc właśnie takich zamierzam użyć!

„Wypasione” biblioteki

Istnieje wiele bibliotek, które można wykorzystać do sterowania NeoPixeli. Jeśli używam Arduino lub pokrewnego MCU, prawie zawsze korzystam z biblioteki Adafruit (<https://bit.ly/2LDpXPK>). Inną łatwą w użyciu biblioteką Arduino do programowania NeoPixeli (i innych urządzeń), którą poleca wielu moich znajomych, jest FastLED Animation Library (<http://fastled.io/>). Alternatywnie, jeśli korzystam z Teensy MCU od PJRC.com, używam ich biblioteki OctoWS2811 (<https://bit.ly/2YLER2I>) w połączeniu z adapterem OctoWS2811 (<https://bit.ly/2SNyJmJ>).

W przypadku biblioteki Adafruit, zaczynasz od ustalenia łańcucha NeoPixels, w ramach którego określasz, ile pikseli będzie w łańcuchu i które z pinów Arduino chcesz nimi sterować. Dla każdego z NeoPixeli biblioteka zarezerwuje trzy bajty w pamięci SRAM Arduino, co oznacza, że na przykład Arduino Uno z 2 KB pamięci SRAM jest ograniczone do obsługi maksymalnie około 500 NeoPixeli, pozostawiając 512 bajtów pamięci SRAM wolnych na inne rzeczy.

Założmy, że utworzyliśmy ciąg 10 NeoPixeli, które będą ponumerowane od 0 do 9, i nazwaliśmy ten ciąg MyNeos. Założmy teraz, że używamy instrukcji takiej jak `MyNeos.setPixelColor(i, COLOR_HOT_PINK)`, gdzie `i` jest numerem NeoPixela, który chcemy zmienić, a `COLOR_HOT_PINK` jest 24-bitową wartością szesnastkową reprezentującą składowe RGB naszego pożądanego koloru. Ważne jest, aby pamiętać, że w rzeczywistości nie zmienia to wartości NeoPixela w łańcuchu; zamiast tego zmienia wartość w zarezerwowanym obszarze pamięci SRAM Arduino. Dopiero po użyciu polecenia `MyNeos.show()` wszystkie wartości w pamięci zostaną przesłane do fizycznego łańcucha. Nie martw się, zobaczymy proste przykłady wszystkich tych rzeczy w programach, które stworzymy później.

Bity i bajty

Uderzyło mnie, że jest jedna rzecz, która może okazać się myląca, jeśli jesteś nowicjuszem w korzystaniu z NeoPixeli – fakt, że istnieją dwa sposoby określenia koloru, którego chcemy użyć. Zaczniemy od przypomnienia sobie, że każdy NeoPixel zawiera czerwone, zielone i niebieskie diody LED. Ponadto zawiera trzy 8-bitowe generatory PWM, po jednym dla każdej z podrzędnych diod LED. Oznacza to, że każdej diodzie podrzędnej możemy przypisać wartość z zakresu od 0 do 255.

Założmy, że chcemy ustawić siódmy NeoPixel w łańcuchu na kolor, który możemy nazwać intensywnym fioletem (pamiętaj, że ten piksel

będzie w rzeczywistości numerem 6, ponieważ zaczynamy liczyć od 0). Założmy ponadto, że aby uzyskać ten kolor, chcemy, aby składowa czerwona wynosiła 128, składowa zielona wynosiła 0, a składowa niebieska wynosiła 255. W tym przypadku twórcy biblioteki NeoPixel firmy Adafruit zaimplementowali wszystko w taki sposób, że możemy użyć instrukcji takiej jak: `MyNeos.setPixelColor(6,128,0,255)`.

Jednak spryciarze zaimplementowali również takie rozwiązania, że możemy określić połączone komponenty RGB jako pojedynczą 24-bitową wartość. Do tego typu rzeczy najlepiej jest używać systemu szesnastkowego, więc moglibyśmy uzyskać nasz elektrycznie fioletowy kolor za pomocą: `MyNeos.setPixelColor(6,0x8000FF)`, gdzie `0x80` odpowiada 128 w systemie dziesiętnym, `0x00` odpowiada 0 w systemie dziesiętnym, a `0xFF` odpowiada 255 w systemie dziesiętnym.

W wielu przypadkach lepiej jest użyć tego 24-bitowego podejścia, ponieważ pozwala nam to robić takie rzeczy, jak wstępne definiowanie koloru za pomocą czegoś takiego jak `#define COLOR_ELECTRIC_VIOLET 0x8000FFU`, gdzie `U` jest używane do wskazania, że jest to wartość bez znaku. Użycie `U` (lub `u`) jest opcjonalne – kompilator zazwyczaj sam sobie w tym poradzi – ale rzadko szkodzi dać mu okazjonalną wskazówkę i sprawić, że twoje intencje będą jaśniejsze dla kogoś innego czytającego twój kod. Teraz moglibyśmy użyć `MyNeos.setPixelColor(6,COLOR_ELECTRIC_VIOLET)` w naszym programie. Jeśli pobierzesz szkice omówione w dalszej części tej kolumny, zobaczysz, że właśnie to zrobiliśmy.

Ekscytujesz mnie

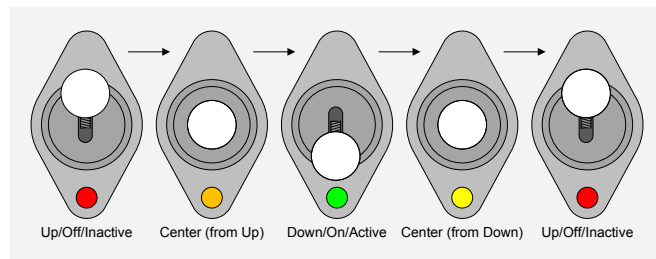
Zanim przejdziemy do mojego nowego projektu, musimy najpierw uporządkować nasze wcześniejsze eksperymenty z przełącznikami. Korzystając z podobnej konfiguracji do tej omówionej w poprzednim odcinku, zamontowałem pojedynczy NeoPixel Flora na płytce prototypowej i napisałem mały szkic, aby kontrolować go za pomocą jednobiegowego przełącznika SPCO.

Podobnie jak w przypadku standardowej trójkolorowej diody LED w poprzednim odcinku, użyjemy koloru czerwonego do wskazania, kiedy przełącznik jest wyłączony/nieaktywny, zielonego do wskazania, kiedy przełącznik jest włączony/aktywny, oraz pomarańczowego lub żółtego, gdy przełącznik znajduje się w pozycji środkowej, aby zapewnić wskazanie jego poprzedniego stanu (**rysunek 3**). Możesz pobrać szkic (plik CB-Jul20-01.txt – dostępny na stronie PE z lipca 2020 r.) i obejrzeć film (<https://bit.ly/3cyCyUF>), aby zobaczyć to wszystko w akcji.

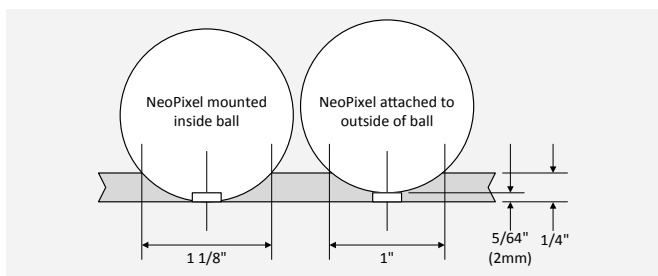
Zestaw pitek

I tak dochodzimy do mojego nowego hobbystycznego projektu. Jak wspomniałem w poprzednim odcinku, postanowiłem zbudować wspólną matrycę opartą na piłeczkach pingpongowych podświetlanych przez NeoPixelse – coś w rodzaju „ściany wideo”, którą można zobaczyć na YouTube (<https://bit.ly/3aG1itl>).

Zdecydowałem również, że moim pierwszym prototypem będzie mały prototyp ping-ponga 12×12=144. Okazało się, że w USA można kupić torbę 144 piłeczek pingpongowych za jedyne 11 USD, ale wiedziałem,



Rysunek 3. Użycie trójkolorowego NeoPixela z przełącznikiem SPCO i dwoma kolorami dla pozycji środkowej



Rysunek 4. Alternatywne mocowania NeoPixel (teoretyczne)

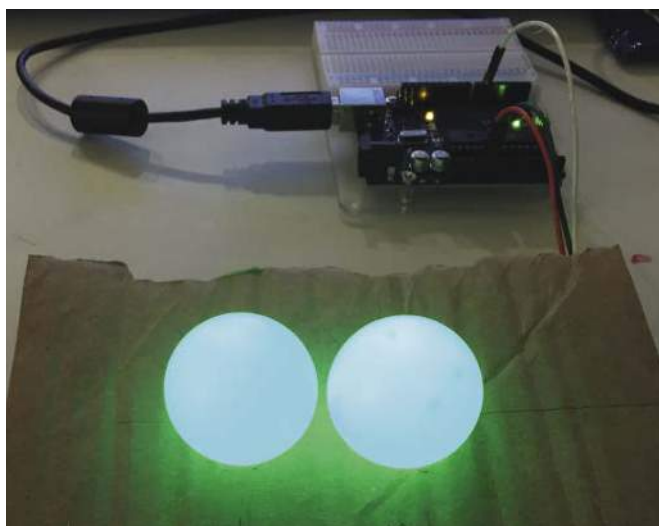
że będę potrzebował kilku zapasowych, więc kupiłem dwie torby, co dało mi w sumie 288 piłeczek pingpongowych.

Oczywiście od razu kusilo mnie, aby „pójść w większe i lepsze” – powiedzmy $15 \times 15 = 225$ matryc – ale zamówiłem już pięć metrów paska NeoPixel 30 pikseli na metr od Adafruit (<https://bit.ly/3dOa5v4>), co da mi 150 NeoPikseli, więc postanowiłem trzymać się pierwotnego planu. Dzięki Bogu, że tak zrobiłem, ponieważ wszystko wymaga znacznie więcej wysiłku i zajmuje znacznie więcej czasu niż pierwotnie planowałem.

Zacząłem od zastanowienia się, w jaki sposób zamierzam przymocować segmenty paska NeoPixel do piłeczek pingpongowych. Pracowałem z kawałkiem sklejkę o grubości 6 mm ($1/4$ cala). Chciałem, aby same paski były zlicowane z powierzchnią drewna. Wiedząc, że NeoPiksele wystają z powierzchni pasków o 2 mm ($5/64$ cala), istnieją dwa oczywiste rozwiązania (rysunek 4).

Pierwszą opcją jest wycięcie otworu w piłeczce pingpongowej i zamontowanie NeoPixel wewnątrz piłki. W tym przypadku, ponieważ piłeczki mają średnicę 38 mm ($1\ 1/2$ cala) (dopuszczam również 1,5 mm ($1/16$ cala) między piłeczkami dla „miejsca do poruszania się”), otwór, który wywiercę w płytce, będzie musiał mieć średnicę 28,6 mm ($1\ 1/8$ cala), aby podstawa piłeczki znajdowała się na równi ze spodem płytki. Drugą opcją byłoby przymocowanie NeoPixeli na zewnątrz kuli. W tym przypadku otwór, który wywiercę w płytce będzie musiał mieć średnicę 25 mm (1 cal).

Szczerze mówiąc, nie sądziłem, że będzie duża różnica między tymi dwoma rozwiązaniami w odniesieniu do ich wyglądu, ale postanowiłem zbudować prototyp przy użyciu kawałka tektury i zestawu piłeczek. Stworzyłem mały filmik pokazujący tę samą sekwencję kolorów wyświetlanych w obu piłeczkach (<https://bit.ly/2WZN6pq>), a szkic, którego użyłem, można pobrać (plik CB-Jul20-02.txt – dostępny na stronie PE z lipca 2020 r.).



Rysunek 5. Alternatywne mocowania NeoPixeli (prototyp)

Powiedziałem sobie, że wyglądają tak samo, ale zapytałem również moją żonę (Gina the Gorgeous) i mojego 25-letniego syna (Joseph the Common-sense Challenged) i oboje stwierdzili, że jedna wygląda lepiej (gładko) i jaśniej niż druga. Choć na **rysunku 5** wyglądają tak samo, to naprawdę widać różnicę, gdy patrzy się na nie w prawdziwym świetle. Oczywiście lepszą opcją była ta, która wymagała ode mnie wycięcia 10 mm ($3/8$ cala) otworów w 144 piłeczkach pingpongowych. „Ojej”, powiedziałem do siebie (lub słowa w tym stylu).

Ale potem zajrzałem do skrzynki z narzędziami. Największe wiertło, jakie miałem, miało średnicę 25 mm (1 cal). „No cóż”, powiedziałem do siebie, „los podjął decyzję za mnie, a w końcu różnica między tymi dwoma rozwiązaniami jest naprawdę niewielka”.

Poszedłem więc i wywierciłem otwory 144×25 mm (1 cal), przeszlifowałem wszystko i pomalowałem deskę na czarno. Kiedy skończyłem, wziąłem jedną piłeczkę pingpongową i włożyłem ją do otworu. Co? Spód piłeczki zrównał się z dolną powierzchnią płyty! Jak to możliwe?

Okazało się, że to, co uważałem za deskę o grubości 6 mm ($1/4$ cala), gdy zobaczyłem ją leżącą w garażu, miało w rzeczywistości tylko 5 mm ($3/16$ cala) grubości. Wróciłem więc do konieczności wycinania otworów o średnicy 10 mm ($3/8$ cala) w moich 144 piłeczkach pingpongowych. Prawdę mówiąc, nie byłem tym zbyt przerażony, ponieważ świadomość, że wybrałem łatwiejszą, ale gorszej jakości opcję, nie dawała mi spokoju.

Budowanie tablicy

Jestem pewien, że podobnie jak ja, przez lata spędziłeś znacznie więcej czasu niż chciałbyś pamiętać na frustrującym zadaniu wiercenia otworów w piłeczkach pingpongowych. Tym razem wymyśliłem coś innego. Najpierw wziąłem piłkę i trzymałem ją pod światło, aby określić, gdzie dwie półkule są połączone, a następnie użyłem trwałego markera, aby zrobić kropkę na środku jednej z półkul (nie chcę, aby linia łączenia była widoczna). Następnie użyłem szablonu, aby zaznaczyć okręgi o średnicy 10 mm ($3/8$ cala) wyśrodkowane na kropce. Na koniec użyłem małych, ostrych, zakrzywionych nożycek do paznokci, aby przebić mały otwór w środku koła i ostrożnie odciąć materiał piłeczki. Następnie powtórzyłem



Rysunek 6. Jestem przekonany, że to najlepszy układ piłeczek pingpongowych na naszej ulicy

ten proces jeszcze 143 razy (piękny, długi wieczór, który już nigdy nie powróci).

Następnym krokiem było przymocowanie piłeczek pingpongowych do drewnianego arkusza. Nie będę was znużał problemami z wyrównaniem piłeczek i przyrządem, który musiałem stworzyć. Wystarczy powiedzieć, że jeśli kiedykolwiek zbuduję wyświetlacz wielkości ściany z piłeczek pingpongowych,

zrobię to przy użyciu podmatryc 8×8. Pomińmy więc zgrzytanie zębami i rozdieranie szat i przejdźmy do części, w której mówię: „Jak zawsze, mój pistolet do klejenia na gorąco okazał się wiernym przyjacielem”.

Szczerze mówiąc, uważam, że efekt końcowy wygląda dość elegancko (**rysunek 6**). Miałem nadzieję, że wszystko uda mi się uruchomić na czas przed napisaniem tego odcinka, ale tak się nie stało. Nadal muszę dołączyć 144 segmenty paska NeoPixel, a następnie przylutować wszystko razem (144×3=432 połączenia lutowane „plus”, „minus”, „dane”), więc nie będziemy mogli zobaczyć tego małego piękna w akcji aż do następnego odcinka.

Seeeduino XIAO

W przeszłości pokusiłbym się o użycie 8-bitowego Arduino Nano do tego projektu, ponieważ jest on stosunkowo mały i dyskretny (<https://bit.ly/2Lfk2jy>). Z drugiej strony, Nano ma tylko zegar 16 MHz, 32 KB pamięci Flash i 2 KB pamięci SRAM, co jest nieco ograniczające.

Wtedy jednak los znów się do mnie uśmiechnął, ponieważ ludzie z Seeed Studio opowiedzieli mi o swoim Seeeduino XIAO (<https://bit.ly/3ckK31c>) i nawet wysłali mi jeden egzemplarz do zabawy. Kosztujący zaledwie 4,90 USD i będący wielkości małego znaczka pocztowego (**rysunek 7**), to małe cudo może pochwalić się 32-bitowym procesorem Arm Cortex-M0+ pracującym z częstotliwością 48 MHz, 256 KB pamięci Flash i 32 KB pamięci SRAM. Jedną rzeczą, na którą należy zwrócić uwagę, jest to, że złącze programowania to USB typu C, co oznacza, że będziesz potrzebować kabla USB-A do USB typu C.

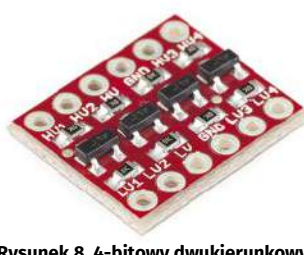
Seeeduino XIAO posiada 11 pinów cyfrowych/analogowych, z których 10 obsługuje PWM, a jeden może zapewnić wyjście prawdziwego przetwornika cyfrowo-analogowego (DAC). Piny te mogą być również wykorzystywane do obsługi interfejsu UART, interfejsu SPI i interfejsu I²C. Ponadto, to małe cudo byłoby świetne do implementacji efektów świetlnych do noszenia. Myślę, że można śmiało powiedzieć, że Seeeduino XIAO pojawi się w wielu moich przyszłych projektach.

Utrzymywanie poziomu

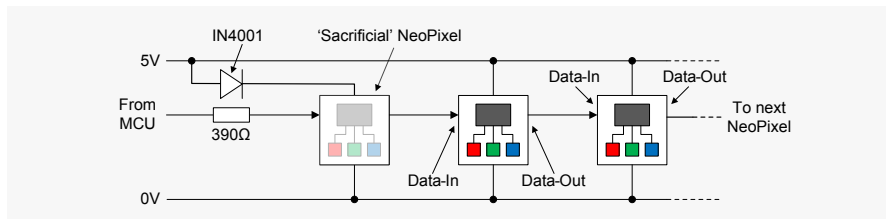
Z drugiej strony, Seeeduino XIAO może być zasilany z tego samego zasilacza 5 V, którego używam do moich NeoPixeli. Jest jednak jedno „ale” lub „pewien szkopuł” w tym wszystkim (czuję się w hojnym nastroju, więc pozwolę ci wybrać metaforę). Piny I/O Seeeduino XIAO używają interfejsu 3,3 V, ale moje NeoPiksele wymagają sygnałów danych 5 V, więc potrzebujemy jakiegoś sposobu na konwersję między nimi.



Rysunek 7. Seeeduino: mój nowy najlepszy przyjaciel



Rysunek 8. 4-bitowy dwukierunkowy konwerter poziomów logicznych SparkFun



Rysunek 9. Tani i sprytny konwerter napięć

W przeszłości odniosłem wiele sukcesów z 4-kanalowym dwukierunkowym konwerterem poziomów logicznych firmy SparkFun (<https://bit.ly/2WHnvRW>). Ta kosztująca zaledwie 2,95 USD płytka typu break-out (moduły dostępne w postaci panelu po procesie v-cut, czyli nafrezowane ostrzem w kształcie litery V, gotowe do depanelizacji w rękach przez rozłamywanie) może być używana do konwersji sygnałów 3,3 V na ich odpowiedniki 5 V i odwrotnie (**rysunek 8**). Można go nawet używać z magistralą I²C, która wymaga rezystorów podciągających, ponieważ BOB ma rezystory podciągające 10 kΩ po obu stronach każdego kanału.

Mając to na uwadze, potrzebuję jedynie przekonwertować sygnał z pojedynczego wyjścia cyfrowego 3,3 V na Seeeduino XIAO. Co więcej, potrzebuję tylko jednokierunkowej konwersji poziomu. Z obu tych powodów konwerter poziomu BOB firmy SparkFun byłby przesadą.

Na szczęście natknąłem się na niesamowity hack na Hackaday.com (<https://bit.ly/35LjlfL>), który zapewnia proste rozwiązanie wymagające tylko jednej diody (IN4001 ogólnego przeznaczenia jest idealny do tego zadania) i NeoPixela „na zmarnowanie” (**rysunek 9**).

Działa to w ten sposób, że arkusz danych NeoPixel określa wartość logiczną 1 jako $0,7 \times V_{cc}$. Tak więc, jeśli zasilamy NeoPixel napięciem 5 V, logiczna 1 będzie wynosić $0,7 \times 5 = 3,5$ V. W rzeczywistości sygnał 3,3 V z Seeeduino XIAO może działać, ale może też nie działać.

Rozwiązaniem jest zasilenie ofiarnego NeoPixela przez diodę IN4001. Ponieważ dioda ta ma spadek napięcia w stanie przewodzenia wynoszący 0,7 V, oznacza to, że pierwszy NeoPixel jest zasilany napięciem $V_{cc} \text{ wynoszącym } 5 - 0,7 = 4,3$ V. To z kolei oznacza, że pierwszy NeoPixel będzie postrzegał sygnał $0,7 \times 4,3 = 3,01$ V jako logiczną 1, więc wyjście 3,3 V z mikrokontrolera idealnie wpasowuje się w te wyliczenia.

Tymczasem sygnał Data-Out 4,3 V z pierwszego NeoPixela jest więcej niż wystarczający do zasilania sygnału Data-In do drugiego NeoPixela w łańcuchu. Oznacza to, że używamy pierwszego NeoPixela w roli konwertera poziomu napięcia. W wielu systemach nadal możemy używać tego poświęconego NeoPiksele do wskazywania czegoś, ponieważ będzie on po prostu nieco słabszy niż pozostałe NeoPiksele. Jednak w przypadku mojego układu 12×12 piłeczek pingpongowych, po prostu dołączę dodatkowy NeoPixel przed pierwszą piłeczką pingpongową i zawsze ustawię go na kolor czarny (wyłączony).

Następnym razem

Do czasu następnego odcinka, będę miał moją matrycę 12×12 uruchomioną i działającą, więc będziemy mogli rozważyć niektóre programy i efekty, które możemy na niej uruchomić. Do tego pamiętnego dnia życzę miłej zabawy.



Komentarze lub pytania?
Napisz do Maxa na adres:
max@CliveMaxfield.com

Sprytne porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania



W poprzednim odcinku (EdW 01/2024) obiecałem, że w tym miesiącu zastanowimy się, w jaki sposób operatory << i >> wykonują swoją magię. Niestety, będziemy musieli odłożyć to na później, ponieważ czytelnik David R Humrich, który pochodzi z Perth w Australii, wysłał mi e-mail z dość interesującym pytaniem dotyczącym użycia nawiasów klamrowych { }.

Zanim zajmiemy się pytaniem Davida, przypomnijmy, że { } może być używane do tworzenia tak zwanych „instrukcji złożonych”. Jest to mechanizm używany przez język programowania C do grupowania wielu instrukcji w coś, co można uznać za pojedynczą instrukcję.

Rozważmy na przykład, co się stanie, gdy zdefiniujemy funkcję o nazwie MyFunction ():

```
void MyFunction ()
{
    // Udawaj, że ten komentarz jest oświadczeniem
    // Udawaj, że ten komentarz jest oświadczeniem
    // Udawaj, że ten komentarz jest oświadczeniem
}
```

Kiedy wywołujemy tę funkcję z innego miejsca w programie, komputer „widzi” wszystkie instrukcje w funkcji jako tworzące jedną logiczną całość.

To samo dzieje się, gdy używamy { } wraz z instrukcją sterującą, taką jak if (). Po pierwsze, założmy, że jeśli warunek jest prawdziwy, chcemy wykonać tylko jedną akcję, w którym to przypadku moglibyśmy napisać to w następujący sposób:

```
if (done == true) fred = fred + 1;
```

W języku C nie ma znaczenia, ile białych znaków użyjemy:

```
if (done == true)
    fred = fred + 1;
```

Zakładamy oczywiście, że zmienne done i fred zostały zadeklarowane w innym miejscu programu. Załóżmy teraz, że chcemy wykonać kilka akcji, jeśli nasz warunek jest prawdziwy. Można to zrobić w następujący sposób:

```
if (done == true) fred = fred + 1;
if (done == true) jane = jane - 1;
if (done == true) bert = fred + jane;
```

Oprócz tego, że wygląda to głupio, jest to nieefektywne, ponieważ wykonujemy ten sam test trzy razy. Ten przykład wzywa nas do użycia instrukcji złożonej w następujący sposób:

```
if (done == true) // Równe powyższemu
{
    fred = fred + 1;
    jane = jane - 1;
    bert = fred + jane;
}
```

W tym przypadku złożone oświadczenie jest zarówno bardziej wydajne, jak i jaśniejsze, co próbujemy osiągnąć.

Powrót do Davida

Wracając do pytania Davida. Zapytał on, co by się stało, gdyby ktoś użył { }, a tym samym utworzył samodzielną instrukcję złożoną w środku funkcji – przez „samodzielną” rozumiemy, że nie jest ona powiązana z instrukcją sterującą, taką jak if () lub for ().

W rzeczywistości jest to całkowicie legalne. Jeśli chcesz, możesz po prostu użyć { }, aby zebrać grupę instrukcji razem, aby wyjaśnić sobie i innym, że uważasz te instrukcje za powiązane. Są one często określane jako „blok”, a korzystanie z tej techniki może być określane jako „programowanie blokowe”. Można również zagnieżdżać { { } } do dowolnego poziomu. Ale naprawdę interesujące jest to, że oprócz instrukcji, bloki mogą również zawierać deklaracje zmiennych.

Dlaczego jest to interesujące? Cieszę się, że o to zapytałeś. W poprzednim odcinku rozmawialiśmy o zmiennych globalnych i lokalnych. Zauważyliśmy, że zmienne globalne są deklarowane poza dowolną funkcją i mogą być widziane i modyfikowane przez dowolną funkcję. Dla porównania, zmienne lokalne są deklarowane wewnątrz funkcji i mogą być widziane i modyfikowane tylko przez funkcję, w której zostały zadeklarowane.

Rozmawialiśmy również o „zakresie” zmiennej, który odnosi się do zakresu, w jakim zmienna może być widoczna. Na przykład, jeśli zadeklarujemy zmienną jako część pętli for (), zakres zmiennej jest ograniczony do tej pętli (EdW 01/2024). Cóż, jeśli zmienna jest zadeklarowana wewnątrz bloku, jej zakresem jest ten blok; to znaczy, że nie może być „widziana” poza blokiem, nawet przez inne części funkcji, w której znajduje się blok. Rozważmy przykład kodu poniżej:

```
int John = 6;
```

```
void MyFunction ()
{
    int jane = 9;

    { // Pierwszy blok
        int bert = jane + John
        // Więcej rzeczy
    }

    { // Drugi blok
        int jack = jane - John;
        // Więcej rzeczy
    }
}
```

Pamiętaj, że używam początkowych wielkich i małych liter odpowiednio dla moich zmiennych globalnych i lokalnych (EdW 12/2023). Tak więc John jest zmienną globalną, której zakresem jest każda funkcja w programie, podczas gdy jane jest zmienną lokalną, której zakres jest ograniczony do MyFunction (). Dla porównania, zakres zmiennej bert jest ograniczony do pierwszego bloku, a zakres zmiennej jack jest ograniczony do drugiego bloku.

Podzielenie dużego programu na kilka mniejszych, dobrze zdefiniowanych funkcji ułatwia testowanie każdej funkcji oddzielnie i ponowne wykorzystanie funkcji w różnych programach. Ponadto, jedną z zalet korzystania z techniki blokowej w celu ograniczenia zakresu zmiennych jest to, że ułatwia ona ponowne wykorzystanie tych bloków w innych funkcjach i innych programach.

Następnym razem

Nie powiem, czym zajmiemy się następnym razem, ponieważ rzeczy szybko się tutaj zmieniają (odrobiłem swoją lekcję). Musimy po prostu poczekać i zobaczyć. Do tego czasu, miłego oglądania! ■

Clive „Max” Maxfield

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

Precyzyjny filtr aktywny drugiego rzędu do subwoofera

W publikacji opisano prostą metodę obliczania precyzyjnego filtra aktywnego drugiego rzędu do subwoofera o regulowanej skokowo częstotliwości podziału. Układ ten może znaleźć zastosowanie w konstrukcji profesjonalnych subwooferów przeznaczonych do użytku w studiach nagraniowych lub w zastosowaniach estradowych. Główną zaletę tego filtra stanowi możliwość zachowania stałego współczynnika dobroci w całym zakresie regulacji częstotliwości podziału.

Wprowadzenie

Na stronie internetowej firmy Analog Devices można znaleźć ciekawe narzędzie wspomagające komputerowo projektowanie filtrów aktywnych opartych o wzmacniacze operacyjne. Aplikacja ta nosi nazwę Analog Filter Wizard: <https://tools.analog.com/en/filterwizard/>.

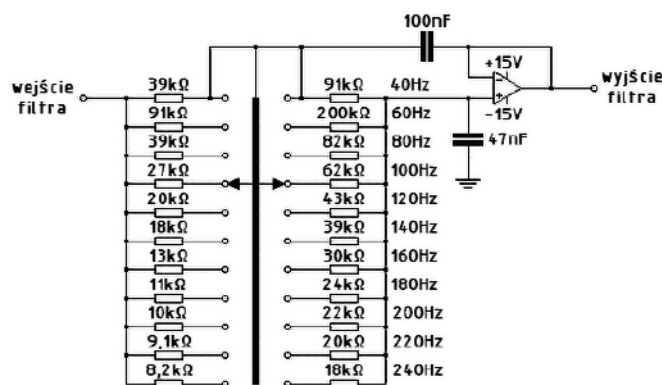
Narzędzie to umożliwia projektowanie aktywnych filtrów dolnoprzepustowych, pasmowoprzepustowych oraz górnoprzepustowych w konfiguracji Sallen Key oraz z wielokrotnym sprzężeniem zwrotnym. Przy pomocy tej aplikacji możemy projektować filtry o różnych częstotliwościach podziału, różnym nachyleniu i różnych współczynnikach dobroci. Narzędzie pozwala nam wykreślać charakterystyki amplitudowo-częstotliwościowe i fazowo-częstotliwościowe filtra. Umożliwia także wykreślenie charakterystyki opóźnienia grupowego filtra w funkcji częstotliwości oraz odpowiedzi filtra na skok jednostkowy. Aplikacja ta dobiera wartości elementów w sposób automatyczny przy uwzględnieniu wybranego szeregu wartości oraz tolerancji zastosowanych elementów RC. Narzędzie to pomoże nam zaprojektować precyzyjny filtr aktywny drugiego rzędu do subwoofera.

Projektowanie precyzyjnego filtra aktywnego drugiego rzędu do subwoofera

Częstotliwość podziału oddzielająca pasmo przepustowe od zaporowego będzie regulowana skokowo w zakresie od 40 Hz do 240 Hz z krokiem co 20 Hz. Będzie to filtr drugiego rzędu o nachyleniu 12 dB/

okt. wykonany w konfiguracji Sallen Key. Aplikacja Analog Filter Wizard podpowiada nam, że najniższy poziom szumów zapewni nam zastosowanie w układzie kondensatorów foliowych o pojemnościach 100 nF oraz 47 nF. Przełączaniu będą podlegać jedynie dwa oporniki wchodzące w skład filtra. Schemat układu wraz z wartościami zastosowanych elementów elektronicznych przedstawiono na **rysunku 1**.

Na początek należy wyznaczyć wartości rezystancji oporników w zależności od jednej z jedenastu częstotliwości podziału. Wartości te ujęto w **tabeli 1**.



Rysunek 1. Schemat precyzyjnego filtra aktywnego drugiego rzędu do subwoofera

Tabela 1. Wartości rezystancji oporników filtra wyznaczone przez aplikację Analog Filter Wizard w oparciu o szereg wartości E24 i tolerancję jednoprocenową

Częstotliwość [Hz]	Rezystancja pierwszego opornika filtra [kΩ]	Rezystancja drugiego opornika filtra [kΩ]
40	39,0	91,0
60	27,0	62,0
80	20,0	43,0
100	16,0	36,0
120	13,0	30,0
140	12,0	27,0
160	10,0	22,0
180	9,1	20,0
200	8,2	18,0
220	7,5	16,0
240	6,8	15,0

Tabela 2. Obliczone wartości rezystancji dodatkowych oporników R_2 dla połączeń równoległych

Częstotliwość [Hz]	Dodatkowy opornik R_2 dla pierwszego opornika filtra [kΩ]	Dodatkowy opornik R_2 dla drugiego opornika filtra [kΩ]
40	---	---
60	87,8	194,6
80	41,1	81,5
100	27,1	59,6
120	19,5	44,8
140	17,3	38,4
160	13,4	29,0
180	11,9	25,6
200	10,4	22,4
220	9,3	19,4
240	8,2	18,0

Tabela 3. Wartości rezystancji dodatkowych oporników R_2 dla połączeń równoległych dobrane z szeregu wartości E24

Częstotliwość [Hz]	Dodatkowy opornik R_2 dla pierwszego opornika filtra [kΩ]	Dodatkowy opornik R_2 dla drugiego opornika filtra [kΩ]
40	---	---
60	91,0	200,0
80	39,0	82,0
100	27,0	62,0
120	20,0	43,0
140	18,0	39,0
160	13,0	30,0
180	11,0	24,0
200	10,0	22,0
220	9,1	20,0
240	8,2	18,0

Z uwagi na to, że jedenastopozycyjny dwusekcyjny przełącznik obrotowy będzie dołączał każdorazowo równolegle do każdego z oporników zakresu 40 Hz dodatkowy opornik, musimy wyznaczyć rezystancje tych dodatkowych oporników. Posłużymy się do tego celu formułą pozwalającą wyznaczyć rezystancję wypadkową dwóch oporników w połączeniu równoległym:

$$R_W = \frac{R_1 \cdot R_2}{R_1 + R_2}$$

Rezystancja wypadkowa jest znana, za rezystancję R_1 przyjmujemy każdorazowo rezystancję jednego z oporników zakresu 40 Hz. Potrzebujemy zatem wyznaczyć formułę pozwalającą obliczyć rezystancję dodatkowego opornika R_2 , który zostanie dołączony równolegle:

$$R_W = \frac{R_1 \cdot R_2}{R_1 + R_2} \quad \left| \cdot (R_1 + R_2) \right.$$

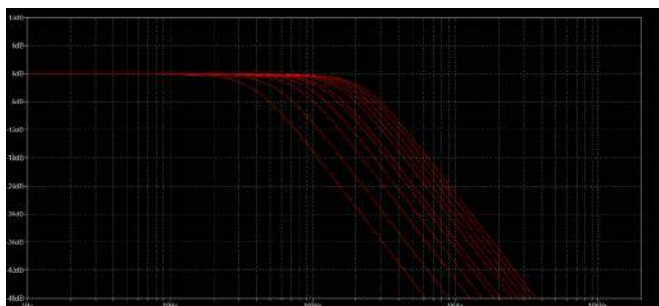
$$R_W \cdot (R_1 + R_2) = R_1 \cdot R_2$$

$$R_W \cdot R_1 + R_W \cdot R_2 = R_1 \cdot R_2$$

$$R_W \cdot R_2 - R_1 \cdot R_2 = -R_W \cdot R_1$$

$$R_2 \cdot (R_W - R_1) = -R_W \cdot R_1 : (R_W - R_1)$$

$$R_2 = \frac{-R_W \cdot R_1}{R_W - R_1}$$



Rysunek 2. Zestaw charakterystyk amplitudowo-częstotliwościowych precyzyjnego filtra aktywnego drugiego rzędu do subwoofera o regulowanej skokowo częstotliwości podziału

Tabela 4. Wartości rezystancji wypadkowych R_W uzyskane po dołączeniu równolegle oporników R_2 dobranych z szeregu wartości E24

Częstotliwość [Hz]	Dodatkowy opornik R_2 dla pierwszego opornika filtra [kΩ]	Dodatkowy opornik R_2 dla drugiego opornika filtra [kΩ]
40	---	---
60	27,3	62,5
80	19,5	43,1
100	16,0	36,9
120	13,2	29,2
140	12,3	27,3
160	9,8	22,6
180	8,6	19,0
200	8,0	17,7
220	7,4	16,4
240	6,8	15,0

Teraz dla obliczonych wartości rezystancji dobieramy wartości, które występują w szeregu wartości E24 (tabela 3).

