

ELEKTRONIKA

dla wszystkich

nr 3/2023 (326) • marzec • www.elportal.pl

Myjka ultradźwiękowa o dużej mocy

DIY PLUS
tylko dla prenumeratorów

PROJEKTY dla elektroników

- ▶ Łatwe w budowie. Aktywne monitory Hi-Fi z opcjonalnymi subwooferami, część 2
- ▶ Latarnia morska „Nocny Opiekun”
- ▶ Programowalny termoregulator, część 2

DIY dla wszystkich

- ▶ Wyświetlacz na matrycy diod LED
- ▶ Podręczny monitor serca EKG
- ▶ Bezprzewodowy Power-Bank
- ▶ Autonomiczny robot-samochodzik omijający przeszkody

TUTORIALE

- ▶ Silniki krokowe w praktyce, część 4: Sterowniki bipolarnych silników krokowych
- ▶ Audio out: PE Mini-monitorowa zwrotnica dla przetworników Wavecor, część 2
- ▶ Powrót do układów tensometrycznych
- ▶ Praktyczny kurs op-ampów
- ▶ Edukacja w EdW dla szkół i uczelni:
Wykład 4 – Półprzewodniki mocy
- ▶ Programowanie wizualne z XOD. Przedstawiamy wizualne programowanie XOD dla Arduino
- ▶ Pokój Nauczycielski



16,90 zł (w tym 8% VAT)
ISSN 1425-1698 Indeks 33362X
0 3 >
9 771425 169238



EP.com.pl

Największy
portal dla
elektroników
konstruktorów

FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

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

www.piekarz.pl



Poznaj najważniejsze na rynku tytuły z segmentów

**Bezpłatna
przesyłka**



Tylko prenumeratorzy
mają dostęp do inspirujących
projektów w zbiorze **DIY PLUS**
na www.elportal.pl

Zaprenumeruj
„Elektronikę dla Wszystkich”,
a dostaniesz bezpłatny
dostęp do archiwalnych
e-wydań EdW!

Nie dotyczy wydań z ostatnich 24 miesięcy.

na start
do 6* wydań gratis

po 5 latach
nieprzerwanej
prenumeraty
do 12* wydań gratis

* Cena prenumeraty rocznej **na start** wynosi 185,90 zł. Przy zamówieniu prenumeraty dwuletniej za 304,20 zł oszczędność wynosi równowartość sześciu wydań „Elektroniki dla Wszystkich”.

Przedłużasz prenumeratę? Aby otrzymać zniżkę lojalnościową, przedłuż prenumeratę po zalogowaniu się do swojego panelu na www.ulubionykiosk.pl, gdzie znajdziesz atrakcyjną ofertę prenumerat, która uwzględnia przysługujące Ci zniżki za lojalność. Po 5 latach nieprzerwanej prenumeraty otrzymasz **rabat 50%** na prenumeratę dwuletnią.

Oferta dotyczy prenumerat drukowanej.

Wszystkie opcje prenumeraty i e-prenumeraty znajdziesz na stronie www.UlubionyKiosk.pl

Po opłaceniu prenumeraty przysłemy Ci kod dostępu do projektów DIY plus na www.elportal.pl

prenumerata@avt.pl

AVT-Korporacja sp. z o.o., ul. Leszczynowa 11, 03-197 Warszawa,
konto 18 1050 1012 1000 0024 3173 1013

eprasa.pl 250f0d7be5

8



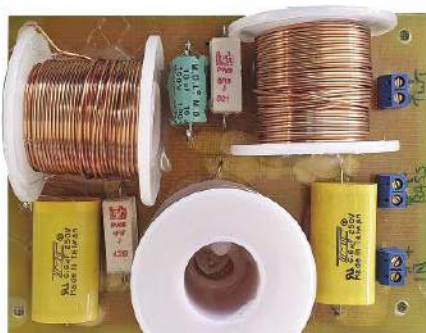
Projekty dla elektroników:

Myjka ultradźwiękowa o dużej mocy, część 1.....	8
Łatwe w budowie. Aktywne monitory Hi-Fi z opcjonalnymi subwooferami, część 2.....	16
Latarnia morska „Nocny Opiekun”.....	26
Programowalny termoregulator, część 2.....	30

Tutoriale:

Silniki krokowe w praktyce, część 4: Sterowniki bipolarnych silników krokowych.....	39
Audio out: PE Mini-monitorowa zwrotnica dla przetworników Wavecor, część 2.....	43
Powrót do układów tensometrycznych.....	49
Praktyczny kurs op-ampów.....	54
Edukacja w EdW dla szkół i uczelni: Wykład 4 – Półprzewodniki mocy.....	58
Programowanie wizualne z XOD.	
Prezentujemy wizualne programowanie XOD dla Arduino.....	66
Pokój Nauczycielski.....	72

26



16



DIY dla wszystkich:

Wyświetlacz na matrycy diod LED.....	76
Podręczny monitor serca EKG.....	81
Bezprzewodowy Power-Bank.....	84
Autonomiczny robot-samochodzik omijający przeszkody.....	88

DIY PLUS

PALPi konsola do gier retro.....	91
Wielokanałowy analizator FFT widma.....	91

Rubryki stałe:

Prenumerata.....	3
Od wydawcy.....	5
Pocztą.....	6

A za miesiąc w marcowym EdW



Mini Wi-Fi LCD Backpack

Wyposażenie modułu LCD Backpack w łączność Wi-Fi stwarza nieograniczone możliwości zastosowań, na przykład do wyświetlania parametrów automatyki domowej. Dostęp do Internetu przez domowe czy biurowe Wi-Fi pozwala zaprzęgnąć płytkę Arduino jako zdalny sterownik w różnych aplikacjach.

Myjka ultradźwiękowa dużej mocy, część 2

Urządzenie do czyszczenia ultradźwiękami niezwykle skutecznie usuwa najbardziej nawet uciążliwe zabrudzenia z różnych przedmiotów – od zastosowań domowych do częściowego przemysłowych. W pierwszej części (EdW 02/2023) opisaliśmy budowę, właściwości i zasadę działania tego urządzenia. Teraz przyszedł czas na opis montażu i uruchomienia.

Łatwe w budowie aktywne kolumny HiFi z opcjonalnymi subwooferami, część 3

Zakończyliśmy budowę kolumn aktywnych, które brzmią doskonale. Może jeszcze poprawić brzmienie dodając subwoofer, a nawet dwa.

Kolumny głośnikowe „Concreto”

To absolutnie unikalna konstrukcja – głośniki zamontowane w bloczkach budowlanych. Intrygujące, prawda?

Plus zwykła porcja intrygujących projektów DIY.

Plus wiele artykułów w Twoich ulubionych cyklach Tutoriali, w tym polecamy nowy kurs holenderskiego autora – „Praktyczny kurs op-ampów”.

**W kioskach
od 1 kwietnia**



Niech moc będzie z Wami

Moc to energia, a energia to życie. Ludzie wiedzą o tym od tysięcy lat, ale dopiero ostatnio dowiedzieli się, że produkcja energii z paliw kopalnych może zabić życie na naszej planecie, a przynajmniej uczynić je nieznośnym. Okazuje się, że najmniej szkodliwa dla środowiska naturalnego jest energia elektryczna pozyskiwana ze źródeł odnawialnych. Świat, już przecież dawno zelektryfikowany, szybko przestawia się na cywilizację totalnie elektryczną. Nie dość, że mamy elektryczne pralki, lodówki, kuchenki, automatyczne linie produkcyjne, to również wszystkie pojazdy, od hulajnogi do autobusów mają być elektryczne.

Wszyscy słyszeliśmy o słynnym sporze Edisona z Teslą o prąd – ma być prąd przemienny (AC), czy stały (DC)? Chociaż w przesyłowych sieciach energetycznych płynie prąd przemienny, to wiele urządzeń wymaga zasilania prądem stałym. Zatem niezbędne jest przetwarzanie AC w DC (prostowniki) i odwrotnie DC w AC (falowniki, na przykład w fotowoltaice). Potrzebne są też rozwiązania umożliwiające sterowanie dużymi prądami. Potrzebne są elementy, które wykonują te zadania. Są to półprzewodnikowe elementy mocy. To gałąź elektroniki przeciwległa do mikroelektroniki. Tranzystor MOS w układzie scalonym pamięci półprzewodnikowej pobiera energię z mocą na poziomie zaledwie nanowatów. Natomiast tranzystor MOS mocy przełącza prądy w obciążeniu o mocy rzędu megawatów.

Tak szeroko rozpięte są skrzydła elektroniki. Elektronika mocy to już pogranicze elektryki, a nawet energetyki, zresztą ta część elektroniki jest często nazywana energoelektroniką. Amatorski pasjonat elektroniki raczej omija energoelektronikę prądów kiloamperowych i napięć kilowoltowych, ale zdarza mu się konstruować układy o mocy setek watów, a to już jest elektronika mocy.

A jeśli nawet Czytelnik EdW nie będzie konstruował układów przeznaczonych dla farm fotowoltaicznych i wiatrowych, to wypada wiedzieć cokolwiek o półprzewodnikowych elementach mocy stosowanych w tych układach. Dlatego wykład w tym wydaniu EdW jest poświęcony temu tematowi.

To temat fascynujący zarówno różnorodnością rodzajów elementów, jak i fantastycznymi perspektywami i rozwojem. O ile mikroelektronika zdaje się dobiegać do kresu dalszych możliwości rozwoju, rozumianego jako ciągły wzrost skali integracji, o tyle perspektywy rozwoju elektroniki mocy rysują się wspaniałe. Nie trzeba uzasadniać, że wszystko co dotyczy poprawy efektywności sterowania produkcją i zużyciem energii elektrycznej znajduje się na topie potrzeb cywilizacyjnych. A mamy w tej dziedzinie wyjątkowe rezerwy rozwoju przez stosowanie nowych materiałów półprzewodnikowych. Jest tu sytuacja całkowicie odmienna od tej w mikroelektronice, gdzie liczy się tylko jeden materiał – krzem. W elektronice mocy krzem jest już mocno wypierany przez tzw. półprzewodniki szerokopasmowe – SiC (węglik krzemu) i GaN (azotek galu), a na horyzoncie pojawia się diament.

Oj będzie się działo!

Wiesław Marciniak

W rubryce „Począ” zamieszczamy fragmenty listów od Czytelników. Szczególnie chętnie publikujemy komentarze do artykułów w bieżących wydaniach EdW oraz propozycje istotne dla planowania przyszłych wydań.

Kolumny głośnikowe

Witam serdecznie!

Jestem zainteresowany zbudowaniem kolumn głośnikowych Senator, opisanych w nr 09/2022 i 10/2022 EdW. Autor artykułu Pan Leo Simpson proponuje wykonanie ścianek kolumn z płyt meblowych dostępnych w Australii. W Polsce trudno o dokładnie taki materiał. I tu moje pytanie i prośba, czy można zbudować kolumny po prostu z płyty MDF powszechnie dostępnej, względnie z drewna bukowego zachowując wymiary. Druga sprawa to czy boki kolumny (front, góra i bok) mają podwójną grubość. Moje drugie pytanie to, czy wynika to tylko ze względów estetycznych, czy też ma to znaczenie akustyczne? Jeśli ma to znaczenie akustyczne, to można od razu użyć grubszego materiału na owe ścianki. Czy można jakoś skontaktować się z autorem aby Go o to zapytać.

Za ewentualne wyjaśnienie z góry dziękuję i Pozdrawiam

Bogusław Kuleba

Oto odpowiedź, pisownia oryginalna

Gen Dobrze!

Nice to hear from you all the way from Poland!

Yes its ok to use MDF because our construction was based on materials from a local “Bunnings” store so please use MDF.

The double sides are for both aesthetic reason AND for superior acoustics so please use the same dimensions if possible.

Regards,

Allan Linton-Smith

Red. To przykład korespondencji z autorami zagranicznymi. Na pytania Czytelników do Autorów, wysłane za pośrednictwem redakcji EdW, otrzymuje się szybko odpowiedź od Autora. Pod tym względem odległa Australia może być bliższa niż sąsiednia dzielnica w Warszawie.

Wykład...

Wykład o tranzystorach przeczytałem z wielkim zainteresowaniem, chociaż 20 lat temu skończyłem technikum elektroniczne i wiem, kto dostał Nobla za wynalazek tranzystora. Jednak jeszcze w szkole, gdy dowiadywałem się o jakimś odkryciu, to zawsze ciekawiło mnie, jak doszło do tego odkrycia, czy wynalazku. Co wiedział odkrywca i jakie rozumowanie doprowadziło go do wyjaśnienia tej czy innej zagadki przyrody. Przecież wybitni uczeni i wynalazcy nie zgadywali innowacyjnych rozwiązań, tylko dochodzili do nich stopniowo w drodze logicznego wyciągania wniosków ze znanych faktów i stawiania hipotez do udowodnienia. Dlatego dużą satysfakcję sprawiło mi czytanie wykładu o tranzystorach, w którym nie tylko śledzimy proces twórczy, ale też dowiadujemy się o emocjach twórców – ludzi jak każdy z nas. Świetna lektura. Proszę o więcej takich wykładów (...) gdy 20 lat temu uczyłem się w technikum, słyszałem, że krzem zostanie zastąpiony w przyszłości przez arsenek galu i azotek galu. Czy ten czas kiedyś nastąpi?

K.G.

Red. W niektórych segmentach elektroniki to się dzieje. GaAs odgrywa ważną rolę w optoelektronice i mikrofalach, a GaN razem z SiC przejmują od Si coraz większą część elektroniki mocy – o tym jest wykład w tym numerze.

Przypominamy, że **Prenumeratory EdW** mają dostęp bezpłatny do e-wydań archiwalnych EdW (dostęp do e-wydań archiwalnych niektórych tytułów AVT jest też bezpłatny dla prenumeratów tych tytułów). Nie dotyczy wydań z ostatnich 24 miesięcy. **Wydawnictwo AVT**

Patronat AVT

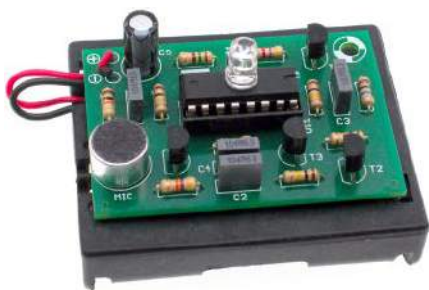
Poniżej prezentujemy listę szkół biorących udział w programie PATRONAT AVT, który jest całkowicie bezpłatny, a szkoły objęte tym patronatem korzystają z różnych benefitów, takich jak bezpłatne prenumeraty, darmowe pakiety próbne kitów AVT, itp. Szkoły, które dopiero teraz dowiadują się o naszej akcji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

- Centrum Edukacji Zawodowej, 82-200 Malbork, De Gaulle'a 75a
- Centrum Edukacji Zawodowej i Biznesu, 66-400 Gorzów Wielkopolski, Pomorska 67
- Gminny Zespół Szkolno-Przedszkolny nr 4 w Więckach, 42-110 Popów, Więcki, Szkolna 1
- Górnośląskie Centrum Edukacyjne im. Marii Skłodowskiej-Curie w Gliwicach, 44-100 Gliwice, Okrzei 20
- Noworudzka Szkoła Techniczna w Nowej Rudzie, 57-401 Nowa Ruda, Stara Droga 4
- Regionalne Centrum Edukacji Zawodowej w Biłgoraju, 23-400 Biłgoraj, Kościuszki 98
- Regionalne Centrum Edukacji Zawodowej w Lubartowie, 21-100 Lubartów, 1 Maja 82
- Szkoła Podstawowa im. Rodziny Bohaterów II Wojny Światowej w Żałakowie, 83-342 Kamienica Królewska, Żałakowo 6
- Techniczne Zakłady Naukowe w Dąbrowie Górniczej, 41-300 Dąbrowa Górnicza, Zawidzkiej 10
- Technikum nr 4 im. Marii Skłodowskiej-Curie, 41-902 Bytom, Katowicka 35
- Zespół Placówek Edukacyjno-Wychowawczych w Gołdapi, 19-500 Gołdap, Wojska Polskiego 18
- Zespół Placówek Oświatowych w Rudniku, 32-440 Sułkowice, Rudnik, Szkolna 55
- Zespół Szkolno-Przedszkolny nr 2 w Wiśle, 43-460 Wisła, Malinka 53
- Zespół Szkolno-Przedszkolny nr 3 w Gliwicach, 44-122 Gliwice, Żwirki i Wigury 85
- Zespół Szkolno-Przedszkolny w Choceniu, 87-850 Chocień, Sikorskiego 12
- Zespół Szkolno-Przedszkolny w Ostroźnicy, 47-280 Pawłowiczki, Ostroźnica, Kościelna 42
- Zespół Szkół Budowlano-Elektrycznych im. Jana III Sobieskiego w Świdnicy, 58-100 Świdnica Śląska, Wątrzybska 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, Tyrankiewiczów 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. Stanisława Staszica, 39-400 Tarnobrzeg, Kopernika 1
- Zespół Szkół nr 1 w Przysietnicy, 36-200 Brzozów, Przysietnica 198
- Zespół Szkół im. ks. dra Jana Zwierza w Ropczycach, 39-100 Ropczyce, Mickiewicza 14
- Zespół Szkół nr 10 im. Prof. Janusza Groszkowskiego w Zabrze, 41-807 Zabrze, Chopina 26
- 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, Zamojskiego 62
- Zespół Szkół Powiatowych im. Stanisława Staszica w Opocznie, 26-300 Opoczno, Kossaka 1a
- Zespół Szkół Publicznych w Szewnie, 27-400 Ostrowiec Świętokrzyski, Szewna, Langiewicza 3
- Zespół Szkół Spożywczych i Hotelarskich w Radomiu, 26-600 Radom, Św. Brata Alberta 1
- Zespół Szkół Techniczno-Informatycznych w Elblągu, 82-300 Elbląg, 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



Najbardziej popularne kity AVT

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



AVT788 Lampka LED reagująca na klaskanie:
klaskacz, włącznik dźwiękowy
<https://sklep.avt.pl/avt788.html>



AVT723 Uniwersalna gra zgrzewnicowa
<https://sklep.avt.pl/avt723.html>



AVT594 Zdalnie sterowany potencjometr
do aplikacji audio
<https://sklep.avt.pl/avt594.html>



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



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



AVT3225 Uniwersalny sterownik silnika krokowego
<https://sklep.avt.pl/avt3225.html>



AVT3200 Uniwersalny timer 0 do 99 min.
<https://sklep.avt.pl/avt3200.html>



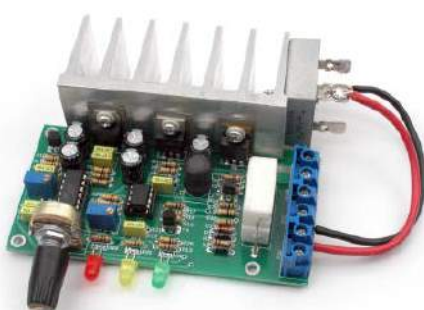
AVT990 Automatyczny włącznik świateł
<https://sklep.avt.pl/avt990.html>



AVT732 Whisper - łowca szeptów. Superczuły
podłuch przewodowy
<https://sklep.avt.pl/avt732.html>



AVT5553 Sterownik zgrzewarki oporowej
<https://sklep.avt.pl/avt5553.html>



AVT3120 Automatyczna ładowarka
akumulatorów ołowianych
<https://sklep.avt.pl/avt3120.html>



AVT3166 Regulator do prostownika
<https://sklep.avt.pl/avt3166.html>



Pełna oferta na: sklep.avt.pl

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



Myjka ultradźwiękowa o dużej mocy, część 1

Ta spora i mocna myjka ultradźwiękowa jest idealna do czyszczenia dużych przedmiotów, takich jak części mechaniczne i delikatne wyroby. Posiada wiele funkcji i jest również dość łatwa w budowie.

Zapewne widziałeś już małe, tanie myjki ultradźwiękowe dostępne w sieci. Są one świetne do czyszczenia przedmiotów takich jak biżuteria, okulary itp.

Ale co zrobić, jeśli chcesz oczyścić coś nieco większego i bardziej zabrudzonego, o nietypowych, skomplikowanych kształtach

albo pragniesz dysponować urządzeniem o szerszej gamie możliwości czyszczących?

Czyszczenie wtryskiwaczy paliwa, starego gaźnika lub innych skomplikowanych części jest zadaniem brudnym, męczącym i czasochłonnym, wymagającym moczenia w cuchnących rozpuszczalnikach, takich jak

benzyna, nafta lub odtłuszcacz, i szorowania różnymi szczotkami w celu oczyszczenia części. Jest to trudne i żmudne zadanie i często nie dociera się do małych zakamarków, które jak na złość odgrywają ważną rolę.

Nasza Myjka Ultradźwiękowa znacznie ułatwia to zadanie. Wystarczy umieścić

Cechy

- Stosowany przetwornik nominalnie 40 kHz, 50 W lub 60 W
- Regulowany poziom mocy
- Wskaźnik poziomu mocy
- Przyciski Start i Stop z sygnalizacją stanu pracy
- Minutnik automatycznego wyłączenia od 20 sekund do 90 minut
- Miękki start
- Wyłączenie i wskazanie błędu przeciążenia lub rozruchu
- Diagnostyka poziomu mocy
- Automatyczna lub ręczna kalibracja przetwornika
- Minimalizacja efektu fali stojącej
- Obsługuje częstotliwość rezonansową przetwornika od 34,88 kHz do 45,45 kHz



„Robocze” elementy naszej Myjki Ultradźwiękowej przed podłączeniem przetwornika do wanny czyszczącej. Obsługa jest dość prosta: włączamy, ustawiamy minutnik i naciskamy przycisk „start”!

elementy w kąpeli z rozpuszczalnika, nacisnąć przycisk i wrócić później, aby wyjąć części lśniąco i czyste. Można w niej czyścić nawet przestrzenie wewnętrzne! Wykorzystuje przetwornik piezoelektryczny o dużej mocy i generator ultradźwiękowy, aby usunąć brud i zanieczyszczenia za pomocą energii ultradźwięków.

W przypadku delikatniejszych części można zmniejszyć moc, aby zapobiec uszkodzeniu czyszczonych elementów.

Napężenie śruby zapewnia, że dyski piezoelektryczne pozostają zawsze ściśnięte, nawet podczas pracy urządzenia, co zapobiega ich rozerwaniu.

Gdy do płytek piezoelektrycznych zostanie przyłożone napięcie, elementy piezoelektryczne generują siły, które ścisną bądź rozsuwają dwie metalowe okładki w takt częstotliwości generatora ultradźwiękowego. Nasza myjka ultradźwiękowa zasila przetwornik piezoelektryczny z częstotliwością bliską

jego nominalnej częstotliwości rezonansowej wynoszącej 40 kHz.

Rysunek 4 pokazuje zależność mocy od częstotliwości dla konkretnego przetwornika ultradźwiękowego, którego używamy. Jego częstotliwość rezonansowa wynosi 40 kHz z tolerancją ± 1 kHz. Stwierdziliśmy, że pod obciążeniem nasz przetwornik ma częstotliwość rezonansową 38,8 kHz.

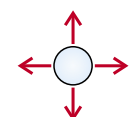
Częstotliwość pracy przetwornika musi być zachowana w ścisłych granicach

Jak to działa?

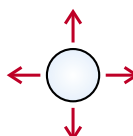
Metalowy pojemnik jest wypełniony rozpuszczalnikiem, wodą dejonizowaną lub zwykłą gorącą wodą oraz detergentem lub środkiem zwilżającym. Przetwornik ultradźwiękowy miesza zawartość wanny; przy wyższych poziomach mocy czoło fali ultradźwiękowej powoduje kawitację, generując mikropęcherzyki próżniowe, które następnie ulegają kolapsowi, tworząc bardzo silne uderzenia czoła cieczy. Pokazuje to rysunek 1.

Gdy czoło fali mija, przywracane jest normalne ciśnienie, a próżniowy pęcherzyk zapada się, wytwarzając falę uderzeniową. Ta fala uderzeniowa pomaga poluzować cząstki brudu na czyszczonym przedmiocie (rysunek 2). Wielkość pęcherzyków zależy od częstotliwości ultradźwięków; są one mniejsze przy wyższych częstotliwościach.

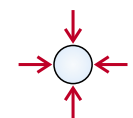
Stosujemy powszechnie dostępny przetwornik ultradźwiękowy Langevina, mocowany śrubami, przedstawiony na rysunku 3. Składa się on z płytek piezoelektrycznych umieszczonych pomiędzy metalowymi elektrodami. Śruba centralna nie tylko utrzymuje zespół razem, ale ma decydujące znaczenie dla zapewnienia, że elementy piezoelektryczne nie ulegną uszkodzeniu podczas emisji ultradźwięków. Śruba jest dokręcana do określonego wcześniej napężenia i blokowana (klejona) w tym położeniu, aby zapobiec jej poluzowaniu.



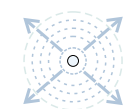
Tworzą się pęcherzyki próżniowe



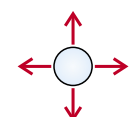
Pęcherzyki rosną przy zmniejszonym ciśnieniu



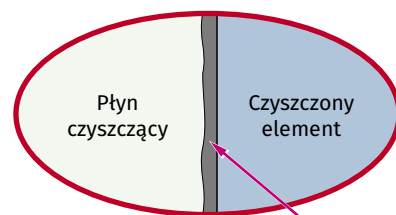
Pęcherzyki kurczą się przy powrocie ciśnienia



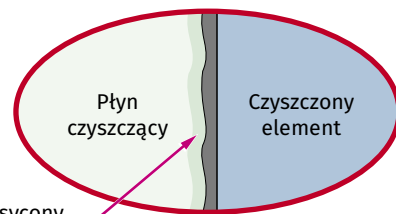
Pęcherzyki zapadają się gwałtownie generując fale uderzeniowe



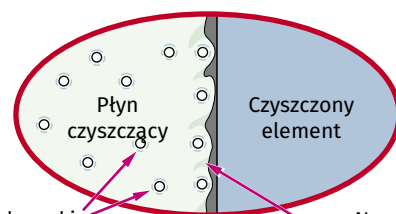
Tworzą się nowe pęcherzyki



Zanieczyszczenia powierzchniowe



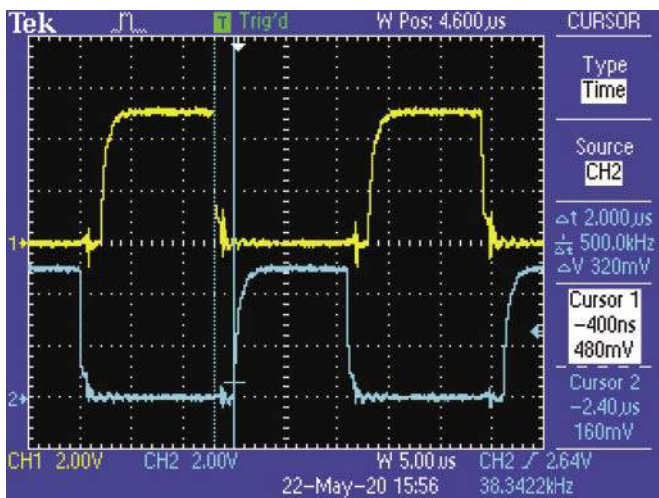
Nasycony płyn czyszczący



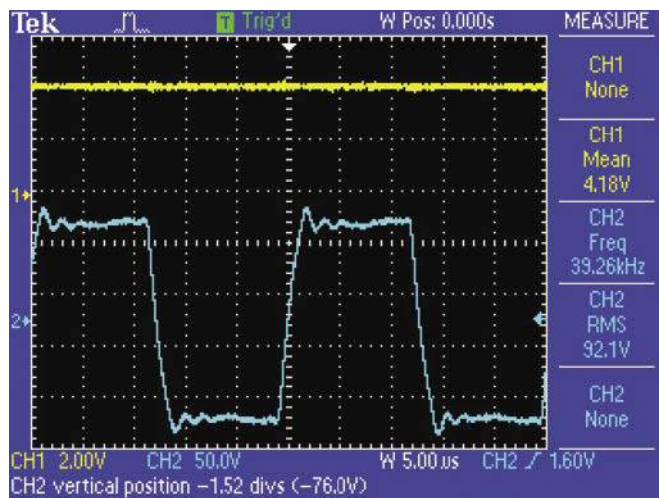
Pęcherzyki kawitacyjne

Nasycony płyn czyszczący

Rysunek 1 i 2. Fale dźwiękowe wytwarzane przez myjkę ultradźwiękową gwałtownie tworzą i zginiatają pęcherzyki w cieczy. Kiedy pęcherzyki zapadają się, generują lokalne fale uderzeniowe. Ta "kawitacja" pobudza warstwę rozpuszczalnika, która ma kontakt z brudem, tłuszczem i zanieczyszczeniami, pomagając je rozbici i szybciej rozpuścić. Można to robić ręcznie – nazywa się to szorowaniem – ale jest to żmudna praca i trudno jest dostać się do zakamarków i wewnętrznych przestrzeni czyszczonych części!



Oscylogram 1. Wysterowanie bramek MOSFET-ów Q1 (górny przebieg, żółty) i Q2 (dolny przebieg, cyjan) mierzone na wyjściach 5 i 6 układu IC1. Pionowe kursory pokazują czas martwy 2 µs, gdy oba MOSFET-y nie są wysterowane. Pokazana sytuacja, gdy Q1 wyłącza się a Q2 włącza się; czas martwy jest taki sam między wyłączeniem Q2 i włączeniem Q1.



Oscylogram 2. Dolny przebieg (cyjan) pokazuje napięcie wyjściowe transformatora przy wysterowaniu przetwornika ultradźwiękowego przy częstotliwości 39,26 kHz. Górny przebieg pokazuje pomiar prądu za pośrednictwem napięcia na wejściu AN11 układu IC1 (TP1). 4,18 V oznacza prąd 2,98 A zasilający uzwojenia pierwotne transformatora przy napięciu zasilania 12 V. Odpowiada to około 35,8 W mocy dostarczanej do przetwornika.

tolerancji, aby utrzymać stały poziom mocy. Mała zmiana częstotliwości w stosunku do punktu rezonansowego spowoduje dość wyraźne zmniejszenie mocy. Dodatkowo, impedancja przetwornika zmienia się w zależności od obciążenia. Podczas pracy bez obciążenia w powietrzu impedancja jest znacznie niższa w porównaniu z sytuacją, gdy przetwornik zasila wannę wypełnioną płynem czyszczącym.

Szczegóły obwodu

Schemat ideowy myjki ultradźwiękowej pokazany jest na rysunku 5. Główną częścią jest mikrokontroler PIC16F1459 (IC1). IC1 przełącza dwa MOSFET-y (Q1 i Q2), które zasilają na przemian uzwojenia pierwotne transformatora T1. T1 wytwarza napięcie krokowe 100 V AC (RMS) do zasilania przetwornika ultradźwiękowego.

IC1 włącza również diodę LED (LED1) sygnalizacji zasilania i diody LED poziomu mocy (LED2-LED6); ponadto nadzoruje potencjometr minutnika (VR1) oraz przełączniki S2 i S3, służące do ręcznego uruchamiania i zatrzymywania pracy myjki.

Na swoim wejściu analogowym AN11, na końcówce 12, układ IC1 monitoruje również prąd płynący przez MOSFET-y Q1 i Q2. Steruje także miękkim startem ładowania głównego kondensatora bocznikującego za pomocą tranzystora Q5 i MOSFET-a Q6.

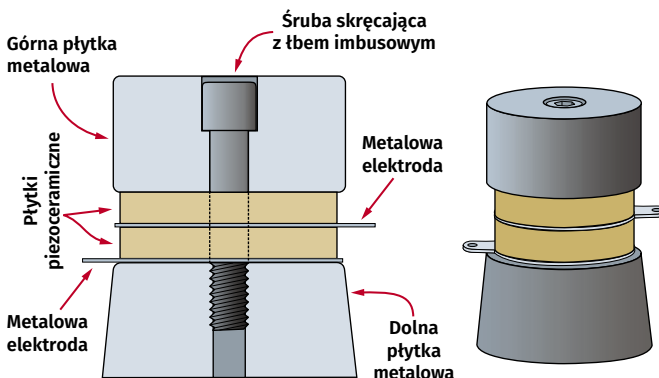
Praca transformatora

Komplementarny generator fal w IC1 jest używany do wysterowania MOSFET-ów Q1 i Q2 w trybie naprzemiennym (push-pull). Transformator ma odczep środkowy, aby umożliwić ten rodzaj zasilania. Generator PWM w IC1 posiada regulowany czas martwy,

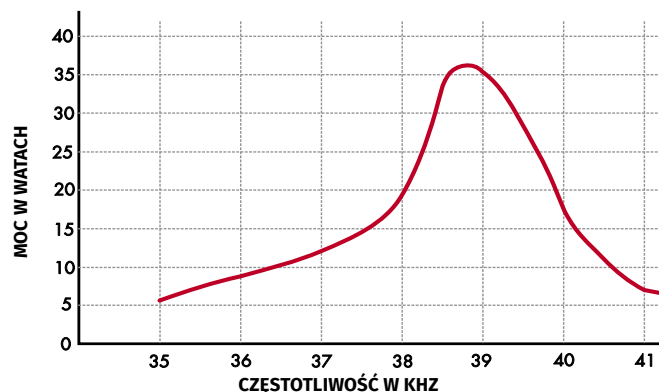
tak, że wyłączenie jednego MOSFET-a następuje zanim drugi MOSFET zostanie włączony (Oscylogram 1). Zapobiega to „zwarciom” podczas jednoczesnego włączenia obu MOSFET-ów, które w przeciwnym razie spowodowałyby przegrzanie tranzystorów.

Wyjścia cyfrowe RC5 i RC4 układu IC1 dostarczają komplementarnych sygnałów wysterowania bramek dla MOSFET-ów Q1 i Q2. Ponieważ napięcia na tych wyjściach wahają się tylko od 0 V do 5 V, używamy MOSFET-ów o bramkach sterowanych poziomami logicznymi TTL. Standardowe MOSFET-y wymagają sygnałów bramki o wartości co najmniej 10 V dla pełnego przewodzenia, ale MOSFET-y poziomu logicznego zazwyczaj przewodzą w pełni przy 4,5 V, a czasami nawet przy niższych napięciach.