Na koniec sprawdzamy jakie wartości rezystancji wypadkowych R_W uzyskamy po dołączeniu oporników R_2 dobranych z szeregu wartości E24. Możemy porównać je z wartościami rezystancji oporników filtra zamieszczonymi w tabeli 1.

Jeśli w handlu nie uda się znaleźć jedenastopozycyjnego dwusekcyjnego przełącznika obrotowego, można go zastąpić przełącznikiem obrotowym dwusekcyjnym o mniejszej liczbie pozycji, ograniczając w ten sposób zakres regulacji częstotliwości podziału. Podczas obliczania parametrów filtra przyjęto krok co 20 Hz. Nic nie stoi jednak na przeszkodzie aby na podstawie informacji zawartych w tym artykule przeliczyć układ dla innych częstotliwości podziału wedle uznania. Autor pozostawia w tym względzie czytelnikom całkowitą dowolność.

Książka o układach elektronicznych do subwooferów aktywnych

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

mgr inż. Tomasz Łysek



Rysunek 3. Okładka książki pt. „Wprowadzenie do projektowania układów elektronicznych subwooferów aktywnych. Poradnik praktyczny”

REKLAMA

Akumulatory



Ogniwa galwaniczne omówione w artykule „Elektrochemia” nazywane są „ogniwami pierwotnymi”. Jeśli ogniwo jest wyczerpane, nie można go ponownie naładować. Ponadto istnieją „ogniwa wtórne”, zwane w codziennym użytku akumulatorami. Akumulatory można ładować, gdy są wyczerpane.

Zasada działania akumulatorów

Czym jest akumulator?

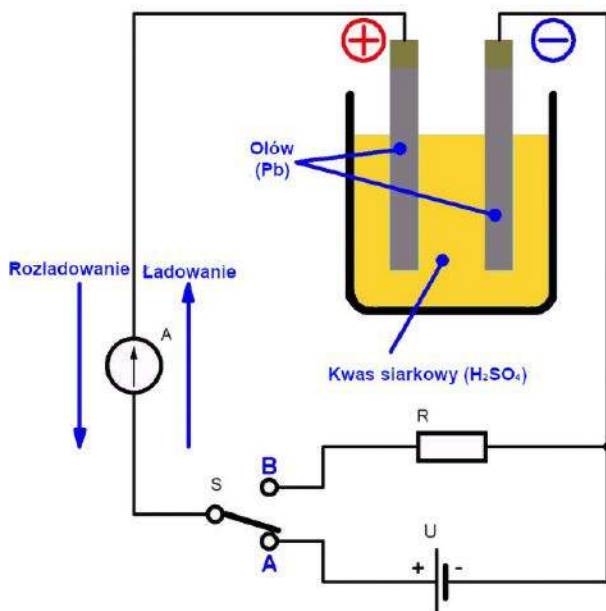
Akumulatory są magazynami energii elektrycznej. Po dodaniu energii elektrycznej do akumulatora jest ona przekształcana w energię chemiczną. Nazywa się to ładowaniem akumulatora. Po naładowaniu energia jest zachowana w związkach chemicznych zawartych w akumulatorze. Jeśli następnie do akumulatora zostanie podłączone obciążenie elektryczne, energia chemiczna zostanie ponownie przekształcona w energię elektryczną. Nazywa się to rozładowaniem akumulatora. Teoretycznie proces ładowania i rozładowywania można powtarzać w nieskończoność.

Pierwszy akumulator

Wynalezienie akumulatora przypisuje się francuskiemu fizykowi o nazwisku Planté. W 1860 roku przeprowadził on eksperyment, który ilustruje poniższy rysunek. Dwie ołowiane elektrody zostały zawieszone w szklanym pojemniku w kąpeli z rozcieńczonym kwasem siarkowym. Oczywiście jest, że taki układ nigdy nie może tworzyć pierwotnego ogniwa elektrochemicznego, ponieważ używane są dwie płytki z tego samego metalu, a potencjał elektryczny nigdy nie może powstać między nimi.

Co ustalił Planté

Można podłączyć dwie płytki ołowiane za pomocą amperomierza A i przełącznika S do ogniwa pierwotnego U lub do rezystora R.



Eksperyment, który doprowadził do wynalezienia akumulatora
(© 2018 Jos Verstraten)

Jeśli ustawisz przełącznik S w pozycji A, zobaczysz, że prąd wypływa z ogniwa pierwotnego do ogniwa wtórnego. Po pewnym czasie prąd ten zmniejsza się, aż spadnie do zera. Jeśli następnie odłączysz ogniwo pierwotne od ogniwa wtórnego, a w jego miejsce podłączysz obciążenie w postaci rezystora, zmieniając ustawienie przełącznika z pozycji A na B, zauważysz, że prąd ponownie przepływa przez obwód, ale teraz w przeciwnym kierunku. Jeśli płynie prąd, musi również istnieć napięcie. Jest bardzo mało prawdopodobne, że napięcie to powstaje w rezystorze R lub w amperomierzu A. Oczywiście wnioskiem jest to, że tylko zbiornik z płytkami ołowianymi i kwasem siarkowym musi być odpowiedzialny za powstanie napięcia w obwodzie. Zauważasz, że prąd po pewnym czasie staje się mniejszy, by w końcu opaść do zera.

Wnioski Plantégo

Planté był w stanie wyciągnąć następujące wnioski:

- Najwyraźniej pojemnik wypełniony rozcieńczonym kwasem siarkowym, w którym zawieszono dwie ołowiane płytki, jest w stanie „przechowywać” energię elektryczną. Konstrukcja przypomina więc nieco butelkę lejdejską (kondensator), która również może to robić.
- Jednak energia elektryczna pozostaje zmagazynowana w pojemniku przez bardzo długi czas. To duża różnica w porównaniu z butelką lejdejską, z której zmagazynowana energia elektryczna znikała bardzo szybko.
- Zmagazynowana energia elektryczna przejawia się w postaci napięcia, które można zmierzyć między dwiema płytkami ołowianymi.
- Zmagazynowana energia elektryczna może być ponownie wykorzystana w dowolnym momencie poprzez podłączenie płytek ołowiowych do odbiornika.
- Proces ten można powtarzać wielokrotnie.

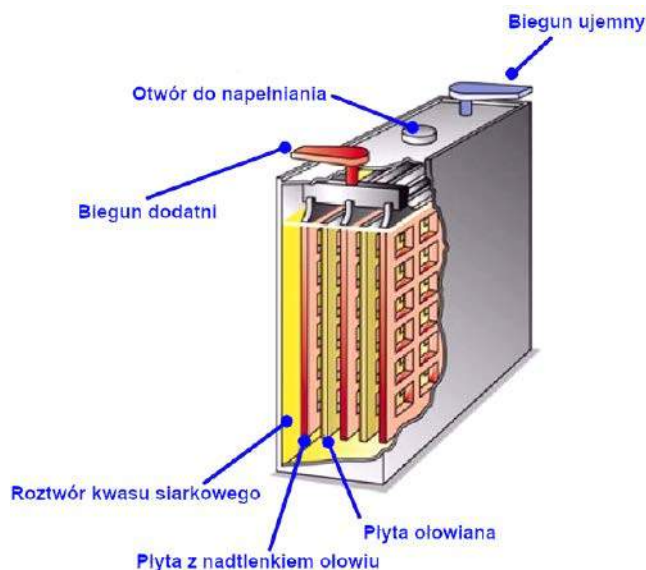
Wspaniały wynalazek z praktycznymi ograniczeniami

W zasadzie Planté dokonał wspaniałego wynalazku. Niestety, w tamtym czasie niewiele można było z nim zrobić. W końcu nie ma sensu magazynować energię elektryczną generowaną przez pierwotne ogniwo elektrochemiczne w ten uciążliwy sposób. Jeśli energia elektryczna była gdzieś potrzebna, znacznie łatwiej było wykonać od razu ogniwo pierwotne, takie jak stos Volty. Ładowanie akumulatora za pomocą ogniwa elektrochemicznego to konwersja energii, która jest bardzo mało użyteczna. Dopiero gdy znacznie później wynalazcy tacy jak Siemens i Gramme wynaleźli dynamo, odkrycie Plantégo zyskało naprawdę istotne praktyczne zastosowanie.

Nowoczesny akumulator kwasowo-ołowiowy

Sto pięćdziesiąt lat udoskonaleń technologicznych

To niewiarygodne, ale nowoczesny akumulator kwasowo-ołowiowy nadal działa na tej samej zasadzie, co pierwsza wersja akumulatora Plantégo. Ponad sto pięćdziesiąt lat rozwoju technologii przyniosło oczywiście niezbędne udoskonalenia. Przykładowo, ogniwo nowoczesnego akumulatora już nigdy nie będzie składać się z dwóch zwykłych płyt ołowianych. Stosowana jest konstrukcja warstwowa, patrz rysunek poniżej, składająca się z wielu arkuszy ołowiu, które są naprzemiennie połączone równolegle. W ten sposób zwiększa się efektywna powierzchnia obu elektrod, bez uszczerbku dla wymiarów ogniwa. Co więcej, nie stosuje się czystego ołowiu. Płyty wykonane są ze stopu ołowiu i antymonu, który jest znacznie bardziej odporny na działanie kwasu siarkowego. Jeden zestaw płyt (biegun ujemny) jest pokryty warstwą bardzo mocno rozdrobnionego czystego ołowiu Pb, drugi zestaw (biegun



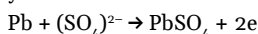
Budowa ogniwa nowoczesnego akumulatora kwasowo-ołowiowego
(© www.aljequestions.nl)

dotatni) jest pokryty cienką warstwą nadtlenku ołowiu PbO_2 . Dzięki takiej konstrukcji zwiększa się żywotność i wydajność akumulatora.

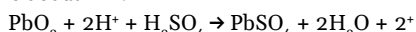
Konwersja energii

Energia elektryczna powstaje w procesie przekształcenia ołowiu Pb i nadtlenku ołowiu PbO_2 w siarczan ołowiu $PbSO_4$, oczywiście w wyniku dysocjacji elektrolitycznej kwasu siarkowego H_2SO_4 . Podczas rozładowywania akumulatora kwasowo-ołowiowego zachodzą następujące procesy.

Na biegunie ujemnym:



Na biegunie dodatnim:



Odkłada się siarczan ołowiu $PbSO_4$. Odwrotna sytuacja ma miejsce podczas ładowania akumulatora. Siarczan ołowiu utworzony na obu płytach jest przekształcany z powrotem w ołów z jednej strony i nadtlenek ołowiu z drugiej. Dlatego to zdysocjowane jony kwasu siarkowego H^+ i $(SO_4)^{2-}$ zapewniają transport elektronów przez roztwór, zarówno podczas ładowania, jak i rozładowywania. Zjawisko to zostało wyraźnie przedstawione na poniższym rysunku.

Kwasowość

Kwasowość, czyli stężenie kwasu siarkowego w akumulatorze kwasowo-ołowiowym, jest ważnym parametrem, ponieważ wskazuje na stan naładowania akumulatora. Wzory chemiczne pokazują, że podczas



Kwasomierz lub areometr, za pomocą którego można sprawdzić stan naładowania akumulatora kwasowo-ołowiowego (© www.gerstaecker.nl)

rozładowywania kwas siarkowy jest przekształcany w siarczan siarki. Im bardziej akumulator jest rozładowany, tym niższe jest stężenie kwasu siarkowego w roztworze. Kwasowość będzie zatem spadać w miarę rozładowywania akumulatora. Na rynku dostępne są niewielkie wskaźniki (patrz rysunek powyżej), za pomocą których można łatwo zmierzyć kwasowość akumulatora kwasowo-ołowiowego. Jeśli dużo pracujesz z akumulatorami, na przykład do zasilania łodzi lub kampera, zalecamy zakup takiego wskaźnika, zwanego areometrem. Kwasowość daje znacznie lepsze wskazanie stanu akumulatora niż zwykły pomiar napięcia na zaciskach.

Areometr

Areometr to nic innego jak miernik, za pomocą którego można zmierzyć ciężar właściwy cieczy. Przy całkowicie rozładowanym akumulatorze ciężar właściwy roztworu kwasu siarkowego wynosi $1,14 \text{ g/cm}^3$. Wartość ta wzrasta do $1,28 \text{ g/cm}^3$, gdy akumulator jest w pełni naładowany (w przypadku akumulatorów bezobsługowych użycie areometru jest niemożliwe z powodu braku korków dla poszczególnych ogniwa – przyp. tłum.).

Ładowanie i rozładowywanie akumulatora kwasowo-ołowiowego

Krzywa ładowania

Akumulatory kwasowo-ołowiowe charakteryzują się wieloma właściwościami, z których najważniejsze są krzywe ładowania i rozładowania. Krzywa ładowania wskazuje, jak napięcie na zaciskach ogniwa wzrasta w funkcji czasu ładowania. W przypadku akumulatora kwasowo-ołowiowego krzywa ładowania ma typowy kształt pokazany na poniższym rysunku. Całkowicie rozładowane ogniwo akumulatora kwasowo-ołowiowego ma napięcie na zaciskach $2,1 \text{ V}$. Podczas ładowania akumulatora napięcie ogniwa wzrasta dość szybko do $2,2 \text{ V}$, a następnie bardzo powoli wzrasta do $2,3 \text{ V}$. Następnie napięcie szybko wzrasta do $2,7 \text{ V}$, gdzie można zauważyć, że w ogniwie występuje duża produkcja gazu, a elektrolit nie jest już przezroczysty, ale mleczny. Napięcie końcowe $2,7 \text{ V}$ jest napięciem końca ładowania. Dalsze ładowanie nie ma sensu, a nawet jest bardzo szkodliwe dla żywotności ogniwa.

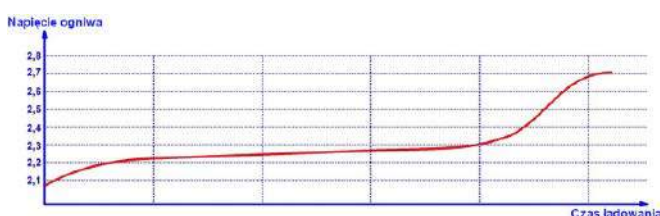
Cechy w pełni naładowanego ogniwa

W pełni naładowane ogniwo ma następujące cechy:

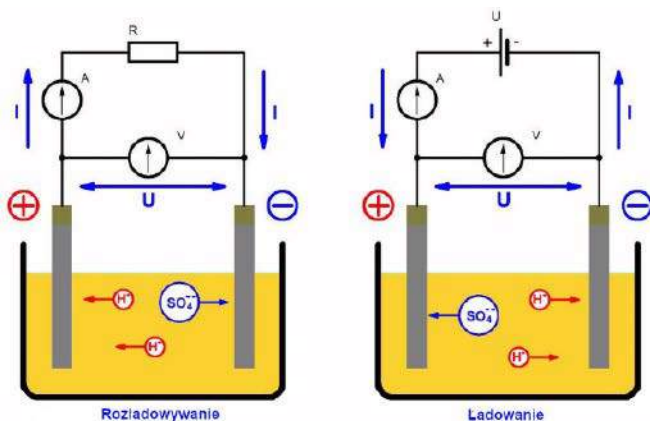
- Napięcie końcowe ogniwa osiąga wartość $2,7 \text{ V}$ i jej nie przekracza.
- Elektrody dodatnie nabierają kasztanowego koloru.
- Elektrody ujemne są szare.
- Następuje obfite wydzielanie się wodoru.
- Elektrolit staje się mleczny.
- Gęstość elektrolitu wzrosła do $1,28 \text{ g/cm}^3$.

Rozładowywanie akumulatora kwasowo-ołowiowego

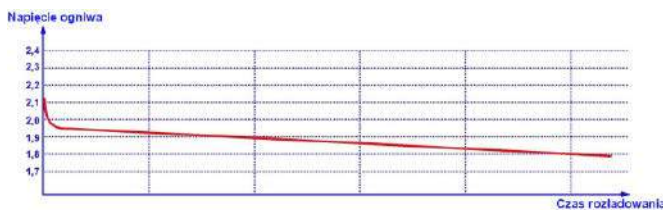
Gdy prąd ładowania zostanie wyłączony, napięcie końca ładowania $2,7 \text{ V}$ spadnie bardzo szybko do około $2,1 \text{ V}$. Jeśli ogniwo akumulatora jest obciążone, napięcie ogniwa spadnie bardzo szybko do $1,9 \text{ V}$. Następnie



Krzywa ładowania ogniwa akumulatora kwasowo-ołowiowego
(© 2018 Jos Verstraten)



Transport jonów przez roztwór podczas ładowania (po prawej) i rozładowywania (po lewej) akumulatora kwasowo-ołowiowego (© 2018 Jos Verstraten)



Krzywa rozładowania ogniwa akumulatora kwasowo-ołowiowego
(© 2018 Jos Verstraten)

napięcie spada stopniowo do 1,8V. Nazywa się to napięciem końca rozładowania ogniwa. Dalsze rozładowywanie nie ma sensu, a nawet jest niebezpieczne dla żywotności ogniwa. Dane te można podsumować jako krzywą rozładowania, narysowaną na poniższym rysunku.

Charakterystyka akumulatora

Wprowadzenie

Podobnie jak ogniwa pierwotne, ogniwa wtórne również posiadają charakterystyczne parametry. Część z nich pokrywa się z parametrami ogniw pierwotnych, ale istnieją również pewne różnice.

SEM

Jest to napięcie, które można zmierzyć na dwóch biegunach nieobciążonego ogniwa. W przypadku ogniw ołowiowych napięcie to wynosi 2,7 V przy pełnym naładowaniu i 1,8 V przy pełnym rozładowaniu. Wartość SEM jest całkowicie niezależna od rozmiaru ogniwa.

Rezystancja wewnętrzna

Rezystancja wewnętrzna ogniw wtórnych jest bardzo niska. Wartość rezystancji wewnętrznej zależy w dużym stopniu od powierzchni płyt. Typowy akumulator kwasowo-ołowiowy ma rezystancję wewnętrzną wynoszącą zaledwie 10 mΩ! Oczywiście jest zatem, że zwarcie w pełni naładowanego akumulatora skutkuje bardzo dużym prądem zwarciovym, który natychmiast powoduje rozgrzanie do czerwoności nawet dość grubych miedzianych przewodów!

Prąd ładowania

Prąd ładowania to nominalna wartość prądu ładowania i rozładowania. Zasadniczo do ładowania stosuje się maksymalny prąd 1 A na kg

ołowiu. Przy rozładowywaniu można tę wartość podwoić. Prąd ładowania można też wyznaczyć opierając się na pojemności (patrz następny punkt). Prądy ładowania i rozładowania są wtedy ustawione jako równe jednej dziesiątej pojemności. Akumulator kwasowo-ołowiowy o pojemności 48 Ah można zatem ładować i rozładowywać prądem 4,8 A. Przekroczenie tej wartości przez dłuższy czas spowoduje wypaczenie elektrod (maksymalny prąd rozładowania może być ponad sto razy większy od wartości nominalnej, pod warunkiem, iż nie trwa to zbyt długo, przykładowo akumulator 4 Ah może dostarczyć przez kilka sekund do 50 A bez uszkodzenia – przyp. tłum.).

Pojemność

Pojemność to ilość energii elektrycznej, wyrażona w amperogodzinach, którą w pełni naładowany akumulator może dostarczyć od momentu rozpoczęcia rozładowywania do momentu, gdy napięcie ogniwa spadnie do 1,8 V. Pojemność jest iloczynem prądu rozładowania i liczby godzin rozładowania. Można zatem rozładowywać akumulator o pojemności 500 Ah przez maksymalnie 10 godzin prądem o natężeniu 50 A. Pojemność zależy od ilości aktywnego materiału płyt w akumulatorze, a więc w rzeczywistości od powierzchni płyt, ale także od kwasowości elektrolitu.

Wydajność

Wydajność to stosunek energii pochłoniętej podczas ładowania do energii uwolnionej podczas rozładowywania. Oczywiście więcej energii jest pochłanianie niż uwalnianie, więc sprawność musi być w każdym przypadku niższa niż 1. Jednak akumulatory mogą mieć bardzo wysoką sprawność, wartości między 0,85 a 0,95 nie są czymś wyjątkowym. Gdyby nie fakt, że większość akumulatorów wykorzystuje substancje chemiczne, które nie są zbyt przyjazne dla środowiska, akumulatory te byłyby idealnymi magazynami energii dla ogniw fotowoltaicznych (w praktyce 99% akumulatorów kwasowo-ołowiowych podlega recyklingowi, czyniąc je najczęściej recyklingowanymi produktami na świecie – przyp. tłum.). ■

Jos Verstraten



Akumulatory

Rozwiązanie znajdziesz na www.elportal.pl/quizy

1. Pierwszy akumulator stworzył:

- ☐ Planté;
- ☐ Volta;
- ☐ Ampère;

2. Pierwszy akumulator używał płyt zrobionych z:

- ☐ miedzi i cynku;
- ☐ miedzi i ołowiu;
- ☐ tylko ołowiu.

3. Akumulator ten używał kwasu:

- ☐ azotowego;
- ☐ siarkowego;
- ☐ solnego.

4. Maksymalne napięcie na ogniwie akumulatora kwasowo-ołowiowego może wynosić:

- ☐ 2,7 V;
- ☐ 2,3 V;
- ☐ 1,8 V.

5. Minimalne napięcie na ogniwie akumulatora kwasowo-ołowiowego może wynosić:

- ☐ 2,7 V;
- ☐ 2,3 V;
- ☐ 1,8 V.

6. Areometr pozwala zmierzyć:

- ☐ pojemność akumulatora;
- ☐ stopień naładowania akumulatora;
- ☐ stopień zużycia akumulatora.

7. Areometr robi to przez pomiar:

- ☐ przejrzystości elektrolitu;
- ☐ lepkości elektrolitu;
- ☐ gęstości elektrolitu.

8. Prąd ładowania powinien wynosić:

- ☐ 1/2 pojemności akumulatora;
- ☐ 1/10 pojemności akumulatora;
- ☐ 1/20 pojemności akumulatora.

9. Mimo że akumulator kwasowo-ołowiowy jest znany od ponad 160 lat, na początku nie był używany, ponieważ:

- ☐ elektrolit był zbyt drogi;
- ☐ ryzyko eksplozji wodoru było zbyt duże;
- ☐ jedynym źródłem energii i tak były ogniwa elektrochemiczne.

10. Współczesne akumulatory uzyskują lepszą wydajność i żywotność, gdyż jedna z elektrod w ogniwie pokryta jest:

- ☐ nadtlakiem ołowiu;
- ☐ siarczanem ołowiu;
- ☐ tlenkiem siarki.

Aneng AN870, multimetr 19999 – test multimetru



Dzięki wyświetlaczowi 4½ cyfry, rzeczywistej dokładności $\pm 0,05\%$ przy pomiarze napięcia DC i cenie poniżej 130 złotych, ten duży multimetr jest więcej niż wart obszernego testu.

Wprowadzenie do multimetru Aneng AN870

Alternatywne marki, nazwy i ceny

Ten multimetr jest oferowany pod różnymi markami i numerami typów:

- Aneng AN870
- Zetek ZT219 (prawdziwy producent?)
- Richmeters RM219

Ceny różnią się znacznie, a mianowicie od 17,54 € do 52,14 € (poziom cen styczeń 2021). Należy jednak w tym miejscu poczynić komentarz. Multimetr jest identyczny niezależnie od oznaczeń dostawców. To, co różni się dość znacznie, to liczba sond pomiarowych i akcesoriów, które są dołączone przez różnych dostawców. Niektóre tanie oferty zawierają tylko jeden zestaw przewodów pomiarowych. My kupiliśmy nasz egzemplarz jako Aneng AN870 na stronie Banggood za cenę 30,54 € z bardzo obszernym zestawem przewodów pomiarowych (ten sam model dostępny jest też na portalu Allegro.pl, AliExpress.com, oraz w różnych specjalistycznych sklepach internetowych – przyp. tłum.).

Co można zmierzyć za pomocą AN870?

Multimetr ten posiada bardzo rozbudowany zestaw funkcji pomiarowych, które można wybrać za pomocą dziesięciopozycyjnego przełącznika obrotowego oraz przycisków „SELECT” i „Hz/%”:

- napięcie stałe: 19,999 mV, 199,99 mV
- napięcie stałe: 1,9999 V, 19,999 V, 199,99 V, 1,000.0 V
- napięcie zmienne: 19,999 mV, 199,99 mV.
- napięcie zmienne: 1,9999 V, 19,999 V, 199,99 V, 750.0 V
- prąd stały: 199,99 μ A, 1 999,9 μ A
- prąd stały: 19,999 mA, 199,99 mA
- prąd stały: 1,9999 A, 19,999 A
- prąd zmienny: 199,99 μ A, 1,999,9 μ A
- prąd zmienny: 19,999 mA, 199,99 mA
- prąd zmienny: 1,9999 A, 19,999 A
- rezystancja: 199,99 Ω , 1,9999 k Ω , 19,999 k Ω , 199,99 k Ω
- rezystancja: 1,9999 M Ω , 19,999 M Ω , 199,99 M Ω
- pojemność kondensatora: 9,999 nF, 99,99 nF, 999,9 nF, 9,999 μ F, 99,99 μ F, 999,9 μ F, 9,999 mF
- częstotliwość: 99,99 Hz, 999,9 Hz, 9,999 kHz, 99,99 kHz, 999,9 kHz, 9,999 MHz
- wypełnienie: 99%
- temperatura: $+1,000^{\circ}\text{C}$

Pomiary są wykonywane trzy razy na sekundę.

Oczywiście można również testować diody, badać przewodność do 50 Ω , wykrywać fazę i punkt neutralny napięcia sieciowego oraz lokalizować przewody pod napięciem w ścianach.

AN870 posiada również następujące funkcje:

- pomiar względny,



Paczka z zawartością (© 2021 Jos Verstraten)

- wskazanie wartości minimalnej i maksymalnej,
- funkcja zamrożenia wskazań.

Opakowanie miernika

Multimetr AN870 jest dostarczany w pozbawionym oznaczeń producenta czy dystrybutora, wytrzymałym kartonowym pudełku zawierającym czarne etui, w którym znajduje się miernik i wiele akcesoriów. Dołączona jest także 24-stronicowa instrukcja obsługi w doskonałym języku angielskim. Etui ma wymiary 22 cm na 13 cm i jest zamykane na zamek błyskawiczny.

Co znajduje się w etui?

Wiele, patrz zdjęcie poniżej. Oprócz dwóch w pełni izolowanych przewodów pomiarowych o długości około 80 cm, opakowanie zawiera również dwa prawie tak samo długie przewody pomiarowe do samodzielnego montażu. Są one wyposażone po obu stronach w izolowane złącza z gwintowanymi otworami, do których można wkręcić różne sondy pomiarowe, wtyki bananowe i zaciski krokodylkowe. W opakowaniu znajduje się łącznie czternaście takich wkręcanych elementów.

Wreszcie, po lewej stronie zdjęcia widać termoparę typu K do pomiaru temperatury. Spiralny kabel zakończony jest dwoma dość amatorsko wyglądającymi wtykami bananowymi.

Sam multimetr AN870

Ten multimetr jest dość dużym urządzeniem o wymiarach 180 mm na 90 mm na 45 mm. Z bateriami urządzenie waży 368 gramów. Właściwy miernik jest umieszczony w elastycznej plastikowej osłonie, która jest dostarczana w kolorze czerwonym lub zielonym. Z tyłu tej osłony znajduje się cienka plastikowa płytka, która służy jako podpórka



Akcesoria, które są dostarczane wraz z AN870 (© 2021 Jos Verstraten)



AN 870 w elastycznej plastikowej osłonie (© 2021 Jos Verstraten)

do stawiania miernika na stole. Jest ona dość chwiejna i mogłaby być nieco solidniejsza. Z tyłu miernika znajduje się oddzielna komora baterii, która jest przymocowana do miernika za pomocą jednej śruby. AN870 jest zasilany dwiema standardowymi bateriami AA 1,5 V. Z tyłu osłony znajdują się dwa uchwyty, w które można wpiąć sondy pomiarowe. Nie jest to jednak zbyt przydatne! Miernik należy wyjmować z osłony tylko wtedy, gdy konieczna jest wymiana jednego z dwóch bezpieczników.

Wyświetlacz i przyciski sterujące

Miernik posiada monochromatyczny wyświetlacz LCD o wymiarach 62 mm na 38 mm z „staromodnymi” siedmiosegmentowymi cyframi o wysokości 21 mm. Wyświetlacz nie ma stałego podświetlenia, co jest poważną wadą w przypadku pomiarów w słabo oświetlonych pomieszczeniach. Można jednak podświetlić wyświetlacz za pomocą białej diody LED po lewej stronie, naciskając przycisk „HOLD”, przez ponad dwie sekundy. Rezultat tego działania jest jednak minimalny. Pokrętko ma średnicę 44 mm i jest łatwe w obsłudze. Jednak pokrętko obraca się dość sztywno, więc trzeba trzymać miernik w lewej ręce, gdy obsługuje się pokrętko prawą ręką.

AN870 posiada sześć przycisków z następującymi funkcjami:

- **SELECT** Wybór między prądem stałym a zmiennym; między mV, °C lub °F; między rezystancją, pojemnością lub ciągłością oraz między napięciem zmiennym, częstotliwością lub wypełnieniem sygnału.
- **HOLD** Naciśnij krótko, aby zamrozić wyświetlacz i długo, aby włączyć lub wyłączyć podświetlenie wyświetlacza.
- **RANGE** Naciśnij krótko, aby włączyć ręczny wybór zakresu. Każde krótkie naciśnięcie tego przycisku powoduje wybranie następnego zakresu. Naciśnij dłużej niż dwie sekundy, aby włączyć automatyczny wybór zakresu.
- **REL** Włącza tryb względnego. Bieżący odczyt jest zapisywany w pamięci miernika jako względny punkt zerowy. Wartość ta jest automatycznie odejmowana od wszystkich nowych odczytów. Idealny na przykład do kompensacji rezystancji przewodów pomiarowych podczas pomiaru niskich wartości rezystancji.
- **MAX/MIN** Po pierwszym naciśnięciu tego przycisku na wyświetlaczu pojawi się tylko maksymalna zmierzona wartość. Po drugim



Przyciski sterujące AN870 (© AliExpress)

naciśnięciu wyświetlana jest tylko minimalna zmierzona wartość. Naciśnięcie przycisku przez ponad dwie sekundy powoduje wyjście z trybu MAX/MIN.

- **Hz/%** Gdy przełącznik obrotowy znajduje się w pozycji V~, przełącza się na pomiar częstotliwości lub wypełnienia przebiegu na wejściu.

Tryb automatycznego wyłączania zasilania

Po włączeniu AN870 poprzez obrócenie przełącznika obrotowego z pozycji 'OFF', miernik automatycznie przejdzie w tryb 'Auto Power Off'. Tekst 'APO' pojawi się w lewym górnym rogu wyświetlacza, a miernik wyłączy się po piętnastu minutach bezczynności. Na minutę przed wyłączeniem miernik wyemituje pięć sygnałów dźwiękowych. Jeśli chcesz, aby AN870 był włączony w sposób ciągły, musisz nacisnąć i przytrzymać przycisk „SELECT” podczas włączania miernika.

Elektronika w AN870

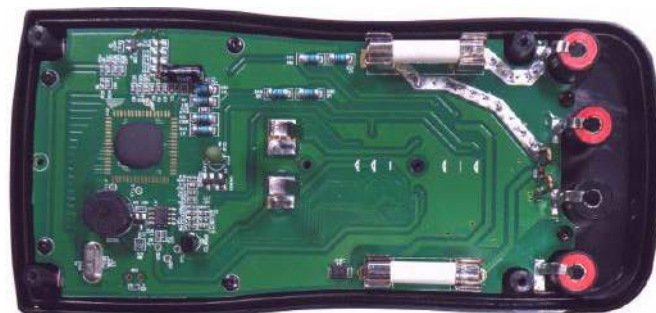
Płytką drukowaną

Po zdjęciu (z pewnym wysiłkiem) osłony, z tyłu widoczne będą cztery małe śruby krzyżakowe. Po odkręceniu tych śrub można otworzyć obudowę miernika. Płytką drukowaną jest precyzyjnie wykonana i zawiera wszystkie części.

To, co od razu rzuca się w oczy, to dwa duże szklane bezpieczniki wypełnione piaskiem, o wymiarach 6,3 mm × 32 mm, które są bezpośrednio podłączone do dwóch złączy pomiarowych 20 A i mA/μA na froncie. Bezpieczniki te mają napięcie przebicia 250 VAC i zapewniają, że w razie ich przypadkowego przepalenia nie dojdzie do przebicia napięcia w bezpieczniku. Jeden bezpiecznik ma wartość 200 mA, a drugi 20 A. Pomiędzy dwoma bezpiecznikami jest miejsce na uchwyt baterii, można zobaczyć styki na środku płytki drukowanej. Rezystory bocznikowe do pomiarów prądu mają wartość 0,01 Ω dla zakresu 20 A (R33) i 1,0 Ω dla zakresu mA/μA (R24). Rezystor 0,01Ω i bezpiecznik 20A są podłączone do złączy na płycie czołowej za pomocą szerokich ścieżek PCB po obu stronach płytki. Co więcej, ścieżki te są pokryte grubą warstwą cyny. Złącze do pomiaru napięcia połączone jest do układu przez dwa rezystory 5 MΩ połączone szeregowo (R29 i R30). Wejście napięciowe miernika jest zabezpieczone przed nadmiernymi napięciami za pomocą termistora PTC (PTC1) i dwóch tranzystorów (Q3 i Q4). Na płycie drukowanej znajduje się pamięć EEPROM typu P24C02A (IC1). Pamięć ta służy prawdopodobnie do przechowywania danych kalibracyjnych. W końcu nowoczesne multi-metry nie są już regulowane za pomocą potencjometrów montażowych, ale za pomocą wartości cyfrowych, których procesor używa do korygowania zmierzonych wielkości. Q5 to napięcie referencyjne 1,2 V typu ICL8069. Dwa tranzystory Q1 i Q2 są prawdopodobnie odpowiedzialne za sterowanie brzęczykiem i podświetleniem wyświetlacza.

Złącza pomiarowe

Cztery złącza pomiarowe zostały skonstruowane w najprostszy sposób, przez wygięcie kawałka metalu w cylindryczny kształt, mając nadzieję, że sprężystość takiej konstrukcji zagwarantuje dobry



Płytką drukowaną multimetru AN870 (© 2021 Jos Verstraten)

kontakt z wtykami przewodów pomiarowych w dłuższej perspektywie. Naszym skromnym zdaniem mogłoby to być nieco bardziej profesjonalne! Oczywiście styki te są wpasowane w plastikowe izolatory, gdy płytka drukowana i obudowa są ze sobą skręcone, ale długoterminowa stabilność jest nadal wątpliwa.