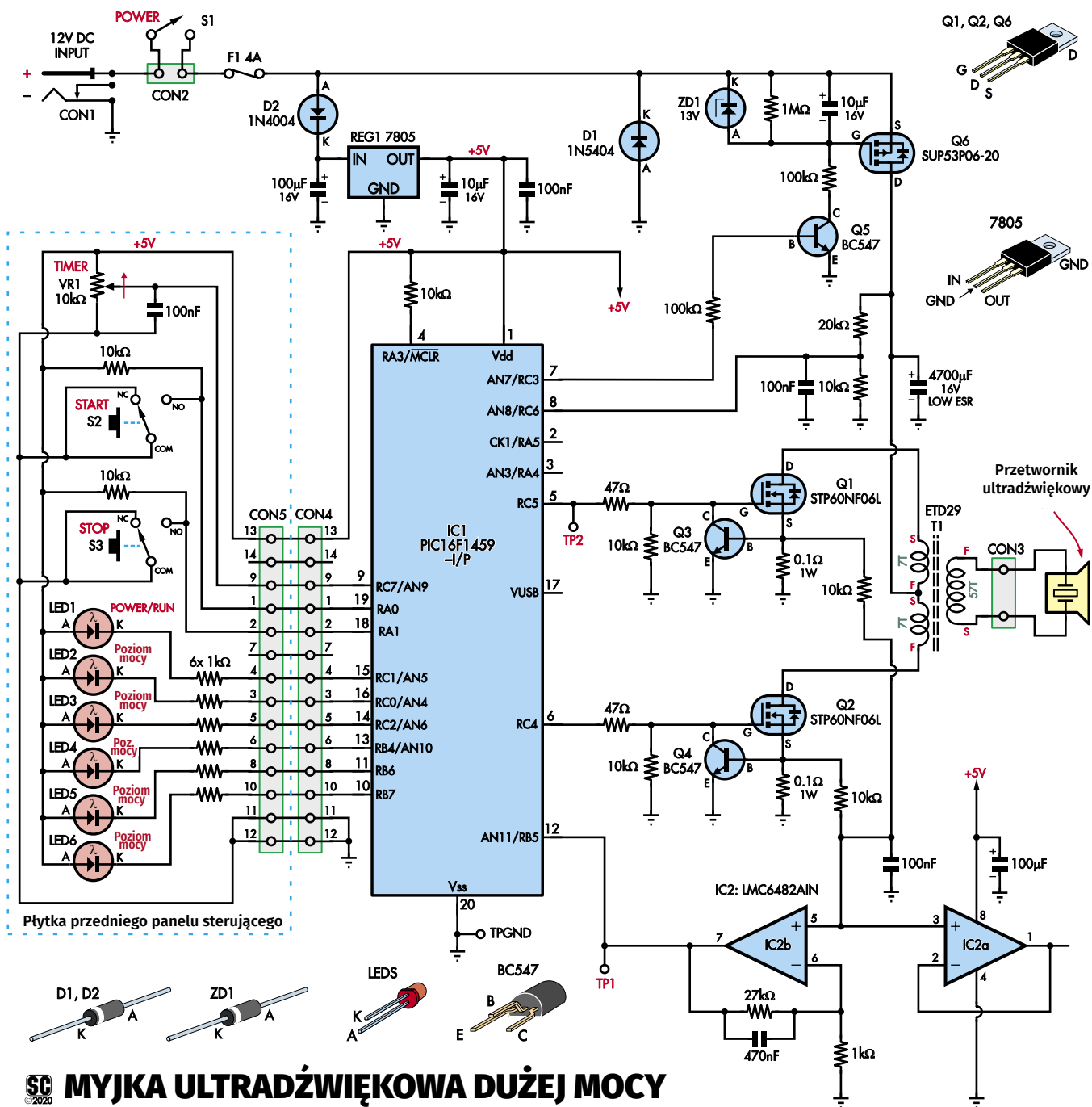
W przypadku stosowanych przez nas MOSFET-ów STP60NF06L



Rysunek 3. Przedstawia budowę przetwornika ultradźwiękowego, którego używamy. Dwie płytki piezoelektryczne (ceramiczne) są umieszczone pomiędzy dwiema półkami korpusu, z elektrodami umożliwiającymi podanie napięcia na elementy piezoelektryczne. Kompresja piezoceramiki spowodowana naprężeniem pochodzącym od śruby mocującej całość jest krytyczna dla zapobieżenia przedwczesnemu uszkodzeniu przez wibracje ultradźwiękowe.



Rysunek 4. Krzywa częstotliwość vs moc dla przetwornika w naszym prototypie. Większość przetworników z nominalnym rezonansem 40 kHz powinna mieć podobną charakterystykę, ale dokładna częstotliwość rezonansu będzie się różnić, podobnie jak stromość zboczy. Dlatego też, nasza myjka posiada automatyczną procedurę kalibracji, aby znaleźć ten rezonans; ustawienie 100% mocy uruchamia przetwornik na częstotliwości bliskiej rezonansowi, podczas gdy niższe ustawienia mocy są na wyższych częstotliwościach.



SC MYJKA ULTRADŹWIĘKOWA DUŻEJ MOCY

Rysunek 5. Kompletny schemat ideowy myjki ultradźwiękowej. IC1 wytwarza komplementarne sygnały sterujące bramkami MOSFET-ów Q1 i Q2, które z kolei zasilają uzwojenie pierwotne transformatora T1 w trybie push-pull. W rezultacie na CON3 pojawia się około 100 V AC. Prąd jest monitorowany przez dwa rezystory bocznikujące 0,1 Ω przy źródłach Q1 i Q2, przez wzmacniacz IC2b podłączony do wejścia analogowego AN11 układu IC1; moc jest obliczana na podstawie tego pomiaru i pomiaru napięcia na wejściu analogowym AN8.

rezystancja włączenia (pomiędzy drenem a źródłem) wynosi 14 mΩ przy 30 A i napięciu bramki 5 V. Są one przystosowane do pracy ciąglej z prądem 60 A i posiadają zabezpieczenie przed przepięciami, które blokuje napięcie dren-źródło na poziomie 60 V.

Q1 i Q2 przewodzą naprzemiennie, i z kolei zasilają oddzielne połowki uzwojenia pierwotnego transformatora T1, ze środkowym odczepem podłączonym do zasilania +12 V. Kiedy MOSFET Q1 jest włączony, prąd płynie w jego sekcji uzwojenia pierwotnego transformatora. Q1

pozostaje włączony przez mniej niż 25 μs (zakładając częstotliwość pracy 40 kHz), a następnie jest wyłączany.

Każdy z MOSFET-ów jest wyłączany na dwie mikrosekundy przed włączeniem przeciwnego. Przykładowo włączony Q2 zasilą następnie prądem swoją sekcję uzwojenia pierwotnego transformatora T1 i pozostaje włączony przez taki sam czas jak uprzednio Q1. Oba MOSFET-y zostają ponownie wyłączone przez dwie mikrosekundy, zanim Q1 zostanie powtórnie włączony. Przerwa,

kiedy oba MOSFET-y są wyłączone, jest „czasem martwym” i powoduje, że wyłączenie MOSFET-ów zajmuje trochę czasu.

Bez czasu martwego, oba MOSFET-y przez krótki czas byłyby włączone jednocześnie. Spowodowałoby to ogromne skoki prądu zwarciowego, nie tylko powodując przegrzanie MOSFET-ów, ale również generując duże skoki prądu pobieranego z kondensatora filtrującego i zasilacza DC.

Naprzemienne włączanie MOSFET-ów generuje falę prostokątną AC w uzwojeniu

wtórny transformatora T1. Przy przekładni transformatora 8,14:1 (57 zwojów wtórnego i 7 zwojów pierwotnego), i napięciu 12 V AC na uzwojeniu pierwotnym, uzwojenie wtórne dostarcza około 98 V AC do przetwornika piezoelektrycznego.

Fale stojące

Praca myjki ultradźwiękowej ze stałą częstotliwością w pobliżu rezonansu jest efektywna, ponieważ impedancja przetwornika jest w tych warunkach prawie czysto rezystancyjna. Nie jest to jednak idealne rozwiązanie dla minimalizacji fal stojących w kąpielni myjącej. Gdy częstotliwość pozostaje stała, fale stojące mogą przybierać na sile.

Fale te są spowodowane odbiciami w fazie od czyszczonych części i ścian zbiornika. Może to spowodować uszkodzenie delikatnych części.

Nasza myjka ultradźwiękowa ma możliwość zmniejszenia mocy do użytku z delikatnymi częściami, ale nawet większe części mogą mieć w sobie delikatne sekcje, zwłaszcza w cienkościennych wnękach.

Aby uniknąć powstawania fal stojących, częstotliwość może być zmieniana w czasie, aby zapobiec jednakowej, niezmienną fazie fali, która skutkowałaby interferencjami (wzmocnieniem i wygaszaniem fali) i zakłóceniami w różnych miejscach wanny. Jak pokazuje wykres zależności mocy od częstotliwości, zmiana częstotliwości nawet o niewielką wartość drastycznie zmienia moc. To nie jest optymalne rozwiązanie, jeśli częstotliwość jest zmieniana w sposób ciągły, ponieważ zmniejsza to moc czyszczenia.

Zamiast tego, pracujemy z przetwornikiem na stałej częstotliwości przez 10 sekund, następnie przez krótki czas pracujemy na różnych częstotliwościach, po czym wracamy do częstotliwości maksymalnej mocy na kolejne 10 sekund. Wystarczy to, aby gęstość mocy w poszczególnych miejscach myjki zmieniła się.

W międzyczasie częstotliwość zmienia się w małych krokach 37,5 Hz w zakresie 2,4 kHz (64 kroki zmiany częstotliwości) przez około 400 ms. Oznacza to, że moc jest zmniejszana tylko przez około 4% czasu pracy. Zmiana częstotliwości zmienia fazę drgań ultradźwiękowych w wannie, dając czas na wygaśnięcie fal stojących pojawiających się w okresie stałej częstotliwości, co zapobiega ich narastaniu do szkodliwego poziomu.

Zabezpieczenie nadprądowe

Zabezpieczenie nadprądowe dla MOSFET-ów działa na dwa sposoby. Oba opierają się na wykrywaniu natężenia prądu poprzez pomiar napięcia na rezystorach 0,1 Ω pomiędzy źródłami Q1 i Q2 a masą.



Przetwornik 40 kHz jest dostępny zarówno w Australii, jak i w Internecie. Zauważ jednak, że jeśli kupisz w sieci, to musisz się upewnić, że dostaniesz typ 40 kHz – dostępne są inne częstotliwości i wyglądają one dość podobnie. (Patrz panel na stronie 15).

Pierwsza metoda wykorzystuje tranzystory NPN Q3 i Q4. Mają one złącza baza-emiter połączone z rezystorami 0,1 Ω .

Detekcja nadmiaru prądu występuje, gdy napięcie na rezystorze 0,1 Ω przekroczy około 0,5 V, czyli gdy przez Q1 lub Q2 przepływa prąd większy niż 5 A. Takie napięcie baza-emiter wystarcza, aby tranzystor Q3 lub Q4 zaczął wtedy przewodzić.

Płynący prąd kolektor-emiter zmniejsza napięcie na bramce sąsiedniego MOSFET-a. Ma to wpływ na zwiększenie rezystancji włączenia MOSFET-a, co następnie zmniejsza płynący prąd. Zabezpieczenie to jest zabezpieczeniem szybko działającym, aktywnym w każdym cyklu przełączania MOSFET-ów.

W tym samym czasie napięcia na dwóch rezystorach 0,1 Ω są uśredniane przez parę rezystorów 10 k Ω i filtrowane przez kondensator 100 nF. To uśrednione napięcie jest następnie podawane na nieodwracające wejście na końcówce 5 wzmacniacza operacyjnego IC2, który wzmacnia sygnał 28 razy (27 k Ω ÷ 1 k Ω +1). Uśrednianie efektywnie zmniejsza o połowę wartość mierzonego napięcia, ponieważ tylko jeden z MOSFET-ów: Q1 lub Q2, jest w danym momencie włączony.

Ogólne wzmocnienie jest zatem 14-krotne. Napięcie wyjściowe z końcówki 7 układu IC2b jest mierzone na wejściu analogowym AN11 układu IC1 (końcówka 12) – patrz Oscylogram 2.

Napięcie to jest przekształcane na postać cyfrową i przetwarzane przez IC1. Jeżeli napięcie to utrzymuje się na poziomie 4,9 V lub wyższym przez okres 160 ms, zasilanie przetwornika ultradźwiękowego jest wyłączane.

Napięcie to reprezentuje średnią wartość 350 mV mierzoną na każdym rezystorze 0,1 Ω , lub średni przepływ prądu 3,5 A. Oblicza się to jako $(4,9 \text{ V} \div 14) \div 0,1 \Omega$.

Błąd przeciążenia prądowego jest sygnalizowany przez miganie diod LED2, LED4 i LED6 na wyświetlaczu poziomu mocy na przednim panelu. Gdy to nastąpi, konieczne jest wyłączenie i ponowne uruchomienie zasilania w celu wznowienia czyszczenia. Jeśli problem będzie się utrzymywał, konieczne będzie znalezienie przyczyny.

Sterowanie mocą

Prąd mierzony na wejściu AN11 jest również wykorzystywany do sterowania mocą podawaną do przetwornika ultradźwiękowego. Maksymalna moc znamionowa przetwornika wynosi 50 W, ale nie jest to moc ciągła. Zalecana moc ciągła to 43 W. My ograniczamy moc do bardziej bezpiecznego poziomu 36 W. Przy zasilania 12 V prąd wymagany dla tej mocy wynosi 3 A.

Podczas pracy prąd jest monitorowany na wejściu AN11, a napięcie na przetworniku jest również próbkowane, poprzez dzielnik rezystancyjny, na wejściu analogowym AN8 (końcówka 8). Dzięki temu mikroprocesor może obliczyć moc podawaną do transformatora w miarę regulacji częstotliwości, tak aby utrzymać tę moc na wymaganym poziomie.

Zegar taktowania instrukcji IC1 pochodzi z jego wewnętrznego oscylatora, a zatem częstotliwości wyjściowe PWM również pochodzą z tego oscylatora. Może być on regulowany w małych krokach za pomocą rejestru OSCTUNE. W ten sposób można zmieniać częstotliwość wewnętrznego oscylatora w zakresie 12% w 128 krokach. W przypadku sterowania przetwornikiem ultradźwiękowym o częstotliwości 40 kHz, pozwala to na uzyskanie zakresu regulacji 4,8 kHz w krokach co 37,5 Hz.

Rozdzielczość kroku 37,5 Hz jest wystarczająco mała, aby wysterować przetwornik ultradźwiękowy na pożądanym poziomie mocy. Rejestr OSCTUNE nie posiada jednak wystarczającego zakresu częstotliwości, aby zapewnić nam możliwość wysterowania przetwornika ultradźwiękowego będącego w rezonansie poza zakresem od 37,6 kHz do 42,4 kHz.

Aby poszerzyć zakres pracy, urządzenie kalibruje się automatycznie (można to również zainicjować ręcznie). W ten sposób zostaje znaleziona przybliżona częstotliwość rezonansowa przetwornika przy użyciu zgrubnej regulacji. Dokładne dostrojenie odbywa się następnie za pośrednictwem systemu OSCTUNE; umożliwia to zastosowanie wielu różnych przetworników.

Ta zgrubna kalibracja jest wykonywana za pomocą rejestru PR2, który ustawia okres, a tym

samą częstotliwość generacji PWM. Dla naszego układu zapewnia to kroki o częstotliwości około 540 Hz. Ograniczamy zakres regulacji zgrubnej od 34,88 kHz do 45,45 kHz. Zakres ten odpowiada wszystkim przetwornikom, które mają nominalny rezonans 40 kHz.

Tak więc rezonans przetwornika jest ustalany z krokiem 540 Hz poprzez regulację PR2, a znaleziona wartość najbliższa rezonansowi przetwornika jest przechowywana w pamięci nieulotnej flash. Jak łatwo zauważyć, błąd częstotliwości rezonansowej ustalonej przy pomocy PR2 jest mniejszy od połowy kroku, czyli <270 Hz. Dokładne dostrojenie następuje przy pomocy rejestru OSCTUNE w krokach co 37,5 Hz, co wymaga maksymalnie 7 kroków (maksymalnie 262,5 Hz). Błąd określenia częstotliwości rezonansowej przetwornika nie przekracza zatem 20 Hz i ma niewielki wpływ na jego pracę. Rejestr OSCTUNE może następnie zmieniać częstotliwość o co najmniej 2,1 kHz powyżej i 2,1 kHz poniżej wartości ustawionej początkowo przez rejestr PR2 (2,1 kHz $\approx$ 2,4 kHz – 270 Hz), ponieważ z nominalnego zakresu przestrajania przy pomocy rejestru OSCTUNE wynoszącego $\pm 2,4$ kHz, maksymalnie ok. 300 Hz używane jest na dokładne dostrojenie do częstotliwości rezonansowej.

Różne poziomy mocy są dostępne poprzez regulację częstotliwości przetwornika. Największa moc występuje przy częstotliwości najbliższej rezonansowi, natomiast niższe poziomy mocy wykorzystują częstotliwość powyżej rezonansu, przy której przetwornik wytwarza mniejszą moc.

Dostępnych jest dziewięć poziomów mocy, w zakresie od 100% (36 W) do 10% (około 3,6 W). W zależności od charakterystyki przetwornika, najniższy poziom mocy może być niedostępny.