Dość niezgrabna konstrukcja złączy pomiarowych (© lygte-info.dk, pod redakcją Josa Verstratena)

Aneng AN870 na stanowisku testowym

Pomiary napięcia stałego na zakresach mV

AN870 ma dwa zakresy mV: 19,999 mV oraz 199,99 mV. Dokładność wg specyfikacji wynosi $\pm(0,05\% + 3)$. Do generowania tak małych napięć używamy cyfrowo regulowanego zasilacza podłączonego do dzielnika napięcia 1:100 zbudowanego na rezystorach. Nasz najnowszy nabytek, multimetr Fluke 8842A jest używany jako miernik referencyjny. Jest on połączony równolegle z AN870. Zasilanie jest regulowane do momentu, aż zmierzone napięcia na Fluke będą jak najbardziej zbliżone do pożądanych wartości testowych 10 mV, 20 mV, 50 mV i 100 mV. Wyniki są podsumowane w poniższej tabeli i mówią jasnym językiem. Doskonała wydajność AN870!

Pomiary napięcia stałego na zakresach V

AN870 ma cztery zakresy: 1,9999 V, 19,999 V, 199,99 V oraz 1000,0V. Dokładność wg specyfikacji dla tych pomiarów również wynosi $\pm(0,05\% + 3)$. Poniższa tabela porównuje wskazanie na AN870 ze wskazaniem na Fluke 8842A jako miernika wzorcowego. Wyniki są nieco gorsze niż przy pomiarze napięcia na zakresie mV, ale przy średnim odchyleniu 0,068% nie ma powodu do niezadowolenia. Te gorsze wyniki mogą oczywiście wynikać również z faktu, iż przy większości napięć testowych aktywne były tylko cztery z pięciu cyfr.

Zakresy prądu stałego do 2 A

AN870 ma 3x2 zakresy prądu: 199,99 μ A i 1,999,9 μ A; 19,999 mA i 199,99 mA; 1,9999 A i 19,999 A. Dokładność wg specyfikacji wynosi

$\pm(0,5\% + 3)$. AN870 jest podłączony szeregowo z 8842A do zasilacza, który jest ustawiony jako źródło prądu stałego. Ponieważ nasz zasilacz może dostarczyć maksymalnie tylko dwa ampery, a Fluke może mierzyć tylko do 1,99999 A, pomiary są ograniczone do tej wartości. Wyniki zostały podsumowane w poniższej tabeli. Przy średnim błędzie wynoszącym 0,45%, testowany egzemplarz jest zgodny ze specyfikacją.

Pomiar dużych prądów

AN870 może, według producenta, mierzyć prądy do 20A. W instrukcji ani na urządzeniu nie ma informacji, które wskazywałyby na konieczne ograniczenie czasu trwania takich pomiarów. Różne recenzje AN870 twierdzą, że ścieżki na płytce drukowanej i bocznik 20A nie są w stanie wytrzymać tak dużych prądów przez długi czas. Mamy zasilacz 12VDC, który może dostarczyć 30A, idealny do przesyłania tak wysokiego prądu przez AN870! Zasilacz jest obciążony szeregowym połączeniem miernika i dwoma równolegle połączonymi rezystorami 1 Ω /100W. Miernik jest otwarty, a temperatura jest mierzona w sześciu punktach obwodu pomiarowego 20A za pomocą odsłoniętej termopary i pasty przewodzącej ciepło.

Początkowy prąd 20A dość szybko spada do wartości 15A. Jest to wynikiem dodatniego współczynnika temperaturowego rezystorów drutowych. Wyniki tego testu zostały podsumowane na poniższym rysunku. Po dziesięciu minutach temperatura ustabilizowała się i pomiar został przerwany. Zmierzone wartości nie są alarmujące. Trzeba przyznać, że pomiar ten został wykonany przy otwartym mierniku AN870, więc generowane ciepło może być rozproszone znacznie łatwiej niż w przypadku miernika zamkniętego. Test ten udowadnia jednak, że bez problemu można mierzyć tym miernikiem prądy o natężeniu przekraczającym 10A (ale raczej nie przez dwie minuty lub więcej przy zamkniętej obudowie – przyp. tłum.). Warto zauważyć, iż miernik będzie emitował przerywany sygnał dźwiękowy przy pomiarach powyżej 10A.

Zakresy napięcia zmiennego mV

Dostępne są dwa zakresy, 19,999mV i 199,99mV, z dokładnością wg specyfikacji wynoszącą $\pm(0,3\% + 3)$. Ponieważ w praktyce takie zakresy pomiarowe są używane głównie do pomiaru napięć w obwodach małosygnałowych, wykonujemy ten test przy częstotliwości 1kHz. Problem polega na tym, że nie mamy dobrego miernika referencyjnego dla tak małych napięć przemiennych. Nasz cyfrowo regulowany generator funkcyjny UTG9005C-II firmy Uni-T ma dokładność ustawień tylko $\pm 3\%$. Dlatego nie należy traktować wyników podsumowanych w poniższej tabeli jako bezwzględnego wskazania dokładności AN870, ale raczej jako wskazanie, że miernik zapewnia wiarygodne wyniki w tym obszarze.

Zakres napięć zmiennych V

AN870 ma cztery zakresy: 1,9999 V, 19,999 V, 199,99 V oraz 750,0 V z dokładnością $\pm(0,3\% + 3)$. Do testowania tych zakresów używamy autotransformatora (wariaka) i naszego multimetru VC650BT firmy Voltcraft. Ma on jednak dokładność porównywalną z AN870 i dlatego nie

Dokładność pomiaru małych napięć statycznych (© 2021 Jos Verstraten)

Ustawione napięcie	Napięcie zmierzone przez Fluke 8842A	Napięcie zmierzone przez Aneng AN870	Procentowe odchylenie
10 mV	10,0485 mV	10,044 mV	0,044%
20 mV	20,060 mV	20,06 mV	0,00%
50 mV	50,012 mV	50,02 mV	0,016%
100 mV	100,000 mV	100,02 mV	0,02%

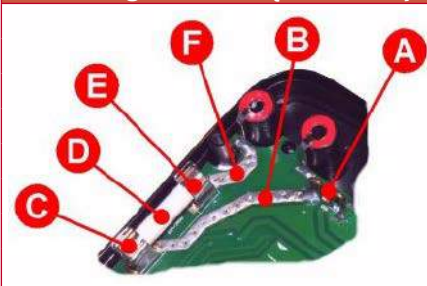
Dokładność podczas pomiaru większych napięć DC (© 2021 Jos Verstraten)

Ustawione napięcie	Napięcie zmierzone przez Fluke 8842A	Napięcie zmierzone przez Aneng AN870	Procentowe odchylenie
2,50 V	2,4999 V	2,502 V	0,084%
5,00 V	5,0134 V	5,014 V	0,012%
7,50 V	7,5027 V	7,508 V	0,070%
10,00 V	10,0061 V	10,015 V	0,089%
20,00 V	20,003 V	20,02 V	0,085%
30,00 V	30,008 V	30,03 V	0,073%

Dokładność pomiaru prądu stałego (© 2021 Jos Verstraten)

Ustawiony prąd	Prąd zmierzony przez Fluke 8842A	Prąd zmierzony przez Aneng AN870	Procentowe odchylenie
10 mA	10,005 mA	9,985 mA	0,20%
20 mA	19,522 mA	19,50 mA	0,18%
50 mA	55,829 mA	55,96 mA	0,23%
100 mA	100,849 mA	100,95 mA	0,10%
200 mA	200,82 mA	199,20 mA	0,81%
500 mA	500,63 mA	497,2 mA	0,69%
1,00 A	1,00005 A	992,3 mA	0,77%
2,00 A	1,99772 A	1,9843 A	0,67%

Pomiar nagrzewania się miernika podczas pomiaru prądu stałego 15 A (© 2021 Jos Verstraten)



Punkt pomiarowy	Temperatura po		
	2. minutach	5. minutach	10. minutach
A	135°C	145°C	147°C
B	75°C	79°C	80°C
C	64°C	72°C	73°C
D	64°C	72°C	73°C
E	68°C	79°C	80°C
F	60°C	71°C	72°C

może być używany jako miernik referencyjny. Dlatego porównaliśmy tylko pomiary w poniższej tabeli bez obliczania odchylenia procentowego.

Szerokość pasma dla sygnałów sinusoidalnych

Instrukcja nie wspomina o tym ważnym parametrze. Zmierzyliśmy szerokość pasma dla sygnału sinusoidalnego o amplitudzie $1 V_{RMS}$. Wyniki przedstawiono w poniższej tabeli. Jeśli zdefiniujemy szerokość pasma jako częstotliwość, przy której odczyt spadł o 3dB lub do współczynnika 0,707, to AN870 wydaje się mieć pasmo pomiarowe o szerokości około 3 kHz.

Dokładność pomiaru wartości skutecznej

Podobnie jak wszystkie nowoczesne mierniki, AN870 mierzy wartość skuteczną napięcia przemiennego. Jest to wartość napięcia przemiennego, które generuje taką samą moc cieplną w rezystorze jak napięcie stałe o tej samej wartości. Poniższa tabela pokazuje, jak AN870

wyświetla wartość skuteczną trzech różnych kształtów sygnału o wartości skutecznej $1 V_{RMS}$ (pomiar taki powinno się przeprowadzić w odniesieniu do wskazań V_{RMS} na oscyloskopie cyfrowym – przyp. tłum.).

Pomiar częstotliwości

AN870 ma sześć zakresów pomiarowych od 99,99 Hz do 9,999 MHz w pełnej skali z dokładnością $\pm(0,1\% + 2)$. Nie jest łatwo sprawdzić zakres pomiarowy częstotliwości multimetru. Zakres ten zależy od wielu czynników, takich jak amplituda sygnału i jego kształt. Aby dać jakieś wyobrażenie, dokonaliśmy pomiaru z sygnałem sinusoidalnym o amplitudzie $1 V_{RMS}$. Wyniki zostały podsumowane w poniższej tabeli.

Pomiar wypełnienia sygnału

Wypełnienie sygnału lub stosunek długości czasów włączenia do wyłączenia sygnału prostokątnego wskazuje, jaki procent czasu trwania okresu stanowi sygnał „H”. Symetryczna fala prostokątna ma zatem wypełnienie wynoszące 50%. AN870 mierzy tę wartość w jednym zakresie od 1% do 99% z dokładnością $\pm(0,1\% + 2)$. Przetestowaliśmy to przy częstotliwościach 1 kHz i 1 MHz, patrz tabela poniżej. Sygnał miał amplitudę 1 V.

Pomiar rezystancji

AN870 ma siedem zakresów rezystancji, od 199,99 Ω do 199,99 M Ω . Dokładność różni się w zależności od zakresu i wynosi od $\pm(0,2\% + 3)$ w najczęściej używanych zakresach do $\pm(5,0\% + 5)$ w najwyższym zakresie. Ponieważ nasz Fluke 8842A mierzy również rezystancję bardzo dokładnie, włączyliśmy ten miernik ponownie i uznaliśmy jego odczyt za 100%. Wyniki zostały podsumowane w poniższej tabeli. Skompensowaliśmy rezystancję przewodów pomiarowych poprzez pomiar

Pomiar małych sygnałów sinusoidalnych o częstotliwości 1 kHz (© 2021 Jos Verstraten)

Sygnał 1kHz sinus z generatora UTG9005C-II	Napięcie zmierzone przez Aneng AN870
10 mV	10,130 mV
20 mV	20,02 mV
50 mV	50,11 mV
100 mV	100,10 mV
199 mV	199,41 mV

Pomiar większych napięć sinusoidalnych o częstotliwości 50 Hz (© 2021 Jos Verstraten)

Napięcie ustawione na autotransformatorze	Napięcie zmierzone przez Fluke 8842A	Napięcie zmierzone przez Aneng AN870
50 V	50,28 V	50,67 V
100 V	100,36 V	100,73 V
150 V	150,28 V	150,78 V
200 V	200,13 V	200,6 V
250 V	250,19 V	250,6 V

Określanie szerokości pasma przy sinusie $1 V_{RMS}$ (© 2021 Jos Verstraten)

Sygnał $1 V_{RMS}$ sinus z generatora UTG9005C-II	Napięcie RMS zmierzone przez Aneng AN870
50 Hz	1,0008 V_{RMS}
1 kHz	1,0065 V_{RMS}
2 kHz	0,9950 V_{RMS}
3 kHz	0,7389 V_{RMS}
4 kHz	0,4321 V_{RMS}
5 kHz	0,2054 V_{RMS}

Dokładność pomiarów wartości skutecznej (© 2021 Jos Verstraten)

Sygnał $1 V_{RMS}$ z generatora UTG9005C-II	Napięcie RMS zmierzone przez Aneng AN870
Sinus	1,0008 V_{RMS}
Prostokąt 50%	0,9290 V_{RMS}
Trójkąt	0,9970 V_{RMS}

Dokładność pomiarów częstotliwości (© 2021 Jos Verstraten)

Sygnał $1 V_{RMS}$ sinus z generatora UTG9005C-II	Częstotliwość zmierzona przez Aneng AN870
10 Hz	9,99 Hz
100 Hz	99,99 Hz
1 kHz	1,001 kHz
10 kHz	9,999 kHz
100 kHz	99,99 kHz
1 MHz	0,9999 MHz
5 MHz	4,999 MHz

względny za pomocą AN870 i przy użyciu metody czteroprzewodowej (Kelvina – przyp. tłum.) z 884,2 A. Przy średnim odchyleniu wynoszącym 0,12%, testowany egzemplarz wypadł lepiej, niż można było się spodziewać.

Pomiar kondensatorów

AN870 mierzy pojemność kondensatorów w siedmiu zakresach od 9,999 nF do 9,999 mF. Dokładność wynosi $\pm(2,0\% + 5)$ w najczęściej używanych zakresach i spada do $\pm(5,0\% + 20)$ w skrajnych zakresach. W przypadku braku dokładnego sprzętu pomiarowego, musimy zadowolić się naszym zestawem kondensatorów z gwarantowaną tolerancją $\pm 1,0\%$. Wyniki zostały podsumowane w poniższej tabeli. W każdym razie odchylenia przy niższych wartościach nie są spowodowane pasywnymi pojemnościami zbyt długich kabli pomiarowych. W takim przypadku zmierzone wartości powinny być wyższe niż wartości kondensatorów.

Pomiar temperatury

Jest kilka rzeczy, na które należy zwrócić uwagę w związku z dostarczoną sondą termopary, patrz rysunek poniżej. Po pierwsze, uważamy, że spiralny przewód jest bardzo niewygodny. Po drugie, nie jesteśmy zadowoleni ze sposobu, w jaki przewód jest podłączony do dwóch wtyczek bananowych. Bardzo cienkie przewody są przykręcone do wtyczek bananowych bez żadnej formy odciążenia. To nie przetrwa zbyt długo!

Wreszcie, sonda wygląda ładnie, ale jest wyjątkowo nieodpowiednia do sprawdzania temperatury komponentu elektronicznego. Lepsza by była tak zwana „odsłonięta termopara” (patrz wstawka), która pochłania znacznie mniej ciepła i mierzy temperaturę, na przykład tranzystora, znacznie dokładniej i szybciej.

AN870 ma pojedynczy zakres temperatur od -20°C do $+1,000^{\circ}\text{C}$ z dokładnością $\pm(2,5\% + 5)$. Nie mamy jeszcze dokładnego miernika temperatury, więc porównaliśmy wyniki AN870 z wynikami naszego TM-902C o identycznej dokładności. W tym teście obie termopary zostały sztywno połączone ze sobą i zanurzone w podgrzewanej wodzie. Poniższa tabela porównuje wyniki obu mierników.

Bezdotykowe wykrywanie napięcia sieciowego

Funkcja NCV przełącznika obrotowego umożliwia bezdotkowy pomiar napięcia sieciowego. W stanie spoczynku na wyświetlaczu pojawia się EF. Trzymając miernik górną częścią przy ścianie, można wykryć położenie przewodu pod napięciem w ścianie. Jeśli AN870 wykryje pole elektromagnetyczne, miernik wyemituje sygnał dźwiękowy, a na wyświetlaczu pojawi się kilka kresek. W miarę zbliżania się do przewodu pod napięciem na wyświetlaczu pojawi się więcej kresek, a miernik będzie emitował szybsze sygnały dźwiękowe. Dokładność tej funkcji nie jest jednak szczególnie wysoka.

Funkcja ta umożliwia również, przynajmniej zgodnie z instrukcją obsługi, wykrywanie fazy w gniazdku ściennym. Podłącz czerwone złącze V/ Ω /Hz jednym przewodem pomiarowym do styków gniazda ściennego. Jeśli przewód znajduje się w otworze zerowym, na wyświetlaczu pozostanie wskazanie EF. Jeśli przewód znajduje się w otworze fazowym, na wyświetlaczu pojawią się cztery kreski, a AN870 wyda szybki sygnał dźwiękowy.

Jeśli chodzi o tego typu pomiary, wolimy polegać na starym testerze napięcia z neonówką!

Nasza opinia na temat Aneng AN870

Nie ma zbyt wielu powodów do krytykowania tego miernika. Producent mógł nieco bardziej profesjonalnie wykonać cztery gniazda pomiarowe i zrobić podpórkę z nieco solidniejszego plastiku. Zmierzone przez nas parametry prawie w całości odpowiadają wartościom podanym przez producenta. Uważamy, że AN870 to wspaniały multimetr dla każdego hobbysty elektroniki! ■

Jos Verstraten



Dołączona do miernika termopara (© AliExpress)

Pomiar wypełnienia przebiegu prostokątnego (© 2021 Jos Verstraten)

Wypełnienie przebiegu prostokątnego 1 V _{pp} z generatora UTG9005C-II	Wypełnienie zmierzone przez Fluke 8842A dla 1 kHz	Wypełnienie zmierzone przez Fluke 8842A dla 1 MHz
1%	1,0%	---
10%	10,0%	3,4%
50%	50,0%	48,5%
90%	90,0%	91,9%
99%	99%	---

Dokładność pomiaru rezystancji (© 2021 Jos Verstraten)

Rezystor testowy 0,1%	Opór zmierzony przez Fluke 8842A	Opór zmierzony przez Aneng AN870	Procentowe odchylenie
10 Ω	10,001 Ω	10,00 Ω	0,01%
100 Ω	100,014 Ω	100,19 Ω	0,47%
1 k Ω	1,00007 k Ω	1,0012 k Ω	0,11%
10 k Ω	9,9962 k Ω	10,007 k Ω	0,11%
100 k Ω	100,082 k Ω	100,21 k Ω	0,12%
1 M Ω	1,00182 M Ω	1,0040 M Ω	0,21%

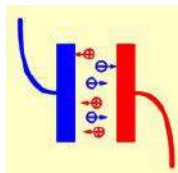
Dokładność pomiaru pojemności kondensatorów $\pm 1\%$ (© 2021 Jos Verstraten)

Kondensator $\pm 1\%$	Pojemność zmierzona przez Aneng AN870
100 pF	85 pF
470 pF	462 pF
1 nF	999 pF
4,7 nF	4,719 nF
10 nF	9,992 nF
47 nF	46,97 nF
100 nF	100,7 nF
1 μF	1,004 μF

Porównanie pomiarów temperatury (© 2021 Jos Verstraten)

Temperatura zmierzona przez TM-902C	Temperatura zmierzona przez Aneng AN870
20,0 $^{\circ}\text{C}$	20 $^{\circ}\text{C}$
40,0 $^{\circ}\text{C}$	40 $^{\circ}\text{C}$
60,0 $^{\circ}\text{C}$	59 $^{\circ}\text{C}$
80 $^{\circ}\text{C}$	79 $^{\circ}\text{C}$
100 $^{\circ}\text{C}$	99 $^{\circ}\text{C}$

Elektrochemia



Elektrochemia to dział chemii fizycznej badający elektryczne aspekty reakcji chemicznych, a także w mniejszym stopniu własności elektryczne związków chemicznych. Dzięki jej odkryciom mogły powstać baterie i akumulatory, oraz różne procesy używane zarówno przez przemysł, jak i samą elektronikę.

Czym jest elektrochemia?

Od energii chemicznej do elektrycznej i z powrotem

Elektrochemia zajmuje się głównie dwoma zjawiskami. Po pierwsze, konwersją energii chemicznej w energię elektryczną, która stanowi podstawę akumulatorów i baterii. Ponadto, konwersją energii elektrycznej w energię chemiczną, co można podsumować terminami elektroliza i galwanotechnika. Dzięki tym technikom możliwe jest rozłożenie wodnych roztworów soli na ich atomy, ale także pokrycie kawałka metalu cienką warstwą innego metalu. Szczególnie ten ostatni proces ma wiele zastosowań; część złotej biżuterii tak naprawdę jest wykonana ze srebra, na które naniesiono cienką warstwę złota. Galwanizacja jest również szeroko stosowana w elektronice. Jedną z metod nakładania cyny na warstwy miedziane, obok najbardziej popularnego procesu mechanicznego (HASL), oraz chemicznego (chemical tin plating) jest właśnie galwanizacja (electroplating). W elektrocynowaniu cyna jest nakładana za pomocą prądu elektrycznego. Jony cyny w roztworze reagują z powierzchnią miedziową pod wpływem prądu, co prowadzi do osadzania się warstwy cyny na powierzchni PCB. Procesy galwaniczne wykorzystuje się również podczas metalizacji otworów w płytkach dwu i wielowarstwowych (płytki PCB stanowi katodę i umieszczana jest w elektrolicie, który zawiera sole miedzi, po czym na skutek przepływu prądu elektrycznego na ściankach otworów jak również na całej jej powierzchni osadza się (dodatkowa) warstwa miedzi). Złącza krawędziowe płytek drukowanych są w ten sam sposób pokrywane cienką warstwą złota. To samo dzieje się ze stykami lepszych typów przełączników i przełączników (i różnych złączy – przyp. tłum.). Krótko mówiąc, elektrochemia ma wiele nieoczekiwanych zastosowań.

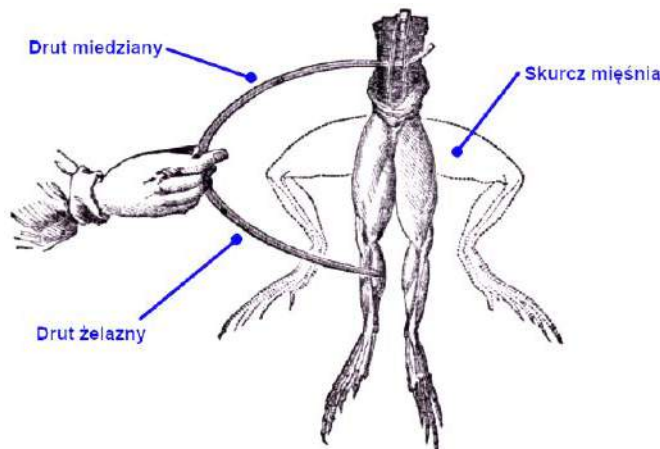
Historia rozwoju elektrochemii

Elektryczność i chemia

Od XVIII wieku wiadomo było, że elektryczność i chemia mają ze sobą wiele wspólnego. Przy okazji doświadczeń z elektrycznością statyczną zaobserwowano, iż elektryczność jest w stanie powodować kurczenie się mięśni. Fakt ten był używany jako uzasadnienie używania elektryczności statycznej w celach terapeutycznych w latach 50tych XVIII wieku. Rażenie prądem z maszyny elektrostatycznej miało leczyć kurcze mięśni i paraliż.

Luigi Galvani

Włoski fizjolog i anatom, Luigi Galvani zajmował się badaniami nad działaniem i rolą nerwów. Eksperymentując z ciałami martwych żab zaobserwował on, iż mięśnie kurczą się pod wpływem prądu. Pozwoliło mu to zbudować prosty, choć czuły instrument wykrywający elektryczność – galwanoskop żabi. W przyrządzie tym mięsień z nogi świeżo zabitej żaby miał przyłączone przewody na obu końcach. Gdy przyłączyło się te przewody do źródła elektryczności, mięsień się kurczył, co potwierdzało jej obecność. Przy okazji tego odkrycia Galvani zaobserwował też, iż mięsień żaby może się kurczyć bez obecności zewnętrznego źródła elektryczności. Rysunek poniżej pokazuje jego słynny



Słynny eksperyment Galvaniego, który stanowi podstawę elektrochemii (© Wikimedia Commons)

eksperyment z nogami żaby, w którym do końców mięśni podłączone były druty wykonane z różnych metali. Po połączeniu ze sobą wolnych końców, mięśnie się kurczą.

Galvani słusznie wywnioskował z tego, że elektryczność jest w jakiś sposób obecna w zamkniętym układzie miedzi, żelaza i mięśni. Galvani jednak wyciągnął z tego błędny wniosek, według którego w zwierzętach istnieje do tej pory nieznaną siłą vitalną, którą nazwał „elektrycznością zwierzęcą”, która nie jest tożsama z „naturalną” elektrycznością (pioruny) czy „sztuczną” elektrycznością wytwarzaną przez człowieka z pomocą maszyn elektrostatycznych. Według jego teorii mózg wytwarza „elektryczny fluid”, który nerwami trafia do mięśni, powodując ich skurcze. Ponadto tkanki według Galvaniego zachowują się analogicznie do butelek lejdejskich (prymitywnych kondensatorów) i gromadzą ten „elektryczny fluid”. Teoria ta była generalnie przyjęta przez środowisko naukowe, z jednym wyjątkiem.

Alessandro Volta

Rodak Galvaniego, Alessandro Volta nie uznawał tego wyjaśnienia. Według niego mięsień był niczym więcej niż instrumentem używanym do określenia obecności elektryczności. Elektryczność była wynikiem kontaktu dwóch różnych metali. W końcu, gdyby eksperyment Galvaniego został przeprowadzony z dwoma przewodami wykonanymi z tego samego metalu, nic by się nie stało. Volta zaczął pracować nad swoją wizją eksperymentalnie i wkrótce odkrył, że do wytworzenia elektryczności nie jest potrzebny żaden materiał zwierzęcy. Jeśli umieścić dwa przewodniki wykonane z różnych metali w roztworze przewodzącym, między dwoma przewodnikami pojawiła się różnica napięcia. W ten sposób Volta był w stanie stworzyć pierwszą prymitywną baterię, która nosi nazwę „ogniwo Volty”, znaną również jako „stos Volty”.

Ogniwo Volty

Ogniwo Volty składało się ze szklanego naczynia wypełnionego rozcieńczonym roztworem kwasu siarkowego. W naczyniu tym zawieszono miedzianą i cynkową płytkę. Volta był w stanie wykazać, że między dwiema płytkami wytworzyło się napięcie około 1,1V. W innym eksperymencie Volta stworzył swój stos, który składał się z wielu dysków ułożonych naprzemiennie: dysk miedziany, dysk z tektury nasączonej solanką lub innym elektrolitem, dysk cynkowy, i ponownie dysk miedziany, tekturowy i cynkowy, i tak dalej. Stos ten nie tylko

zapewniał dużo wyższe napięcie, niż pojedyncze ogniwo, ale też prąd elektryczny w ilości wystarczającej do prowadzenia innych eksperymentów. Stanowiło to poważny problem dla ówczesnych naukowców. Koncepcja „elektronu” była wciąż nieznana, ale wiadomo było, że prąd elektryczny musi składać się z „czegoś”. To „coś” musiało więc również przepływać przez przewodzący roztwór stosu Volty. Minie wiele czasu, zanim zjawisko to zostanie naukowo wyjaśnione.

Zasada działania ogniwa elektrycznego

Struktura materii

Aby zrozumieć zasadę działania ogniwa elektrycznego, nauka musiała nieco lepiej poznać strukturę materii. Gdy ludzie osiągnęli punkt, w którym wiedzieli, że substancje składają się z cząsteczek i atomów oraz, że te atomy mają dodatnie jądro i ujemny ładunek wokół tego jądra przechowywany w elektronach, mogli zacząć badać i wyjaśniać zjawiska zademonstrowane przez Voltę, Galvaniego i im współczesnych badaczy.

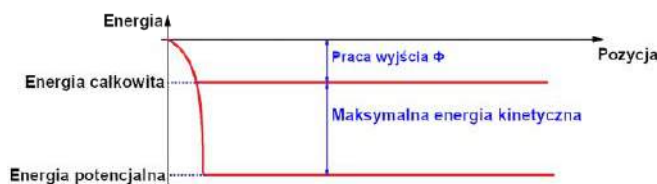
Dysocjacja elektrolityczna

W 1884 roku szwedzki chemik Svante Arrhenius wykazał, że jeśli pewne substancje chemiczne zostaną rozpuszczone w wodzie, roztwór ten może przewodzić prąd elektryczny. Takie substancje zostały nazwane „elektrolitami”, słowem złożonym z greckich pojęć, które można luźno przetłumaczyć jako „ciekłe substancje o właściwościach elektrycznych”. Arrhenius podejrzewał, że przewodzenie to powstało, ponieważ atomy substancji podzieliły się na oddzielne części dodatnie i ujemne, które mogą swobodnie poruszać się w cieczy. Te naładowane cząstki zostały nazwane „jonami”. Ponieważ sam atom jest elektrycznie obojętny, było jasne, że ładunek wszystkich dodatnich jonów w cieczy musi być dokładnie równy ładunkowi wszystkich ujemnych jonów w cieczy, tak aby całość pozostała elektrycznie obojętna. Zjawisko polegające na tym, że substancje, które same w sobie nie są przewodnikami elektrycznymi, stają się nimi po rozpuszczeniu w wodzie, zostało nazwane przez Arrheniusa „dysocjacją elektrolityczną”.

Zachowanie wolnych elektronów w metalach

Drugim zjawiskiem, które należy wyjaśnić, jest zachowanie elektronów w metalach. Jeśli czytałeś poprzednie artykuły z tej serii, wiesz, jak elektrony w substancjach przewodzących, takich jak blok metalu, mogą swobodnie przemieszczać się od atomu do atomu. Dlatego mówi się o „swobodnych elektronach”. Wolny elektron przeskakuje niejako z atomu na atom, tworząc w ten sposób ślad jonów dodatnich. Są to atomy, które utraciły elektron i przez to stały się naładowane dodatnio. Te „dziury” są wypełniane przez inne wolne elektrony, dzięki czemu sam blok metalu jest elektrycznie neutralny.

Ruch jest jednak formą energii. Poruszające się swobodnie elektrony w przewodniku mają zatem pewną energię kinetyczną ruchu. Wielkość tej energii zależy od właściwości przewodnika. W zasadzie jeśli elektron posiada wystarczająco dużo energii, może całkowicie opuścić sieć wiązań atomów. Jednak energia kinetyczna wolnych elektronów w metalu jest zbyt mała, aby umożliwić elektronom ucieczkę z metalu. Bryła metalu tworzy niejako pułapkę dla wolnych elektronów, z której nie ma ucieczki. Jest to możliwe tylko wtedy, gdy dodatkowa energia jest w jakiś sposób dostarczona wolnym elektronom lub jeśli energia

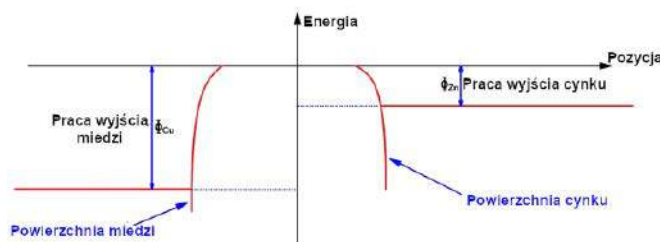


Energia swobodnego elektronu w funkcji jego położenia względem krawędzi metalu (© 2017 Jos Verstraten)

niezbędna do opuszczenia metalu jest w jakiś sposób obniżona. W typowym metalu deficyt energii do ucieczki jest dość mały, zaledwie kilka elektronowoltów (eV). Jednak wolne elektrony pozostają w metalu, ponieważ nie ma możliwości pozyskania tej dodatkowej energii. Energię tę nazywa się „pracą wyjścia” i oznacza się grecką literą Φ . Poniższy rysunek przedstawia graficzną reprezentację energii wolnego elektronu w kawałku metalu. Wyraźnie widać, że energia ta dąży do zera, gdy zbliża się do krawędzi metalu.

Zjawisko Seebecka

Fakt, że wolne elektrony w różnych metalach mają różne energie, ma bardzo daleko idące konsekwencje. Załóżmy, że dwa metale są połączone ze sobą bardzo sztywno, na przykład cynk i miedź. Oba metale mają różne wartości Φ , co skutkuje różnicą energii w miejscach, w których atomy cynku i miedzi stykają się ze sobą. Wyjaśnia to poniższy rysunek.



Zjawisko Seebecka wynika z różnicy energii między wolnymi elektronami w dwóch metalach (© 2017 Jos Verstraten)

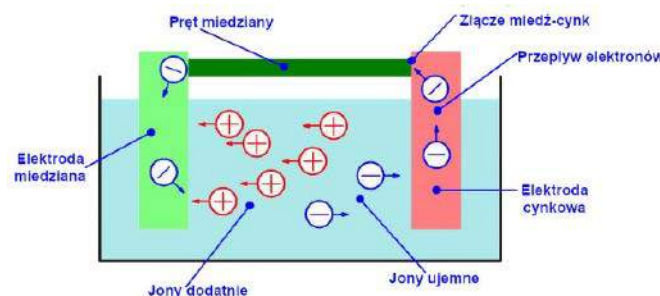
Powstaje różnica energii około 1eV, w wyniku czego wolne elektrony w cynku mogą teraz łatwiej uciec z siatki wiązań cynku i przenieść się do siatki wiązań miedzi. Przepływ elektronów z cynku do miedzi następuje na styku obu metali. Po krótkim czasie od połączenia dwóch metali w układzie ustala się równowaga. Występuje wtedy niedobór ładunku ujemnego na powierzchni cynku i niedobór ładunku dodatniego na powierzchni miedzi. Całość wygląda wtedy trochę jak naładowany kondensator. Pomiędzy dwoma metalami powstaje różnica ładunków, a tym samym różnica napięć.

Ogniwo Volty

Dzięki teoretycznym pracom Arrheniusa i opisanym postępom w myśleniu o wolnych elektronach, możliwe było naukowe wyjaśnienie zasady działania ogniwa Volty. Nowoczesna wersja ogniwa Volty składa się, patrz rysunek poniżej, z naczynia wypełnionego elektrolitem, którym jest rozcieńczony kwas siarkowy, w którym zanurzone są miedziana i cynkowa płytki. Płytki te nazywane są „elektrodami”.

Wyjaśnienie działania ogniwa Volty

Założmy, że obie płytki są połączone miedzianym prętem. Następnie na płytce cynkowej powstaje punkt, w którym miedź i cynk są ze sobą ściśle połączone – „złącze miedź-cynk”. Spowoduje to ucieczkę elektronów z cynku do miedzi. Elektrony te mogą jednak swobodnie rozprzestrzeniać się po całej powierzchni miedzi



Budowa nowoczesnego ogniwa Volty (© 2017 Jos Verstraten)

za pośrednictwem miedzianego pręta, dzięki czemu część miedzianej płytki, która jest zanurzona, również wykazuje nadwyżkę elektronów. Jednocześnie płytka cynkowa zanurzona w cieczy będzie wykazywać nadmiar jonów dodatnich. Występuje zatem różnica napięć między dwoma płytkami. Będziesz teraz musiał poradzić sobie z opisaną już dysocjacją elektrolityczną. Atomy kwasu siarkowego obecne w wodzie są poddawane różnicy napięć i dzielą się na jony dodatnie i ujemne. Jony dodatnie są jednak przyciągane przez elektrony na płycie miedzianej. Jony ujemne z cieczy są przyciągane do płytki cynkowej. Jony dodatnie z cieczy łączą się z elektronami ujemnymi w płycie miedzianej, tworząc atomy. Jony dodatnie absorbują elektrony z płytki miedzianej. To samo dzieje się po stronie cynku. Tutaj jony ujemne z cieczy łączą się z cynkiem, dzięki czemu jony ujemne oddają elektrony płycie cynkowej. W rezultacie układ nie osiąga równowagi, elektrony mogą nadal przepływać przez układ, a mianowicie:

- z cynku do miedzi w miejscu styku tych dwóch metali,
- z miedzi do jonów dodatnich w cieczy,
- od jonów ujemnych w cieczy do płytki cynkowej.