Wskaźniki LED

Diody LED 2–6 wskazują, który z dziewięciu poziomów mocy jest wybrany, przy czym dioda LED2 świeci się, aby wskazać najniższy poziom mocy. Następny stopień w górę to świecenie diod LED2 i LED3, następnie LED3 i tak dalej, aż do momentu, gdy świeci się tylko dioda LED6, wskazująca najwyższy poziom mocy.

Poziom mocy reguluje się poprzez przytrzymanie przełącznika Start. Przełącznik przechodzi przez dziewięć możliwych poziomów aż do maksimum, po czym ponownie opada. Przełącznik może być zwolniony przy żądanym ustawieniu poziomu. Podczas regulacji poziomu mocy przetwornik nie jest zasilany.

Dioda LED On/Run (LED1) pokazuje, kiedy do układu jest podłączone zasilanie. Ta dioda LED sygnalizuje również działanie

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

- 1 dwustronna płytka drukowana o kodzie 04105201, 103,5×79 mm.
- 1 dwustronna płytka drukowana o kodzie 04105202, 65×47 mm
- 1 etykieta/naklejka na panel, 115×90 mm (patrz tekst)
- 1 obudowa z odlewu aluminium, 115×90×55 mm (Jaycar HB5042)
- 1 ultradźwiękowy przetwornik tubowy 50/60 W 40 kHz (impedancja rezonansowa 10...20 Ω) [patrz tekst]
- 1 zasilacz impulsowy 12 V DC 60 W lub podobny [Jaycar GH1379, Altronics MB8939B] lub
- 1 akumulator 12 V (10 Ah lub większy) z podwójnym przewodem 5 A+
- 1 karkas (cewka) transformatora EPCOS ETD29 13-końcówkowy, B66359W1013T001 (T1) [RS Components 125-3669, element14 1422746].
- 2 rdzenie ferrytowe EPCOS ETD29 N97, B66358G0000X197 (T1) [RS komponenty125-3664, element14 1422745]
- 2 klipsy/zaciski do rdzeni ferrytowych EPCOS ETD29, B66359S2000X000 lub równoważne (T1) [elementy RS 125-3668, element14 178507]
- 1 przełącznik mini przechylny SPST 6A (S1) [Altronics S3210, Jaycar SK0984]
- 2 przełączniki przyciskowe chwilowe SPDT (S2, S3) [Altronics S1393]
- 2 zaślepki do przełączników S2 i S3 [Altronics S1403]
- 1 gniazdo zasilania 5 A do montażu na płytce PCB, 5,5/2,5 mm ID (CON1) [Jaycar PS0520, Altronics P0621A]
- 1 wtyk zasilania 5 A, 5,5/2,5 mm ID [Jaycar PP0511, Altronics P0165] (opcja)
- 1 pionowe 2-stykowe gniazdo z zaciskami śrubowymi (CON2) [Jaycar HM3112+HM3122]
- 1 złącze śrubowe 2-stykowe do montażu na płytce drukowanej raster 5,08 (CON3) [Jaycar HM3130, Altronics P2040A]
- 1 wtyk IDC Z-WS14 prosty do druku, 14-kołkowy (CON4) [Altronics P5014]
- 1 gniazdo Z-FC14 zaciskane na taśmie AWG28, 14 żył (do CON4) [Altronics P5314]
- 1 wtyczka przejściowa 14 żył IDC (CON5) do lutowania na PCB [Altronics P5162A]
- 2 oprawki bezpieczników 3AG do montażu na płytce drukowanej (F1)
- 1 bezpiecznik 4A 3AG (F1)
- 1 potencjometr liniowy 10 k Ω 16 mm (VR1)
- 1 pokrętko pasujące do potencjometru
- 1 20-stykowa podstawa DIL IC wąska (dla IC1)
- 1 8-stykowa podstawa DIL IC (dla IC2)
- 3 podkładki silikonowe i tuleje izolacyjne do obudów TO-220
- 4 przyklejane gumowe nóżki

Części obudowy przetwornika

- 1 złączki PVC DWV (spust, odpływ i odpowietrzenie) o długości 50 mm; końcówka (Holman DWVF0192) i adapter (Holman DWVF0022) lub
- 1 rura PCV o długości 40 mm i średnicy 50 mm
- 1 przepust kablowy dla kabla 3...6,5 mm
- Neutralny utwardzalny uszczelniający silikonowy (np. dach i rynna)
- Żywica epoksydowa (np. JB Weld) do klejenia metalu

Części do prób

- 1 odcinek długości 100 mm drutu miedzianego ocynowanego 0,7 mm
- 4 gwintowane podkładki dystansowe M3 o długości 9 mm
- 4 śruby M3×6
- dodatkowy odcinek emaliowanego drutu miedzianego o średnicy 0,63 mm
- Kable, przewody i osprzęt
- 1 wkręt M3×6 (do REG1)
- 3 śruby M3×9 (dla Q1, Q2 i Q6)
- 4 nakrętki sześciokątne M3
- 1 przepust kablowy dla kabla o średnicy 3...6,5 mm
- 1 odcinek długości 800 mm emaliowanego drutu miedzianego DNE o średnicy 1 mm (T1 uzwojenie pierwotne)
- 1 odcinek długości 3,6 m emaliowanego drutu miedzianego DNE o średnicy 0,63 mm (T1 uzwojenie wtórne)
- 1 odcinek długości 1 m kabla o podwójnej powłoce 0,75 mm lub przewodu w podwójnej powłoce lub przewodu osemkowego (do podłączenia przetwornika)
- 1 odcinek długości 160 mm przewodu przytępczeniowego 5 A (1 mm²)
- 1 odcinek długości 200 mm 14-żyłowego kabla taśmowego IDC
- 8 kołków PC
- 1 odcinek długości 30 mm 5 mm rurki termokurczliwej (dla połączeń S1)
- 1 rolka elektrycznej taśmy izolacyjnej

Półprzewodniki:

- 1 mikrokontroler PIC16F1459-I/P zaprogramowany za pomocą wsadu 0410520A.hex (IC1)
- 1 CMOS dual op amp LMC6482AIN (IC2)
- 1 7805 regulator liniowy 5 V 1 A (REG1)
- 2 MOSFET-y N-kanatowe STP60NF06L wyzwalane poziomem logicznym TTL (Q1, Q2)
- 3 tranzystory NPN BC547 (Q3-Q5)
- 1 MOSFET P-kanatowy SUP53P06-20 (Q6)
- 1 dioda Zenera 13 V 1 W (ZD1)
- 1 dioda 3 A 1N5404 (D1)
- 1 dioda 1 A 1N4004 (D2)
- 6 diod LED 3 mm (czerwonych lub zielonych) (LED1-LED6)

Kondensatory:

- 1 kondensator elektrolityczny 4700 μ F 16 V low-ESR
- 2 kondensatory elektrolityczne 100 μ F 16 V MB PC
- 2 kondensatory elektrolityczne 10 μ F 16 V MB PC
- 1 kondensator poliestrowy 470 nF MKT
- 4 kondensatory poliestrowe 100 nF MKT

Rezystory: (0,25 W, 1%, jeśli nie podano inaczej)

- 1 szt. 1 M Ω 2 szt. 100 k Ω 1 szt. 27 k Ω 1 szt. 20 k Ω 8 szt. 10 k Ω 7 szt. 1 k Ω
- 2 szt. 47 Ω 2 szt. 0,1 Ω 1 W (SMD 6432/2512-rozmiar; Panasonic ERJ1LWKF10CU lub podobny) [RS Components 566-989]



Ta fotografia pokazuje, jak będzie wyglądać ukończona myjka ultradźwiękowa, gdy w przyszłym miesiącu zajmiemy się budową i testowaniem. Pokażemy również, jak skonfigurować wanny do mycia ultradźwiękowego przy użyciu tanich pojemników „kuchennych”.

urządzenia. Dioda LED gaśnie podczas kalibracji przetwornika, a następnie zapala się po znalezieniu wymaganej wartości dla PR2.

Trwa to kilka sekund, chyba że coś jest nie tak, np. brak podłączonego przetwornika.

Po uruchomieniu, dioda LED1 świeci tylko wtedy, gdy przetwornik jest zasilany przy wymaganym ustawieniu mocy; pełni ona funkcję wskaźnika „blokady”.

Po naciśnięciu przełącznika Stop, zasilanie przetwornika przestaje działać, diody poziomu mocy gasną, a włącza się dioda zasilania. Dioda LED1 gaśnie po wyłączeniu głównego źródła zasilania przez S1 lub po odłączeniu lub wyłączeniu samego zasilania.

Minutnik

VR1 jest elementem sterującym pracą minutnika. Napięcie z jego styku podawane jest na wejście analogowe AN9 układu IC1 (końcówka 9) i zmienia się w zakresie od 0 V do 5 V. Odpowiada to zakresowi pracy minutnika od 20 sekund do 90 minut.

Minutnik rozpoczyna pracę po naciśnięciu przełącznika Start. Po upływie wybranego okresu zasilanie przetwornika wyłącza się.

Przełączniki S2 i S3 podłącza się odpowiednio do wejść RA0 i RA1 układu IC1. Wejścia te są utrzymywane w stanie wysokim (na poziomie 5 V) przez rezystory podciągające 10 kΩ. Zamknięcie przełącznika jest wykrywane po jego naciśnięciu, ponieważ wejście jest wtedy zwierane do 0 V.

Zauważ, że używamy przełączników przyciskowych, które mają styki wspólne (C), normalnie zamknięte (NC) i normalnie otwarte (NO).

Styki na przełączniku są ułożone w linii, ze wspólnym stykiem na jednym końcu, stykami NO w środku i stykami NC na drugim końcu. Zazwyczaj oznacza to, że do poprawnego działania przełącznika należy go odpowiednio zorientować na płytce drukowanej.

Jednak zaprojektowaliśmy płytkę drukowaną w taki sposób, że obie orientacje przełączników będą działać, jeśli końcówki C (wspólna) i NC (normalnie zwarta) zostaną połączone ze sobą na płytce.

Zasilanie

Zasilanie 12 V DC dla układu jest doprowadzane przez CON1. Wymagane są minimum 4 A. Jeśli używamy akumulatora 12 V, powinien on mieć pojemność 10 Ah lub więcej.

Zasilanie jest przełączane przez S1, który jest podłączony do PCB za pomocą zacisku śrubowego (gniazda CON2). Zasilanie doprowadzane jest następnie do regulatora 5 V (REG1) poprzez diodę D2 zabezpieczającą przed odwrotną polaryzacją. Liniowy regulator REG1 dostarcza 5 V wymagane przez IC1 i IC2.

12 V DC doprowadzane jest również do MOSFET-a Q6 przez bezpiecznik F1. Ten MOSFET jest używany jako przełącznik soft-startu do powolnego ładowania dużego kondensatora bocznikującego 4700 µF low-ESR. Bez miękkiego startu, ładowanie kondensatora 4700 µF spowodowałoby znaczny prąd udarowy. Mogłoby to spowodować przepalenie bezpiecznika lub wyłączenie zasilacza impulsowego 12 V.

Przy pierwszym włączeniu zasilania, Q6 jest wyłączony i kondensator 4700 µF nie jest ładowany. Po naciśnięciu przełącznika Start, wyjście RC3 układu IC1 przechodzi w stan wysoki (5 V), co powoduje włączenie tranzystora Q5. Napięcie bramki P-kanalowego MOSFET-a Q6 zaczyna spadać do 0 V, ponieważ kondensator 10 µF ładuje się poprzez rezystor 100 kΩ podłączony do kolektora Q5.

W miarę jak MOSFET zaczyna przewodzić, kondensator 4700 µF powoli się ładuje. Po pół sekundy ładowanie bramki zostaje przerwane przez wyłączenie Q5 i po 250 ms opóźnienia. Następnie za pomocą wejścia analogowego AN8 układu IC1 mierzone jest napięcie na kondensatorze 4700 µF.

Jeśli napięcie na kondensatorze jest poniżej 9 V (3 V na AN8), wszystkie diody poziomu migają dwa razy na sekundę. To wskazuje, że albo kondensator 4700 µF ma upływność, albo jest zwarcie w obwodzie powodujące rozładowanie kondensatora. Należy wtedy odłączyć zasilanie i znaleźć usterkę.

Jeśli nie ma błędu, Q5 jest ponownie włączany, aby kontynuować ładowanie bramki Q6. Potrzeba jednej sekundy, aby napięcie bramki spadło 7,5 V poniżej napięcia źródła, w tym czasie Q6 jest prawie w pełni włączony. Po kilku kolejnych sekundach napięcie bramki będzie bardzo bliskie 0 V, utrzymując pełne 12 V pomiędzy bramką a źródłem. Dioda Zenera ZD1 chroni bramkę przed przepięciem ograniczając napięcie bramka-źródło do -13 V.

Ochronę przed odwrotną polaryzacją dla sekcji zasilania układu zapewnia 4A bezpiecznik F1, dioda D1 oraz integralne diody pasytywne w MOSFET-ach Q1 i Q2. Przy odwrotnym podłączeniu zasilania diody te przewodzą prąd, skutecznie ograniczając napięcie zasilania na poziomie -0,7 V i chroniąc kondensator elektrolityczny 4700 µF przed nadmiernym napięciem wstecznym. Prąd ten szybko przepali bezpiecznik i odetnie odwrotne podłączone zasilanie.

Wanna

Przetwornik ultradźwiękowy musi być przymocowany do zewnętrznej części odpowiedniego pojemnika. Może on być wykonany ze stali nierdzewnej, aluminium lub tworzywa sztucznego, tak aby drgania ultradźwiękowe były skutecznie sprzężone z płynem. Sztywniejsze materiały sprzęgają fale ultradźwiękowe z mniejszymi stratami.

Idealnie byłoby, gdyby wanna miała płaski bok lub podstawę, do której można przymocować przetwornik. Materiał musi być również kompatybilny z żywicą epoksydową używaną do przyklejenia przetwornika do wanny. Metale są najbardziej odpowiednim materiałem.

Znaleźliśmy serię pojemników „gastronorm”, które są idealne, w sklepie z zaopatrzeniem kuchennym. Są to rodzaje pojemników na żywność, które często można zobaczyć w bufetach. Pasują do stołów parowych, które utrzymują jedzenie w cieple, i są dostępne w różnych kształtach i rozmiarach, z kilkoma dobrymi opcjami w pobliżu idealnej objętości 4 L.

Możesz dostać je wykonane ze stali nierdzewnej, poliwęglanu lub polipropylenu; dwie pierwsze opcje są najlepsze. My kupiliśmy nasze (na zdjęciu) na www.nisbets.com.au (mają sklepy w NSW, Vic, Qld i ACT). Prawdopodobnie nie będziesz miał problemu z zakupem zbliżonych na lokalnym rynku.

Polecamy pojemnik o głębokości 150 mm (pojemność 4 L), tacę o głębokości 100 mm (pojemność 3,7 L) lub tacę o głębokości 100 mm (pojemność 2,5 L), w zależności od Twoich wymagań.

Taca o głębokości 150 mm jest głęboka i prostokątna, podczas gdy taca o głębokości 100 mm jest bardziej kwadratowa i płytka. Inne pojemniki mają pośrednie wymiary.

Można również zakupić pojemniki ze stali nierdzewnej lub przezroczystego lub czarnego poliwęglanu, które pasują do wszystkich tych elementów, co byłoby dobrym pomysłem, jeśli czyszysz za pomocą rozpuszczalnika o silnym zapachu (zwłaszcza jeśli planujesz pozostawić rozpuszczalnik w łaźni, gdy go nie używasz).

Większe rozmiary wanien z większą ilością płynu będą miały słabszy efekt czyszczenia niż mniejsze pojemniki z niewielką ilością płynu.

Płyn używany w kąpeli może być wodą z kranu z kilkoma kroplami detergentu jako środka zwilżającego. Inne płyny, które można stosować to woda dejonizowana, alkohol (spirytus metylowy, alkohol izopropylowy itp.), aceton lub podobne rozpuszczalniki.

Skuteczność czyszczenia znacznie wzrasta, gdy płyn jest podgrzany. Wypełnienie około czterech litrów jest idealne dla mocy dostępnej z przetwornika ultradźwiękowego.



Gdybym wiedział, że przyjdiesz, upiekłbym ciasto... To niektóre z pojemników ze stali nierdzewnej, jakie znaleźliśmy w sklepie z artykułami kuchennymi i które byłyby idealne do tego projektu. Wybierz rozmiar i głębokość, które najlepiej pasują do Twojego zastosowania.

W przypadku głębszych pojemników można by je napełnić mniejszą ilością płynu do czyszczenia mniejszych przedmiotów.

Jednak trzeba będzie ponownie skalibrować urządzenie po każdej zmianie poziomu płynu, a może się okazać, że wyłączy się ono z mniejszą ilością płynu w zbiorniku ze względu na spadek impedancji przetwornika, a moc dostarczana przekroczy 40 W.

Takie podejście wymagałoby pewnych eksperymentów, aby móc je z powodzeniem stosować. Procedura rekalkibracji zostanie opisana później. Należy również zauważyć, że trzeba by zamontować przetwornik dość nisko na zbiorniku (lub pod dnem), aby umożliwić stosowanie różnych poziomów płynu.

Podsumowanie

W następnym miesiącu przedstawimy szczegóły budowy, w tym sposób nawijania transformatora T1, etapy montażu płytki drukowanej, okablowanie, hermetyzację przetwornika, przygotowanie obudowy i montaż końcowy.

Opiszemy również procedury testowania i kalibracji oraz podamy kilka wskazówek dotyczących najbardziej efektywnego wykorzystania myjki ultradźwiękowej.

Pozyskiwanie części do myjki ultradźwiękowej ...

Jak zwykle, dwie płytki PCB, zaprogramowany mikrokontroler do tego projektu i niektóre inne części, w tym MOSFET-y, można zamówić w SILICON CHIP ONLINE SHOP – szczegóły na stronach:

<https://www.siliconchip.com.au/Shop/7/5629>

Postanowiliśmy również zaopatrzyć się w przetwornik ultradźwiękowy, ponieważ nie jest on lokalnie łatwo dostępny. Jaycar sprzedawał przetwornik o mocy 50 W (Cat AU5556),

ale według ich strony internetowej, został on wycofany z produkcji.

Nasze przetworniki mają moc 50 W i są przeznaczone do pracy przy częstotliwości 40 kHz. Powinny być dostępne w magazynie do czasu ukazania się drugiej i ostatniej części tego artykułu w przyszłym miesiącu (Cat SC5629 @ \$54.90).

Możesz zdobyć pozostałe części do tego projektu od zwykłych dostawców: tzn. Jaycar i/lub Altronics; lub element14 albo RS Components dla bardziej specjalistycznych części.

Możesz również uzyskać prawie wszystkie części z Digi-Key; oferują oni bezpłatną ekspresową dostawę do Polski dla zamówień powyżej 200 PLN.

Części PVC do obudowy przetwornika są dostępne w sklepach z narzędziami, takich jak Bunnings, podczas gdy pojemniki do kąpeli są dostępne w Nisbets, jak opisano powyżej. ■

John Clarke

Adaptacja do wydania
polskiego – Andrzej Nowicki

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

REKLAMA

PIPEK DRĘCZYCIEL

AVTEDU625

sklep.avt.pl

Łatwe w budowie

Aktywne monitory Hi-Fi z opcjonalnymi subwooferami, część 2

W zeszłym miesiącu przedstawiliśmy ten nowy, fantastyczny, aktywny system głośników, który świetnie wygląda i brzmi, a do tego jego budowa nie jest droga. Nie potrzeba też ekstremalnych umiejętności czy specjalistycznych narzędzi. Poprzednio opisaliśmy projekt obudowy, wybór przetworników i wydajność. W tej drugiej części podamy kompletne szczegóły dotyczące montażu głośników głównych – monitorów.

Zdecydowaliśmy się na wykonanie tych głośników ze sklejk, ponieważ przy odrobinie staranności będzie ona wyglądała naprawdę ładnie. Jest to wciąż stosunkowo tani i sztywny materiał, a do tego dość łatwy w obróbce. Warto jednak pamiętać, że nawet drobny błąd może mieć fatalne konsekwencje.

Jak wyjaśniliśmy w zeszłym miesiącu, używamy dwóch rodzajów przetworników Altronicsa (lub trzech, jeśli zbudujemy opcjonalne subwoofery). Mają one przystępną cenę, dając jednocześnie doskonałe rezultaty.

Wszystko, co jest wymagane do osiągnięcia wysokiej wierności brzmienia, to dwie starannie zaprojektowane zwrotnice, których konstrukcja jest opisana poniżej.

Opiszemy również, jak zamontować użyte do zasilania tych aktywnych kolumn dwa „płytkowe” wzmacniacze klasy D (jeśli budujesz subwoofery, dla samych monitorów wystarczy jeden). Można je również wykorzystać w innych konstrukcjach głośnikowych. Są to gotowe moduły wzmacniaczy, więc montaż jest dość prosty.

Budowa zwrotnic pasywnych

Dla każdego systemu stereo należy zamontować dwie zwrotnice pasywne, tj. po jednej na każdą obudowę monitora.

Zwrotnica pasywna zamontowana jest na płycie drukowanej o kodzie 01101201



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://bit.ly/3XgpjiP>
Materiały dodatkowe są również dostępne na stronie edw.elportal.pl: <https://bit.ly/3XgqCtj>

i wymiarach 137×100 mm. Podczas dopasowywania elementów należy korzystać ze schematu montażowego płytek drukowanych (rysunek 12), oraz poglądowego zdjęcia.

Są tylko trzy kondensatory, z których żaden nie jest spolaryzowany, i trzy rezystory 5 W. Wszystkie powinny być oznaczone swoimi wartościami, więc po prostu wlutuj je tam, gdzie pokazano.

Chociaż nie jest to krytyczne, dobrze jest umieścić korpusy rezystorów kilka milimetrów nad górną częścią PCB, aby umożliwić obieg powietrza chłodzącego.

Następnie należy zamontować trzy identyczne dwudrożne terminale śrubowe, z otworami na przewody w kierunku na zewnątrz krawędzi płytki drukowanej. Pozostają jeszcze trzy duże dławiki z rdzeniem powietrznym, które po prostu „zrobiono” z całych szpul emaliowanego drutu miedzianego.

Aby wykonać te cewki, usuń naklejki znajdujące się na każdym z końców szpuli. W środku szpul zobaczysz 1–2 cm odcinki

drutu. Wyciągnij je, a następnie zeskrób dokładnie emalię ostrym nożem.

Przedłużamy ten przewód lutując do odsłoniętego końca kilka centymetrów pocynowanego drutu miedzianego. Wszystkie zakupione przez nas szpule miały 1–2 cm odcinki drutu wewnątrz szpuli, co nie wystarczy, aby osiągnąć do PCB.

Następnie wywierć mały otwór z boku plastikowej szpuli, aby umożliwić przeciągnięcie drugiego końca miedzianego drutu i zabezpieczyć go klejem na gorąco lub uszczelniaczem. Upewnij się, że drut jest mocno naprężony na szpuli, więc nic nie może się luzować i przesuwać. Tę końcówkę drutu również należy oczyścić z emalii.

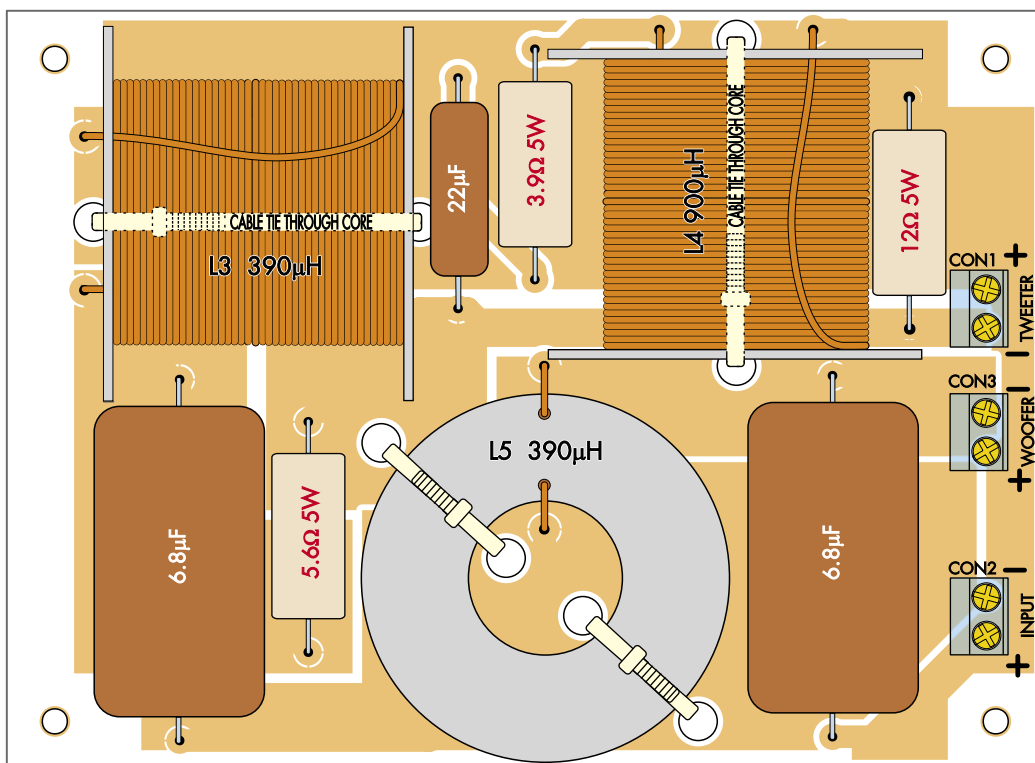
My zdecydowaliśmy się pomalować drut na szpuli lakierem elektroizolacyjnym. Jest on twardy i trzyma zwoje drutu mocno razem, więc nic nie może brzęczeć lub drgać. Jest to jednak czynność opcjonalna.

Po przylutowaniu wyprowadzeń cewek do ich umocowania na płycie drukowanej

Rysunek 12. Budowa zwrotnic pasywnych nie jest zbyt trudna, ponieważ potrzebne są tylko trzy cewki, trzy kondensatory, trzy rezystory i trzy bloki zacisków śrubowych. Cewki (które są pełnymi szpulami drutu miedzianego!) są nieporęczne i ciężkie, więc upewnij się, że są odpowiednio unieruchomione za pomocą opasek kablowych (jak na rysunku tutaj) lub uszczelnacza akrylowego/silikonowego.

Wczesny prototyp PCB jest pokazany poniżej, w rozmiarze nieco mniejszym niż rzeczywisty. Istnieją pewne różnice w komponentach pomiędzy zdjęciem a schematem montażowym po prawej stronie. Postępuj zgodnie ze schematem! Potrzebne są dwie zwrotnice – po jednej na każdą obudowę.

Zwróć też uwagę na nasze uwagi dotyczące kleju termotopliwego – choć ładnie trzyma cewki, to może zmięknąć, a nawet puścić, jeśli cewki się nagrzeją. Stąd przepis na użycie opasek kablowych, jak pokazano po prawej.



użyliśmy kleju termotopliwego, ale po wykonaniu tej czynności zdaliśmy sobie sprawę, że nie był to najlepszy pomysł. Klej termotopliwy jest dość skuteczny przy montażu takim jak ten, ale może wyglądać trochę niechlujnie, jeśli pracujesz niestarannie. I może puścić, jeśli płyta staje się zbyt gorąca podczas użytkowania.

Jeśli go używasz, uważaj, żebyś nie poparzył nim sobie boleśnie palców.

Przyklejenie cewek za pomocą neutralnie utwardzanego uszczelnacza silikonowego jest prawdopodobnie lepszym rozwiązaniem, ale najlepiej użyć opasek kablowych.

Tak więc końcowa płyta ma otwory, które pozwolą Ci przywiązać cewki za pomocą

odpowiednich dużych opasek kablowych, jak pokazano na rysunku 12.

Montaż obudów monitorów

Rysunek 13 (na następnej stronie) pokazuje jak pociąć dwa arkusze sklejk o wymiarach 600×1200 mm i grubości 15 mm na kawałki, które będą potrzebne do budowy dwóch głośników półkowych (monitorów). Element oznaczony jako „Subwoofer 2 front” jest potrzebny tylko wtedy, gdy zamierzasz zbudować subwoofery; w przeciwnym razie możesz go zostawić jako część ścianek.

Pokazane są również rysunki wycięć. Wybierając głębokość głośnika równą 297 mm, można przeciąć jeden arkusz przez środek, aby uzyskać spód, boki i wierzch dla pary głośników przy minimalnym koszcie i komplikacji.

Pozwala to na pozostawienie 6 mm na cięcie. Jeśli twoje kawałki będą nieco szersze niż 297 mm, nie będzie to stanowiło większego problemu, o ile wszystkie będą tej samej wielkości.

Użyliśmy materiału o grubości 15 mm o strukturze pięciowarstwowej ze sklepu Bunnings. Jest on dostępny w większości magazynów z narzędziami. Wybraliśmy go na podstawie ceny i dostępności.

Jeśli zależy Ci na naprawdę gładkim wykończeniu, korzystnie będzie wybrać drewno wyższej klasy, stąd nasza sugestia na liście części, aby użyć sklejk szklanych. Odradzamy eksperymenty z ręcznym



Zwróć uwagę na nasz komentarz dotyczący użycia kleju termotopliwego (widocznego na fotografii): opaska kablowa zaciskająca przez otwór cewki (jak pokazano na Rysunku 12 powyżej) jest znacznie bardziej pewna (zwłaszcza, gdy cewka się nagrzeje!).



Wytnij wzmocnienia jak pokazano na rysunku i użyj kleju oraz 3-4 śrub, aby przymocować je do paneli bocznych. Pokazany jest tutaj panel boczny subwoofera, ale dla paneli bocznych monitorów zasada jest ta sama.

- Sprawdź, czy najdłuższe wkręty do drewna, które posiadasz, mają odpowiednią długość, aby przejść przez materiał usztywniający i skrócić go z panelami głośnikowymi, nie wychodząc przy tym na wylot. Stwierdziliśmy, że wkręty 28 mm były w sam raz dla naszego materiału o grubości 15 mm. Należy również pamiętać, że w przypadku wkrętów mocujących głośniki od zewnątrz, należy wywiercić otwory pilotażowe i szafować je, aby łby wkrętów znalazły się w jednej płaszczyźnie.
- Potrzebujesz kleju i śrub. My użyliśmy standardowego kleju poli (winylooctanowego) (PVA), w Polsce dostępnego po nazwą Wikol. W trakcie budowy należy mieć pod ręką trochę akrylowego wypełniacza, którym można wypełnić wszelkie szczeliny.
- Do usztywnienia użyliśmy ścinków skleiki (prętów) 15×15 mm. Są one wystarczająco duże, aby umożliwić przykręcenie elementów, nie zajmując zbyt wiele objętości

wewnętrznej. Jeśli twoje wzmocnienie ma nieco inny wymiar, nie martw się niewielką zmianą objętości.

Po wycięciu wszystkich elementów, złoż je w następujący sposób:

Boki

Zacznij od przymocowania pionowego wzmocnienia do wewnętrznej strony paneli bocznych. Upewnij się, że pozostawiłeś szczeliny z przodu i z tyłu każdego panelu bocznego, nieco szersze niż grubość panelu przedniego i tylnego.

Najłatwiej jest to osiągnąć trzymając kawałek odciętego materiału obok usztywnienia podczas jego mocowania.

Na górnej i dolnej krawędzi przykręć od środka 2 listwy wzmocnienia krawędziami równo z górną i dolną krawędzią płyty bocznej. Wzmocnienie przykręć bardzo lekko od środka przy górnej i dolnej krawędzi, tak, aby po przykręceniu panelu górnego i dolnego wzmocnienie przyciągnęło górny/dolny panel do paneli bocznych. Pozwoli to na zminimalizowanie szczelin z boku obudowy.

Panele dolne

Wywierć otwory pilotażowe i przykręć panel dolny do boków, wkrętami wchodzącymi przez dno od zewnątrz w panele boczne. Jeśli planujesz zamontować głośniki w taki sposób,



że panel dolny będzie widoczny, prawdopodobnie będziesz chciał przykręcić panel dolny od wewnątrz.

Jeśli mocujesz panel dolny od wewnątrz skrzynki, wywierć otwory we wzmocnieniach (pewną orientację da Ci fotografia po lewej na dole strony), przez które przejdą śruby. Upewnij się, że te otwory we wzmocnieniu ustawione są pod kątem, aby wiertło pracowało prawidłowo, a końcówka śrubokręta prawidłowo wchodziła w nacięcia we łbie wkręta. Jeśli wywiercisz otwory pilotażowe idealnie pionowo, bardzo trudno będzie przykręcić potem panel dolny/górny.



Różne etapy procesu budowy. Zdjęcie po lewej stronie pokazuje wewnętrzne uszczelnienie w trakcie budowy skrzynki, aby zapewnić jej szczelność. Ważne jest również, aby utrzymać wyrównanie i prostokątny styk paneli (jak pokazano na środkowym zdjęciu); zdjęcie po prawej pokazuje śruby używane do zapewnienia stałej szczeliny wokół panelu przedniego i „naprostowania” obudowy (patrz tekst).



Użycie zacisku stolarskiego podczas wiązania kleju zapewni, że prostokątne narożniki pozostaną właśnie takie – pod kątem prostym.

Panele górne

Dokładnie wyrównaj panel górny z bocznymi i przytwierdź od środka wkrętami z tyłu skrzynki.

Następnie sprawdź wyrównanie paneli względem siebie. Jeśli podczas montażu przesunęły się, teraz jest ostatni moment, aby to skorygować!

W razie potrzeby nie bój się wywiercić nowego otworu w usztywnieniu, co pozwoli na użycie nowej śruby, następnie usuń pierwszą śrubę i użyj drugiego otworu, aby spróbować montażu ponownie. Nikt nigdy nie zobaczy drugiego otworu mocującego po wewnętrznej stronie obudowy, ale bez problemu zobaczy krzywo skrócone panele.

Gdy tył górnego panelu jest już prawidłowo wyrównany i szczelny, przykręć przód i środek panelu.

Panele przednie

Zanim przymocujesz panel przedni, musisz mieć pomysł na wykończenie obudów. Ja zabezpieczałem pudełka, więc mogłem po prostu zamontować panel przedni. Na tym etapie można go przykleić i przykręcić. Ponownie, zrób to od wewnątrz wkrętami przechodzącymi przez listwy wzmacniające. Wypełniacz akrylowy, kiedy już stwardnieje, jest doskonałym klejem. Użyłem go do przyklejenia frontu.

Pamiętaj, że klej będzie stanowił główne spoiwo, a śruby są po to, by utrzymać wszystko razem, dopóki klej nie stwardnieje.