Dlatego ogniwo Volty nadal wytwarza napięcie i prąd, nie tylko przez krótki czas, ale przez dość długi czas.

Wieczysty ruch?

Wygląda na to, że ogniwo Volty to prawdziwe perpetuum mobile. Magiczne urządzenie, którego naukowcy poszukiwali przez wieki, a które produkowałoby energię lub ruch w nieskończoność. Volta również przez chwilę myślał, że wynalazł takie cudowne urządzenie. W końcu: kto może zaprzeczyć, że opisany prąd elektronowy nie płynie wiecznie przez cieć i miedź/cynk? Niestety, tak nie jest i aby to zrozumieć, należy sięgnąć do chemii. Kwas siarkowy jest substancją, której jedna cząsteczka składa się z siedmiu atomów:

- dwóch atomów wodoru H,
- jednego atomu siarki S,
- czterech atomów tlenu O.

Zgodnie z konwencjami chemicznymi, wzór kwasu siarkowego można zapisać jako H_2SO_4 . Prawdą jest, że cztery atomy tlenu mają ściśle wiązanie z jednym atomem siarki. Dwa atomy wodoru są znacznie luźniej związane z całością. W rezultacie jeśli rozcińcemy kwas siarkowy wodą, atomy wodoru mogą bardzo łatwo opuścić wiązanie molekularne. W wyniku dysocjacji elektrolitycznej dwa atomy wodoru opuszczają cząsteczkę, ale ich elektrony pozostaną. W rezultacie powstają opisane już jony dodatnie, które w rzeczywistości składają się z jąder wodoru. Jony ujemne składają się z pozostałych atomów, tj. tlenu i siarki. Konsekwencją dysocjacji elektrolitycznej jest rozpad cząsteczki kwasu siarkowego na dwa dodatnie jony wodorowe H^+ i ujemny jon SO_4^- . Dodatnie jony wodoru migrują do miedzianej płytki i absorbują nadmiar elektronów. W rezultacie ponownie powstaje neutralny atom wodoru. Płytkę miedzianą szybko pokrywa się małymi pęcherzykami gazu, które składają się z czystego wodoru. Jony ujemne migrują do płytki cynkowej i uwalniają tam nadmiar elektronów. Jednak grupa SO_4 natychmiast wiąże się chemicznie z cynkiem płytki, tworząc siarczan cynku $ZnSO_4$. Płytkę cynkową jest niejako „zjadana” i powoli, ale pewnie staje się cieńsza. Proces generowania napięcia nie trwa wiecznie, ale do momentu, gdy cały kwas siarkowy zostanie przekształcony w wodór z jednej strony i siarczan cynku z drugiej. W tym momencie w roztworze nie mogą już powstawać żadne jony, a elektrony nie mogą być przenoszone z miedzi na płytkę cynkową za pośrednictwem opisanego skomplikowanego szlaku elektrochemicznego. Napięcie zostaje utracone, a ogniwo Volty można ożywić jedynie poprzez ponowne wlanie kwasu siarkowego do wody lub/i wymianę płytki cynkowej.

Polaryzacja

W praktyce jednak napięcie ogniwa spada znacznie szybciej. W pewnym momencie nie tylko widoczne, ale także niezliczone niewidoczne mikroskopijne pęcherzyki wodoru całkowicie odcięły miedzianą płytkę od elektrolitu. Płytkę miedzianą jest wtedy niejako odizolowana i nie może być mowy o transporcie elektronów. Zjawisko to nazywane jest „polaryzacją” ogniwa.

Podsumowanie

Jak zatem można wyjaśnić, że ogniwo elektrochemiczne generuje napięcie? Ten złożony proces można przedstawić w następujący sposób:

- Dwa metale o różnych funkcjach roboczych muszą być ze sobą połączone.
- Jednocześnie oba metale muszą być umieszczone w izolowanym pojemniku, w którym znajduje się elektrolit.
- Elektrony przepływają z jednego metalu do drugiego mogą, dzięki dysocjacji elektrolitycznej, przepływać z powrotem z drugiego metalu do pierwszego poprzez elektrolit.
- Elektrony są transportowane przez roztwór elektrolitu za pośrednictwem jonów, które powstają w wyniku rozpadu cząsteczek substancji chemicznej rozpuszczonej w wodzie.
- Jednak ze względu na transport elektronów przez jony, jony te wytrącają się na obu elektrodach.
- Jony te utworzą wiązania chemiczne z metalem jednej z płytek (lub z obiema płytkami), tworząc nowe substancje chemiczne.
- W rezultacie stężenie substancji pierwotnie rozpuszczonej w wodzie powoli, ale zdecydowanie maleje.
- W pewnym momencie cała oryginalna substancja została chemicznie przekształcona i transport elektronów przez jony w roztworze ustaje.
- Ogniwo nie dostarcza już wtedy napięcia elektrycznego.

Szereg napięciowy

Różnica napięć powstaje między dwiema metalowymi płytkami w wyniku dysocjacji elektrolitycznej. Testy wykazały, że wielkość tego napięcia zależy od rodzaju użytych metali. Można więc sporządzić tabelę, w której wszystkie metale są uszeregowane według wielkości napięcia generowanego przez nie w ogniwie Volty. Tabela taka nazywana jest szeregiem napięciowym. Jak zawsze przy definiowaniu wielkości napięcia, należy gdzieś zdefiniować potencjał odniesienia. W tym przypadku używany jest wodór, który jest utożsamiany z potencjałem 0 V. Zakres napięć najbardziej znanych metali i innych substancji podano w poniższej tabeli. Korzystając z tego szeregu napięć, można łatwo obliczyć,

Zakres napięć różnych substancji stosowanych w ogniwach i bateriach (© 2017 Jos Verstraten)

Substancja	Napięcie	Substancja	Napięcie
Złoto	+1,50 V	Kobalt	-0,29 V
Platyna	+0,86 V	Kadm	-0,40 V
Srebro	+0,80 V	Żelazo	-0,44 V
Rtęć	+0,79 V	Chrom	-0,56 V
Węgiel	+0,74 V	Cynk	-0,76 V
Miedź	+0,34 V	Mangan	-1,10 V
Bizmut	+0,28 V	Glin	-1,67 V
Antymon	+0,14 V	Magnez	-2,40 V
Wodór	0,00 V	Sód	-2,71 V
Ołów	-0,13 V	Potas	-2,92 V
Cyna	-0,14 V	Lit	-2,96 V
Nikiel	-0,23 V		

w jaki sposób Volta uzyskał napięcie 1,1 V ze swojego ogniwa. Napięcie miedzi względem wodoru wynosi +0,34 V. Napięcie cynku względem wodoru wynosi -0,76 V. Wystarczy dodać do siebie te dwa napięcia, aby poznać napięcie, które powstaje w ogniwie utworzonym z płytek cynkowych i miedzianych: 1,1 V. Należy teraz zauważyć, że rodzaj elektrolitu i jego stężenie również odgrywają pewną rolę w napięciu wyjściowym, które może dostarczyć ogniwo.

Tworzenie ogniwa elektrochemicznego o największym możliwym napięciu wyjściowym

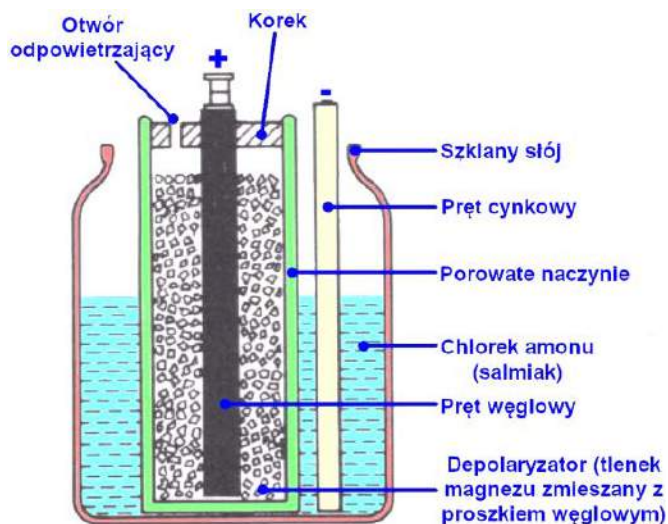
Aby stworzyć ogniwo elektrochemiczne o maksymalnym napięciu wyjściowym, należy wybrać dwa metale, które są jak najbardziej oddalone od siebie w zakresie napięcia i wybrać elektrolit odpowiedni dla tych metali. Ale oczywiście inne czynniki, takie jak cena, toksyczność, dostępność i długoterminowa stabilność również odgrywają bardzo ważną rolę. Zasadniczo możliwe są niezliczone kombinacje płytek i elektrolitów. Oczywiście nie wszystkie kombinacje będą praktycznie użyteczne. Najważniejsze jest znalezienie kombinacji, która:

- jest tania,
- nie zawiera substancji toksycznych,
- jest niezawodna,
- zapewnia dość wysokie napięcie ogniwa,
- nie cierpi zbyt mocno z powodu polaryzacji,
- nie wytwarza łatwopalnych ani toksycznych gazów,
- łatwa i tania w produkcji.

Ogniwo Volty nie kwalifikuje się do miana „praktycznie użytecznego ogniwa” z różnych powodów. Zaszczyt stworzenia takiego ogniwa przypada Georges’owi Leclanché.

Ogniwo Leclanché

Ten Francuz stworzył w 1865 roku ogniwo, którego skład pokazano na poniższym rysunku. Od tamtej pory ogniwo to znane jest jako „ogniwo Leclanché” i nadal stanowi podstawę wielu tanich baterii, które można kupić w każdym supermarkecie. Ogniwo składa się ze szklanego słoika wypełnionego rozcieńczonym roztworem chlorku amonu (salmiaku) jako elektrolitem. Jedną elektrodę tworzy pręt cynkowy, a drugą pręt węglowy. Pręt węglowy nie jest jednak zawieszony bezpośrednio w elektrolicie, lecz w porowatym naczyniu wypełnionym mieszaniną proszku węglowego i tlenku magnezu. Substancje te zapobiegają polaryzacji ogniwa. Po raz kolejny na biegunie dodatnim powstaje wodór. Jednak tlenek magnezu powoduje, iż ten wodór wchodzi w reakcję chemiczną, dzięki



Oryginalna wersja ogniwa Leclanché (© 2017 Jos Verstraten)

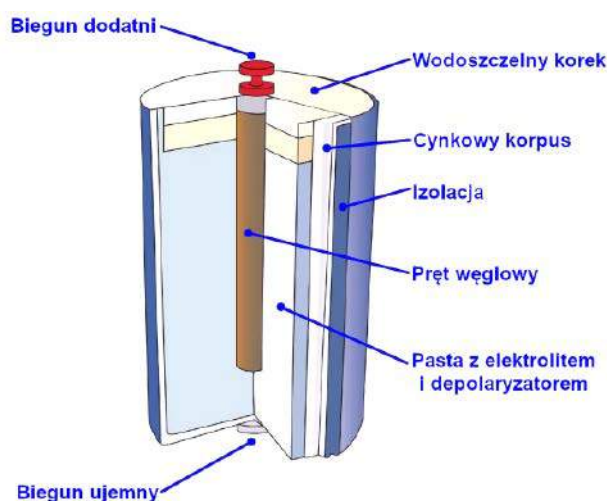
czemu zostaje związany i nie może zaizolować pręta węglowego mikroskopijnymi pęcherzykami gazu. Z szeregu napięć można wywnioskować, że ogniwo Leclanché wytwarza napięcie $+0,74 \text{ V} + -0,76 \text{ V} = 1,5 \text{ V}$. Ogniwo Leclanché było używane wszędzie przez bardzo długi czas. Jediną wadą ogniwa jest to, że depolaryzator działa dość wolno. Jeśli z ogniwa pobierany jest duży prąd, wokół pręta węglowego powstaje tak dużo wodoru, że tlenek magnezu nie jest w stanie go przetworzyć. Dlatego ogniwo polaryzuje się powoli, powodując spadek napięcia. Jeśli jednak ogniwo nie będzie obciążane przez jakiś czas, tlenek magnezu będzie kontynuował swoją pracę i przekształcał powstały wodór, dzięki czemu ogniwo ponownie dostarczy pełne napięcie 1,5 V.

Nowoczesne ogniwa Leclanché

Nowoczesne ogniwa działające zgodnie z tą starożytną zasadą są stosowane na przykład w bateriach płaskich o napięciu 4,5 V oraz powszechnie znanych bateriach AA i AAA (oraz C/R14 i D/R20 – przyp. tłum.) o napięciu 1,5 V. Płaska bateria zawiera trzy ogniwa Leclanché połączone szeregowo, co wyjaśnia napięcie wyjściowe 4,5 V (bateria PP3/6F22 używa sześciu małych ogniw połączonych szeregowo do uzyskania napięcia 9 V, jej niewielki rozmiar jest przyczyną niskiej pojemności – przyp. tłum.). Oczywiście dostarczenie ogniwa z cieczą jako elektrolitem jest dość niepraktyczne. Można jednak użyć wodnistej pasty, która ma właściwość niepełnienia. Budowę takiego nowoczesnego ogniwa pokazano na poniższym rysunku. Kubek cynkowy tworzy teraz zarówno obudowę ogniwa, jak i biegun ujemny. Pręt węglowy jest ponownie używany jako biegun dodatni. Elektrolit jest tworzony przez zmieszanie roztworu salmiaku ze środkami wiążącymi, do czego wcześniej używano trocin. W ten sposób tworzy się dość stałą pastę, która nadal jest wystarczająco wodnista, aby zachodziła dysocjacja elektrolityczna. Takie ogniwo nazywane jest „suchym”, ponieważ elektrolit występuje w postaci pasty.

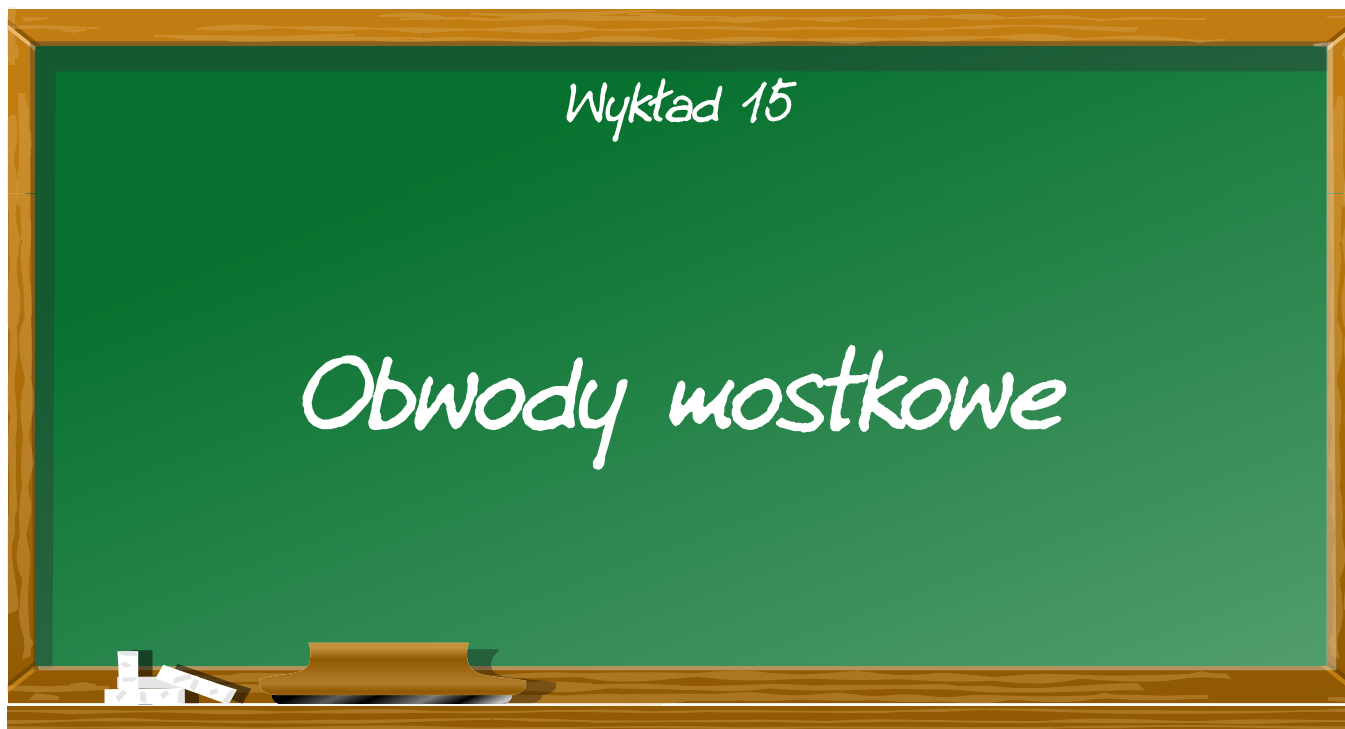
Czytelnikom zainteresowanym głębszym poznaniem omawianych w artykule zagadnień pragnę polecić dwie książki polskiego chemika i popularyzatora nauki, Stefana Sękowskiego: „Galwanotechnikę domową” i „Elektrochemię domową”. Obie książki nie tylko dokładnie zgłębiają poruszane tu tematy, ale demonstrują też praktyczne zastosowania prezentowanej w nich wiedzy. Elektroników-amatorów może zainteresować na przykład proces cynowania z prądem lub bez prądu – przyp. tłum. ■

Jos Verstraten



Budowa nowoczesnego „suchego” ogniwa Leclanché, znanego z baterii 1,5 V AA i AAA (i innych rozmiarów) (© 2017 Jos Verstraten)

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.



Wprowadzenie do koncepcji obwodu mostkowego

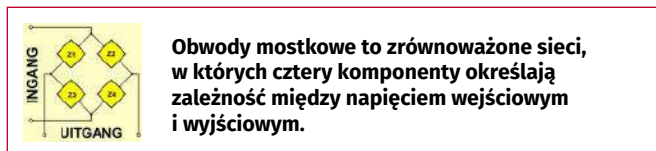
Podstawowy schemat mostka

Poniższy rysunek przedstawia ogólny schemat obwodu mostka. Obwód jest zwykle rysowany w kształcie rombu (lewy schemat), ale można go również narysować tradycyjnie, używając tylko poziomych i pionowych linii (prawy schemat). Cztery bloki Z1, Z2, Z3 i Z4 mogą reprezentować wszelkiego rodzaju komponenty elektroniczne, takie jak rezystory, kondensatory, cewki, diody, tranzystory lub przełączniki.

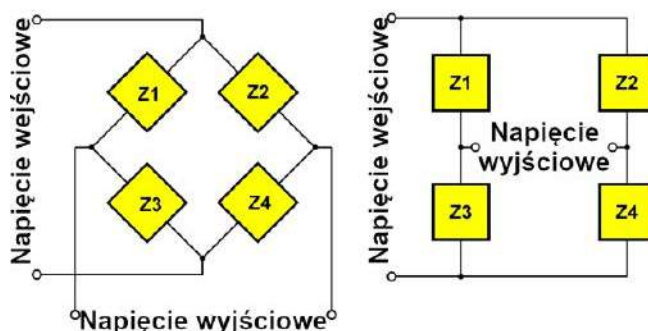
Pływające wejścia i wyjścia

W większości układów elektronicznych wejście ma stronę „gorącą” i „uziemioną”. To samo dotyczy wyjścia. Zarówno wejście, jak i wyjście mają zatem jeden wspólny punkt (zazwyczaj wspólną masę). Ułatwia to przetwarzanie sygnałów w takim obwodzie. Pomiar w takim obwodzie jest również prosty, ponieważ większość przyrządów pomiarowych ma również wejście „gorące” i „uziemiające”.

Inaczej jest w przypadku obwodu mostkowego. Jedną z podstawowych właściwości obwodu mostkowego jest to, że zarówno wejście, jak i wyjście pływają w odniesieniu do masy obwodu. Na schemacie wyraźnie widać, że na wejściu i wyjściu nie ma wspólnego połączenia. Dlatego podczas pomiarów w obwodach mostkowych za pomocą uziemionych przyrządów należy zwrócić szczególną uwagę na tę podstawową właściwość obwodu mostkowego. Na przykład, jeśli podłączysz generator fali sinusoidalnej do wejścia mostka, a oscyloskop do wyjścia, to z definicji zwiernasz jeden z czterech bloków Z1, Z2, Z3 lub Z4!



Obwody mostkowe to zrównoważone sieci, w których cztery komponenty określają zależność między napięciem wejściowym i wyjściowym.



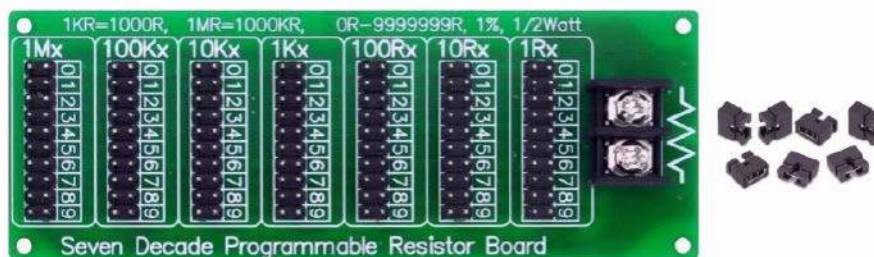
Ogólny schemat mostu (© 2022 Jos Verstraten)



Dwie dekady rezystorowe (© 2022 Jos Verstraten)

Mostek w stanie równowagi

Szereg obwodów mostkowych służy do dokładnego pomiaru wartości rezystora, kondensatora lub cewki indukcyjnej. W takich mostkach mierzona część jest umieszczana, na przykład, w bloku Z3, a blok Z1 jest zastąpiony, na przykład, bardzo dokładnie ustawioną rezystancją. Zaciski wyjściowe są podłączone do bardzo czułego miernika. Zaciski wejściowe są podłączone do źródła napięcia. W zależności od zastosowania może to być napięcie stałe lub przemienne. Regulowany rezystor Z1 jest następnie obracany, aż między dwoma zaciskami wyjścia nie będzie napięcia lub sygnału. Mówi się wtedy, że mostek jest „w równowadze” lub „zbalansowany”. Na podstawie wartości rezystancji w Z1 można następnie obliczyć wartość mierzonego elementu w Z3 za pomocą zwykle prostego wzoru.



Tania alternatywa dla klasycznej dekady rezystorowej (© AliExpress)

Dekada rezystorowa

Do takich pomiarów „mostka w równowadze” zawsze potrzebna jest dekada rezystorowa. Za pomocą takiego przyrządu można bardzo dokładnie ustawić rezystancję, czasami nawet z dokładnością do $\pm 0,1\%$. Poniższe zdjęcie przedstawia dwie wersje takiego urządzenia. W lewej należy utworzyć żądaną rezystancję za pomocą przełączników suwakowych, aby włączyć lub wyłączyć kombinacje wartości 1-2-3-4, 10-20-30-40 itp. Prawa wersja składa się z szeregu dziesięciopozycyjnych przełączników obrotowych, które pozwalają szybko wybrać żądaną rezystancję.

Takie dekady rezystorowe są dość drogie. Tania alternatywa jest sugerowana na zdjęciu powyżej. Za pomocą takiej „dekady rezystorowej na PCB” można zrobić to samo, ale przy nieco większym wysiłku i mniejszej dokładności. Taka płytko drukowana składa się z siedmiu sekcji, a każda z tych sekcji składa się z kolei z dziewięciu rezystorów o takiej samej rezystancji, połączonych szeregowo, na przykład 1 k Ω i takiej samej liczby podwójnych listew kołkowych PCB (goldpin – przyp. tłum.). Za pomocą zworki można zerwać jedną z par kołków na listwie, na przykład włączając pięć z dziewięciu rezystorów i wybierając całkowitą rezystancję 5 k Ω .

Uwaga! Na rynku dostępne są również takie dekady z kondensatorami zamiast rezystorów.

Galwanometr do mostków DC

Drugim niezbędnym przyrządem podczas wykonywania pomiarów „mostków w równowadze” jest bardzo czuły miernik, który może mierzyć zarówno dodatnie, jak i ujemne napięcia lub prądy. Można oczywiście użyć do tego multimetru cyfrowego. Ponieważ jednak trzeba wyregulować mostek do minimalnej wartości prądu na wyjściu, a prąd ten może być zarówno dodatni, jak i ujemny, przydatne jest użycie miernika igłowego z punktem zerowym pośrodku skali. Wtedy można bardzo wyraźnie zobaczyć, jak podczas regulacji mostka igła miernika powoli, ale pewnie osiąga środek skali. Taki miernik nazywany jest „galwanometrem” i można go kupić za jedno do dwóch euro na znanych chińskich platformach sprzedażowych. Pamiętaj, aby kupić taki, który jest wewnętrznie chroniony przed przeciążeniem przez dwie diody!

Alternatywa dla mostków napięcia przemiennego

Równoważenie mostka za pomocą galwanometru jest możliwe tylko wtedy, gdy mostek ten jest zasilany napięciem stałym. Istnieją jednak również mostki, które muszą być zasilane napięciem zmiennym, na przykład w celu pomiaru cewki lub kondensatora. Nie można wtedy użyć galwanometru i należy wtedy użyć czułego miernika prądu przemiennego lub napięcia przemiennego. Tanim rozwiązaniem jest użycie czułych słuchawek. Następnie należy zasilć mostek napięciem przemiennym o częstotliwości na przykład 1 kHz i wyregulować mostek tak, aby dźwięk ze słuchawek był minimalny. Rozpoczynając pomiar od małego napięcia, upewnij się, że nie przeciążysz uszu!

Rodzaje obwodów mostkowych

Podczas naszego obszernego przeglądu literatury fachowej, będącego źródłem tego artykułu, napotkaliśmy nie mniej niż piętnaście obwodów mostkowych. Niektóre z nich znasz już pod inną nazwą i często używasz ich, nie podejrzewając, że należą do rodziny układów mostkowych. Inne są przestarzałe, bardzo rzadkie lub/i stosowane w skrajnych przypadkach, i rzadko lub nigdy nie spotkasz się z nimi w swojej codziennej praktyce elektronicznej.

W kolejnych rozdziałach omówimy:

- Mostek Graetza,
- Mostek H,
- Wzmacniacz mostkowy,
- Mostek Wheatstone’a,
- Mostek Kelvina, zwany też mostkiem Thomsona,
- Mostek Maxwella,
- Mostek Wiena,
- Mostek Sauty’ego,
- Mostek Scheringa,
- Mostek Murraya,



Galwanometr jest przydatny podczas wykonywania pomiarów zrównoważenia mostka z napięciem stałym (© AliExpress)

- Mostek Carey-Fostera,
- Mostek Andersona,
- Mostek Hay'a,
- Mostek Fontany,
- Mostek kratowy.

Mostek Graetza

Funkcja mostka Graetza

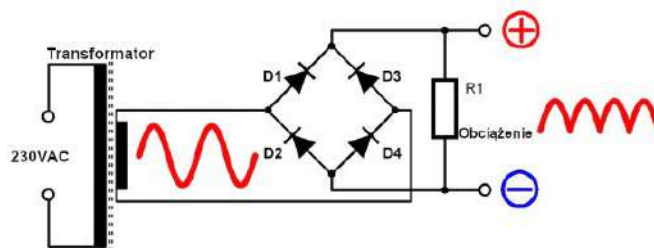
Znasz ten obwód mostkowy aż za dobrze pod nazwą „mostek prostowniczy”. Jest to dobrze znany obwód, który służy do prostowania wtórnego napięcia przemiennego transformatora sieciowego.

Wynalezienie mostka Graetza

Obwód ten został po raz pierwszy użyty w grudniu 1895 roku przez polskiego inżyniera elektryka Karola Pollaka. Dwa lata później został on również wynaleziony i opatentowany przez niemieckiego fizyka Leo Graetza, bez żadnej wiedzy o polskim obwodzie. Dlatego obwód ten wszedł do historii pod jego nazwiskiem.

Budowa mostka Graetza

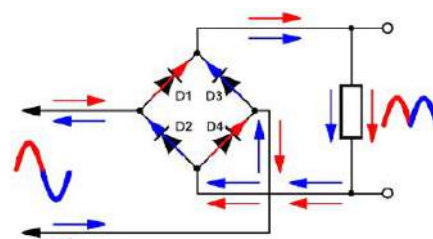
Poniższy rysunek przedstawia skład i działanie mostka Graetza w obwodzie prostownika. Obwód składa się z czterech identycznych diod od D1 do D4. Są one połączone parami równolegle między dwoma przyłączami uzwojenia wtórnego transformatora. Dwa równoległe łańcuchy składają się z dwóch diod połączonych przeciwsośnie. Dwie diody D1 i D3, które są połączone ze sobą katodami, zasilają dodatni biegun wyprostowanego napięcia. Dwie diody D2 i D4, które są połączone ze sobą anodami, zasilają ujemny biegun napięcia stałego.



Mostek Graetza (© 2022 Jos Verstraten)

Przepływ prądu

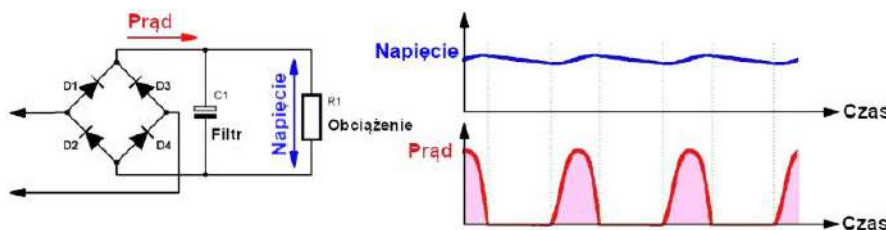
Poniższy rysunek przedstawia przepływ prądu przez mostek Graetza przy dodatnim (czerwonym) i ujemnym (niebieskim) półokresie napięcia na transformatorze. Przy dodatnim półokresie diody D1 i D4 zaczynają przewodzić, a przy ujemnym półokresie przewodzą diody D3 i D2. W obu przypadkach prąd płynie w tym samym kierunku przez rezystor obciążenia. Stąd mówi się, że na wyjściu mostka występuje napięcie **stałe**, choć nie jest ono stałe, lecz cyklicznie zmienia się od zera do maksimum i znowu do zera sto razy na sekundę (przy zasilaniu transformatora z sieci 50 Hz – przyp. tłum.).



Przepływ prądu przez obwód (© 2022 Jos Verstraten)

Wyglądanie

Mostek Graetza dostarcza napięcie stałe na wyjściu, ale napięcie to nie może być używane do zasilania układów. Dlatego mostek jest zakończony dużym kondensatorem elektrolitycznym C1. Służy on jako zbiornik ładunku i utrzymuje napięcie wyjściowe mniej więcej na stałym poziomie wartości szczytowej napięcia sinusoidalnego minus spadek napięcia na dwóch diodach prostowniczych w stanie przewodzenia. Z oscylogramów na poniższym rysunku wynika, że diody nie zawsze przewodzą prąd. Elementy te zaczynają przewodzić tylko wtedy, gdy napięcie dostarczane przez transformator staje się większe niż napięcie na kondensatorze. Wtedy przez diody płynie prąd (fragmenty zakreślone na różowo) by w pełni naładować kondensator C1. W tym momencie napięcie na kondensatorze rośnie, co pokazuje ślad niebieski.



Wyglądanie wyprostowanego napięcia transformatora (© 2022 Jos Verstraten)

Mostek Graetza w praktyce

Prostownik można zmontować z czterech oddzielnych diod, na przykład typu 1N4007. Istnieje jednak kilka wariantów gotowych mostków Graetza. Poniższe zdjęcie daje wyobrażenie o dostępnych wersjach.

Mostek H

Funkcja mostka H

Obwodu tego można użyć do odwrócenia kierunku prądu płynącego przez silnik elektryczny. Dlatego też obwód ten jest bardzo często wykorzystywany w układach stosowanych w robotyce. Mostek ten nie zawdzięcza swojej nazwy osobie, ale sposobowi rysowania go na schematach – przypomina swoją formą literę H.



Kilka przykładów mostków Graetza (© 2007 Rainer Knäpper, Free Art License)

Skład i działanie mostka H

Poniższy rysunek przedstawia podstawowy układ mostka H. Cztery przełączniki S1, S2, S3 i S4 mogą być dowolnymi elementami przełączającymi: przekaźnikami, tranzystorami bipolarnymi, tranzystorami FET i MOSFET. Ze względu na ich doskonałe właściwości przełączania, obecnie stosuje się głównie tranzystory MOSFET. Komponenty te są sterowane czterema sygnałami U1 do U4, które są binarne. Mają one wartość „L” lub „H”. Przy „H” odpowiedni przełącznik jest zamknięty. Możliwe są następujące kombinacje:

- U1 i U4 'H': silnik obraca się zgodnie z ruchem wskazówek zegara,
- U2 i U3 'H': silnik obraca się w kierunku przeciwnym do ruchu wskazówek zegara,
- U1 i U3 'H': silnik jest zwarty i hamuje,
- U2 i U4 'H': silnik jest zwarty i hamuje.

Brak lub tylko jeden sygnał „H”: silnik znajduje się w położeniu neutralnym.

UWAGA: Niebezpieczeństwo zwarcia!

Będzie jasne, że sytuacje:

- U1 i U2 „H”
- U3 i U4 „H”

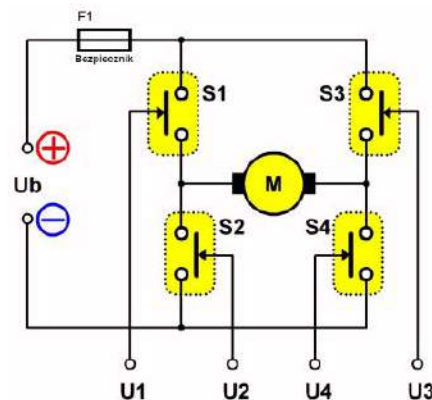
nigdy nie mogą wystąpić. Te kombinacje powodują zwarcie między plusem i minusem napięcia zasilania. Dlatego na schemacie znajduje się bezpiecznik F1. Oczywiście w praktyce mostek H nie jest sterowany czterema sygnałami, ale tylko dwoma. Zazwyczaj dekodowanie od dwóch do czterech sygnałów odbywa się w układzie scalonym. Dzięki pewnemu wewnętrznemu układowi logicznemu w praktyce można zapobiec wyżej wspomnianym sytuacjom zwarciovym.

Mostek H z tranzystorami bipolarnymi

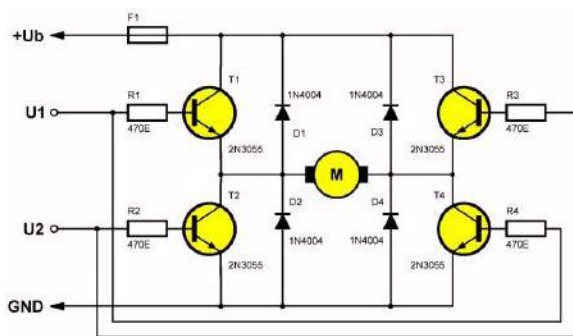
Poniższy rysunek przedstawia schemat, w którym silnik M jest sterowany czterema tranzystorami mocy 2N3055 „końmi roboczymi elektroniki”. Gdy U1 ma wartość „H”, tranzystor T1 jest doprowadzany do stanu nasycenia przez rezystor R1, a T4 przez rezystor R4. Prąd przepływa od lewej do prawej przez silnik. Gdy U2 przyjmuje wartość „H”, to samo dzieje się z tranzystorami T2 i T3. Prąd płynie wtedy od prawej do lewej strony silnika. W tym obwodzie nie może wystąpić sytuacja U1 = U2 = „H”, ponieważ wtedy cztery tranzystory będą przewodzić i wystąpią dwa zwarcia między +Ub i GND. Schemat na czterech tranzystorach bipolarnych typu NPN, stanowiący przedruk z oryginału, wystarcza do przedstawienia zasady działania mostka H natomiast nie oferuje najlepszej efektywności sterowania silnikiem. Na pozycjach T1 i T3 należałoby zastosować tranzystory typu PNP, albo zmienić sposób sterowania mostkiem (tranzystory PNP, odwrotnie niż ma to miejsce w przypadku tranzystorów NPN, „załącza się” stanem niskim na bazie) albo dodać dodatkowe tranzystory odwracające logikę sterowania tranzystorami PNP.

Mostek H z nowoczesnymi tranzystorami N-MOSFET

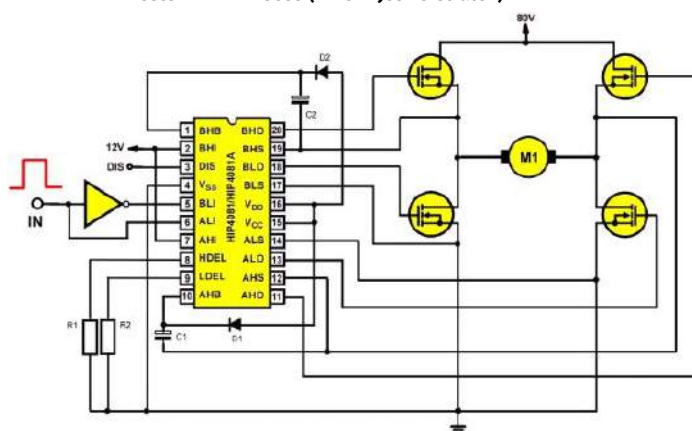
Oczywiste jest, że tranzystory N-MOSFET są idealnymi komponentami do budowania takiego mostka. Niezwykle niska rezystancja przewodzenia takich tranzystorów zapewnia niewielkie straty mocy i niewielką dodatkową powierzchnię chłodzenia. Jednakże, jeśli chcesz zbudować taki mostek z czterech tranzystorów N-MOSFET, będziesz miał problem z tym, że dwa z tych tranzystorów muszą sterować obciążeniem po stronie wysokiej mostka, a dwa z nich muszą sterować po stronie niskiej. Półprzewodniki strony wysokiej muszą otrzymywać napięcie sterujące na bramce, które jest wyższe niż napięcie źródła. Na szczęście opracowano również kompletne układy scalone sterowników do tego celu, takie jak HIP4081A firmy Renesas, które można kupić za około 5,00 € i które można stosować przy napięciach zasilania do +15 V. Schemat zalecany przez producenta pokazano na poniższym rysunku. Wewnętrzna budowa układu zapewnia, że MOSFETy high side i low side wyłączają się i włączają z niewielkim opóźnieniem względem siebie. W ten sposób zapobiega się bardzo dużym prądom udarowym.



Podstawowy obwód mostka H (© 2022 Jos Verstraten)



Mostek H z 4x2N3055 (© 2022 Jos Verstraten)



Mostek H ze sterownikiem HIP4081A i 4xN-MOSFET (© 2022 Jos Verstraten)

Wzmacniacz mostkowy

Funkcja wzmacniacza mostkowego

Dzięki takiemu układowi można, przynajmniej w teorii, czterokrotnie zwiększyć maksymalną moc generowaną przez dany głośnik przy danym napięciu zasilania.

Podstawowy schemat wzmacniacza mostkowego

Poniższy rysunek przedstawia połączenie między głośnikiem a wzmacniaczem mostkowym mocy. W tym układzie głośnik znajduje się pomiędzy dwoma identycznymi stopniami wyjściowymi. Główną zaletą tej konfiguracji jest to, że maksymalna moc jest wysyłana do głośnika z dostępnego napięcia zasilania. Dlatego też ten obwód mostkowy można znaleźć głównie we wzmacniaczach mocy zasilanych niskim napięciem, takich jak wzmacniacze samochodowe, które muszą zadowolić się zasilaniem 12 V lub wzmacniacze zasilane z portu USB 5 V. Wadą jest to, że do dwóch stopni wyjściowych trzeba wysłać sygnały, które mają przeciwne fazy. Jeśli sygnał U1 na jednym wejściu wzrasta dodatnio, to sygnał U2 na drugim wejściu musi spaść ujemnie o tę samą wartość. Trzeba więc wstawić dodatkowy stopień, który obróci fazę sygnału wejściowego o 180 stopni.

Głośnik nie do masy!

Drugą wadą jest to, że głośnik nie jest podłączony do masy. Oba połączenia tego elementu przenoszą napięcie sygnału względem masy. W codziennej praktyce nie będzie to przeszkadzać. Jest to problem, jeśli chcesz dokonać pomiaru w układzie mostkowym. W przypadku zwykłego wzmacniacza z głośnikiem podłączonym do masy, można podłączyć generator sinusoidalny do wejścia i oscyloskop do wyjścia. W końcu oba instrumenty mają również jedno połączenie wspólne do masy. Jeśli zrobisz to samo ze wzmacniaczem mostkowym, wszystko pójdzie nie tak. Sondą oscyloskopu zwierasz jedno z dwóch połączeń głośnikowych do masy (patrz wstęp: „Pływające wejścia i wyjścia”).

Maksymalna moc w normalnym wzmacniaczu

Głośniki mają impedancję 4, 8 lub 16 Ω . Maksymalna moc, jaką można wygenerować w takim głośniku, zależy od maksymalnego skutecznego napięcia sygnału, jakie wzmacniacz może wysłać do głośnika. Maksymalne efektywne napięcie sygnału zależy od dostępnego napięcia zasilania. Załóżmy, że masz głośnik o impedancji 4 Ω , jak pokazano na poniższym rysunku. Wzmacniacz mocy ma napięcie zasilania ± 20 V. Aby działał prawidłowo, na tranzystorach mocy powstanie spadek napięcia około 4 V. Oznacza to, że amplituda U_{amp} napięcia przemiennego na głośniku wynosi maksymalnie 16 V. Wartość tę należy przeliczyć na napięcie skuteczne:

$$U_{eff} = \frac{U_{amp}}{1,41}$$

$$U_{eff} = \frac{16V}{1,41}$$

$$U_{eff} = 11,35V$$

Moc generowana w rezystorze jest dana wzorem:

$$P = \frac{(U_{eff})^2}{R}$$

$$P = \frac{(11,35)^2}{R}$$

$$P = 32,2W$$

Większa moc nie jest możliwa, niezależnie od rodzaju obwodu!

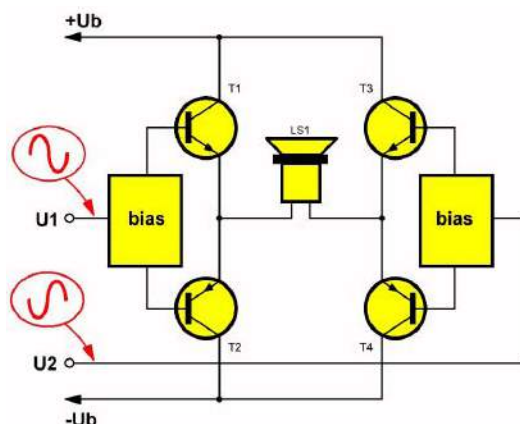
Maksymalna moc we wzmacniaczu mostkowym

Teraz obliczymy, jaką moc można wygenerować w tym głośniku za pomocą wzmacniacza mostkowego o identycznych napięciach zasilania. Schemat został przedstawiony na poniższym rysunku. Wyraźnie widać, że napięcie między obydwoma połączeniami głośnika jest teraz dwukrotnie wyższe niż w poprzednim przykładzie. Dzieje się tak, ponieważ połączenie głośnika (na schemacie kolor zielony) nie jest teraz podłączone do masy, ale również przenosi sygnał w odniesieniu do tej masy. Oba sygnały są w przeciwfazie. Jeśli czerwone złącze głośnika ma napięcie +16 V, to zielone złącze ma napięcie -16 V. Maksymalne napięcie na głośniku wynosi więc 32 V. Odpowiada to wartości skutecznej 22,69 V. Można teraz obliczyć moc rozpraszaną w głośniku: 128,7 W!

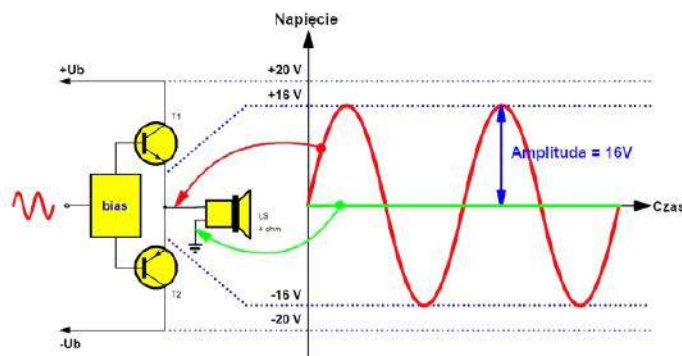
To rzeczywiście, jak stwierdzono we wstępie, cztery razy więcej niż w przypadku tradycyjnego wzmacniacza mocy.

Wzmacniacz mostkowy z dwoma układami TDA2030

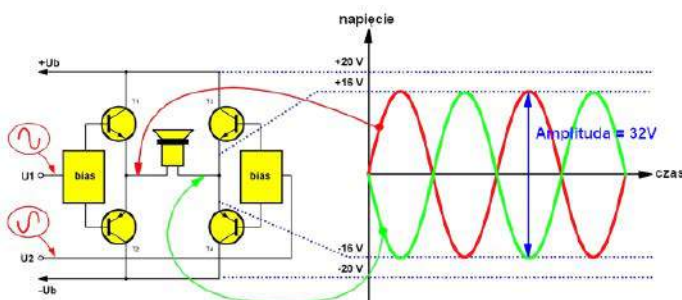
Dzięki wzmacniaczom różnicowym na wejściu tego układu scalonego nie jest to takie trudne. W końcu można przełączyć jeden układ TDA2030 jako wzmacniacz nieodwracający, a drugi jako wzmacniacz odwracający.



Podstawowy schemat mostka wzmacniacza (© 2022 Jos Verstraten)



Moc głośników w „normalnym” wzmacniaczu (© 2022 Jos Verstraten)



Moc głośników w mostku wzmacniacza (© 2022 Jos Verstraten)

Jeśli oba sąysterowane tym samym sygnałem, wyjścia wygenerują dwa sygnały przesunięte w fazie o 180°, które można wykorzystać do zasilania głośnika. Schemat został przedstawiony na poniższym rysunku. Sygnał wejściowy trafia do układu IC1, który jest podłączony jako wzmacniacz nieodwracający. Niewielka część sygnału w fazie jest podawana na wejście odwracające układu IC2 za pośrednictwem rezystora R7. Oba wzmacniacze mają identyczne elementy sprzężenia zwrotnego:

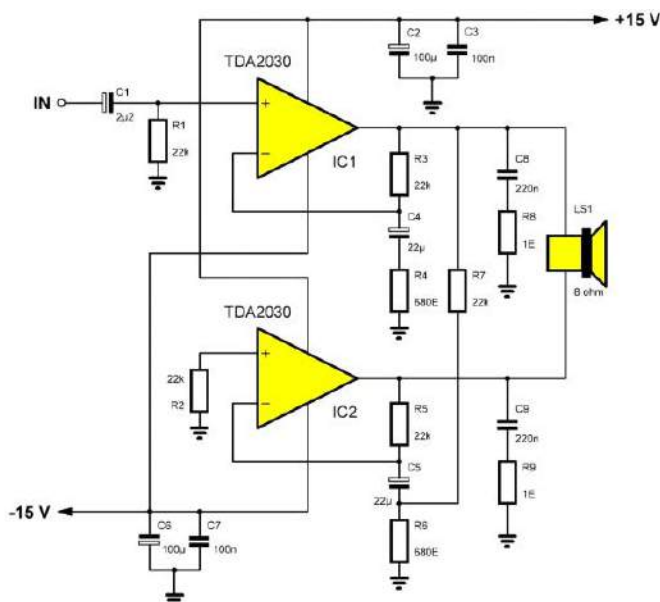
- IC1: R3, C4, R4,
- IC2: R5, C5, R6

i dlatego mają takie samo wzmocnienie. Sygnały, które są w przeciwnej fazie pojawiają się na dwóch wyjściach.

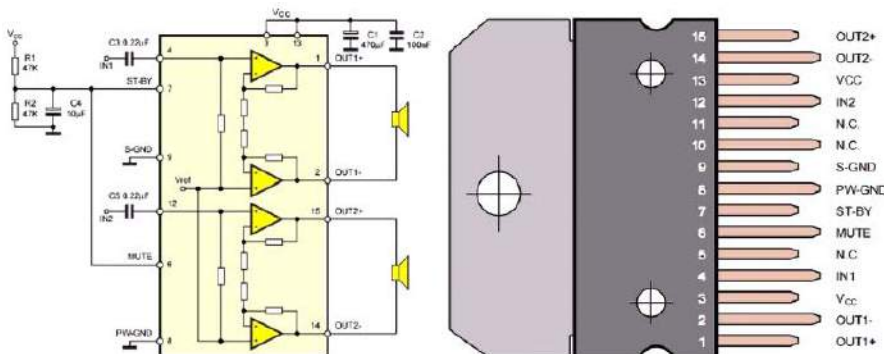
Wzmacniacz mostkowy z układem TDA7266

Oczywiście opracowano również kompletne obwody mostkowe w jednym układzie scalonym. Dobrym przykładem jest układ TDA7266. Ten układ scalony zawiera dwa wzmacniacze mostkowe o maksymalnej mocy 2×7 W dla głośników 8 Ω przy zniekształceniach 10%. Maksymalna moc spada do około 2×2 W przy ograniczeniu dopuszczalnych zniekształceń do 1%. Układ scalony można zasilać napięciem z zakresu od 3 V do 18 V. Poniższy rysunek przedstawia najprostszą aplikację układu TDA7266 wraz z obudową i opisem wyprowadzeń.

TDA7266 jest obecnie dostępny, w tym od Reichelt za 3,30 €. Co ciekawe, kompletny moduł tego wzmacniacza można kupić od różnych dostawców na AliExpress za 2,10 €. Jak pokazuje poniższe zdjęcie, dołączony radiator jest zbyt mały, by skutecznie odprowadzać generowane ciepło, ale jest to oczywiście łatwe do zmiany dla doświadczonego majsterkowicza.



Wzmacniacz mostkowy z 2×TDA2030 (© 2022 Jos Verstraten)



Wzmacniacz mostkowy z 1×TDA7266 (© 2022 Jos Verstraten)

Mostek Wheatstone'a

Funkcja mostka Wheatstone'a

Mostek ten umożliwia pomiar wartości nieznanego rezystora poprzez umieszczenie go po jednej stronie mostka, dwóch znanych rezystorów po dwóch pozostałych stronach i dekadę rezystorowej po czwartej stronie. Dostosowując tę dekadę, aż mostek znajdzie się w równowadze, można obliczyć wartość nieznaną rezystancji za pomocą prostego wzoru.

Wynalezienie mostka Wheatstone'a

Ta metoda pomiaru rezystancji została po raz pierwszy opracowana w 1833 roku przez Samuela Christie, ale do praktycznego użytku została wprowadzona przez Charlesa Wheatstone'a w 1843 roku.

Budowa mostka Wheatstone'a

Poniższy rysunek przedstawia budowę mostka Wheatstone'a. Rezystory R1 i R3 są dokładnie znanymi rezystorami, na przykład o tolerancji ±0,1%. Rezystor R2 to regulowana i czytelna dekada rezystorowa. Rezystancja Rx jest nieznaną rezystancją do zmierzenia.

Mostek Wheatstone'a należy zasilać stabilnym napięciem stałym Ub. W drugiej przekątnej umieść czuły galwanometr M1. Po zamknięciu przełącznika S1 igła galwanometru znajdzie się w jednym z końców skali. Należy obracać przełącznikami rezystorów dekadę, aż igła miernika znajdzie się jak najbliżej pozycji zerowej na środku skali.

Obliczanie nieznannej rezystancji

Gdy igła galwanometru znajduje się w środku, przez przyrząd nie przepływa prąd. Jest to możliwe tylko wtedy, gdy napięcia w punktach B i D są równe. Korzystając z prawa Ohma można stwierdzić, że tak jest, jeśli:

$$\frac{R2}{R1} = \frac{Rx}{R3}$$



Tani moduł z układem TDA7266 (© AliExpress)

Z tego równania wynika:

$$R_x = \frac{R_2}{R_1} \cdot R_3$$

Wartości R_1 i R_3 są znane, wartość R_2 odczytaj z przełączników dekady rezystorowej a następnie oblicz wartość R_x .

Praca z mostkiem Wheatstone'a

Za pomocą tego obwodu można bardzo dokładnie mierzyć rezystancję, zwłaszcza jeśli używa się bardzo czułego galwanometru, na przykład o zakresie od $-100 \mu A$ do $+100 \mu A$. Aby zachować żywotność tego galwanometru, należy oczywiście postępować z pewną ostrożnością. Podłącz mostek do regulowanego źródła zasilania U_b i rozpocznij pomiar z napięciem 0 V. Igła galwanometru znajduje się wtedy oczywiście w pozycji środkowej. Teraz ostrożnie zwiększ nieco napięcie U_b . Igła miernika przesunie się w lewo lub w prawo. Zwiększaj napięcie, aż igła znajdzie się prawie na końcu skali. Następnie przełącz rezystory w dekadzie, aż igła ponownie znajdzie się mniej więcej pośrodku. Następnie zwiększ napięcie zasilania, aż igła ponownie znajdzie się prawie na końcu skali. Następnie dostosuj wartość dekady, aż igła ponownie znajdzie się mniej więcej pośrodku. W ten sposób można zwiększyć czułość pomiaru do niespotykanych dotąd poziomów. Jedyną rzeczą, na którą należy zwrócić uwagę, jest to, aby napięcie zasilania nie stało się tak wysokie, że dwa prądy nagrzewają rezystory. Dokładność może być wtedy ograniczona przez współczynniki temperaturowe tych komponentów. Zbyt duży prąd mógłby też doprowadzić do uszkodzenia galwanometru.

Mostki Wheatstone'a i pomiary czujników

Mostek Wheatstone'a jest idealnym obwodem do normalizacji sygnału wyjściowego sensorów wielkości nieelektrycznych, np. temperatury. Wyjaśni to przykład. Załóżmy, że chcesz zmierzyć temperaturę w piekarniku za pomocą czujnika PT100. Ma on prawie liniową zależność rezystancji w bardzo dużym zakresie temperatur. Jedyną wadą jest to, że rezystancja takiego czujnika wynosi 100Ω przy $0^\circ C$. Załóżmy, że chcemy, by napięcie w V odpowiadało dokładnie takiej samej liczbowej wartości temperatury w $^\circ C$. Jak przekonwertować rezystancję 100Ω na napięcie 0 V?

Można to zrobić za pomocą mostka Wheatstone'a, mierzonego za pomocą wzmacniacza różnicowego, patrz rysunek poniżej. Mostek tworzą rezystory od R_1 do R_5 . Za pomocą potencjometru R_5 można wyregulować mostek w taki sposób, aby między punktami A i B było dokładnie 0V w temperaturze dokładnie $0^\circ C$. Gdy robi się zimniej, napięcie w punkcie A staje się niższe niż w punkcie B. Gdy robi się cieplej, napięcie w tym punkcie staje się wyższe niż w punkcie B. Oba punkty trafiają do wejść wzmacniacza różnicowego, zbudowanego wokół wzmacniacza operacyjnego 1. Obwód ten wzmacnia różnicę napięcia między punktami A i B przez współczynnik, który można ustawić za pomocą potencjometru regulacyjnego R_8 .

Mostek Kelvina, zwany również mostkiem Thomsona

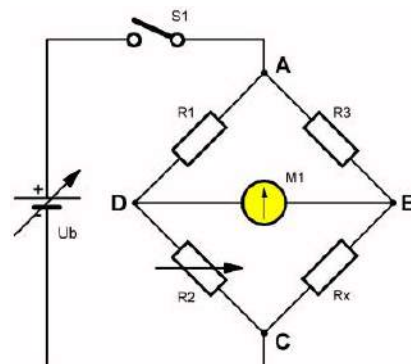
Skąd to zamieszanie z nazwiskami?

Thomson i Kelvin to dwa nazwiska tej samej osoby. William Thomson, który żył w latach 1824–1907, był brytyjskim fizykiem. Za swoje wielkie zasługi dla nauki został pasowany na rycerza Sir Williama Thomsona w 1866 roku i mianowany pierwszym baronem Kelvin (rzeka w hrabstwie Ayr) w 1892 roku. W tym samym roku został pierwszym brytyjskim naukowcem, który został wyniesiony do Izby Lordów, stając się Lordem Kelvinem.

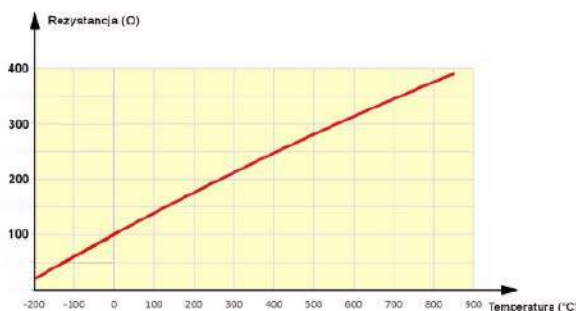
Funkcja mostka Kelvina

Mostek Kelvina umożliwia dokładny pomiar rezystancji o wartości niższej niż 1Ω . Przy tak małych rezystancjach rezystancja przewodów łączących ma znaczący wpływ, którego nie można zignorować. Należy zatem wykluczyć wpływ tych pasożytniczych rezystancji i to właśnie robi mostek Kelvina.

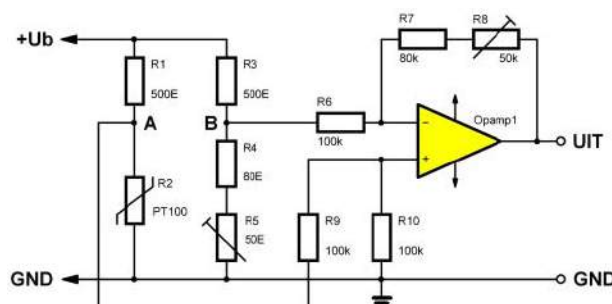
Poniższy rysunek przedstawia podstawowy schemat takiego mostka. Rezystancja R_x jest ponownie nieznaną rezystancją do zmierzenia. Jest ona podłączona do źródła napięcia szeregowo z rezystorem R_s . Rezystancja ta jest tego samego rzędu wielkości co nieznaną rezystancję. Punkty P1, P2, P3 i P4 znajdują się jak najbliżej połączeń R_s i R_x .



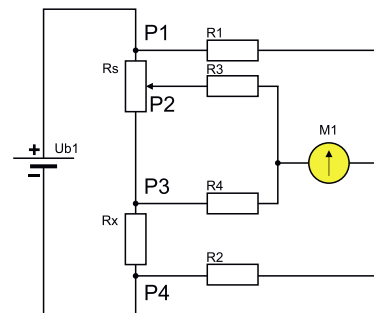
Mostek Wheatstone'a (© 2022 Jos Verstraten)



Charakterystyka czujnika PT100 (© 2022 Jos Verstraten)



Za pomocą mostka Wheatstone'a można „znormalizować” napięcie wyjściowe czujnika (© 2022 Jos Verstraten)



Standardowy schemat mostka Kelvina (© 2022 Jos Verstraten)

Jak działa mostek Kelvina

Rezystory R_s , R_1 , R_2 i R_x tworzą standardowy mostek Wheatstone'a. Ponieważ punkty P1 i P4 znajdują się blisko komponentów, nie przeszkadzają rezystancje pasożytnicze między źródłem zasilania a dwoma rezystorami R_s i R_x . Tego samego nie można jednak powiedzieć o rezystancjach pasożytniczych między R_s i R_x . Aby to zrekomensować, wprowadzono dwie nowe gałęzie z rezystorami R_3 i R_4 . To zbyt skomplikowane, aby w pełni wyjaśnić działanie mostka Kelvina matematycznie w tym i tak już bardzo obszernym artykule. W praktyce należy upewnić się, że stosunek między rezystorami R_1 i R_2 jest równy stosunkowi między rezystorami R_3 i R_4 . Mostek jest równoważony najlepiej jak to możliwe, na przykład poprzez wprowadzenie rezystancji R_s w postaci skalibrowanych rezystorów, które są umieszczane w mostku jeden po drugim, aż galwanometr wskaże najlepszą możliwą równowagę. Po zrównoważeniu można obliczyć nieznaną rezystancję za pomocą wzoru:

$$R_x = R_2 \cdot \frac{R_s}{R_1}$$

Aby zmaksymalizować dokładność, ważne jest, aby na R_s i R_x przyłożyć możliwie największe napięcie. Pomiary są zwykle wykonywane przy maksymalnym prądzie, który może przepływać przez te rezystory.

Mostek Kelvina w praktyce

Dzięki profesjonalnie zaprojektowanemu mostkowi Kelvina możliwy jest dokładny pomiar rezystancji rzędu jednej setnej oma. Wygląd takiego urządzenia pokazano na poniższym rysunku. Rezystor R_s z poprzedniego schematu został zastąpiony dziewięcioma rezystorami $0,01 \Omega$ połączonymi szeregowo. Oczywiście jest, że całkowite wykluczenie rezystancji pasożytniczych jest tutaj nie lada zadaniem! Ta kombinacja jest również połączona szeregowo z metalowym prętem R_{s10} , którego całkowita rezystancja jest dokładnie równa $0,01 \Omega$. Możesz przesuwając styk tam i z powrotem po tym pręcie, tak jak w przypadku reostat z drutem. Ponieważ krzywa rezystancji jest liniowa na całej długości pręta, można obliczyć dokładną wartość rezystancji, porównując położenie suwaka na pręcie z całkowitą długością pręta. Jeśli pręt ma długość 100 cm , a suwak znajduje się 25 cm od początku, to rezystancja między początkiem a suwakiem jest równa jednej czwartej całkowitej rezystancji lub $0,0025 \Omega$. Nad galwanometrem znajduje się potencjometr R_5 , który umożliwia regulację czułości tego przyrządu, zapobiegając w ten sposób uderzeniu igły w ogranicznik na końcu skali.

QJ57 firmy Tinsley

Urządzeniem działającym na zasadzie mostka Kelvina jest QJ57 firmy Tinsley. Za pomocą tego urządzenia można mierzyć rezystancje w zakresie $\text{m}\Omega$. Cena takiego urządzenia wynosi około $750,00 \text{ €}$.

Mostek Maxwella

Funkcja mostka Maxwella

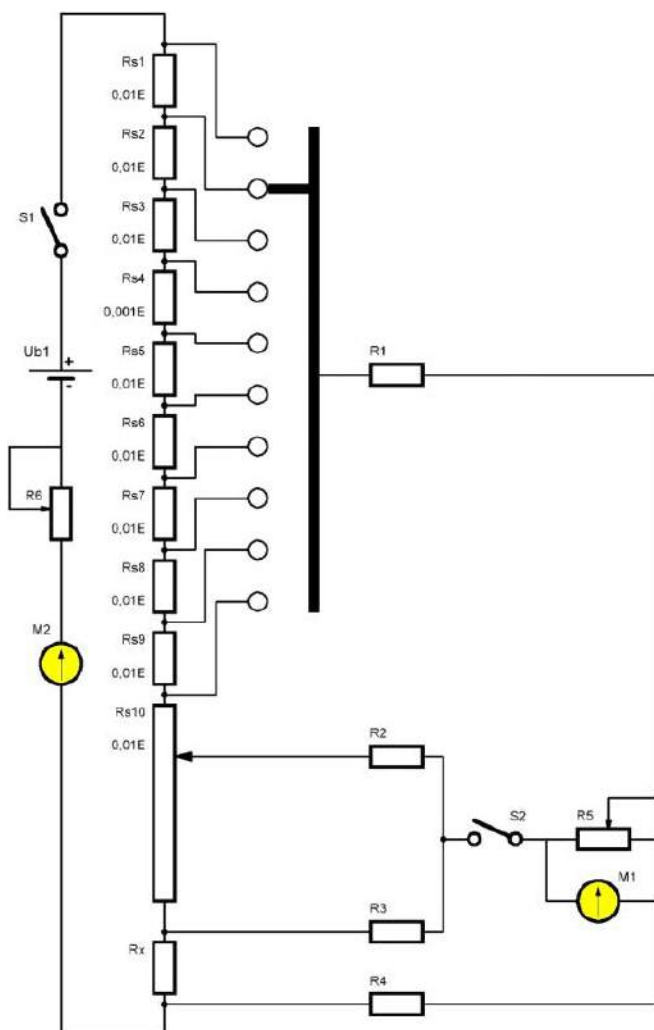
Za pomocą tego obwodu można mierzyć cewki i inne indukcyjności. Jest to odmiana mostka Wheatstone'a, w której jeden z rezystorów jest zastąpiony kondensatorem, a mostek nie jest zasilany napięciem stałym, lecz napięciem zmiennym. Mostek ten jest czasami nazywany również mostkiem Maxwella-Wiena.

Wynalezienie mostka Maxwella

Obwód ten został wynaleziony przez światowej sławy Jamesa C. Maxwella, który po raz pierwszy opisał swój mostek w 1873 roku.

Budowa mostu Maxwella

Mostek jest zbudowany zgodnie z poniższym schematem. Bez wątpienia rozpoznasz strukturę mostka Wheatstone'a, w którym dwie gałęzie są zastąpione przez mierzoną cewkę i kondensator o znanej wartości. Rezystancja R_x jest wewnętrzną rezystancją cewki L_x .



Schemat mostka Kelvina do pomiaru bardzo małych rezystancji (© 2022 Jos Verstraten)



Mostek Kelvina QJ57 firmy Tinsley (© AliExpress)

Kolejną zasadniczą różnicą jest to, że nie można teraz użyć zasilacza prądu stałego do zasilania obwodu, ale należy polegać na napięciu przemiennym. Częstotliwość tego sygnału nie ma znaczenia, więc można na przykład użyć napięcia wtórnego transformatora sieciowego. Konsekwencją tego jest to, że standardowy galwanometr nie jest już użyteczny, ale musisz dołączyć do mostka miernik napięcia lub prądu przemiennego. Oczywiście możliwe jest również zastosowanie wspomnianych już czułych słuchawek.

Obliczanie nieznannej cewki L_x

Rezystory R_1 i R_4 są znanymi dokładnymi rezystorami. Skalibrowane dekady są używane dla R_2 i C_1 . Zasada pomiaru opiera się na fakcie, iż zarówno kondensator, jak i cewka indukcyjna powodują przesunięcie fazowe między napięciem na elemencie a prądem przepływającym przez niego. Jednak te przesunięcia fazowe są przeciwne! W przypadku cewki indukcyjnej prąd będzie opóźniony w stosunku do napięcia, a w przypadku kondensatora prąd będzie wyprzedzał napięcie.

Chodzi o to, aby zrównoważyć mostek zarówno pod względem amplitudy, jak i fazy. W pierwszym przypadku służy do tego rezystor R_2 , a w drugim kondensator C_1 . Jeśli mostek jest w równowadze, można odczytać wartość kondensatora C_1 i rezystora R_2 . Wartość cewki jest wtedy określona wzorem:

$$L_x = R_1 \cdot R_4 \cdot C_2$$

a wartość rezystancji wewnętrznej cewki przez:

$$R_x = \frac{(R_1 \cdot R_4)}{R_2}$$

Mostek Wiena

Funkcja mostka Wiena

Mostek Wiena jest przeznaczony do dokładnego pomiaru wartości częstotliwości sygnału sinusoidalnego w zakresie audio. Sygnał ten jest również napięciem zasilania mostka.

Wynalezienie mostka Wiena

Mostek ten został opracowany przez niemieckiego fizyka Maxa Wiena w 1891 roku, aby pomóc w jego badaniach nad siłą i propagacją dźwięków o różnych częstotliwościach.

Budowa mostka Wiena

Mostek Wiena zawiera cztery rezystory i dwa kondensatory. Dwa regulowane rezystory muszą być równe. To samo dotyczy dwóch kondensatorów.

Obliczanie nieznannej częstotliwości

Warunkiem prawidłowego działania tego mostka jest to, aby oba kondensatory były dokładnie takie same. W praktyce dwa identyczne kondensatory o wartościach 1/10/100/1000 są podłączone do obwodu, na przykład za pomocą przełącznika obrotowego 2×4-pozycyjnego. Dwa rezystory R_1 i R_3 również muszą być identyczne. Można w tym miejscu użyć potencjometru stereo i zapewniając mu skalibrowaną skalę na froncie urządzenia. Mostek ten musi być również wyposażony w przyrząd pomiarowy zdolny do pomiaru napięcia przemiennego w zakresie częstotliwości, dla którego został zaprojektowany. Mostek jest równoważony poprzez najpierw wybór kondensatorów za pomocą przełącznika obrotowego, a następnie przez regulowanie potencjometru stereo do uzyskania minimalnego napięcia na mierniku. Po zrównoważeniu mostka w ten sposób można obliczyć częstotliwość sygnału wejściowego za pomocą wzoru:

$$f = \frac{1}{2\pi RC}$$

gdzie:

π = słynna liczba pi, więc 3,141592653...