Zanim ostatecznie skręcisz wkręty, posłuż się łatą stolarską i sprawdź, czy wszystkie kąty obudowy są proste. Prawo Murphy'ego mówi bowiem, że skoro mogą być krzywe, to będą. Istnieje kilka sposobów, aby naprostować pudełko przed skręceniem go razem. Jedną z opcji jest użycie opasek zapadkowych, lub zacisków stolarskich, ale opiszę sprytniejszą alternatywę.

Umieść duże śruby w szczelinie między panelami bocznymi a panelem przednim. Wkręć

je na tyle daleko, aby uzyskać równe szczeliny po każdej stronie oraz na górze/dole.

Zakładając, że panel przedni jest prostokątny, bo został wycięty profesjonalnym sprzętem – co naszym zdaniem jest rozsądnym założeniem – cała skrzynka będzie teraz na pewno prostokątna. Zobacz fotografię dalej.

Kiedy obudowa będzie już ładna i prosta, dokręć ostatnie śruby i pozwól klejowi wyschnąć przez kilka godzin.

Panele tylne

Proponuję pomalować wewnętrzną stronę wspornika montażowego czarną farbą. Nie przykręcaj panelu tylnego, dopóki wszystko nie będzie gotowe. Jeśli przypuszczasz, że możesz potrzebować ponownie dostać się do środka, połóż piankowy pasek uszczelniający na wewnętrznych wspornikach-wzmocnieniach. Dzięki temu będziesz miał duże szanse na przeprowadzenie poprawek wewnątrz obudowy.

Pospieszaliśmy się i za wcześniej przykleiliśmy tylny panel. Prace montażowe przez małe otwory na głośnik niskotonowy były naprawdę uciążliwe!

Dodawanie usztywnienia

Użyj ścinków, aby zrobić usztywnienie, które będzie przebiegać poziomo przez głośnik. Umieść go mniej więcej na środku paneli bocznych i użyj wypełniacza akrylowego, aby przykleić go na miejscu. Nie obawiaj się, gdy akryl stwardnieje, usztywnienie będzie wystarczająco mocne, aby wytłumić rezonans tych paneli.

Wycinanie otworów

Kiedy karkasy obudów są już gotowe, nadszedł czas na wycięcie otworów. Zobacz szczegóły na zdjęciu obok. Na tylnym panelu, na jednym głośniku będziesz miał duży otwór na wzmacniacz, a na drugim mały na zaciski głośnikowe. Wytnij wszystkie otwory przed wykończeniem głośników. Wymiary i miejsca

wycięć są przedstawione na rysunkach arkusza cięcia (rysunek 13). Wykończenie głośników to tak naprawdę kwestia gustu i wyboru materiałów.

Port bass-reflexu

Port wykonany jest z rury PVC o średnicy 40 mm i długości 10,5 cm. Ma ona średnicę wewnętrzną około 38 mm. Jej długość jest umiarkowanie ważna, więc postaraj się zmieścić ją w granicach ± 3 mm.

Kiedy jeszcze możesz dostać się do środka obudowy, użyj palca, aby nałożyć warstwę wypełniacza akrylowego od wewnątrz wokół portu, kiedy będzie on już włożony przez przedni panel. Jeśli otwór na port jest trochę za duży, aby go zabezpieczyć, użyj wypełniacza, i pozwól mu zastępnąć.

Wykończenie skrzynki

Zauważ, że na końcu krawędzie sklejkі zostały zaokrąglone pod kątem 45°. W ten sposób podkreśliłem fakt, że w tym miejscu znajduje się słój końcowy. Pierwszy raz zrobiłem to podczas eksperymentu, aby sprawdzić, czy ma to wpływ na rozproszenie na krawędziach panelu przedniego. Byłem znacznie mniej przekonany co do wpływu na rozproszenie w porównaniu z wpływem frezowania na estetykę.

Frezowanie to zostało wykonane przy użyciu taniego freza i uchwytu kąтового 45° z łożyskiem krawędziowym. Krawędzie wykonałem w dwóch przejściach, pierwsze mniej więcej do połowy ostatecznej głębokości.



Gotowa obudowa, wyszlifowana i gotowa do bejcowania, lakierowania lub malowania.



Wycinanie otworów wyrzynarką może nie jest najbardziej estetyczną pracą, ale krawędzie są zakryte przez mocowanie głośników, więc nie musi to być dzieło sztuki! Porada – użyj wycinarki trepanacyjnej z wiertłem pilotującym!

Bejcowanie/lakierowanie

Aby zabejcować i polakierować obudowę, należy szlifować, szlifować i jeszcze raz szlifować. Zawsze szlifuj wzdłuż włókien. Jeśli będziesz szlifował w poprzek włókien, papier ścierny rozerwie włókna w drewnie, a bejcowanie/barwienie podkreśli to w sposób, którego na pewno nie zaakceptujesz. Jakkolwiek dziwnie to brzmi, oznacza to, że szlifierki tarczowe są praktycznie bezużyteczne.

Na początek stosuj papier ścierny o ziarnie 120. Używaj go bez nadmiernej oszczędności i wyrzucaj, gdy tylko się zatrze. Potem wystarczają dwie warstwy lakieru. Idealnie byłoby przeszlifować skrzynkę papierem ściernym o granulacji 240 (lub 400) po nałożeniu pierwszej warstwy lakieru, a następnie oczyścić z pyłu. Szlifuj tak, aby uzyskać piękne, gładkie wykończenie. Kolejna warstwa lakieru wyjdzie gładka, dając lustrzane wykończenie.

Teraz podłącz głośnik niskotonowy i wysokotonowy za pomocą drutu o średnicy co najmniej 1 mm. Nie trzeba też przesadzać, wystarczy nie używać za cienkiego drutu.

Podczas montażu głośnika niskotonowego i wysokotonowego, umieść pasek taśmy piankowej wokół krawędzi otworu montażowego, aby utworzyć ładne uszczelnienie pomiędzy przetwornikiem a panelem przednim.

Nie używaj wypełniacza akrylowego lub silikonowego. Widziałem tak wykonane kolumny, co sprawia, że głośniki są niemożliwe do naprawy!

Głośniki przykręć za pomocą 15 mm wkrętów do drewna, ale jeszcze nie teraz. Uważam, że znacznie łatwiej jest wywiercić otwory pilotażowe przy użyciu wiertła 1,5–2 mm. Dzięki temu wyrównanie jest łatwe i zmniejsza się szansa, że coś się wyślizgnie, gdy zaczniesz przykręcać głośniki.

Aby zamontować zwrotnicę w każdej skrzynce, najpierw przytnij parę czerwono-czarnych przewodów o dużej obciążalności (lub kabla w kształcie ósemki), tak aby sięgały od zacisków wejściowych zwrotnicy wewnątrz obudowy, przez otwór z tyłu, z około 10 cm zapasem długości na zewnątrz skrzynki.

Odizoluj przewody na obu końcach i podłącz każdą parę do zacisków wejściowych zwrotnicy pasywnej.

W tym miejscu najłatwiej jest również przymocować przewody idące ze zwrotnicy do głośnika niskotonowego i wysokotonowego, upewniając się, że są one na tyle długie, że wychodzą z przodu obudowy, aby można je było przymocować do zacisków głośników przed montażem.

Teraz przyklej piankę do spodu płytek drukowanych zwrotnic i przykręć je do dolnych paneli obudów, używając wkrętów 10–12 mm. Podłącz przewody głośnika niskotonowego i wysokotonowego zgodnie z oznaczeniami na płycie drukowanej.

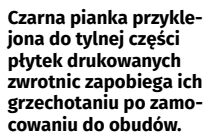
Wnętrza głośników należy wyłożyć warstwą waty poliakrylonitrylowej, jak pokazano na rysunku. Umocuj ją na miejscu za pomocą zszywek tapicerskich. Jeśli nie masz pistoletu do zszywek, użyj gwoździ 40 mm i wbij je na głębokość 10 mm, a następnie zagnij je, aby utrzymać watę na miejscu. Wata jest dostępna w sklepach rzemieślniczych, takich jak Lincraft. Kup gruby materiał i nie oszczędzaj na ilości.

Teraz wkręć śruby głośników, uważając, aby śrubokręt nie poślizgnął się i nie trafił w stożek. Śruby z Ibm Philipsa ułatwiają wkręcanie, ale najważniejsze jest to, żeby się nie poślizgnąć!

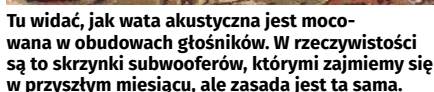
Strona stolarska
– w tym lakierowanie bezbarwnym – jest już zakończona. Pozostaje tylko wbudować wzmacniacz i zwrotnice, zamontować głośniki i ... zrelaksować się muzyką!



Dwa głośniki Altronics wybrane do tego projektu: po lewej stronie znajduje się głośnik wysokotonowy C-3019, a po prawej (oczywiście w innej skali!) głośnik nisko/średniotonowy C-3038.

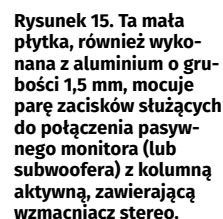
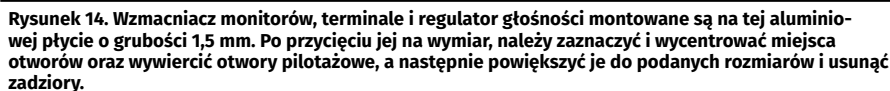


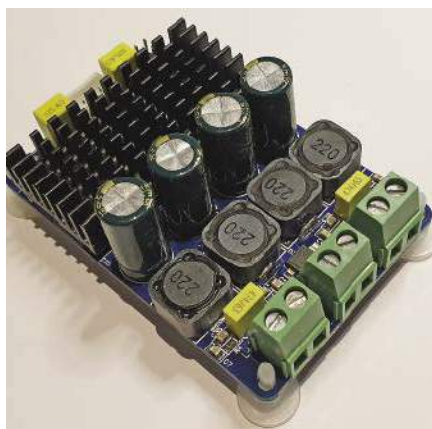
Przed wierceniem otworów należy je wycentrować. Jeśli nie masz punktaka, użyj dużego gwoźdźdza i młotka. Pomoże Ci to uzyskać otwory dokładnie tam, gdzie chcesz. Następnie wywierć otwory pilotażowe



Otwory montażowe muszą być na tyle duże, aby zmieściły się w nich śruby, których użyjesz

Pamiętaj, że ta płyta montowana jest na zewnątrz obudowy głośnika, więc uważaj, aby ładnie ją wykończyć, a najlepiej pomalować aluminium, aby je zabezpieczyć.





Zbliżenie na wzmacniacz stereo TDA7498 80W/Channel Class-D, który kupiliśmy na eBayu za mniej niż 20 \$ – wliczając w to koszty wysyłki! Nie masz najmniejszej szansy na zbudowanie takiego w podobnej cenie. Dodaj zasilacz 24 V DC i wzmacniacz jest gotowy do rock'n'rolla (lub klasyki, lub swingu, lub orkiestry czy jazzu...!) Identyfikacyjny wzmacniacz znajdziesz na znanym portalu jako: TDA7498 Dual-Channel Digital Audio Stereo Power Amplifier Board Class D 2x100W DC 8-32V, przykładowo: (<https://bit.ly/3XS1csc>)

Skoro już przy tym jesteśmy, to teraz jest też dobry moment na wycięcie i wiercenie małej aluminiowej płytki, która będzie przymocowana do tylnej części drugiego (pasywnego) głośnika półkowego. Szczegóły pokazane są na **rysunku 15**.

Jeśli zamierzasz pomalować metalowe płytki (dobrze pasuje kolor czarny), zrób to teraz.

Malowanie nie jest absolutnie konieczne, jeśli zamierzasz przykleić pełnowymiarową etykietę panelu tylnego, w opcji opisanej poniżej, chociaż nadal może to być dobry pomysł, aby zapobiec korozji.

Przygotowanie modułu wzmacniacza TDA7398

Przed montażem wzmacniacza należy sprawdzić mocowanie radiatora do układu scalonego wzmacniacza.

Podczas rutynowego przeglądu wzmacniacza (kupiliśmy wiele egzemplarzy, zanim zdecydowaliśmy się na ten moduł) zauważyliśmy,

że mocowanie radiatora było czasami trochę symboliczne. W rzeczywistości, niektóre radiatory były zamontowane bez pasty termicznej, a niektóre były całkiem luźne!

Aby to sprawdzić, należy odkręcić dwie śruby, które mocują radiator do PCB. Znajdują się one od spodu. Zdejmij radiator i jeśli jest na nim pasta termiczna, usuń ją czystą chusteczką. Jeśli pasty nie ma, podziękuj swojej przeczności, że to sprawdziłeś!

Teraz nałóż trochę świeżej pasty termoprzewodzącej, a następnie dodaj 3 mm ząbkowaną podkładkę przeciwwstrząsową pod leń każdej ze śrub montażowych, jeśli się tam nie znajdowały.

Nasze moduły zostały dostarczone bez nich i jesteśmy pewni, że po kilku latach spędzonych w obudowie głośnika, radiator poluzowałby się.

Kiedy wymienisz śruby, nie dokręcaj ich zbyt mocno. Wywierają one nacisk na układ scalony wzmacniacza, przyciskając do niego radiator i nieco wyginając płytkę drukowaną.

To działa, o czym świadczy sposób montażu wielu radiatorów procesorów komputerowych, ale to jest mały układ, więc trzeba z nim uważać.

Tak więc gdy już pierwsza śruba „chwyci”, zrób dodatkowy obrót. Następnie zrób to samo z drugą. Powtarzaj tę czynność, aż poczujesz, że radiator jest dociskany do wzmacniacza IC.

Dodaj jeszcze po pół obrotu lub dokręcaj tak długo, aż poczujesz, że całość tworzy sztywny zespół, a obciążenie jest przejmowane przez PCB. Pozwól PCB trochę się wygiąć; w ten sposób tworzy się sprężyna, która będzie trzymać radiator ciasno na wzmacniaczu.

Gdy wszystko jest już razem, warto dodać trochę czerwonej farby lub lakieru do paznokci na łby śrub, aby zablokować je mocno i zapobiec ich odkręcaniu się pod wpływem wibracji.

Dopasowanie elementów do płyty

Teraz wzmacniacz jest gotowy do zamontowania na płycie bazowej. Najpierw jednak

powinieneś zastanowić się, jak zamierzasz oznaczyć złącza z tyłu płyty.

W naszym przypadku przekopaliśmy się przez dolną szufladę w kuchni i znaleźliśmy maszynkę do etykietowania – wygniatań liter w samoprzylepnej taśmie PCW.

Wykonała ona świetną robotę, tworząc etykiety na tylną płytę wzmacniacza. Etykiety przylgną do płyty po kilku miesiącach, gdy zapomnisz, która wtyczka do czego służy!

Jeśli zamierzasz przykleić etykiety, możesz to zrobić później. Jako inną opcję przygotowaliśmy projekt nakładki, który można pobrać w formacie PDF ze strony internetowej Silicon Chip i wydrukować na folii transparentnej (jako lustro, więc tusz trafia do środka) lub na etykiecie samoprzylepnej, przymocowanej do zewnętrznej strony tylnego panelu, jeśli wolisz takie rozwiązanie.

Jeśli zamierzasz przykleić etykietę na całym panelu, musisz to zrobić przed zamontowaniem innych elementów. Po przyklejeniu, wytnij ostrym nożem otwory dla poszczególnych komponentów i jesteś gotowy do dalszej budowy.

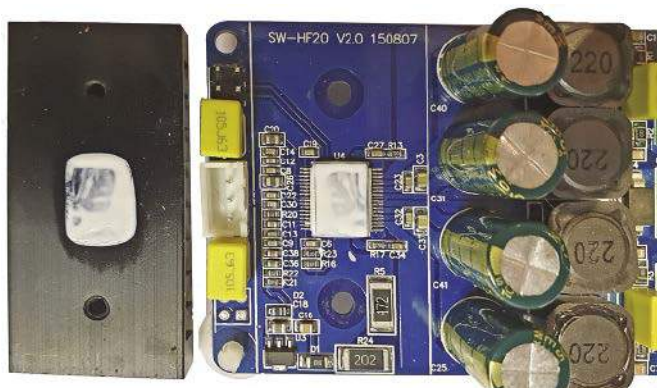
Teraz dopasuj do płytki wzmacniacz, zaciski wejściowe, zaciski wyjściowe dla drugiego głośnika i złącze zasilania.

Pozycje montażowe elementów oraz informacje o okablowaniu pokazane są na **rysunku 16** oraz na zdjęciu pod nim.

Zacznij od zamocowania zacisków wejściowych, wyjściowych i zasilania oraz potencjometru głośności, a następnie podłącz je zgodnie z rysunkiem. Upewnij się, że po zamontowaniu wzmacniacza płytowego wejścia przewodów zacisków głośnikowych będą skierowane do góry. Jeśli tego nie sprawdzisz, możesz się później pomylić!

Ostrożnie dokręć nakrętkę gniazda DC, ponieważ gwint jest aluminiowy. Dokręć ją, ale uważaj, żeby nie przesadzić. Gdy montujesz potencjometr głośności, dokręć jego nakrętkę.