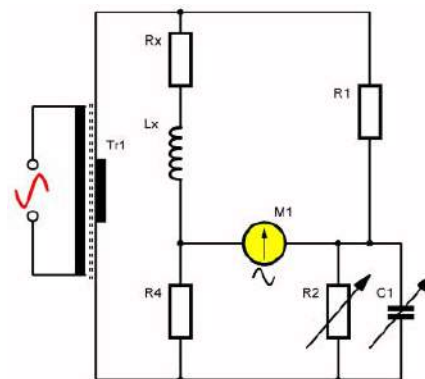
$R = R_1 = R_3$

$C = C_1 = C_2$

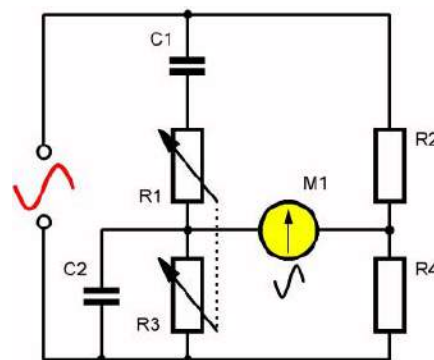
Mostek Wiena w obwodach oscylatora

Oczywistym jest, iż mostek Wiena jest całkowicie przestarzałym przyrządem do pomiaru częstotliwości. Cyfrowe mierniki częstotliwości stabilizowane rezonatorami kwarcowymi są znacznie łatwiejsze w obsłudze i mają znacznie wyższą dokładność (oraz zakres mierzonych częstotliwości – przyp. tłum.). Jednak mostek Wiena jest nadal bardzo często używany w generatorach fali sinusoidalnej. Jak widać na poniższym schemacie, mostek Wiena składa się z równoległego połączenia rezystora i kondensatora (R_2/C_2) oraz szeregowego połączenia identycznych elementów (R_1/C_1). Jeśli oba rezystory i oba kondensatory są tej samej wielkości, częstotliwość oscylacji można obliczyć ze wzoru:

$$f_o = \frac{1}{2\pi RC}$$



Zasada działania mostka Maxwella (© 2022 Jos Verstraten)

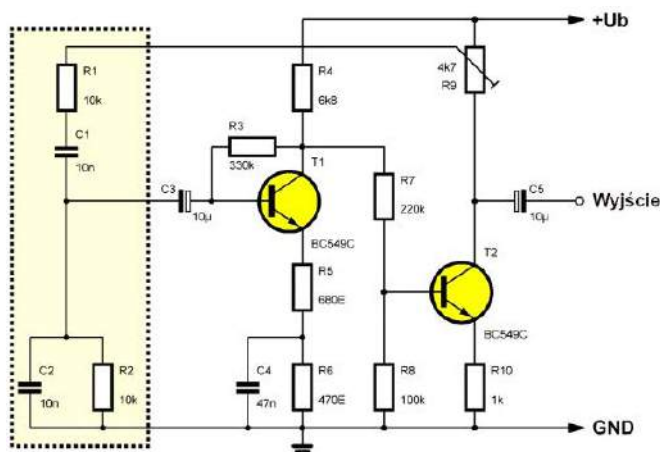


Zasada działania mostka Wiena (© 2022 Jos Verstraten)

Obwód ten wymaga użycia dwóch stopni wzmacniających, ponieważ przesunięcie fazowe o 180° spowodowane przez pierwszy stopień musi zostać odwrócone przez drugi stopień. Wejście i wyjście wzmacniacza muszą być zatem w fazie! Aby spełnić ogólny warunek amplitudy oscylatora, całkowite wzmocnienie musi być większe niż trzy. Jest to konsekwencją podziału napięcia utworzonego przez sieć mostka Wiena między wyjściem a wejściem, co osłabia sygnał sprzężenia zwrotnego.

Za pomocą rezystora R9 można ustawić współczynnik sprzężenia i sprawić, że obwód spełni warunek amplitudy (wzmocnienie większe niż 3). Aby ustabilizować wartość współczynnika wzmocnienia dokładnie na tych trzech, obwód musi być wyposażony w szereg sprzężeń zwrotnych. Przedstawiony obwód (jeden z wielu) zawiera zarówno napięciowe, jak i prądowe ujemne sprzężenie zwrotne. Rezystor R3 stabilizuje pierwszy stopień wzmacniacza, rezystory R5 i R10 zapewniają prądowe ujemne sprzężenie zwrotne w obwodach emiterowych dwóch półprzewodników.

Częstotliwość obwodu można zmieniać poprzez równoczesną zmianę dwóch rezystorów lub dwóch kondensatorów mostka. W praktyce wiele identycznych kondensatorów jest zwykle włączanych do sieci za pomocą podwójnego przełącznika. Jeśli użyjesz szeregu kondensatorów, z których każdy następny jest 10 razy większy od poprzedniego, możesz ustawić częstotliwość obwodu za pomocą tego przełącznika w sposób dekadowy. Dwa stałe rezystory są zastąpione potencjometrem stereo, za pomocą którego można ustawić częstotliwość na dowolną wartość w każdej dekadzie.



Mostek Wiena zastosowany w oscylatorze sinusoidalnym (© 2022 Jos Verstraten)

Mostek Sauty'ego

Pomiar kondensatorów

Ten mostek jest bezpośrednio oparty na mostku Wheatstone'a i jest przeznaczony do pomiaru wartości kondensatora. Mostek jest (oczywiście) zasilany napięciem przemiennym, więc należy użyć w nim miernika napięcia przemiennego. Wartość nieznanego kondensatora obliczymy ze wzoru:

$$C_x = C1 \cdot \frac{R2}{R1}$$

Mostek Scheringa

Pomiar pojemności i ESR kondensatorów

Ten mostek, pokazany na poniższym rysunku, pozwala określić nie tylko pojemność kondensatora, ale także jego ESR czyli równoważną rezystancję szeregową. Na schemacie te dwie wielkości są reprezentowane przez C_x i R_x . Mostek jest zrównoważony poprzez regulację $R1$ i $C1$. W stanie równowagi można obliczyć wartość dwóch mierzonych wielkości za pomocą poniższych wzorów:

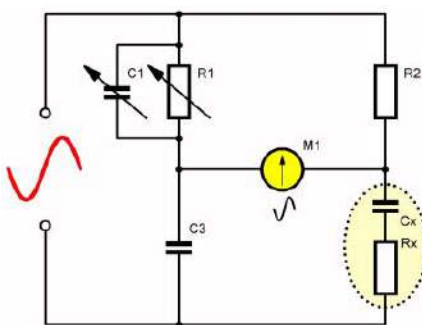
$$C_x = \frac{R1 \cdot C3}{R2}$$

oraz:

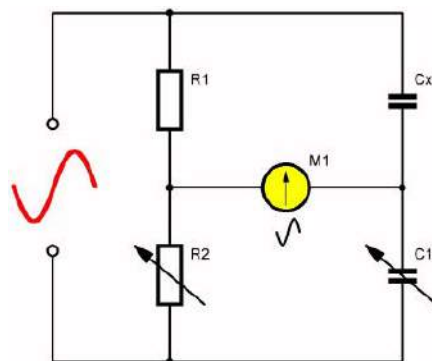
$$R_x = \frac{C1 \cdot R2}{C3}$$

Mostek Scheringa w praktyce

Takich urządzeń nie znajdziemy w większości warsztatów elektronicznych. Są jednak producenci, którzy wytwarzają takie egzotyczne produkty, czego dowodem jest poniższy model ASICO.



Mostek Scheringa (© 2022 Jos Verstraten)



Mostek Sauty'ego (© 2022 Jos Verstraten)



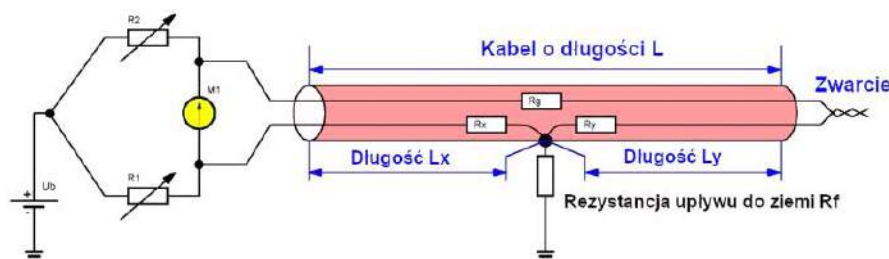
Laboratoryjna wersja mostka Scheringa (© IndiaMART)

Mostek Murray'a

Określanie upływności w kablach

Za pomocą tego mostka można zlokalizować miejsce, w którym podziemny kabel wykazuje przebicie do ziemi, tzw. upływność. Podstawowy schemat został przedstawiony na poniższym rysunku. Dwużyłowy kabel o całkowitej długości L ma dobrą żyłę i żyłę, która ma upływ do ziemi. Kabel jest reprezentowany przez rezystory R_g , R_x i R_y . R_g jest rezystancją żyły o całkowitej długości L kabla, R_x i R_y są rezystancjami żyły do i po wystąpieniu upływu do ziemi. R_f to rezystancja upływu prądu do ziemi z uszkodzonej żyły. Dwie żyły kabla

są połączone ze sobą na jednym końcu kabla. Powoduje to zwarcie kabla. Obie żyły kabla są podłączone na drugim końcu za pomocą rezystorów R_1 i R_2 do źródła napięcia stałego, które dostarcza wysokie napięcie. Drugi biegun tego źródła zasilania połączony jest z ziemią. W ten sposób powstał mostek z czterema rezystorami i źródłem napięcia. Oczywiście w drugiej przekątnej należy umieścić galwanometr. Mostek jest zrównoważony poprzez zmianę rezystancji R_1 i R_2 . Jeśli mostek jest w równowadze, można obliczyć długość L_x za pomocą wzoru:



Mostek Murray'a umożliwia określenie lokalizacji zwarcia doziemnego w kablu (© 2022 Jos Verstraten)

$$L_x = 2 \cdot L \cdot \frac{R_1}{(R_1 + R_2)}$$

Mostek Careya-Fostera

Pomiar bardzo niskich wartości rezystancji

Mostek Careya-Fostera został opisany w 1872 roku przez Georga Careya Fostera. Jest to wariacja na temat tradycyjnego mostka Wheatstone'a do pomiaru bardzo niskich wartości rezystancji. Podstawowa konfiguracja jest pokazana na poniższym rysunku. R_x to mierzona rezystancja. R_1 , R_2 i R_y to trzy rezystory w przybliżeniu tego samego rzędu wielkości co R_x , ale których wartości są znane z dużą dokładnością. Po lewej i prawej stronie znajdują się dwie duże metalowe powierzchnie kontaktowe, których rezystancja jest praktycznie zerowa. Długi drut oporowy o niskiej impedancji i o znanej długości L jest umieszczony pomiędzy tymi powierzchniami (A-B). Do tego drutu można przykręcić zacisk C. Całość jest zasilana przez źródło napięcia stałego U_b , które dostarcza niskie napięcie, ale jest w stanie dostarczyć bardzo duży prąd. Czarne okręgi reprezentują punkty, w których różne komponenty są połączone ze sobą z możliwie najniższą impedancją.

Pomiar rozpoczyna się od przesunięcia styku C nad drutem oporowym do momentu zrównoważenia mostka Careya-Fostera. Igła galwanometru M_1 znajduje się na środku skali. Teraz należy odczytać długość L_x i obliczyć ją jako wartość procentową w stosunku do całkowitej długości L . Współczynnik ten należy nazwać %x. Następnie zmienić pozycje R_x i R_y i określić nowy punkt równowagi



Mostki

Rozwiązanie znajdziesz na www.elportal.pl/quizy

1. Z ilu bloków składa się ogólna forma mostka?

- ☐ z trzech;
- ☐ z czterech;
- ☐ z pięciu.

2. Do pomiarów mostkowych przy zasilaniu napięciem stałym najlepiej nadaje się:

- ☐ multimetr cyfrowy;
- ☐ galwanometr;
- ☐ mikrowoltomierz analogowy.

3. Mostkiem Wheatstone'a pozwala dokładnie mierzyć:

- ☐ rezystancję;
- ☐ reaktancję;
- ☐ indukcyjność cewki.

4. Mostek Wiena pozwala mierzyć:

- ☐ częstotliwość;
- ☐ indukcyjność cewki;
- ☐ pojemność kondensatora.

5. Mostek Wiena pozwala także na:

- ☐ pomiar ESR;
- ☐ generowanie stabilnego przebiegu sinusoidalnego;
- ☐ generowanie stabilnego napięcia doniesienia.

6. Mostek Maxwella używa regulowanych kondensatora i rezystora by wyznaczyć:

- ☐ indukcyjność i rezystancję nieznaną cewki

- ☐ pojemność i ekwiwalentną rezystancję szeregową nieznanego kondensatora;
- ☐ częstotliwość rezonansową i dobroć nieznanego obwodu LC.

7. Mostek Scheringa używa regulowanych kondensatora i rezystora by wyznaczyć:

- ☐ indukcyjność i rezystancję nieznaną cewki;
- ☐ pojemność i ekwiwalentną rezystancję szeregową nieznanego kondensatora;
- ☐ częstotliwość rezonansową i dobroć nieznanego obwodu LC.

8. Mostek Sauty'ego używa regulowanych kondensatora i rezystora by wyznaczyć:

- ☐ indukcyjność i rezystancję nieznaną cewki;
- ☐ pojemność nieznanego kondensatora;
- ☐ ekwiwalentną rezystancję szeregową nieznanego kondensatora.

9. Mostek Murray'a pozwala określić:

- ☐ grubość izolacji;
- ☐ rezystancję izolacji;
- ☐ miejsce przebicia izolacji podziemnego kabla.

10. Mostek Thomsona służy do pomiaru:

- ☐ bardzo małych pojemności nieznaną kondensatorów;
- ☐ bardzo małych indukcyjności nieznaną cewek;
- ☐ bardzo małych rezystancji.

styku C. Teraz odczytaj długość L_y i przekształć ją w liczbę procentową w stosunku do całkowitej długości drutu oporowego. Nazwij ten stosunek $\%y$. Oblicz rezystancję drutu w procentach i nazwij ją $\%total$. Rezystancja R_x wynika ze wzoru:

$$R_x = (\%total \cdot (\%y - \%x)) + R_y$$

Na zakończenie tej długiej historii...

Najbardziej używane mostki

Na przestrzeni dziejów zaprojektowano więcej układów mostkowych. Nie przetrwały one jednak próby czasu i nie można się już z nimi spotkać. Wymieniamy je tutaj bez wyjaśnień i schematów, aby uzupełnić tę historię.

Mostek Andersona

Jest to obwód wywodzący się z mostka Maxwella, opracowany przez Alexandra Andersona w 1891 roku i przeznaczony do pomiaru wartości i rezystancji szeregowej cewki. Jednak wzory do obliczania tych dwóch wielkości są skomplikowane, a mostek jest bardzo nieprzyjazny dla użytkownika ze względu na dość kłopotliwą regulację.

Mostek Hay'a

Jest to również obwód wywodzący się z mostka Maxwella, który został zaprojektowany do pomiaru bardzo wysokich indukcyjności. Ten mostek również posiada dość skomplikowane wzory do obliczania wartości indukcyjności własnej.

Mostek Fontany

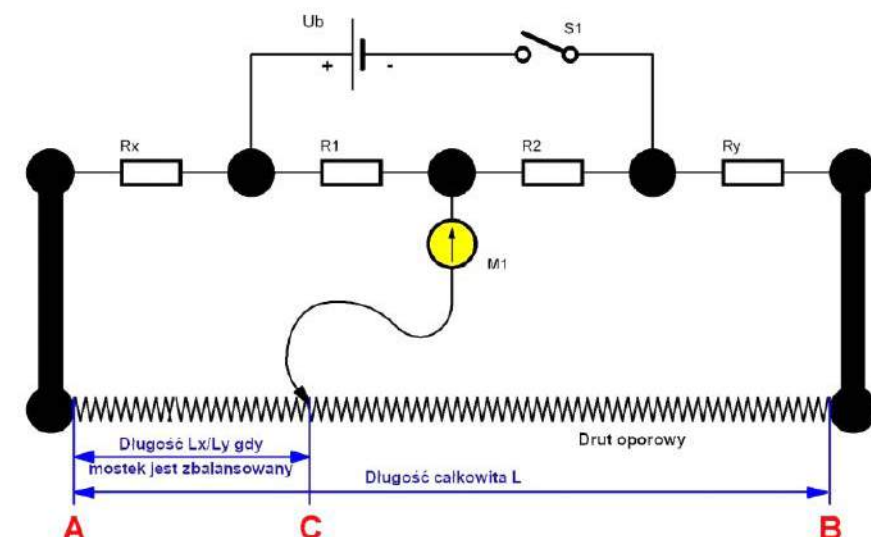
Dość skomplikowany obwód opisany w 2003 roku przez Giorgio Fontanę, który umożliwia zaprojektowanie konwertera napięcia na prąd. Wyjątkowe jest to, że obwód ten działa również dla napięcia przemiennego i ma dużą szerokość pasma.

Mostek kratowy

Obwód ten jest korektorem fazy wynalezionym w 1927 roku przez Otto Zobela. Tłumienie i impedancja filtra są niezależne od częstotliwości, ale to samo nie dotyczy różnicy faz między wejściem a wyjściem. Wzrasta ona wraz ze wzrostem częstotliwości sygnału. Takie mostki były używane do kompensacji przesunięć fazowych między lewym i prawym kanałem sygnału stereo. ■

Jos Verstraten

REKLAMA



Podstawowy projekt mostka Careya-Fostera (© 2022 Jos Verstraten)

Publikujemy dla projektantów i programistów elektroniki. Odwiedź

ELPORTAL.pl

Znajdziesz nas również na Facebooku: facebook.com/ElportalPL

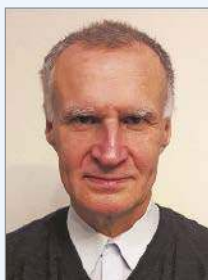
Uczmy się na cudzych błędach

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.

Prosty automatyczny wyłącznik światła w łazience

Szukanie po omacku wyłącznika światła w łazience jest czasem kłopotliwe. Czy nie można zrobić tak, aby światło zaświecało się automatycznie po otwarciu drzwi? Taki „automatyczny wyłącznik” można zrealizować na wiele sposobów, zależnie od własnych preferencji i przeznaczonego na ten cel budżetu. Prezentowany tu projekt stawia na funkcjonalność i prostotę. Istotną częścią takiego projektu jest wybór czujnika. Można zastosować jakiś czujnik inteligentny jak czujkę ruchu lub fotosensor. W tym projekcie wybór padł na prosty kontaktron, do którego zbliżamy magnesik.

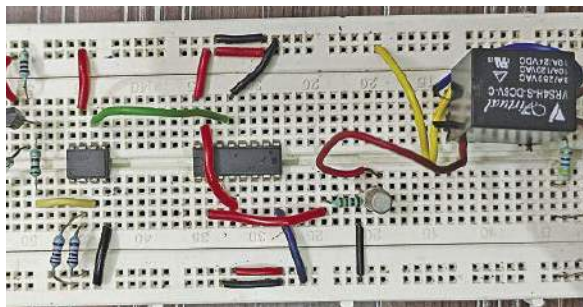
Kontaktron należy zamontować na framudzie drzwi, a magnesik na skrzydle drzwi, tak aby przy zamkniętych drzwiach oba elementy były blisko siebie. Kontaktron powinien być typu NO (Normal Open), a jego styki ulegają zwarceniu po zbliżeniu magnesiku.

Mimo prostoty części logicznej układu trzeba mieć na uwadze, iż mamy tu do czynienia z napięciem sieci energetycznej 230 VAC. Najważniejszym priorytetem jest bezpieczeństwo zarówno przy pracach montażowych jak i podczas codziennego użytkowania. Jeśli nie masz doświadczenia w tym zakresie, należy zasięgnąć konsultacji wykwalifikowanego elektryka.

Na **rysunku 1** pokazano prototyp zmontowany przez autora. Schemat ideowy jest natomiast na **rysunku 2**.

W układzie wykorzystano następujące podzespoły: transformator sieciowy X1, mostek prostowniczy BR1, stabilizator 5 V LM7805 (IC1), kontaktron typu NO, niewielki magnesik oraz tranzystor pnp BC557 (T1), wzmacniacz operacyjny 741 (IC2), licznik dekadowy 4017 (IC3), tranzystor npn 2N2219 (T2), przełącznik z cewką 5 V (RL1) i ponadto niewielką liczbę prostych rezystorów i dwa kondensatory.

Zaproponowany układ zasilany jest pojedynczym napięciem 5 V. Napięcie to pozyskano przy użyciu transformatora sieciowego obniżającego napięcie do wartości 9 VAC (z obciążalnością do 500 mA). W celu pozyskania wartości DC wykorzystano mostek



Rysunek 1. Prototyp wykonany przez autora

Wykaz elementów:

Półprzewodniki:

IC1: stabilizator 5 V LM7805
IC2: wzmacniacz operacyjny 741
IC3: licznik dekadowy 4017
BR1: mostek prostowniczy 1 A
LED1: dioda LED 5 mm
T1: tranzystor pnp BC557
T2: tranzystor npn 2N2219

Rezystory: (wszystkie 0,25 W/±5%)

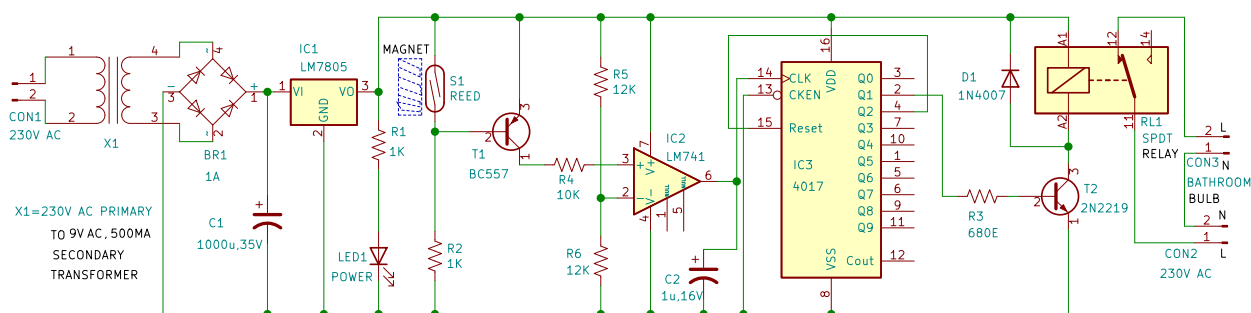
R1, R2: 1 kΩ
R3: 680 Ω
R4: 10 kΩ
R5, R6: 12 kΩ

Kondensatory:

C1: 1000 μF/35 V elektrolityczny
C2: 1 μF/16 V elektrolityczny

Inne:

CON1-CON3: złącze 2-pinowe
S1: miniaturowy kontaktron w obudowie szklanej, typu NO
RL1: przełącznik 5 V z pojedynczymi stykami
X1: transformator sieciowy: uzwojenie pierwotne 230 VAC, wtórne 9 VAC/0,5 A niewielki magnes



Rysunek 2. Schemat ideowy układu

prostowniczy BR1 i kondensator filtrujący C1. Tak wstępnie odfiltrowane napięcie stałe jest dalej stabilizowane trzynóżkowym stabilizatorem LM7805. Obecność napięcia 5 V jest potwierdzona świeceniem diody LED1.

Działanie układu jest proste. Gdy drzwi łazienki są zamknięte, magnesik powinien być na tyle blisko kontaktronu, że jego styki są zwarte. Wtedy tranzystor T1 jest wyłączony, gdyż jest zwarta jego baza z emiterem. W konsekwencji, potencjał wejścia nieodwracającego wzmacniacza operacyjnego utrzymuje się poniżej potencjału referencyjnego na wejściu odwracającym (tu dzielnik rezystancyjny ustala połowę napięcia zasilania). W tej sytuacji wyjście WO n.6 utrzymuje napięcie bliskie dolnego zasilania. To oznacza stabilny stan niski na wejściu zegarowym licznika dekadowego 4017. Otwarcie drzwi skutkuje oddaleniem magnesu od kontaktronu. Kontaktron rozwiera styki, co odblokowuje bazę tranzystora pnp. Przewodzenie w obwodzie kolektor-emiter tranzystora T1 skutkuje podniesieniem napięcia na wejściu nieodwracającym WO powyżej referencyjnego na wejściu „minus”. Wyjście WO zmienia stan z niskiego na wysoki, co daje aktywne zbocze zegara naliczające licznik. Jeśli wcześniej licznik był w stanie gdy Q0=1, teraz „jedynek” przesuwają się na wyjście Q1. To włącza tranzystor T2 i w konsekwencji przełącznik RL1. Styki tego przełącznika zamykają obwód zasilania oświetlenia łazienki. Zamknięcie drzwi nie generuje żadnej reakcji. Magnesik zbliża się do kontaktronu. Jego styki ponownie się zamykają, co wyłącza tranzystor T1. Potencjał wejścia nieodwracającego WO obniża się poniżej napięcia panującego na wejściu odwracającym. Wyjście n.6 WO zmienia stan z wysokiego na niski. Sygnał zegarowy licznika 4017 wykazuje zbocze opadające, a to jest nieaktywne i dlatego nie wywołuje żadnej reakcji. Stan wysoki utrzymuje się na wyjściu Q1 i światło nadal świeci. Kolejne zbocze narastające (aktywne) wystąpi przy następnym otwarciu drzwi. Zakładamy, że oznacza to chęć wyjścia z łazienki. Kolejność sygnałów jest taka jak wcześniej: rozwarcie styków kontaktronu, włączenie tranzystora T1 i przełączenie wzmacniacza operacyjnego działającego tu jako analogowy komparator. Teraz narastające zbocze zegara przesuwają jedynkę logiczną z wyjścia Q1 na Q2. A to natychmiast zeruje licznik, czyli stan „1” przyjmuje wyjście Q0. W każdym razie, zmiana stanu Q1 na niski wyłącza tranzystor T2, przełącznik i w konsekwencji światło gaśnie (*lepiej byłoby, jeśliby światło zgasiło dopiero po zamknięciu drzwi wychodząc z łazienki – przypis red.*)

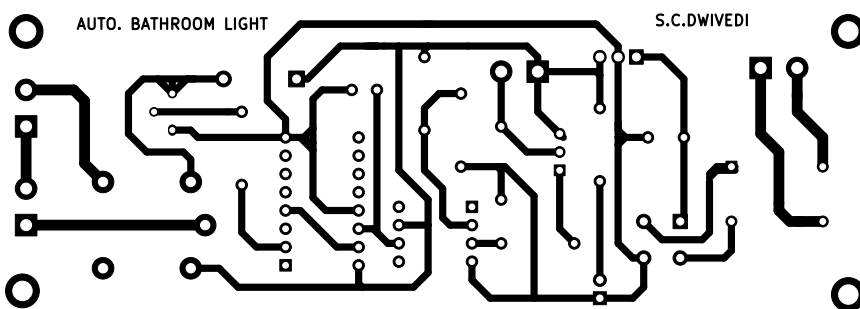
Krótko mówiąc logika układu jest następująca: otwarcie drzwi do łazienki zaświeca światło; wchodzisz, zamykasz drzwi – światło świeci; po „załatwieniu swoich spraw” wychodzisz, światło gaśnie w momencie otwarcia drzwi; zamykasz drzwi, światło pozostaje zgaszzone, układ czeka aż ktoś znów zechce wejść do łazienki.

Konstrukcja i instalacja układu

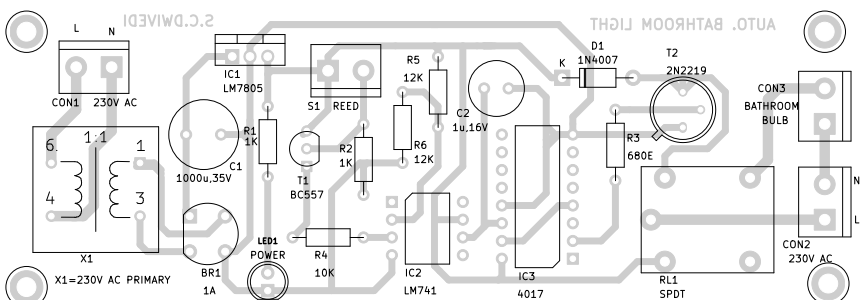
Na potrzeby projektu przygotowano jednostronną płytkę PCB, co pokazano na **rysunku 3**. W papierowym wydaniu pisma, rozmiary płytki powinny odpowiadać naturalnej skali 1:1. Na **rysunku 4** widzimy ułożenie elementów na PCB, co powinno ułatwić prace montażowe. Po uzbrojeniu PCB, należy układ umieścić w odpowiednio przygotowanej obudowie. Z przodu należy zamontować diodę LED1 oraz złącze CON2. Z tyłu należy wyprowadzić przewody łączące żarówki oświetlenia.

Kluczową częścią projektu jest montaż czujnika otwarcia/zamknięcia drzwi. Czyli kontaktronu i zbliżanego doń małego magnesu. Kontaktron powinien być typu „Normal Open”, tak aby bez obecności magnesu styki kontaktronu były rozwarne, a po zbliżeniu magnesu ulegały zwarceniu. Po zmontowaniu elementów na PCB, należy teraz podłączyć przewody od kontaktronu. Jeśli nie popełniliśmy żadnego błędu podczas montażu, układ powinien być gotowy do pracy i nie wymaga żadnej regulacji ani strojenia. ■

Suresh Dwivedi



Rysunek 3. Projekt druku płytki PCB



Rysunek 4. Rozmieszczenie elementów na PCB

Ten artykuł był wcześniej opublikowany na łamach „EFY”, maj 2023 (efymag.com)

Uwagi i poprawki

Od Red. EdW: Układ jest bardzo prosty, mimo to nie uniknięto tu kilku istotnych błędów. Ponadto, często bywają nieporozumienia, co znaczy klucz otwarty, a co zamknięty. W tekście zmieniliśmy, iż kontaktron powinien być typu NO, choć autor kilkakrotnie pisze NC, jak również w spisie materiałów czytamy – small normally closed type glass reed switch. Skorygowaliśmy też opis, kiedy switch się zamyka a kiedy otwiera. Zakładamy, że klucz otwarty to styki rozwarne, a zamknięty – zwarte. Kontaktron Normal Open to taki który ma styki rozwarne bez obecności pola magnetycznego, a po zbliżeniu magnesu styki ulegają zwarcu.

Powiedzieliśmy, projekt jest koncepcyjnie prosty, a najwięcej kłopotów mogą sprawić prace montażowe. I tu też pozwoliliśmy sobie zmienić tekst, gdyż autor sugeruje aby magnes montować w futrynie drzwi, a kontaktron na skrzydle drzwiowym. Taki montaż stwarzałby o wiele więcej kłopotów. Ale zajmijmy się samym schematem ideowym.

Pomijamy również błąd, iż autor w tekście odwołuje się, iż wyjściem wzmacniacza operacyjnego jest nóżka 1. Wiele wzmacniaczy operacyjnych i komparatorów faktycznie zachowuje taki standard, że jeśli wejściem plus i minus są wyprowadzenia n.3 i n.2, to wyjście jest na nóżce 1. Tu autor zastosował „wiekowy” LM741 z możliwością regulacji offsetu napięcia niezrównoważenia, co angażuje nóżki 1 i 5. W obudowie ośmio-nóżkowej mieści się tylko jeden WO, i wyjście ma on na n.6. Tutaj, wymogi co do parametrów WO są praktycznie „żadne”, offsetu także kompensować nie ma potrzeby, i 741 można zastąpić innym dowolnym, (np. 4558, 393, 358, TL082), który będzie miał wyjście na n.1. Ale, błędy na schemacie ideowym rysunek 2 są poważniejsze. Kontaktron zwiera bazę-emiter tranzystora T1, a wejście nieodwracające WO podłączone jest praktycznie wprost do kolektora tego tranzystora. Otwarcie drzwi skutkuje rozwarciem styków S1 i wtedy T1 na wyjściu + WO wymusza stan wysoki. Po zamknięciu drzwi baza z emiterem T1 zostaje zwarta. Wtedy na kolektorze nie jest wypracowany stan niski, ale stan wysokiej impedancji. Jak to zinterpretuje WO? LM741 ma wejścia w postaci pary różnicowej tranzystorów npn, więc odcięcie wejścia „plus” powinno być odczytane jako stan niski. Ale taki projekt to „szkolny błąd”

i nie tylko dlatego, że inny typ wzmacniacza operacyjnego (np. z wejściem w postaci pary różnicowej tranzystorów pnp) nie będzie działał poprawnie. Z kolektora T1 koniecznie trzeba dać jakiś rezystor do masy (który zapewni ewidentny stan niski, gdy tranzystor jest wyłączony). R4 między kolektorem tranzystora T1 i wejściem nieodwracającym WO jest praktycznie niepotrzebny. Można go przepiąć do masy i kolektor połączyć wprost z wejściem wzmacniacza. To uleczy najpoważniejszy błąd, ale jest ich więcej. Na wyjściu WO zastosowano kondensator C2=1 μ F. Można się domyślać, iż załatwi on problem drgania styków mechanicznych. Wzmacniacz i licznik 4017 są elementami na tyle szybkimi, że o tym problemie trzeba koniecznie pomyśleć. Ale wpinanie dość dużej pojemności na aktywne wyjście WO to nie jest dobry sposób. Można to na co najmniej kilka sposobów załatwić od strony wejścia, lub na wyjściu jeśli zastosujemy układ z wyjściem typu otwarty kolektor (oczywiście trzeba wtedy dodać jakiś rezystor do plusa zasilania). WO LM741 pracuje tu w roli analogowego komparatora, więc nawet lepiej jeśliby go zastąpić komparatorem np. LM393. Ale to tylko skromna uwaga, bo jest znacznie gorzej. Komparator czy WO jest tu w ogóle niepotrzebny. Można z niego zrezygnować bez żadnego uszczerbku dla działania układu. Licznik 4017 nalicza na zboczu narastającym i jest to wejście typu Schmitta. Nie ma zatem ograniczeń co do stromości zbocza, może więc przyjąć sygnał wprost z kolektora T1, a R4 jak wcześniej, należy przepiąć do masy. Z rozwiązania w obrębie licznika 4017, też trudno być zadowolonym. To dekadowy licznik Johnsona. Jednak tutaj zlicza „modulo dwa”. Czyli pracuje jak zwykły przerzutnik typu T, licząc do dwóch (0, 1, 0, 1 itd). Jeśli ktoś ma w szufladzie nadmiar elementów 4017, to można go wykorzystać, ale też nie jest to dobre. Reset zeruje licznik, czyli wystawia jedynkę logiczną na wyjście Q0, które nie jest wykorzystane. Kolejne narastające zbocze zegara przenosi „1” na wyjście Q1, które jest właściwym i jedynym wykorzystanym wyjściem licznika. Kolejny impuls zegarowy przenosi „1” na Q2, które natychmiast zeruje licznik. Czyli „parzysty” impuls zegarowy przenosi jedynkę logiczną nie na Q2, ale na Q0. Tak ma być, i w ten sposób licznik dekadowy stał się licznikiem binarnym. Ale takie działanie to jest hazard, którego w układach sekwencyjnych

należy unikać. To jest błąd konstrukcyjny nawet jeśli układ zadziała dobrze. Czas impulsu Reset jest bliżej nieokreślony. Nic nie gwarantuje, że wyzerują się wszystkie przerzutniki licznika Johnsona w układzie scalonym 4017. Czas impulsu zerującego jest nieokreślony, a należałoby go przytrzymać niezależnie od zaniku jedynki logicznej na wyjściu Q2. Ten prosty wymóg wymagałby elementu opóźniającego w torze sprzężenia zwrotnego między Q2 a Reset. To jest jednak dość kłopotliwe. Poprawne rozwiązanie z przerzutnikiem monostabilnym można by uznać za „zbyt drogie” rozwiązanie. Można „iść na skróty” i w pętli sprzężenia zwrotnego wpiąć jakiś kondensator, ew. plus diodę. Ale to nie załatwi sprawy na poziomie „teoretycznej poprawności”. Na poziomie „praktyki” można nic nie robić, bo układ powinien działać zgodnie z oczekiwaniami. Jednak, zgoda na hazard w układzie sekwencyjnym jest po prostu błędem.

Autor reklamuje swój układ jako prosty, i faktycznie takim jest. Ale mógłby być jeszcze prostszy i równocześnie lepszy. Można wyrzucić nie tylko wzmacniacz operacyjny, ale i licznik dekadowy. A w to miejsce zastosować najprostszego przerzutnik typu D zapętlony z „zanegowanego Q” na D, co uczyni przerzutnik T zliczający do dwóch. Przerzutnik ten powinien reagować na wpis narastającym zboczem sygnału zegarowego, tak samo jak zastosowany tu licznik Johnsona. Natomiast o drganie styków kontaktronu trzeba koniecznie zadbać, bo byłoby bardzo denerwującym, jeśliby trzeba kilkakrotnie zamykać i otwierać drzwi, aby światło w łazience się zaświeciło. A potem to samo, aby zgasło. I tak, układ zaświeca i gasi światło w reakcji na co drugie otwarcie drzwi. Licznik „nalicza” na otwarcie, nie na zamknięcie drzwi. Pożądanym jest, aby tak właśnie było. Ale i tak nie wie, kiedy do łazienki wchodzi, a kiedy wychodzi. Na tę niedogodność trzeba się zgodzić, bo jej likwidacja wymagałaby „sztucznej inteligencji” (np. jakiegoś drugiego czujnika informującego w którą stronę idziesz). Jeśli licznik modulo 2 „się pomyli”, to co najwyżej będzie potrzebny dodatkowy ruch zamknięcia i otwarcia drzwi.

Tak czy inaczej, układ jest bardzo prosty i uleczenie dostrzeżonych błędów nie jest trudne. Wydaje się, że sukces lub porażka będzie zależeć bardziej od konstrukcji mechanicznej. Od sposobu montażu kontaktronu i magnesiku na drzwiach łazienki. Wydaje się

też, że wygodniejsza byłaby trochę inna sekwencja działania: światło zaświeca się w reakcji na otwarcie drzwi; następnie zamykamy drzwi – bez reakcji; kolejne otwarcie drzwi – bez reakcji; i światło gaśnie dopiero po następnym zamknięciu drzwi. To wymagałoby jedynie nieznacznej rozbudowy tego prostego układu. Trzymając się schematu z rysunku 2, należałoby dodać sumę logiczną wyjścia Q1 licznika 4017 i sygnału sprzed tego licznika, czyli stan logiczny wyjścia WO. Można do tego zaangażować jakąś bramkę logiczną, ale można i prościej. Taką sumę logiczną załatwią dwie diody. Katody obu diod podłączyć do rezystora R3. Anodę jednej – na wyjście Q1 IC3, a anodę drugiej na wyjście 6 IC2. W tak poprawionym układzie,

nadal aktualna jest uwaga o możliwości rezygnacji ze wzmacniacza operacyjnego i/lub zastąpieniu licznika dekadowego przerzutnikiem typu T. Można mieć obawę, czy tak utworzona suma logiczna nie jest źródłem kolejnego hazardu na zboczu narastającym zegara (kiedy w tym samym momencie ulegają zmianie oba składniki sumy; jeden zmienia się ze stanu wysokiego na niski, a drugi w przeciwnym kierunku). Gdyby nawet tak było, to to niekorzystne zjawisko łatwo jest usunąć wpinając niewielki kondensator na wyjście tak utworzonej sumy. Jeśliby skorzystać z opcji utworzenia sumy logicznej z dwóch diod, to kondensator należałoby wpiąć między obie katody i masę. Tutaj jednak, nawet ten zabieg nie jest potrzebny,

gdyż sygnał sprzed licznika (przejście z zera na jedynkę logiczną; zbocze narastające) wyprzedza (o czas propagacji) zmianę stanu wyjścia Q1 z wysokiego na niski. Reasumując. Wniesiona poprawka jest bardzo „tania”. A bardziej komfortowo jest, jeśli światło zaświeci się w reakcji na otwarcie drzwi; następnie zamknięcie drzwi nie wywoła żadnej reakcji; kolejne otwarcie drzwi też światła nie gasi; natomiast zamknięcie drzwi – światło zgasi. Natomiast, już najbliższe otwarcie drzwi – światło znów zaświeci.

Analizując schemat z rysunku 2 warto też zauważyć, iż błędem mniejszego kalibru jest brak kondensatora elektrolitycznego na wyjściu stabilizatora LM7805.

REKLAMA

Wydawnictwo AVT nawiąże współpracę redakcyjną z osobami dobrze operującymi terminologią elektroniki i słowem pisanym. Propozycja szczególnie interesująca dla nauczycieli elektroniki, autorów artykułów, skryptów i książek.

Aplikacje prosimy kierować na adres: redakcja@elportal.pl



Tani system automatyki z użyciem uP PIC16F676



Obecnie na rynku dostępnych jest bardzo wiele systemów szeroko rozumianej automatyki, zarówno domowej jak i przemysłowej. Często umożliwiają zdalne sterowanie w podczerwieni, wykorzystują łączę Bluetooth, wybieranie tonowe DTMF, domowe sieci Wi-Fi, łączność radiową RF, a nawet zdalne sterowanie głosem. W bieżącym projekcie proponujemy tanie i proste rozwiązanie, które można wykorzystać w automatyce domowej jak i w procesach przemysłowych. Możliwości systemu są jednakże ograniczone do zdalnego włączenia bądź wyłączenia dowolnego odbiornika energii elektrycznej.

Zdalne sterowanie pilotem w podczerwieni jest bardzo rozpowszechnione i nie warto w tym zakresie „odkrywać Ameryki”. Funkcjonuje kilka uzgodnionych protokołów wg których nadajnik komunikuje się z odbiornikiem. Najpopularniejsze to: protokół opracowany przez firmę NEC, Philips, JVC, SIRC (Sony Infrared Remote Control) lub RC5 który zdominował zdalne sterowanie w branży RTV. W bieżącym projekcie wykorzystamy standard NEC-a. **Od Red. EdW:** Można się domyślać, iż o takim wyborze zadecydowały względy marketingowe; pilot który pokazuje zdjęcie na **rysunku 1** jest mały poręczny i przede wszystkim bardzo tani równocześnie

umożliwia wysłanie 21 różnych kodów, co w tym przypadku w zupełności wystarcza.

Zatem nadajnik mamy gotowy, odbiornik chcemy wykonać na popularnym mikrokontrolerze z rodziny PIC. Dokładniej, wykonać musimy dekodery, gdyż sam odbiornik podczerwieni jest również dostępny za bardzo skromną cenę. To TSOP1738. Jednak, żeby zdekodować sygnał pilota, musimy dokładnie poznać jego strukturę. I od tego zaczniemy.

Strukturę ramki w której zakodowane są bity informacji pokazano na **rysunku 2**. Logika kodowania zera i jedynki logicznej zawiera się w czasie odstępu wysyłanych impulsów.

Logiczne 0 to impuls o czasie trwania 562,5 mikrosekundy za którym następuje przerwa, odstęp czasowy o takiej samej szerokości. To w sumie daje czas transmisji zera logicznego 1,125 ms. Logiczną „1” również rozpoczyna impuls (burst) o szerokości 562,5 μ s. Jednak po nim następuje przerwa dłuższa = 1,6875 ms. To daje czas transmisji jedynki równy 2,25 ms. Jednak, aby odbiornik „się nie pogubił”, na początku każdej ramki nadajnik wysyła sekwencję startową. Zatem, naciśnięcie dowolnego przycisku na pilocie skutkuje następującą sekwencją impulsów:

1. 9-cio milisekundowy „burst” rozpoczynający transmisję każdej ramki,

co odpowiada 16-tu okresom podstawowej jednostki impulsu (562,5 μ s) odliczanym przy przesyłaniu każdego bitu

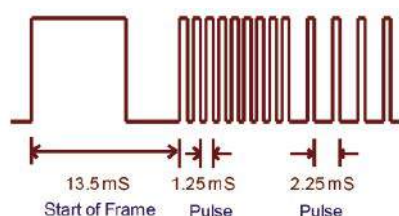
2. Przerwa 4,5 milisekundowa
3. 8-mio bitowy adres odbiornika
4. Powtórzony adres odbiornika w postaci zanegowanej
5. 8-mio bitowy rozkaz
6. 8-mio bitowy rozkaz zanegowany
7. Impuls (burst) kończący ramkę wiadomości, impuls o czasie trwania 562,5 μ s (czas przyjęty jako „jednostka” w protokole NEC-a).

Zatem, w ramce mieszczą się 4 bajty, przy czym adres i komenda są powtórzone w postaci „wprost” i zanegowanej. W każdym bajcie kolejność bitów ustalona jest od bitu najmniej znaczącego (LSB) do bitu najbardziej znaczącego (MSB). Przykładową ramkę pokazano na **rysunku 3**, adres 00-hexadecymalnie (00000000b) i rozkaz ADh (10101101b).

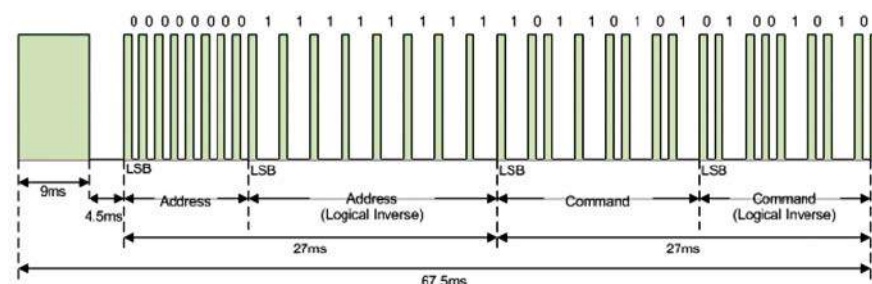
W sumie transmisja pełnej ramki trwa 67,5 milisekundy. 27 ms dla 16-tu bitów adresu (wprost plus zanegowany) i 27 ms dla transmisji 16 bitów rozkazu. Warto zauważyć, iż czas każdej ramki jest stały mimo różnego czasu dla transmisji zera i jedynki logicznej. Wynika to z tego, iż każdy „bit logiczny” transmitowany jest dwukrotnie. W postaci „wprost” i zanegowanej. Trzydziesto-dwu bitowy skład ramki



Rysunek 1. Pilot VIRE firmy NEC



Rysunek 2. Struktura ramki w której bity kodowane są w odstępie impulsów



Rysunek 3. Przykładowa ramka w formacie protokołu NEC

Tabela 1.			
Adres	Adres w negacji	Rozkaz	Rozkaz w negacji
LSB-MSB (0-7)	LSB-MSB (8-15)	LSB-MSB (16-23)	LSB-MSB (24-31)

Tabela 2			
Nr	Wartość dziesiętne	Wartość heksadecymalnie	Przycisk pilota
1	33441975	1FF448B7	OFF
2	33446055	1FE458A7	MODE
3	33444215	1FE47887	MUTE
4	33446255	1FE4807F	RESUME
5	33449935	1FE440BF	PREVIOUS
6	33442575	1FE4C03F	NEXT
7	33441775	1FE420DF	EQ
8	33444415	1FE4A05F	VOL-
9	33448095	1FE4609F	VOL+
10	33440735	1FE4E01F	0
11	33447695	1FE410EF	RPT
12	33440335	1FE4906F	U/SD
13	33444015	1FE450AF	1
14	33448695	1FE4D827	2
15	33446855	1FE47807	3
16	33445855	1FE430CF	4
17	33448495	1FE4B04F	5
18	33443615	1FE400FF	6
19	33442175	1FE4708F	7
20	33444815	1FE4F00F	8
21	33442375	1FE49867	9

protokołu NEC pokazano w tabeli 1. W tabeli 2 zebrano pełną listę kodów odpowiadających wszystkim przyciskom pilota z rysunku 1.

Mikrokontroler PIC16F676

PIC16F676 to ośmiobitowy mikrokontroler Microchipa wykonany w technologii CMOS. PIC16F676 zawiera w sobie pamięć flash i zamknięty jest w obudowie 14-to nóżkowej. Jednostka centralna CPU zawiera zredukowaną listę instrukcji RISC. Mikrokontrolery

PIC są często najlepszą opcją w systemach „wbudowanych” (embedded) i/lub w automatyce przemysłowej. Ten niewielki układ scalony oprócz właściwego mikrokontrolera zawiera wiele obwodów peryferyjnych, co często stanowi idealne rozwiązanie dla wielu indywidualnych projektów w których mikrokontroler stanowi jedynie centralną jednostkę sterującą. Podstawowe cechy chipa PIC16F676 są następujące:

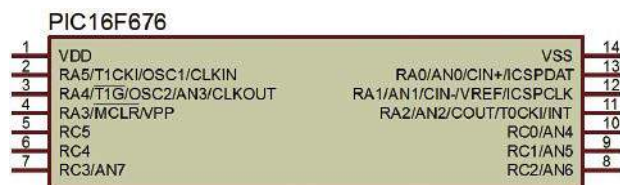
- wewnętrzna pamięć Flash ułatwia i przyspiesza proces opracowania i wykonania całego systemu z mikrokontrolerem
- mikrokontroler ten dostępny jest w obudowach PDIP, SOIC i TSOP; wszystkie o niewielkiej liczbie 14-tu wyprowadzeń,
- w PIC16F676 pamięć programu obejmuje przestrzeń 1,7 kB plus 64 bajty pamięci RAM oraz 128 B EEPROM-u,
- w wielu zastosowaniach cenną cechą jest obecność przetwornika analogowo-cyfrowego. To przetwornik 10-cio bitowy dostępny z aż ośmiu kanałów analogowych. Cecha ta jest szczególnie użyteczna gdy w systemie jest wiele czujników których dane trzeba przetworzyć na postać cyfrową,
- lista cech świadczących o użyteczności mikrokontrolerów PIC jest bardzo długa i trudno tu wszystkie wymienić. Oprócz przetworników ADC jest też analogowy komparator często niezbędny w systemach o charakterze analogowym. O popularności mikrokontrolerów PIC przesądza też cechy ułatwiające ich programowanie jak i możliwość załadowania programu do pamięci Flash bez konieczności wymontowywania mikrokontrolera z systemu. Na rysunku 4 przedstawiono opis wyprowadzeń układu PIC16F676 gdzie widać, że niewielka ilość wyprowadzeń okupiona jest tym, iż poszczególnym

nóżkom mogą być przydzielone różne funkcje. Wykorzystanie tych cech upraszcza hardware systemu z mikrokontrolerem, który jest z jednej strony uniwersalny, a także integruje wiele obwodów peryferyjnych systemu.

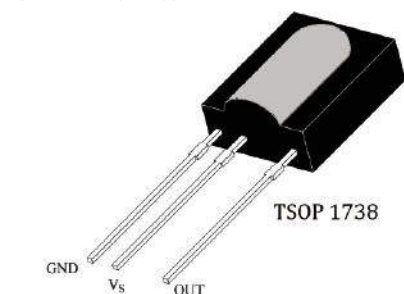
Opis układu i jego działanie

Bieżący projekt wykorzystuje pilot zdalnego sterowania IR NEC oraz mikrokontroler PIC po stronie odbiorczej. Przewidziano możliwość włączania/wyłączania dowolnych odbiorników zasilanych z sieci energetycznej, jak oświetlenie, wentylatory i tym podobne.

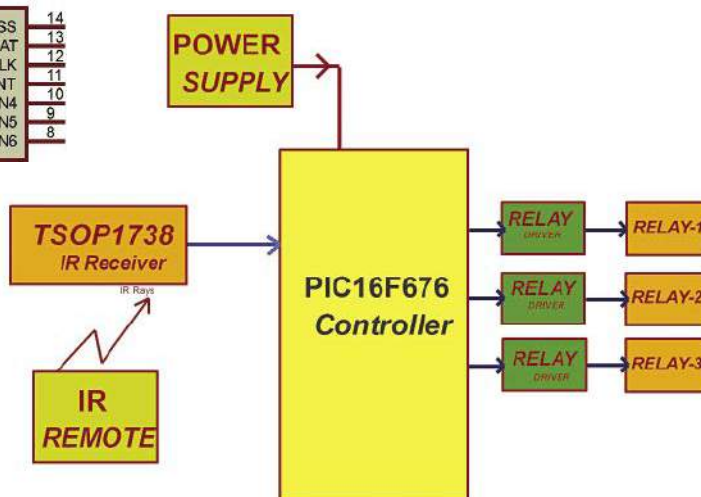
W tabeli 2 zebrano listę kodów wysyłanych z nadajnika-pilota. Każdy przycisk pilota jest opisany, jaką funkcję mu przydzielono. Jednak, o tym nie decyduje nadajnik ale odbiornik. Funkcja każdego przycisku pilota jest taka, jak interpretuje go mikrokontroler po stronie odbiorczej. Zatem, opis pilota można traktować stricte jako umowny. Ponieważ tu chcemy włączać i wyłączać obciążenia dużej mocy, niezbędne będą przełączniki i ich drivery, które zadowolą się niewielkim prądem jakim można obciążyć wyjścia mikrokontrolera. Od strony nadajnika oczekujemy tylko określonego kodu. Odbiornik ma go rozpoznać i wykonać odpowiednią reakcję. Jednak łączność odbywa się za pośrednictwem modulacji światła w zakresie podczerwieni. Większość nadajników pracuje na częstotliwości nośnej w zakresie 30 do 40 kHz. Wykorzystany pilot VIRE NEC wykorzystuje nośną 38 kHz. Zatem, jako front-end po stronie odbiorczej potrzeba odbiornika działającego na tej częstotliwości. Dostępna jest szeroka gama odbiorników TSOP17XX. Tu wykorzystano TSOP1738, którego wygląd i opis wyprowadzeń pokazano na rysunku 5.



Rysunek 4. Opis wyprowadzeń mikrokontrolera PIC16F676



Rysunek 5. Odbiornik podczerwieni TSOP1738



Rysunek 6. Schemat blokowy systemu zdalnego sterowania



Mikrokontroler PIC16F676 i odbiornik TSOP1738 wymagają zasilania 5-cio woltowego. Tu wykorzystano transformator sieciowy z symetrycznym uzwojeniem wtórnym

Działanie systemu jest proste. Naciśnięcie dowolnego przycisku pilota skutkuje wysłaniem sekwencji kodu przydzielonego danemu przyciskowi. Kodowanie zawarte jest w strumieniu impulsów modulujących częstotliwość



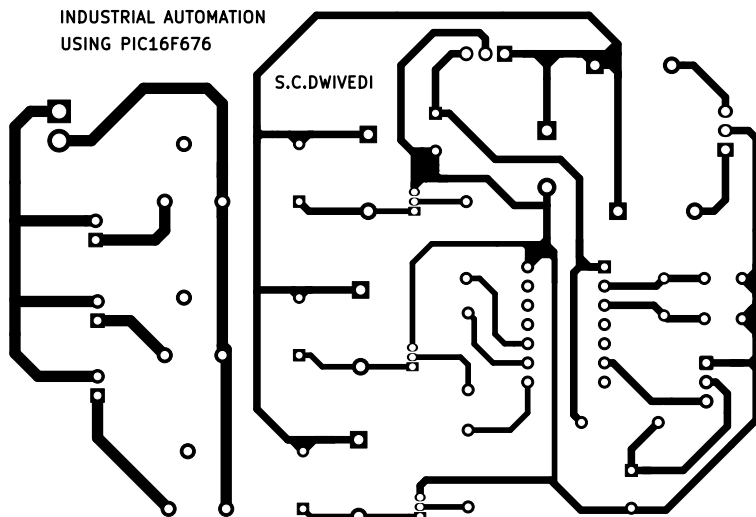
nośną 38 kHz, która z kolei moduluje strumień światła w podczerwieni. Transmisja jest odporna na zakłócenia dzięki wąskopasmowej charakterystyce nadajnika i odbiornika wokół nośnej 38 kHz. Na wejście RC4 mikrokontrolera trafiają już impulsy zgodne z protokołem pokazanym na rysunkach 2 i 3. Zadaniem mikrokontrolera jest ich zdekodowanie i odzy-

Projekt wg schematu z rysunku 7 można rozbudować. Tu jednak wykorzystano tylko 3 wyjścia cyfrowe. Przydzielono im przyciski pilota oznaczone numerami 2, 4 i 6. Naciśnięcie „2” włącza przełącznik RL1, „4” steruje RL2 i przycisk „6” – RL3.

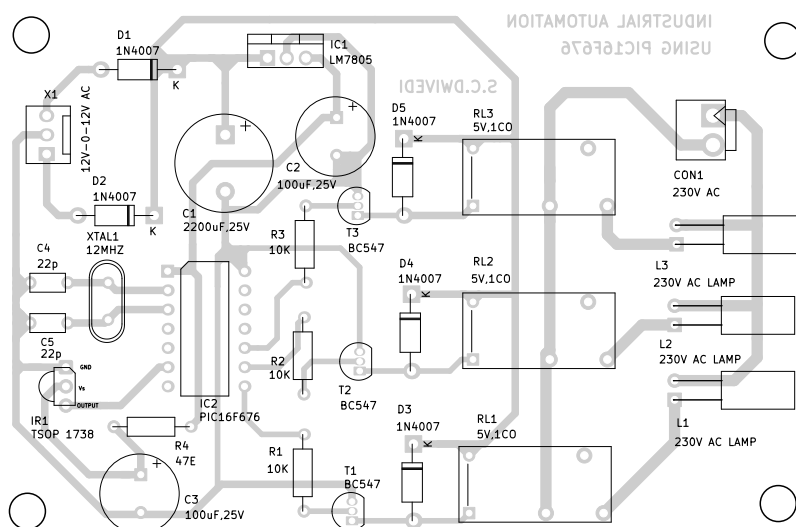
Od Red. EdW: Jak widać nie przydzielono osobnego kodu dla włączenia i wyłączenia przekaźnika. To znaczy, że kolejne naciśnięcia tego samego przycisku, na przemian



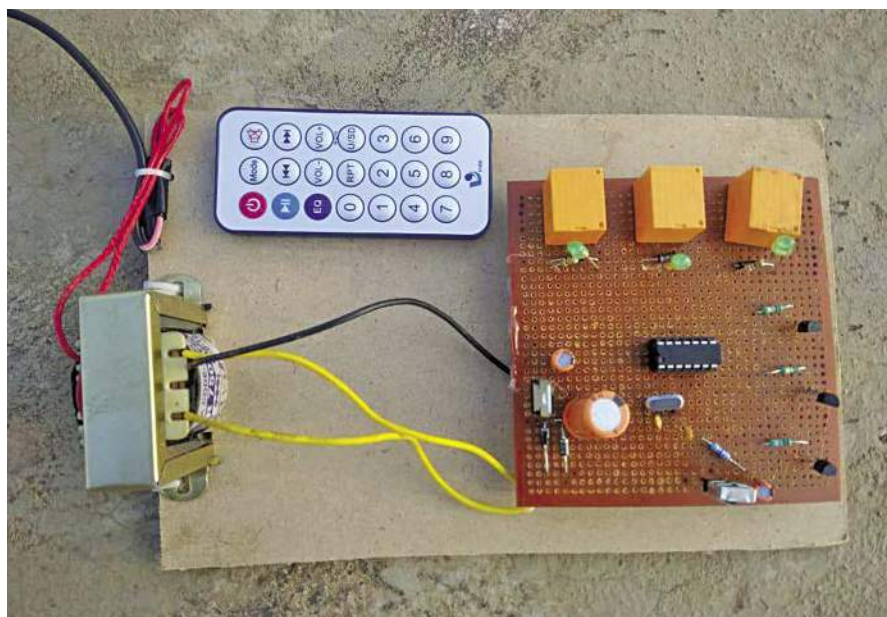
INDUSTRIAL AUTOMATION
USING PIC16F676



Rysunek 9. Płyta PCB „naturalnych rozmiarów”



Rysunek 10. Schemat montażowy elementów na PCB



Rysunek 11. Prototyp wykonany przez autora

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

Półprzewodniki:

IC1: LM7805 – stabilizator +5 V
IC2: PIC16F676 – mikrokontroler
T1-T3: BC547 – tranzystor npn
D1-D5: 1N4007 – dioda prostownicza
IR1: TSOP1738 – moduł odbiornika IR

Rezystory: (wszystkie 0,25 W $\pm 5\%$)
R1-R3: 10 k Ω

Kondensatory:

C1: 2200 μ F/25 V elektrolityczny
C2, C3: 10 μ F/25 V elektrolityczny
C4, C5: 22 pF ceramiczny

Pozostałe:

X1: transformator 230 VAC pierwotne – 12 V-0-12 V 1 A uzwojenie wtórne
CON1: złącze 2-pinowe
L1-L3: żarówki 230 VAC (dowolna moc)
RL1-RL3: przekaźnik 5 V
Pilot zdalnego sterowania NEC

włączają i wyłączają dany przekaźnik. Czyli układ zachowuje się jak przerzutnik typu T. Taki (lub inny) sposób działania leży wyłącznie po stronie programowej projektu.

Oprogramowanie

Układ działa pod kontrolą programu załadowanego do pamięci Flash mikrokontrolera. Program napisano w języku C i skompilowano kompilatorem Mikro C PRO w wersji 7.2.0. Program w języku C jest krótki i zwarty. Po wygenerowaniu kodu źródłowego, został on załadowany do pamięci programu MCU przy pomocy programatora K150.

W programie nie przewidziano wykorzystania przerw ani zaawansowanych trybów pracy mikroprocesora. Pin RC4 zadeklarowano jako cyfrowe wejście i procesor czyta je tak samo jakby był tam podłączony zwykły push-button. Każda zmiana stanu z wysokiego na niski lub odwrotnie, uruchamia procedurę zabezpieczenia przed drganiem styków i uruchamia timer. Odliczony czas przed kolejną zmianą stanu logicznego na linii RC4 jest podstawą decyzji jaki bit został wysłany. W każdym przypadku zmierzony czas trwania impulsu zapisywany jest w tablicy dla późniejszego porównania z wzorcem. Podstawowy impuls trwa 562,5 mikrosekundy, zaś odstęp jest różny gdy kodowana jest jedynka lub zero logiczne. Program w ten sposób kompletuje wszystkie 32 bity całej ramki. Wszystkie 4 bajty zapisywane są w tablicy utworzonej w RAM-ie. Następnie dla komparacji z wzorcem wykorzystywany jest tylko trzeci bajt (reszta może być wykorzystana dla kontroli poprawności transmisji). Możliwości takiego systemu są większe aniżeli wykorzystany tu tryb włącz/wyłącz jakiegoś urządzenia.

REFIKAMA

Robot sterowany gestami i klaśnięciem rąk

Bieżący projekt to w istocie „dwa w jednym”. Umożliwia zdalne sterowanie robotem w bardzo nietypowy sposób – prostymi gestami lub klaśnięciem rąk. Ruchy robota są ograniczone do czterech kierunków: w przód, do tyłu, w lewo i w prawo. Przy okazji projektu wyjaśnimy, jak można sterować urządzeniem za pośrednictwem prostych gestów lub przy pomocy sygnałów dźwiękowych.

W części odbiorczej urządzenia zainstalowano przełącznik dla wyboru trybu pracy. Czyli, czy układ ma reagować na dźwięki czy na gesty. Robot taki może realizować poważne zadania lub może być ciekawą zabawką, jeśli będzie podążał za Tobą np. w trakcie jogi lub innych ćwiczeń. Można również robota skonstruować tak, aby reagował na różne dźwięki w jego otoczeniu.

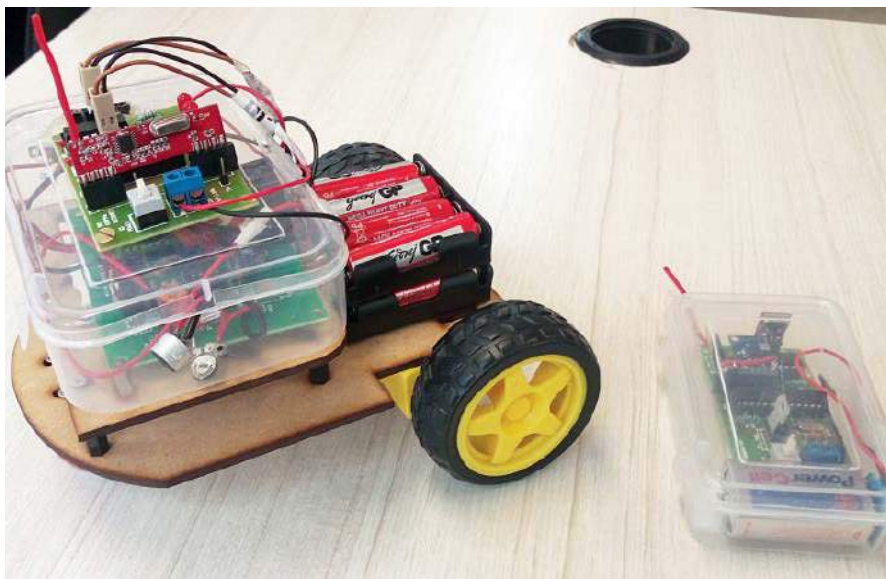
W schemacie wyróżnimy dwie oddzielne sekcje dla obu trybów pracy. Obydwie sekcje mają też oddzielne czujniki. W sekcji, która ma reagować na klaśnięcia czujnikiem jest mikrofon elektretowy. W sekcji gestów jako czujnik ruchu zastosowano akcelerometr.

W sekcji reagującej na dźwięki wykonano swoisty obwód przełączający, który jest aktywowany wtedy, gdy układ rozpozna dźwięk odpowiadający klaśnięciu w dłonie. Obwód składa się z czujnika dźwięku, którym jak powiedziano jest mikrofon. Kolejnymi podzespołami jest timer 555 oraz licznik dekadowy. W sekcji sterowanej gestami oprócz części odbiorczej dochodzi nadajnik. Ruch wykrywany jest przez akcelerometr, który współpracuje z mikrokontrolerem ATmega328P (MCU). Na płycie transmittera jest także koder sygnału oraz bezprzewodowy nadajnik radiowy.

Dalszy opis działania podporządkowany jest ogólnej charakterystyce układu, w której są opisane wspomniane sekcje. Sekcja „clap control” zawiera tylko odbiornik, podczas gdy w sekcji „gestów” mamy odbiornik i nadajnik sygnałów. Cały schemat odbiornika został pokazany na **rysunku 1**. Ta płytka jest fizycznie ulokowana w robocie i może być ukryta jeśli konstrukcja mechaniczna robota na to pozwala.

Opis sekcji sterowania dźwiękiem

W tej części zastosowano mikrofon pojemnościowy (MIC1), wzmacniacz operacyjny NE5534 (IC1), timer NE555 (IC2), licznik dekadowy 4017B (IC3), driver silników L293D (IC4), dwa silniki prądu stałego (M1 i M2), a także potencjometr ustalający próg czułości układu rozpoznającego dźwięki charakterystyczne dla klaśnięć rąk oraz niewiele dodatkowych prostych pasywnych elementów.



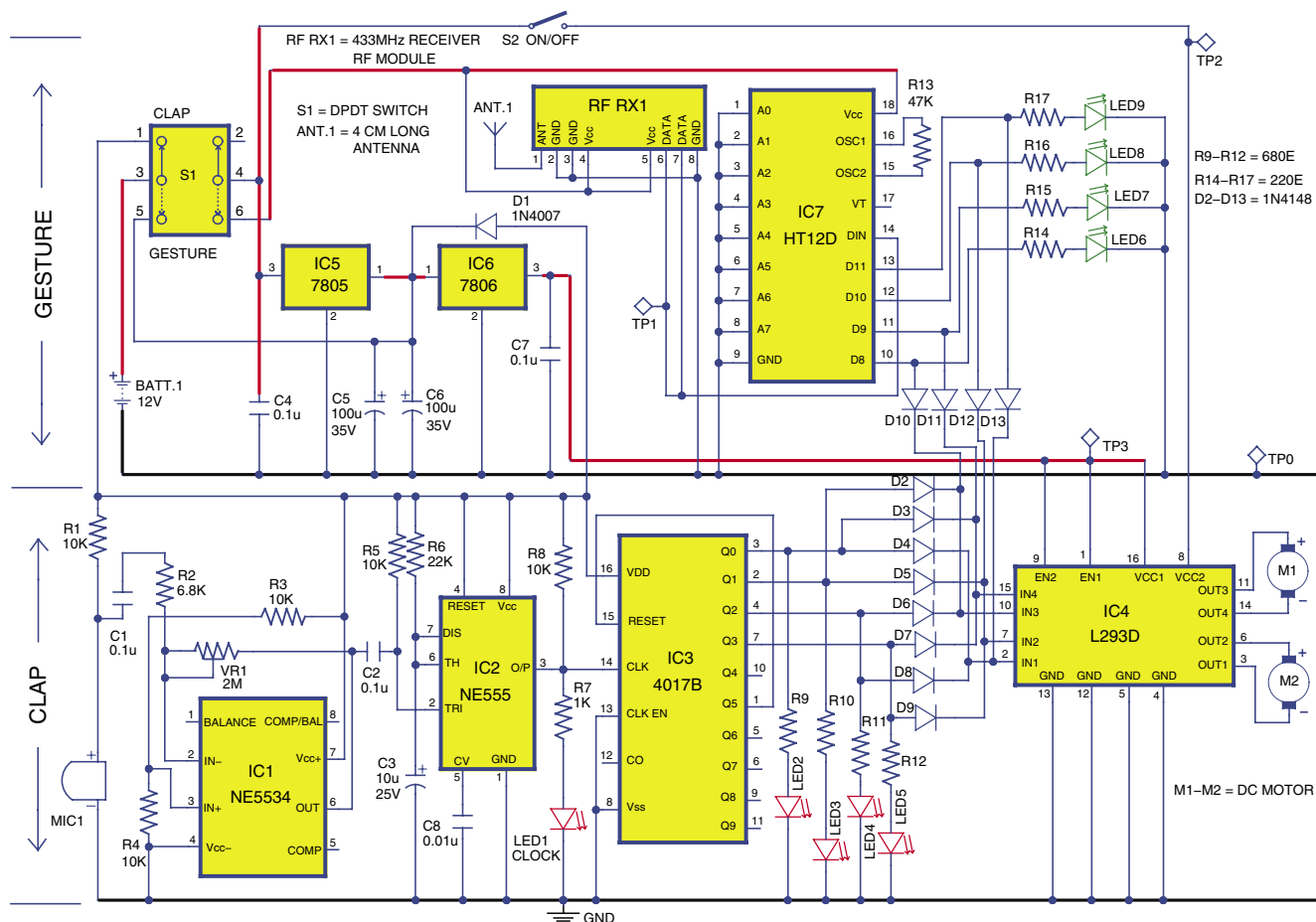
Każdy dźwięk rozpoznany jako klaśnięcie jest komendą wywołującą ruch robota. Klaśnięcia zliczane są „modulo 5” i kolejno wywołują ruch: w przód, do tyłu, skręt w lewo i w prawo. Piąte klaśnięcie zatrzymuje ruch robota. Ponieważ rozpoznawanie dźwięku klaśnięcia rąk może być obciążone błędem i przekłamaniami, w obwodzie zastosowano potencjometr VR1 ustalający wzmocnienie wzmacniacza. Teoretycznie powinien on ustalać odległość z jakiej robot powinien reagować na polecenia głosowe. Układ IC1 pracuje w konfiguracji wzmacniacza odwracającego wzmacniając sygnał wygenerowany przez mikrofon. Timer IC2 pracuje jako przerzutnik monostabilny. Stanowi on opóźnienie i generuje zegar dla licznika IC3. Wyjściem monoflopa jest pin 3 układu IC2 i jest ono bezpośrednio połączone z pinem 14 IC3.

Od Red. EdW. Licznik 4017 nalicza na narastające zbocze impulsu zegarowego, co odpowiada chwili wyzwolenia nonoflopa; taka konfiguracja sprawia, iż licznik nie będzie naliczał szybciej aniżeli wynosi czas przerzutnika monostabilnego; w ten sposób jest „czyszczony” impuls zegarowy i jest to zabieg podobny do zabezpieczenia przed drganiem styków mechanicznych obligatoryjnie stosowanego w klawiaturach współpracujących z „szybką elektroniką”.