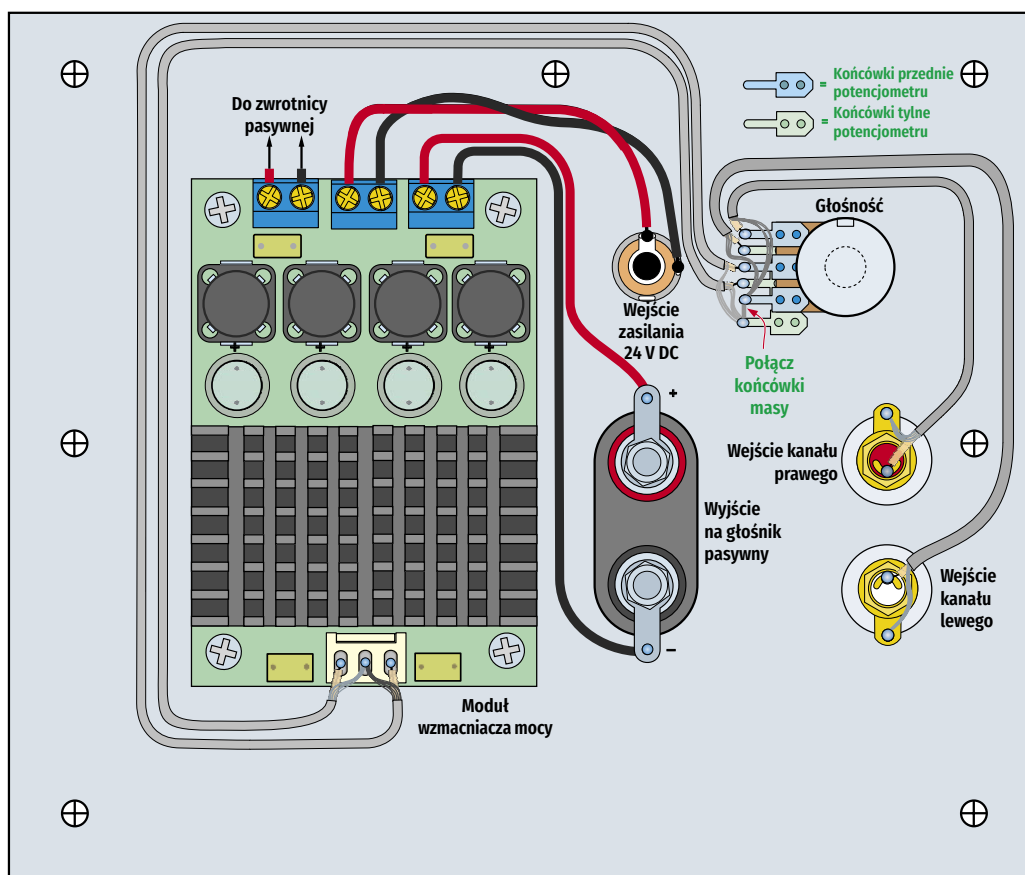
Jeśli przewierciłeś się przez panel dla kołka blokującego potencjometr, uszczelnij otwór



Użyj dużo pasty termoprzewodzącej zarówno na tylnej części radiatora jak i na wzmacniaczu IC...



...i jakiegoś środka zabezpieczającego (lakier do paznokci też działa dobrze), aby zapewnić, że śruby nie poluzują się z czasem.



Rysunek 16. Kiedy płytka wzmacniacza jest już gotowa, przymocuj i podłącz elementy zgodnie z tym, co przedstawia niniejszy schemat montażowy. Wszystkie masy wejść RCA i masa wejścia modułu wzmacniacza są podłączone do zwartych końcówek masy podwójnego potencjometru regulacji głośności – tu na dole po lewej stronie. Styki środkowe gniazda wejściowego RCA podłączone są do oddzielnych zacisków MAX VOL. potencjometru – tu na górze po lewej, natomiast wejścia wzmacniacza wychodzą z odpowiednich zacisków styków ślizgaczy obrotowych potencjometru – tu pośrodku po lewej.

używając odrobiny neutralnego silikonu od wewnątrz. Jeśli go nie masz, użyj trochę akrylowego wypełniacza, którego użyłeś podczas budowy skrzynki.

Podane złącza wejściowe RCA są złączami przelotowymi z integralnymi tulejami izolacyjnymi. Po odpowiednim zamontowaniu, tuleja mieści się wewnątrz otworu 8 mm, izolując gniazdo RCA od panelu.

Moduł wzmacniacza jest montowany po wewnętrznej stronie panelu tylnego na gwintowanych kołkach dystansowych o długości 10–25 mm za pomocą śrub i podkładek przeciwwstrząsowych.

Do okablowania wejściowego i regulacji głośności należy użyć przewodów ekranowanych, a do okablowania zasilania i wyjść głośnikowych – przewodów o dużej obciążalności prądowej lub przewodów typu „8”.

Wejścia wzmacniacza znajdują się na 3-kołkowym, kierunkowym złączu wtykowym o rozstawie 2,54 mm.

Należy usunąć izolację z jednego końca ekranowanego przewodu stereo i zaciśnąć i/lub przylutować dwie wewnętrzne żyły i zewnętrzny ekran do styków wtyczki, jak pokazano na rysunku 16.

Można użyć dwóch oddzielnych, jednożyłowych przewodów ekranowanych lub jednego dwużyłowego przewodu ekranowanego. Ta druga opcja ułatwia nieco budowę.

Zwróć uwagę, że styki tej wtyczki mają dwie części zaciskowe, jedną do kontaktu z gołym przewodem miedzianym i jedną do przytrzymania plastikowej izolacji. Należy upewnić się, że oba są dobrze zaciśnięte. Najlepiej użyć do tego celu narzędzia przeznaczonego specjalnie do tego celu, ale w razie potrzeby można użyć szczypiec do zaciskania.

Najlepiej jest dodać odrobinę lutu na koniec miedzianych drutów po zaciśnięciu

każdego ze złączy, aby wszystko było niezawodnie połączone.

Kiedy wszystkie trzy styki są już gotowe, wepchnij je do gniazda, upewniając się, że złącze masy jest pośrodku. Jeśli styki nie chcą się zatrzasnąć, być może nie zaciśnąłeś izolacji przewodów wystarczająco mocno. Ponadto zatrzasną się one tylko w jednej pozycji względem obudowy gniazda.

Podczas lutowania przewodów do zacisków głośnikowych i gniazda DC będzie

Oto finalna płytka wzmacniacza, gotowa do zamontowania w jednej z obudów (jest to wzmacniacz stereo, więc wystarczy tylko jeden). Kable łączą się z końcówkami wyjściowymi do drugiej obudowy. To zdjęcie porównaj ze schematem montażowym powyżej (rysunek 16).



to o wiele łatwiejsze, jeśli najpierw nałożysz odrobinę topnika na zaciski. Po zakończeniu użyj rurek termokurczliwych, aby mieć pewność, że nic się później nie zewrze.

Zdecydowana większość zasilaczy typu wtyczkowego używa wtyczki o średnicy 5,5/2,5 mm, z minusem na zewnątrz i plusem wewnątrz – rysunek 16 pokazuje okablowanie dla tego przypadku.

Sprawdź swój zasilacz; w mało prawdopodobnym przypadku, gdy jest to typ o odwrotnej polaryzacji wtyczki, zamień miejscami przewody dla gniazda DC.

Powinieneś mieć kompletny wzmacniacz głośnikowy, gotowy do zainstalowania w monitorze lub innym głośniku, który chcesz uczynić „aktywnym”.

Zakończenie montażu głośnika

Tutaj nie pozostało wiele do zrobienia. Jeśli zgodnie z wcześniejszymi zaleceniami podłączyłeś przewody do zacisków wejściowych zwrotnic pasywnych, wystarczy, że podłączysz je do wolnego bloku zacisków na płycie wzmacniacza w głośniku aktywnym, lub przyłączysz je do końcówek po wewnętrznej stronie konektorów w głośniku pasywnym.

Teraz warto zmniejszyć głośność, podłączyć zasilacz, podłączyć źródło sygnału i sprawdzić, czy wszystko działa.

Jeśli tak, to przykręcamy tylne panele do obu obudów i jesteśmy gotowi do odsłuchu!



Przeciwny (zewnątrzny) widok płytki wzmacniacza widocznej na stronie 23. Jest ona przykręcona do wycięcia w obudowie aktywnej (patrz poniżej). Chociaż regulacja głośności może być ustawiana dowolnie, to prawdopodobnie nie będzie zbyt wygodna. Wyobraźmy sobie, że ta regulacja może być „ustawiona do akceptowalnego poziomu i zapomniana”, a głośność regulowana ze źródła sygnału

Dobrym pomysłem jest przyklejenie wokół krawędzi tej samej taśmy piankowej, co w przypadku przetworników, aby dobrze je uszczelnić.

Przy okazji, chociaż uważamy, że równowaga basu/średnich tonów/wysokich tonów w tych głośnikach jest idealna, to gdybyście mieli wrażenie, że są one nieco „mdłe”, można nieznacznie zmienić ustawienie

zwrotnic pasywnych, aby zwiększyć o około 2 dB poziom tonów wysokich.

W tym celu należy usunąć rezystory 12 Ω i zamienić rezystory 5,6 Ω na 4,7 Ω . Jednak może to również prowadzić do zwiększenia zniekształceń, ponieważ głośniki wysokotonowe będą wtedy znacznie słabiej tłumione.

Sugerujemy, abyś najpierw dobrze posłuchał głośników i upewnił się, że naprawdę chcesz dokonać tej zmiany. Jednakże, możesz łatwo wrócić do pierwotnego ustawienia, jeśli po wypróbowaniu tej modyfikacji nie będziesz zadowolony z rezultatu.

W przyszłym miesiącu

Zakończymy ten projekt opisem opcjonalnych subwooferów. Oczywiście, będąc opcjonalnymi, nie muszą one trafić do Twojego pokoju czy na podłogę, możesz przecież użyć głośników tak, jak opisano do tej pory. To zależy od Ciebie... ale subwoofery naprawdę wydobywają z dźwięku to, co najlepsze! ■

Phil Prosser

Adaptacja do wydania polskiego – Andrzej Nowicki



Przód i tył gotowych głośników. Ujęcie z tyłu to głośnik, w którym znajduje się wbudowany wzmacniacz audio.

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



Materiały dodatkowe dostępne są na stronie
Silicon Chip: <http://bit.ly/3JUrox5>
Materiały dodatkowe są również dostępne
na stronie edw.elportal.pl: <https://bit.ly/3XL69rg>

Latarnia morska „Nocny Opiekun”

Latarnia „Nocnego Opiekuna”, jak prawdziwa latarnia morska, na krótko rozświetla ciemność, aby uchronić dziecięce nocne marzenia przed utonięciem w ciemności. Jest to doskonały projekt dla początkujących; jest łatwy w budowie, a przy okazji zapoznasz się z ważnymi elementami teorii układów elektronicznych.

Czy Twoje dzieci lub wnuki zerkają czasami z ciekawością na części elektroniczne i gadżety, z którymi pracujesz na stole warsztatowym?

W takich chwilach warto mieć pod ręką prosty projekt, który zachęci następne pokolenie do zajęcia się tym hobby.

Kiedy moje wnuki planowały ostatnio wizytę, zapytano mnie, czy mógłbym pomóc 8-latkowi zbudować „coś elektronicznego”. Czy to nie brzmi znajomo?

Próbując znaleźć odpowiedni układ dla dzieci ważne jest, aby mogły go zbudować w miarę szybko, zanim stracą zainteresowanie. Równocześnie powinien być na tyle użyteczny, by zyskać aprobatę rodziców.

Od kilku lat na półce nad moim stołem warsztatowym mam obwód migającego światła. Zbudowałem go podczas testowania pewnych pomysłów na dyskretne zasilacze o wysokiej wydajności. Wykorzystujący odpadowe części układ został zbudowany na skrawku płytki prototypowej. Obecnie pozwala mi zużywać do samego końca ogniwa 1,5 V niezdolne do zasilania pilota, latarki czy aparatu fotograficznego. Jest to prosty sposób na wydobywanie ostatniego tchnienia energii z takich prawie martwych baterii.

Pomyślałem, że zamiast budować mrugającą lampkę, mogę uczynić ją nieco bardziej użyteczną i ekscytującą dzięki kilku prostym ulepszeniom. Po pierwsze, zaprojektowałem płytkę drukowaną (PCB), aby



ułatwić dzieciom (i rodzicom, dziadkom lub opiekunom) budowę. Płytkę drukowaną pozwoliła mi naśladować powszechnie rozpoznawalny obiekt i uczynić projekt bardziej atrakcyjnym i interesującym.

Zasugerowała też kilka innych zastosowań, o których będzie mowa później.

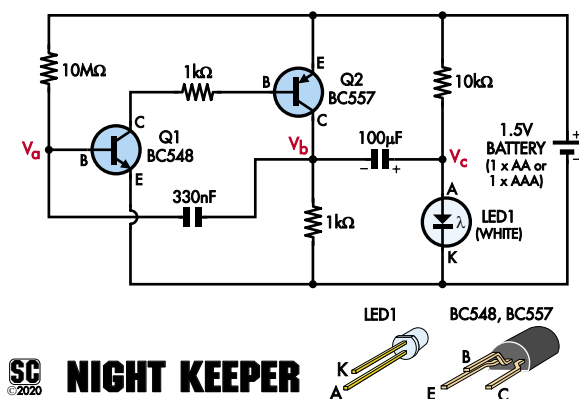
Oto więc „Nocny opiekun”.

Budowa, z pomocą osoby dorosłej, leży w zasięgu możliwości bystrego 10–12-latka. Ponieważ wymagana jest lutownica, niezbędny będzie ścisły nadzór osoby dorosłej i dobrze wentylowane miejsce pracy. Stół kuchenny z podobną wolną przestrzenią roboczą o powierzchni około jednego metra kwadratowego jest idealny; przykryj go wykładziną lub kartonem, aby chronić powierzchnię.

Opis układu

Ten prosty i dobrze znany obwód oscylatora (pokazany na **rysunku 1**) składa się z dwóch tranzystorów, białej diody LED i kilku elementów biernych. Zasilana pojedynczym ogniwem 1,5 V dioda LED miga jasno raz na sekundę przez wiele miesięcy. Nawet prawie wyczerpana bateria może zasiląć diodę przez miesiąc lub dwa.

Dwa tranzystory tworzące serce urządzenia działają jako wysoce wydajny oscylator pompujący. Po pierwszym włączeniu zasilania napięcie na bazie Q1 (Va)



Rysunek 1. „Nocny Opiekun” wykorzystuje dwutranzystorowy oscylator doysterowania pompy ładunkowej opartej na kondensatorze elektrolitycznym 100 μF i złącza diodowym białej diody LED1. Mniej więcej raz na sekundę, w punkcie oznaczonym jako „Vc” wystąpi o około dwukrotności napięcia baterii (około 3 V), zapewniając napięcie wystarczające do jasnego zaświecenia diody LED przez kilkadziesiąt milisekund.

zaczyna powoli rosnąć, ponieważ kondensator 330 nF jest ładowany przez rezystor 10 MΩ z baterii. Gdy Va osiągnie poziom około 0,6 V, złącze baza-emiter Q1, które działa podobnie jak dioda krzemowa, zostaje spolaryzowane w kierunku przewodzenia.

Tymczasem kondensator 100 μF poprzez rezystor 10 kΩ jest szybko ładowany do poziomu zbliżonego do napięcia baterii. Jest to około 1,5 V dla nowego ogniwa. W ten sposób również na diodzie LED1 (Vc) powstaje napięcie bardzo bliskie 1,5 V. Jednak dioda LED1 nie może się jeszcze świecić, ponieważ białe diody LED potrzebują więcej niż 2,5 V do działania.

Gdy tylko Q1 zacznie się włączać, rosnący prąd baza-emiter powoduje, że prąd kolektora rośnie jeszcze szybciej, ponieważ wzmocnienie prądowe tranzystora (beta lub hFE) jest znacznie większe od jedności. To z kolei powoduje, że złącze emiter-baza tranzystora Q2 również zaczyna przewodzić. W momencie, gdy Q2 zaczyna przewodzić, napięcie Vb zaczyna rosnąć ze względu na prąd emiter-kolektor płynący przez Q2 (inaczej mówiąc Q2 zaczyna zwierać Vb do plusa zasilania).

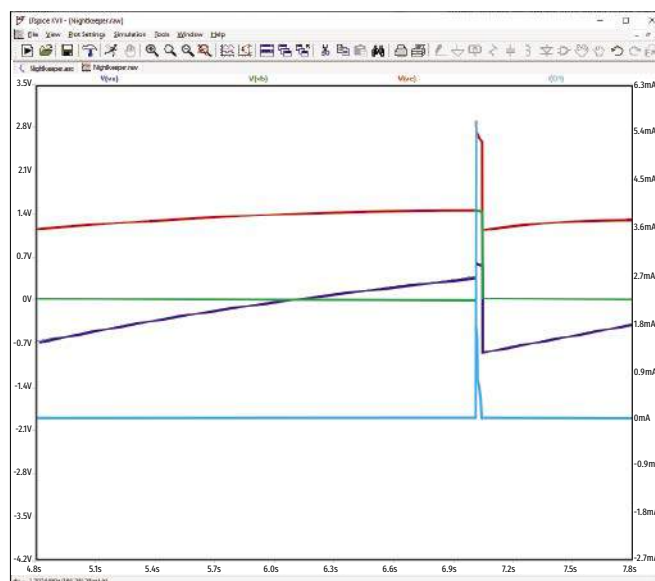
Tranzystor Q2 jeszcze bardziej zwiększa prąd kolektora Q1, w wyniku własnego wzmocnienia prądowego. Wzrastające napięcie Vb powoduje wzrost napięcia Va poprzez jego zatrzaśnięcie, ponieważ wzrost ten jest sprzęgany przez kondensator 330 nF. Wywołuje to szybki efekt „lavinowy” w tranzystorach Q1 i Q2, powodując ich pełne włączenie w wyniku ich łącznego wzmocnienia prądowego.