Licznik IC3 zlicza klaśnięcia i równocześnie jest rejestrem pamiętającym stan pracy robota. Pierwsze klaśnięcie ustawia do stanu wysokiego wyjście Q0 licznika IC3. Dodatkowo wykonano indykację stanu rejestru na diodach LED. Stan aktywny (wysoki) Q0 powinien zaświecić diodę LED2. Reakcją na drugie klaśnięcie powinien być stan wysoki wyjścia Q1. Równocześnie Q0 powinno przejść do stanu

Tabela 1. Sekwencja poleceń dźwiękowych i informacja pokazana na diodach LED

Klaśnięcie	Aktywne wyjście IC3	Ruch robota	Dioda LED
Pierwsze	Q0	Do przodu	LED2 wł.
Drugie	Q1	Do tyłu	LED3 wł.
Trzecie	Q2	W prawo	LED4 wł.
Czwarte	Q3	W lewo	LED5 wł.
Piąte	Q4	Stop	LED2...LED5 wył.



Rysunek 1. Schemat części odbiorczej robota sterowanego kłaśnięciami lub prostymi gestami

nieaktywnego, niskiego. Powinna zaświecić dioda LED3, a LED2 powinna zgasnąć. W ten sam sposób Q2 i Q3 powinny reagować na trzecie i czwarte kłaśnięcie dłońmi. Tą sekwencję wraz ze statusem diod LED pokazuje **tabela 1**.

Diody świecące zastosowano tylko dla wizualizacji stanu robota (i bez nich cały układ też powinien działać). Ważniejszy jest interfejs między rejestrem stanu Q0, Q1, Q2 i Q3 oraz driverem silników. Tutaj między wyjściami IC3, a wejściem IC4, wstawiono diody D2 do D9, które pełnią proste funkcje logiczne. Kiedy Q0 jest w stanie wysokim, robot powinien posuwać się do przodu. Kiedy Q1 przyjmie stan aktywny, robot ma posuwać się w tył. Analogicznie, stan wysoki Q2 ma powodować skręt w prawo, a wysoki poziom na wyjściu licznika Q3 powinien skutkować ruchem w lewo. Licznik ma cztery wyjścia, i jest to wyjście typu „1 z N”, ale licznik skonfigurowano jako modulo-5. Piąte kłaśnięcie resetuje bowiem wszystkie przerzutniki licznika. Czyli Q0, Q1, Q2 i Q3 przyjmują stan niski i robot ma się zatrzymać. Powinny też zgasnąć wszystkie diody LED2 do LED5. Zatem podsumowując – każde kłaśnięcie powinno zaświecać kolejną

z diod LED2 do LED5. Powinien temu towarzyszyć odpowiedni ruch robota. Jeśli ruch nie będzie zgodny z przypisanym wg tabeli 1, może zająć potrzeba zmiany polaryzacji silnika M1 i/lub M2.

Sterowanie robota za pomocą gestów

Ta sekcja robota składa się z nadajnika i odbiornika.

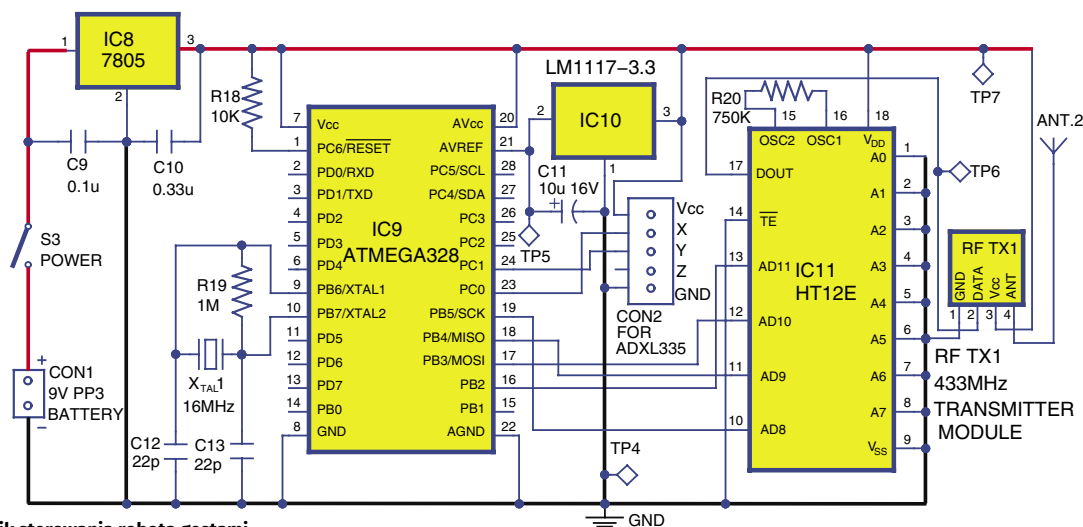
Działanie nadajnika

Schemat nadajnika pokazano na **rysunku 2**. Układ wykonano na bazie następujących elementów: mikrokontroler Atmega328 (IC9), akcelerometr ADXL335, stabilizator napięcia 7805 (IC8), 3,3-woltowy stabilizator LDO LM1117-3,3 (IC10), koder HT12E (IC11), moduł nadajnika radiowego RF 433MHz (RFTX1) i ponadto kilka prostych pasywnych elementów.

Kluczowym elementem jest tu akcelerometr ADXL335. To element czuły na przeciążenia. Mierzy przyspieszenie i rozpoznaje kierunek. W trójwymiarowej przestrzeni „widzi” współrzędne kartezjańskie x, y i z. Współrzędne usytuowane są tak, iż liniowe

przyspieszenie zgodne z kierunkami przypisanymi obudowie tego czujnika dadzą dodatni wynik osobno dla wyjść odpowiadających kierunkom x, y i z. Wyjście ma postać sygnału analogowego, napięcia, którego wartość jest proporcjonalna do przyspieszenia. Zakres miernika to $\pm 3g$ dla każdego wyjścia ($g \approx 9,81 m/s^2$).

Sterowanie gestami jest tu rozumiane jako przyspieszenie dodatnie bądź ujemne wzdłuż kierunków x i y. Cały moduł nadajnika z akcelerometrem powinien zmieścić się w dłoni. Kiedy wykonasz ruch skutkujący przyspieszeniem wzdłuż kierunku x i y, wyjście OUT przypisane tym kierunkom zostanie odczytane przez mikrokontroler Atmega328. Informacja zostanie wysłana przez nadajnik radiowy na module RF TX1 i powinna być odczytana przez odbiornik zamontowany w obudowie robota. W nadajniku jest jeszcze element kodera HT12E, który pośredniczy między mikrokontrolerem a modułem transmittera. Wartość przyspieszenia odpowiadająca gestom nie jest bardzo istotna. Także ilość gestów jest ograniczona do czterech zgodnie z tym co opisano w sekcji sterowania kłaśnięciami rąk.



Rysunek 2. Nadajnik sterowania robota gestami

Wykaz elementów:**Rezystory:** (wszystkie 0,25W, ±5%)

R1, R3-R5
R8, R18: 10 kΩ
R2: 6,8 kΩ
R6: 22 kΩ
R7: 1 kΩ
R9...R12: 680 Ω
R13: 47 kΩ
R14-R17: 220 Ω
R19: 1 MΩ
R20: 750 kΩ
VR1: 2 MΩ – potencjometr montażowy

Kondensatory:

C1, C2, C4
C7, C9: 0,1 μF ceramiczny
C3: 10 μF/25V elektrolityczny
C5, C6: 100 μF/35V elektrolityczny
C8: 0,01 μF ceramiczny
C10: 0,33 μF ceramiczny
C11: 10 μF/16V elektrolityczny
C12, C13: 22 pF ceramiczny

Półprzewodniki:

IC1: NE5534 lub LM741 – wzmacniacz operacyjny
IC2: NE555 – timer
IC3: 4017B – licznik dekadowy
IC4: L293D – driver silników
IC5, IC8: 7805 – stabilizator 5 V
IC6: 7806 – stabilizator 6 V
IC7: HT12D – dekodery
IC9: Atmega328 – mikrokontroler
IC10: LM1117-3.3 – stabilizator 3,3 V
IC11: HT12E – kodery
D1: 1N4007 – dioda prostownicza
D2...D13: 1N4148 – diody sygnałowe
LED1...LED9: diody LED 5 mm

Inne:

S1: przełącznik podwójny DPDT
S2, S3: przełącznik ON/OFF
BATT1: bateria 12 V (8×1,5V AA)
CON1: złącze 2-pinowe
CON2: 5-pinowe złącze typu Berg żeńskie
MIC1: mikrofon pojemnościowy
XTAL1: kwarc 16 MHz
RF-TX1: moduł nadajnika RF 433 MHz
RF-RX1: moduł odbiornika RF 433 MHz
ANT1, ANT2: antena – sztywny drut 4cm
M1, M2: silnik DC 5 V/6 V z przekładnią

Pozostałe:

bateria 9 V PP3
pojemnik baterii 8×1,5 V AA
podstawa pod układ scalony 18-pin (2 szt.)
podstawa pod układ scalony 16-pin (2 szt.)
podstawa pod układ scalony 8-pin (2 szt.)
podstawa 20-pin pod mikrokontroler MCU
obudowa robota
dwa koła na osie silników M1 i M2
koto „mimośrodowe” – „rolka” jako koto przednie
płytki PCB

W nadajniku zastosowano mikrokontroler Atmega328, więc trzeba go oprogramować. Można do tego użyć narzędzia programistyczne platformy Arduino. Program dla robota napisano w języku Arduino. „Świeżą” pamięć mikroprocesora najłatwiej zaprogramować używając oprogramowania IDE na MCU umieszczonym na płytce Arduino Uno.

W pierwszej kolejności trzeba załadować kod bootloadera. W tym celu w IDE wybierz procedurę programowania w systemie (In System Programming) wpisując: File → Examples → ArduinoISP. Przy pomocy bootloadera załadujesz kod źródłowy dla naszego projektu o nazwie *gesture.ino*.

Działanie odbiornika sterowania gestami

Schemat odbiornika pokazuje rysunek 1. W tej sekcji zastosowano: moduł odbiornika radiowego RFRX1 (pracujący na częstotliwości 433 MHz), dekodery HT12D (IC7) oraz fragment wspólny dla obu sekcji, czyli sterownik silnika wraz z dwoma silnikami DC.

Dwubiegunowy przełącznik DPDT (S1) przełącza zasilanie między sekcjami „clap” i „gesture”. Odbiornik RF RX1 jest sparowany z nadajnikiem RF TX1 i odbiera dane rozpoznane przez czujnik akcelerometryczny. Wyjście odbiornika radiowego doprowadzone jest do dekodera i dalej do drivera silnika L293D poprzez diody D10 do D13. Tu także w celu indykacji wizualnej zastosowano diody LED (LED6 do LED9).

W odbiorniku robota istotną rolę pełni także przełącznik S1. Zawiera on dwie pary styków i służy do przełączania trybu pracy. Kiedy przełącznik ten jest w położeniu „clap” tylko sekcja odpowiedzialna za ten tryb jest aktywna i zasilana. Analogicznie jest z trybem odpowiedzialnym za sterowanie robota gestami. Tylko elementy odpowiedzialne za ten tryb są zasilane. W odbiorniku jest jeszcze drugi przełącznik S2. To pojedynczy jednobiegunowy przełącznik, który w położeniu ON zasila driver silników M1 i M2. W węzłach sumowania się sygnałów sterujących driverem zainstalowano aż 8 diod sygnałowych (D2 do D13), które pełnią rolę funkcyj logicznych.

REKLAMA

Certyfikat Underwriters Laboratories

94V-0 E480148 TYPE I

Zakład produkcyjny:

05-650 Warka
ul. M. Ropielewskiej 17
tel. 22 781 63 95
22 781 66 88
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 aluminiowym	Maksymalny wymiar płytek 1w 630 mm	Płyty czelowe FR4	Krótkie terminy
Aktywny kalkulator prototypów na stronie internetowej	Pakowanie Sn lub SnPb inne na życzenie	Trawione szablony SMD	Wykonania super ekspresowe
Maski, opisy montażowe w różnych kolorach			

Konstrukcja układu i jego testowanie

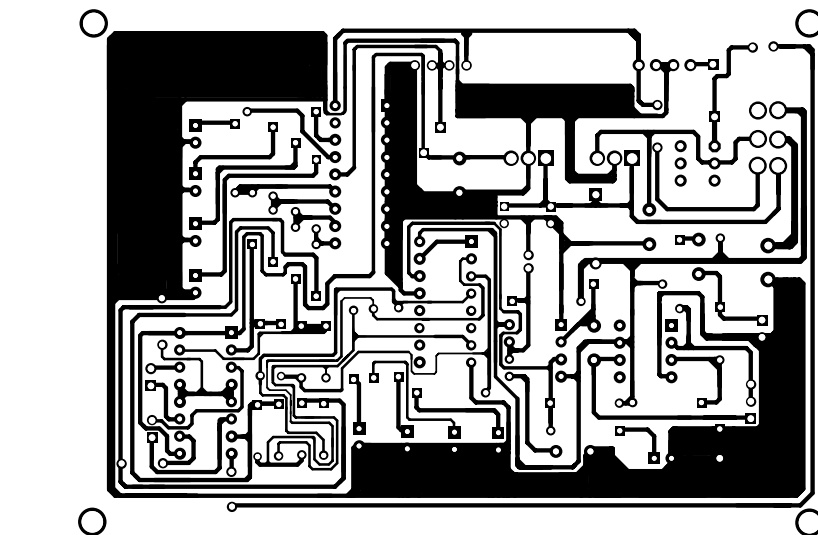
Na **rysunku 3** zamieszczono projekt jednostronnej płytki drukowanej dla odbiornika robota. **Rysunek 4** pokazuje rozmieszczenie elementów na PCB. Odpowiednio, PCB i schemat montażowy dla nadajnika pokazano na **rysunkach 5 i 6**.

Po uzbrojeniu płytek PCB zgodnie ze schematem ideowym, należy je umieścić w odpowiednio przygotowanych obudowach. W dostępnym miejscu zamontuj przełączniki S1 i S2. Przygotowując obudowę należy nie zapomnieć o odpowiednio przygotowanych otworach dla anteny i mikrofonu, tak aby zapewnić nie zakłócającą komunikację radiową i dźwiękową. Na chassis robota zamocuj baterię BATT1. Baterię i silniki M1 i M2 podłącz do odpowiednio przygotowanych złączy na płycie odbiornika.

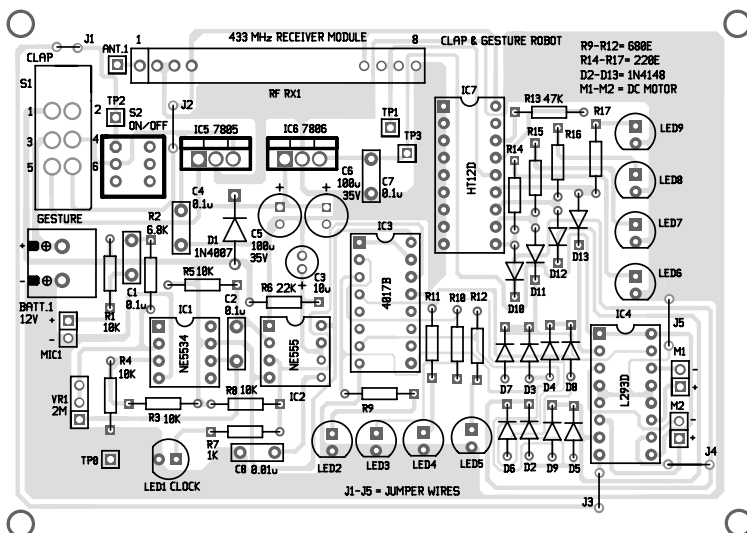
Podobnie należy postąpić z płytą nadajnika. Tu w obudowie należy przewidzieć otwory dla anteny ANT2 i dla przełącznika S3. Źródłem zasilania w nadajniku może być bateria 9 V PP3. Baterię tę podłącz do złącza CON1.

Po wykonaniu prac montażowych można przystąpić do przetestowania pracy robota. Przełącz S1 w położenie „clap mode” i S2 w pozycję ON. Wykonaj kłasknięcia dłońmi w pobliżu mikrofonu. Przy pomocy śrubokręta lekko obracaj suwak potencjometru VR1 ustawiając w ten sposób pożądaną czułość odbiornika. Każde kolejne kłasknięcie powinno skutkować zmianą kierunku poruszania się robota. Piąte kłasknięcie powinno skutkować zatrzymaniem robota. Równocześnie status odbiornika powinien być widoczny na diodach LED2 do LED5.

Następnie przełącz S1 w położenie „gesture mode” oraz przełączniki S2 i S3 w położenie ON. Trzymając w ręku nadajnik wykonaj energiczny ruch w wybranym kierunku. Zwróć uwagę na usytuowanie kierunku nadajnika względem położenia odbiornika. Robot powinien podążać w tym samym kierunku jak ruch nadajnika. W razie potrzeby odwróć biegunowość podłączenia



Rysunek 3. Projekt jednostronnej płytki PCB odbiornika zdalnie sterowanego robota



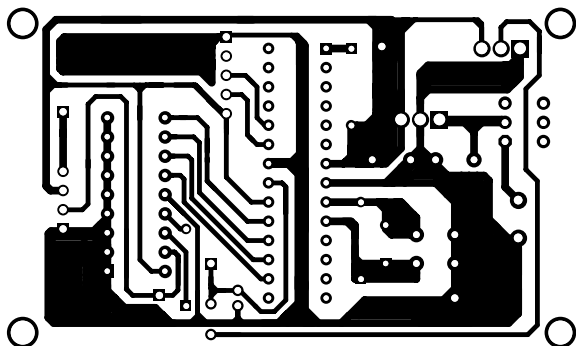
Rysunek 4. Rozmieszczenie elementów na płycie z rysunku 3

silników DC do drivera L293D i/lub zamień miejscami M1 z M2.

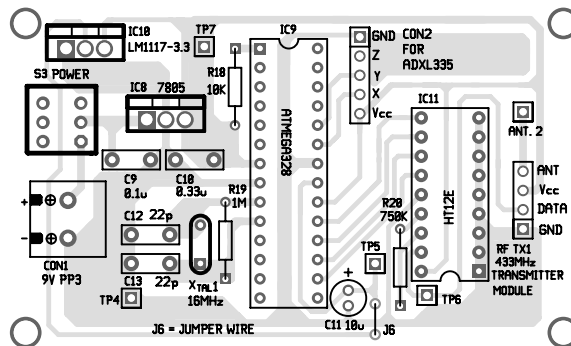
W trybie „gesture mode” istotne jest usytuowanie akcelerometru wzdłuż każdej osi przestrzeni (x, y i z – kierunek pionu). Kalibrację najłatwiej wykonać eksperymentalnie

obserwując ruchy robota w reakcji na ruchy akcelerometru. Zadanie to ułatwia także obecność diod LED (LED6 do LED9).

Tak wykonany robot może wykonywać poważne zadania lub służyć do zabawy. Nadajnik możesz ukryć w kieszeni koszuli lub spodni.



Rysunek 5. Płytkę PCB nadajnika sterowania robota gestami



Rysunek 6. Schemat montażowy płytki PCB nadajnika

Tabela 2. Punkty testowe		
Punkt testowy	Opis	Pozycja przełącznika S1
TP0	masa (GND)	Clap lub Gesture
TP1	ciąg impulsów	Gesture
TP2	5 V jeśli S2 – ON	Clap lub Gesture
TP3	6 V	Clap lub Gesture
TP4	0 V (GND)	Clap lub Gesture
TP5	3,3 V jeśli S3 – ON	Gesture
TP6	5 V jeśli S3 – ON	Gesture

Jeśli układ precyzyjnie doastroisz i będzie wystarczająco czuły, robot będzie podążał za twoimi ruchami bez widocznych gestów transmierem. Może to być zabawne dla twoich przyjaciół, którzy będą się dziwić jak to jest możliwe bez żadnych widocznych więzów. Możesz np. ćwiczyć Jogo z nadajnikiem w kieszeni, a robot będzie naśladował twoje ruchy. To jedynie prosty przykład dla demonstracji możliwości robota, aczkolwiek można wymyślać wiele poważniejszych dla niego zastosowań.

Dla ułatwienia uruchamiania układu, na schemacie przewidziano kilka punktów testowych, które zebrano w tabeli 2. Tabela podaje napięcia w punktach TP0 do TP6 w zależności od trybu (położenia przełącznika S1). Uruchamiając układ zwróć także uwagę na ew. zimne luty, poprawność połączeń, temperaturę

układów scalonych, napięcie baterii i polaryzację podłączenia silników. Zdjęcie prototypu robota wykonanego przez autora pokazuje fotografia tytułowa. ■

Sani Theo

Ten artykuł był wcześniej opublikowany na łamach „EFY”, styczeń 2020 (efymag.com)

Od Red. EdW: W nadajniku zastosowano koder HT12E jako para z dekoderm HT12D pracującym w odbiorniku. To logiczna konstrukcja hardware-u robota. Ale skoro w nadajniku jest też mikrokontroler, to z pewnością można zlecić mu też zadanie, które pełni koder HT12E. Atmega328 jest wydajnym mikroprocesorem z bogatymi peryferiami. Zatem za cenę dodatkowego software-u (a pamięci

jest też dość) można układ uprościć eliminując 18-wyprowadzeniowy układ scalony.

Na krótki komentarz zasługuje też zastosowanie licznika dekadowego 4017. To w istocie licznik Johnsona z dekoderm dziesięciu stanów (jednej dekady). Tutaj ma zliczać do pięciu, po czym ma wrócić do stanu zerowego. Takie działanie ma zapewnić zapełnienie wyjścia Q5 z wejściem Reset. Stan spoczynku, to aktywne wyjście Q4. Gdy pojawi się jedynka logiczna na wyjściu Q5, układ natychmiast wraca do stanu „zero”, kiedy aktywne jest Q0. Niby wszystko w porządku, tak ma być. Prawdopodobnie układ będzie działał zgodnie z oczekiwaniami konstruktora. Mimo to, jest tu niedopatrzanie lub nawet błąd konstrukcyjny. Czas aktywności wyjścia Q5 jest niezdefiniowany i w układzie występuje tzw. hazard. Nie ma pewności, że wszystkie przerzutniki licznika (5 sztuk) zdążą się wyzerować. Gdy zmieni stan przerzutnik najszybszy, wejście Reset wróci do stanu nieaktywnego. Takich zachowań układu należy unikać, nawet jeśli w praktyce problem nie ujawnia się. A więc, co zrobić? Tak całkiem poprawnie, to powinien być krótki monoflop w pętli sprzężenia zwrotnego między Q5 i Reset. Ale praktycznie można to załatwić dwójnikiem RC w tej pętli (o nawet bardzo krótkiej stałej czasowej, na poziomie czasów propagacji w logice licznika 4017).

REKLAMA

UWAGA! Tylko prenumeratorzy czasopism Elektronika dla Wszystkich, Elektronika Praktyczna, Świat Radio oraz Elektronik mogą korzystać z atrakcyjnych rabatów w Sklepie AVT:

- ✓ do 50% na wydania specjalne czasopism Wydawnictwa AVT
- ✓ 20% na kity w wersji A (płytki drukowane do projektów AVT)
- ✓ 10% na pozostałe wersje kitów: (A+, B, C, D)
- ✓ 10% na książki
- ✓ 5% na pozostałe produkty z oferty sklepu

K L U B
AVT
ELEKTRONIKA

Ponadto każdy prenumerator ww. czasopism korzysta z rabatów od 30% do 50% na zakup czasopism z oferty www.UlubionyKiosk.pl

Jak uzyskać rabat? Podczas zamówienia powołaj się na swój numer prenumeraty – otrzymasz go mailowo po zakupie prenumeraty wraz z kartą członkowską Klubu AVT-Elektronika.

Regulamin Klubu AVT-Elektronika znajdziesz na stronie <https://sklep.avt.pl/klub-avt-elektronika>

Elektroniczna maszyna do głosowania w demokratycznych wyborach

Motywacją dla bieżącego projektu jest stworzenie elektronicznej maszyny do głosowania, która może zastąpić tradycyjne głosowanie „papierowe”. Mimo prostoty konstrukcji i niskiego kosztu, maszyna taka wykazuje szereg zalet w porównaniu z ręcznym sposobem głosowania. Proces głosowania przebiega szybciej, a system jest odporny na ew. próby oszukiwania. EVM (Electronic Voting Machine) można zastosować w małych komisjach wyborczych, w szkołach, na uczelni lub w ośrodkach wiejskich. EVM wg bieżącego projektu jest urządzeniem bardzo tanim, a zdecydowanie usprawni proces głosowania i zapewni uczciwość tego procesu.

Założeniem projektu jest wybór jednego (lub więcej) spośród ośmiu kandydatów. Dla każdego kandydata przewidziano jeden z ośmiu przycisków ponumerowanych od S1 do S8. Dziewiąty przycisk oznaczony S9 przewidziano dla wyświetlenia wyniku na wyświetlaczu. Poza tymi przyciskami, w projekcie przewidziano diodę LED oraz buzzer dla indykacji statusu podczas głosowania. Projekt bazuje na płytce Arduino Uno, w której wykorzystano wejścia/wyjścia cyfrowe. Diodę LED podłączono na D3 (pin 18), a dla obsługi buzzera w programie przewidziano wyjście D2 (pin 17). Każdy z głosujących ma prawo wyboru tylko jednego kandydata naciskając jeden z przycisków S1 do S8. Końcowy rezultat ilości zebranych głosów wyświetlany jest po naciśnięciu przycisku S9.

Prototyp układu wykonanego przez autora pokazano na **rysunku 1**.

Opis układu i jego działanie

Schemat ideowy „maszynki EVM” widzimy na **rysunku 2**. Oprócz Arduino Uno wykorzystano tu moduł wyświetlacza LCD z szeregową komunikacją I²C (module 1). Ponadto jest tu 9 przycisków niestabilnych (S1 do S9), dioda LED, jeden rezystor i buzzer (PZ1). Autor zmontował swój projekt na płytce uniwersalnej, potrzebna jest więc także odpowiednia ilość przewodów zakończonych pinami umożliwiającymi wykorzystanie ich jako zworki.

Wykaz elementów:

Półprzewodniki:

Board1: płytka Arduino
LED1: dioda LED 5 mm

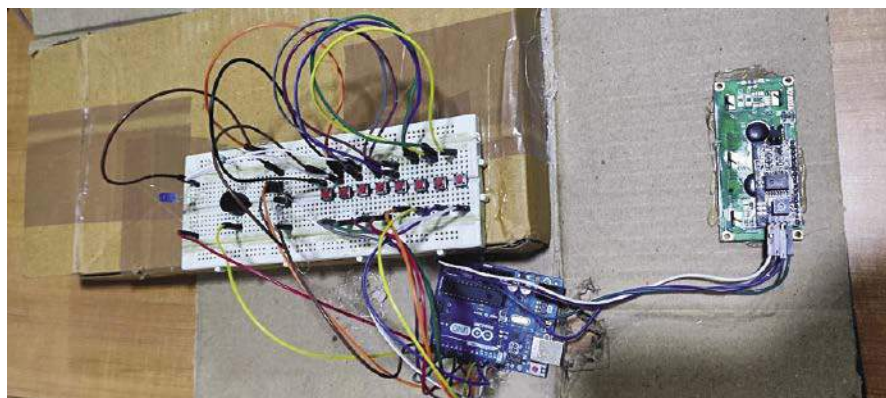
Rezystory: (wszystkie 0,25 W/±5%)
R1: 470 Ω

Inne:

PZ1: buzzer piezoelektryczny
S1-S9: switch „push-to-on”

Ponadto:

uniwersalna płytka „Breadboard”
zworki z przewodami
moduł LCD I²C (RG1602A)



Rysunek 1. Prototyp wykonany przez autora

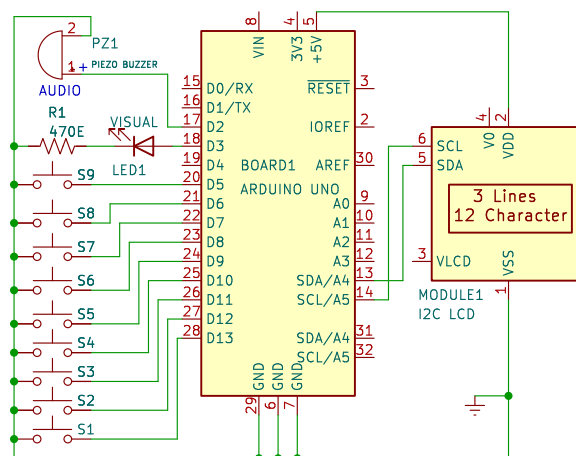
Większa część projektu bazuje na programie. Moduł Arduino zawiera mikrokontroler ATmega 328P, w którym dostępnych jest 14 cyfrowych wejść/wyjść, z których 6 można skonfigurować jako wyjścia PWM. Mikrokontroler zawiera także 6 wejść analogowych, których tu nie wykorzystujemy. Wyjścia A4 i A5 są oprogramowane dla obsługi magistrali I²C. Na płytce Arduino jest także rezonator kwarcowy 16 MHz, złącze USB dla komunikacji z komputerem i/lub dla zasilania, a także złącze typu jack, które można alternatywnie wykorzystać dla zasilania płytki Arduino. Jest tu także przycisk resetu oraz złącze ICSP, które można wykorzystać w procesie zapisania programu do pamięci mikrokontrolera.

Drugim modulem wykorzystanym w tym projekcie jest wyświetlacz znakowy typu LCD. Bazuje on na układzie scalonym HD44780 oraz wykorzystuje transmisję zgodną z protokołem szeregową magistrali I²C. Wykorzystano tu układ scalony PCF8574 w roli adaptera I²C. Ten ośmiobitowy

ekspander dokonuje konwersji szeregowych danych wysyłanych z mikroprocesora na postać równoległą zrozumiałą przez kontroler wyświetlacza. Obsługa układu jest bardzo prosta. Głosowanie polega na wyborze jednego z przycisków przypisanych kandydatom. Naciśnięcie S1 do S8 potwierdzone jest dźwiękiem z buzzera oraz zaświeceniem diody LED1.

Oprogramowanie

Program obsługujący ten projekt nie jest skomplikowany i napisano go z wykorzystaniem



Rysunek 2. Schemat ideowy maszyny do głosowania z wykorzystaniem Arduino Uno

Arduino IDE w wersji 1.8.19. Nie ma jednak problemu z adaptacją programu do innej wersji IDE. Przyciski S1 do S8 podłączono odpowiednio do pinów D13, D12, D11, D10, D9, D8, D7 i D6 płytki Arduino. Switch S9 obsługiwany jest przez wejście D5 portu cyfrowego. Zwarcie tego przycisku powinno wywołać podprogram, który wyświetli rezultat głosowania na wyświetlaczu LCD. Program dla tego projektu jest dostępny pod nazwą Arduino_Code.ino. W celu przegrania szkicu należy wybrać wersję wykorzystanej płytki Arduino oraz numer portu, pod którym Arduino Board jest widziana.

Konstrukcja układu i testowanie działania

W pierwszej kolejności należy wgrać oprogramowanie do mikrokontrolera na Arduino. Teraz zastosowane podzespoły można połączyć wykorzystując uniwersalną płytkę PCB. Finalnie należy przewidzieć obudowę z uwzględnieniem miejsca na Arduino board w wersji Uno. Po sprawdzeniu połączeń

zgodnie ze schematem ideowym można podłączyć zasilanie. Moduł wyświetlacza zasilany jest napięciem +5 V wprost z Arduino.

Status i zliczoną liczbę głosów można sprawdzić na wyświetlaczu. Każdy z głosujących może tylko raz nacisnąć przycisk przypisany jednemu z kandydatów. W układzie przewidziano ośmiu kandydatów, aczkolwiek łatwo jest układ rozbudować dla większej ich ilości. ■

Navpreet Singh Tung

Ten artykuł był wcześniej opublikowany na łamach „EFY”, wrzesień 2023 (efymag.com)

Od Redakcji EdW: Autor nie podaje żadnych szczegółów struktury programu dla tego projektu. Główna część programu polega niewątpliwie na programowej realizacji ośmiu liczników zwyczajnie zliczających ilość naciśnięć każdego z przycisków S1 do S8. Należy się też domyślać, iż układ jest zabezpieczony przed nieuczciwym głosowaniem. Powinien nie pozwolić na naciśnięcie więcej

niż jednego przycisku S1 do S8 w jakimś przewidzianym czasie. Również maszyna nie powinna przyjąć kilkukrotnego naciśnięcia tego samego switcha. To również powinno być łatwo osiągnąć poprzez pętle ustalające programowy timer. Również brak jest informacji jak ten EVM reaguje na naciśnięcie switcha S9. Można się domyślać i oczekiwać, że kolejne naciśnięcia S9 skutkują wyświetleniem numeru kolejnego kandydata i ilości zebranych głosów.

W opisie jest także informacja o statusie wyświetlacza, iż zawiera on dwie linie po 16 znaków. Na schemacie widnieje zaś informacja 3×12 character. Można się domyślać, iż program Arduino potrafi obsłużyć kilka typów wyświetlacza znakowego. Niewiele mówi też informacja, iż o statusie głosowania informuje dioda LED i buzzer. Rozsądnym jest, aby dźwięk buzzera towarzyszył naciśnięciu jednego z przycisków S1 do S8. Natomiast zaświecenie lub zgaśnięcie diody LED powinno sygnalizować gotowość maszyny głosującej do przyjęcia kolejnego głosu.

Prezentujemy 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 formule open source. Prenumeratorów EdW zapraszamy do zapoznania się na www.elportal.pl z niezwykle inspirującymi zasobami rubryki **DIY PLUS**.

DIY PLUS

tylko dla prenumeratorów

Jednokanałowy bezprzewodowy włącznik/wyłącznik PS3

Prezentowana płytka to wielofunkcyjny sprzęt, który zawiera moduł ESP32-Wroom, regulator 3,3 V, przełącznik 5 V, diodę LED zasilania, diodę LED przełącznika i złącze do programowania modułu ESP32. Projekt wymaga napięcia 5 VDC i pobiera 60 mA prądu, gdy przełącznik jest włączony. Projekt może być wykorzystywany w aplikacjach IoT, a przełącznik może być sterowany za pomocą Bluetooth lub Wi-Fi. Sprzęt pomaga użytkownikom rozwijać projekty oparte na przełącznikach jednokanałowych. Przełącznik może być sterowany przez Wi-Fi i bezprzewodową komunikację Bluetooth za pomocą odpowiedniego kodu. Przełącznik jest podłączony do pinu GPIO4 modułu ESP32.

Stereofoniczny procesor audio do telewizor

Ten procesor audio oparty na układzie scalonym PT2315E jest przeznaczony do wszechstronnych zastosowań i obejmuje główną regulację głośności z kompensacją głośności niskich częstotliwości, tłumik wyjściowy głośnika i regulację tonu. Jest to dobre rozwiązanie do domowego przetwarzania sygnału audio. Ze względu na wysokie wymagania dotyczące niezawodności w branży audio, PT2315E poprawia zarówno wydajność audio, jak i zdolność wejściowego prądu udarowego, co powoduje, że PT2315E jest najlepszym rozwiązaniem dla ekonomicznych systemów audio.



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. Leszczynowa 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Adres redakcji:
03-197 Warszawa, ul. Leszczynowa 11
e-mail: edw@elportal.pl, www.elportal.pl

Redaktor merytoryczny:
Mariusz Ciszewski, Paweł Sujko

Dział Reklamy:
Katarzyna Gugala
katarzyna.gugala@elportal.pl, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobański
jakub.sobanski@elportal.pl

Copyright AVT-Korporacja Sp. z o.o., Warszawa, ul. Leszczynowa 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