Proces ten jest zademonstrowany przy pomocy symulacji pokazanej na rysunku 2. Napięcie Va jest pokazane w kolorze cyan, Vb na zielono, a Vc na czerwono. Prąd płynący przez diodę LED1 jest zaznaczony na niebiesko. Widać, że wszystkie trzy napięcia rosną gwałtownie w tym samym czasie, co zbiega się z impulsem prądu diody LED1.

Podczas gdy dioda LED1 świeci, kondensator 330 nF utrzymuje tranzystor Q1 w stanie włączonym i w ten sposób rozładowuje się przez jego złącze baza-emiter. Pozwala to utrzymać Q1 w stanie włączonym przez około 30 ms. Jednak gdy tylko Va spadnie poniżej 0,6 V, Q1 zaczyna się wyłączać. To powoduje, że Q2 również gwałtownie się wyłącza. W rezultacie napięcie Vb spada nagle z 1,5 V do 0 V. Natomiast Va, poprzez kondensator 330 nF, spada następnie z 0,5 V do –1 V.

Takie ujemne napięcie Va wynika z tego, że tuż przed wyłączeniem Q1 i Q2, Va jest na poziomie około 0,5 V, podczas gdy Vb jest na poziomie około 1,5 V. Więc kiedy Vb spada do 0 V, to wskutek sprzężenia przez kondensator również Va spada wg równania: $0,5 \text{ V} - 1,5 \text{ V} = -1 \text{ V}$.

W tym momencie cały cykl zaczyna się od początku. W rezultacie otrzymujemy bardzo wydajny oscylator pompujący, który wytwarza krótki, ale jasny błysk białej diody LED mniej więcej raz



Rysunek 2. Ta symulacja pokazuje jak zmieniają się w czasie napięcia w punktach: Va (cyan), Vb (zielony) i Vc (czerwony) na rysunku 1. Va wzrasta, a następnie wszystkie trzy napięcia nagle wzrastają, co powoduje skokowy impuls prądu przez diodę LED1 (niebieska linia), aż do momentu spadku napięcia i rozpoczęcia procesu od nowa.

na sekundę lub dwie. Jest to w dużej mierze zależne od czasu, jakiego potrzebuje kondensator 330 nF, aby naładować się poprzez rezystor 10 MΩ od –1 V do około 0,6 V.

Zauważ, że chociaż lista części sugeruje typy tranzystorów BC54x i BC55x, możesz również użyć 2N3904, 2N2222 lub 2SC1815 dla tranzystora NPN; oraz 2N3906, 2N2907 lub 2SA1015 dla PNP. Prawie każda para tranzystorów NPN i PNP małej mocy będzie działać, ale pamiętaj, że rozmieszczenia końcówek mogą się różnić.

Budowa

Jeśli wszystkie części są gotowe, budowa „Nocnego Opiekuna” powinna zająć około godziny. Młodsze dzieci mogą potrzebować więcej czasu. Podzielenie budowy na dwie części, w których w jednej krótkiej sesji zamontujesz oporniki i kondensatory, a w drugiej pozostałe części, ułatwi budowę i będzie bardziej odpowiednie dla małych dzieci, które potrafią skupić uwagę tylko przez krótszy czas.

Zakończenie projektu poprzez dodanie uchwytu na baterie i podstawy może być wykonane podczas trzeciej sesji.

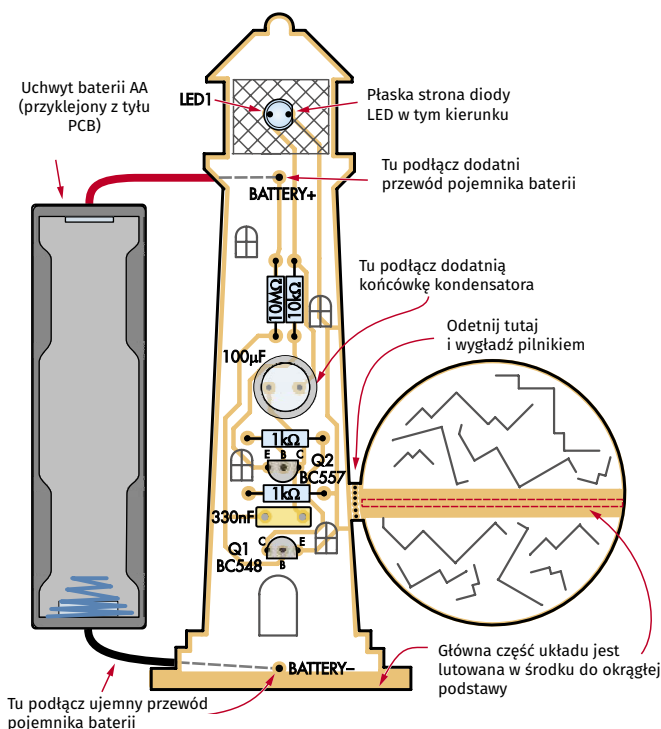
Latarnia „Nocny Opiekun” jest zbudowana na płytce drukowanej o kodzie 08110201 i wymiarach 64×91 mm. Przed rozpoczęciem pracy, odłóż lub odetnij okrągłą podstawę od boku latarni, a następnie spij lub zeszlifuj obie krawędzie na gładko. Dobrym pomysłem jest wykonanie nacięcia wzdłuż krawędzi podziału przed odłamianiem. W tym celu należy kilkakrotnie przejechać ostrym nożem wzdłuż linii przechodzącej przez małe otwory.

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

- 1 płytka drukowana, kod 08110201, 64×91 mm
- 1 tranzystor NPN: BC547, BC548 lub BC549 [Jaycar ZT2154 lub Altronics Z1042]
- 1 tranzystor PNP: BC557, BC558 lub BC559 [Jaycar ZT2164 lub Altronics Z1055]
- 1 biała dioda LED o wysokiej jasności 5 mm [Altronics Z0876E lub Jaycar ZD0190]
- 1 kondensator elektrolityczny 100 μF 16 V [Jaycar RE6130 lub Altronics R5123]
- 1 kondensator 330 nF MKT, ceramiczny lub MLC (kod 0.33, 330n lub 334)
- 1 uchwyt na ogniwa AA lub AAA montowany na płytce drukowanej [AA: Altronics S5029 lub Jaycar PH9203; AAA: Altronics S5051; Jaycar PH9261].
- Klej lub dwustronna taśma piankowa do przymocowania uchwytu ogniwa do tylnych części płytki głównej

Rezystory: (wszystkie 1/4 W, 1% lub 5%)

1 szt. 10 MΩ 1 szt. 10 kΩ 2 szt. 1 kΩ



Rysunek 3. Płytką drukowaną składa się z dwóch części, samej latarni i jej okrągłej podstawy, na której znajdują się zarysy skał! Należy je rozdzielić lub rozciąć przed zamontowaniem pokazanych elementów. Zamiast mocować uchwyt ogniwa za pomocą przewodów (jak pokazano tutaj), zalecamy zamontowanie uchwytu z tyłu płytki.

Odlóż podstawę na bok, a następnie zapoznaj się ze schematem montażowym płytki drukowanej (rysunek 3) i przewodnikiem budowy (rysunek 4), aby zobaczyć, gdzie należy umieścić które części. Wszystkie elementy, z wyjątkiem uchwytu baterii, są montowane na górnej stronie (strona z zarysem elementów i numerami części), a ich wyprowadzenia lutowane są po przeciwnej stronie. Uchwyt baterii jest montowany z tyłu i należy go podłączyć jako ostatni.

Zacznij od zamontowania czterech rezystorów, które można rozpoznać po kolorowych paskach, jak pokazano na rysunku. Rezystory 1% mają zwykle pięć pasków, podczas gdy rezystory 5% mają zwykle cztery. Pokazane są obie możliwości.

Wygnij końcówki każdego rezystora po kolei za pomocą szczypiec lub przyrządu do gięcia, tak aby pasowały do otworów w płytce. Włóż je, jeden po drugim, po kolei, rozsuwając lekko druty, aby utrzymać je na miejscu. Można je zamontować w dowolny sposób. Odwróć płytkę i przylutuj oba wyprowadzenia do pól lutowanych. Następnie odetnij wyprowadzenia na równi z lutem używając ostrych cążek.

Następnie zamontuj kondensator 330 nF. Może to być kondensator poliestrowy, MKT lub ceramiczny, a potem kondensator elektrolityczny. Przylutuj i przytnij wyprowadzenia w ten sam sposób. Upewnij się, że dłuższe wyprowadzenie kondensatora trafi w otwór oznaczony (+) na płytce drukowanej. Pasek na obudowie kondensatora powinien być naprzeciwko symbolu (+).

Teraz nadszedł czas na zamontowanie dwóch tranzystorów. Q1 to tranzystor NPN, natomiast Q2 to tranzystor PNP. Każdy tranzystor musi być zamontowany w odpowiednim miejscu. Z reguły nie są one wsuwane do płytki do końca, a raczej pozostawiane z wyprowadzeniami wystającymi na około 5–10 mm. Ta odległość nie jest krytyczna.

Prawdopodobnie pomocne będzie lekkie podgięcie trzech wyprowadzeń każdego tranzystora przed włożeniem ich do płytki, upewniając się, że płaska strona jest zorientowana jak na rysunku. Po włożeniu końcówek do płytki, rozłóż je jeszcze trochę po tej stronie, aby przytrzymać je

na miejscu przed odwróceniem płytki i przed lutowaniem. Przytnij przewody po zakończeniu lutowania.

Teraz zamontuj białą diodę LED na górze płytki. Ma ona lekko ściętą krawędź z jednej strony. Dioda powinna być włożona w taki sposób, aby pasowała do kształtu nadrukowanego na PCB dla diody. Dłuższe wyprowadzenie anodowe znajdzie się po przeciwnej stronie niż znacznik.

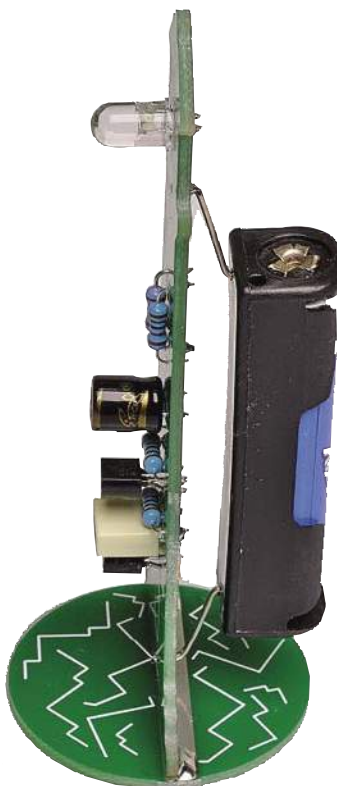
Dokładnie sprawdź, czy wszystkie elementy są prawidłowo rozmieszczone, czy wszystkie wyprowadzenia elementów zostały przylutowane i przycięte. Sprawdź również, czy nie ma odprysków lutu, które mogłyby spowodować zwarcie.

Następnie można zamontować uchwyt na baterie z tyłu płytki. Standardowy uchwyt na ogniwa AA jest na tyle duży, że koniec uchwytu na baterie pozwala na usadowienie latarni na krawędzi półki lub książki, jak pokazano to na zdjęciu. Bateria zapewnia idealną wagę do utrzymania latarni w pionie, przydatną w ciasnych zakamarkach sypialni lub biura.

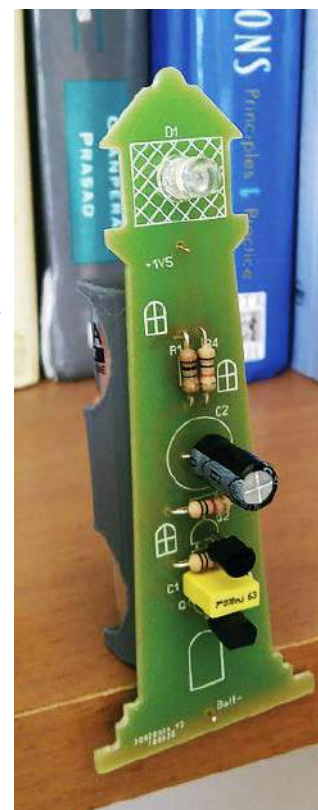
Przewody niektórych uchwytów baterii będą dokładnie pasować do otworów przewidzianych na płytce drukowanej. Przewód dodatni (+) powinien trafić do otworu znajdującego się najbliżej górnej części płytki PCB, obok diody LED. W przypadku innych uchwytów baterii może być konieczne lekkie wygięcie przewodów, aby je dopasować. Można użyć wąskich szczypiec, aby delikatnie wygiąć przewody w odpowiedni kształt, tak aby pasowały do PCB.



Aby połączyć obie płytki razem, najpierw „przyklej” je lutem, a następnie przeprowadź lut wzdłuż ścieżek z cynowanej miedzi na płytce. Będzie to mocne połączenie!



A oto widok z boku pokazujący dwie płytki zlutowane razem i uchwyt baterii na swoim miejscu. OK, trochę oszukiwaliśmy: stwierdziliśmy, że sztywny cynowany drut wystarczy, aby utrzymać pojemnik na miejscu bez kleju czy taśmy.



Nie musisz lutować głównej płytki do podstawy: ciężar uchwytu na baterie AA zapewni, że pozostanie ona na miejscu „wisząc” nad krawędzią półki z książkami.

Najlepiej jest odsunąć uchwyt baterii od tylnej strony płytki drukowanej o około 3 mm. Zapewnia to wystarczającą ilość miejsca, aby przylutować dwa łącza uchwytu baterii do odpowiednich pól na tylnej części PCB.

Mocowanie podstawy

Alternatywnie można dodać okrągłą podstawę PCB. Posiada ona „skalną” nakładkę, która dodaje całości efektowniejszego wyglądu i pozwala postawić „Nocnego Opiekuna” na płaskiej powierzchni. Ta część budowy może wymagać dodatkowej pomocy osoby dorosłej – dwie ręce do trzymania wszystkiego we właściwym miejscu, a dwie pozostałe do trzymania lutownicy i nakładania lutu.

Zacznij od przylutowania dwóch małych „kleksów” lutu na każdym końcu dolnej ocynowanej krawędzi płytki PCB latarni. Podobnie umieść je na ocynowanym pasku znajdującym się na górnej powierzchni okrągłej, bazowej płytki PCB. Główna płytki PCB powinna znajdować się mniej więcej centralnie i pionowo na górze podstawy.

Dotknij lutownicą do dwóch „kleksów” lutu, aby połączyć obie płytki.

W razie potrzeby powtórz tę czynność, przykładając lutownicę na krótko do każdego punktu łączenia, jednocześnie lekko regulując główną płytkę, aż do momentu, gdy główna płytki będzie dokładnie pionowa i wyśrodkowana na podstawie.

Następnie nałóż kolejne krople lutu za pomocą lutownicy wzdłuż połączenia, utrzymując obie płytki w ich ostatecznej pozycji.

Na koniec przeciągnij lutownicą po złączu, aby je wygładzić i nadać mu właściwy wygląd.

Użytkowanie

Czy zauważyłeś, że nie ma przełącznika zasilania? Układ pobiera tak mały prąd, że wyłącznik jest zbędny. Żywotność baterii podczas użytkowania jest podobna do okresu przydatności do użytku. Nowa węglowo-cynkowa bateria AA może zasilать układ przez ponad rok.

Miejmy nadzieję, że twarze młodych budowniczych będą świecić tak jasno jak „Nocny Opiekun” zaraz po włożeniu baterii, kiedy układ zacznie migać.

Zauważ, że zamiast baterii typu AA możesz użyć uchwytu i ogniwa typu AAA. W takim przypadku możesz oczekiwać, że nowe ogniwo wytrzyma coś koło sześciu miesięcy. Czas pracy baterii będzie się różnił w zależności od typu baterii (np. alkaliczna, wysokoenergetyczna itp.) oraz od jej stanu w momencie włożenia (nowa, lekko zużyta, prawie wyczerpana itp.).

Korzystanie z latarni

„Nocny Opiekun” jest przydatną, jasną lampką nocną dla dzieci. Należy jednak pamiętać, że migające światła mogą zakłócać sen, zwłaszcza jeśli są skierowane na twarz.

Ponadto, ze względu na jasność niektórych wysokowydajnych białych diod LED, „Opiekun” nie powinien być umieszczony w miejscu, w którym dioda będzie świecić bezpośrednio w młode i szczególnie wrażliwe oczy. Lepiej jest umieścić latarnię tak, aby światło LED świeciło lekko w górę lub pod kątem prostym, być może na sąsiednią ścianę. Takie układy i tak są generalnie bardziej efektywne do wykorzystania jako lampka nocna.

Starsi konstruktorzy mogą stwierdzić, podobnie jak ja, że „Nocny Opiekun” może być przydatny do lokalizowania przedmiotów w nocy, zarówno dla dzieci, jak i dla dorosłych.

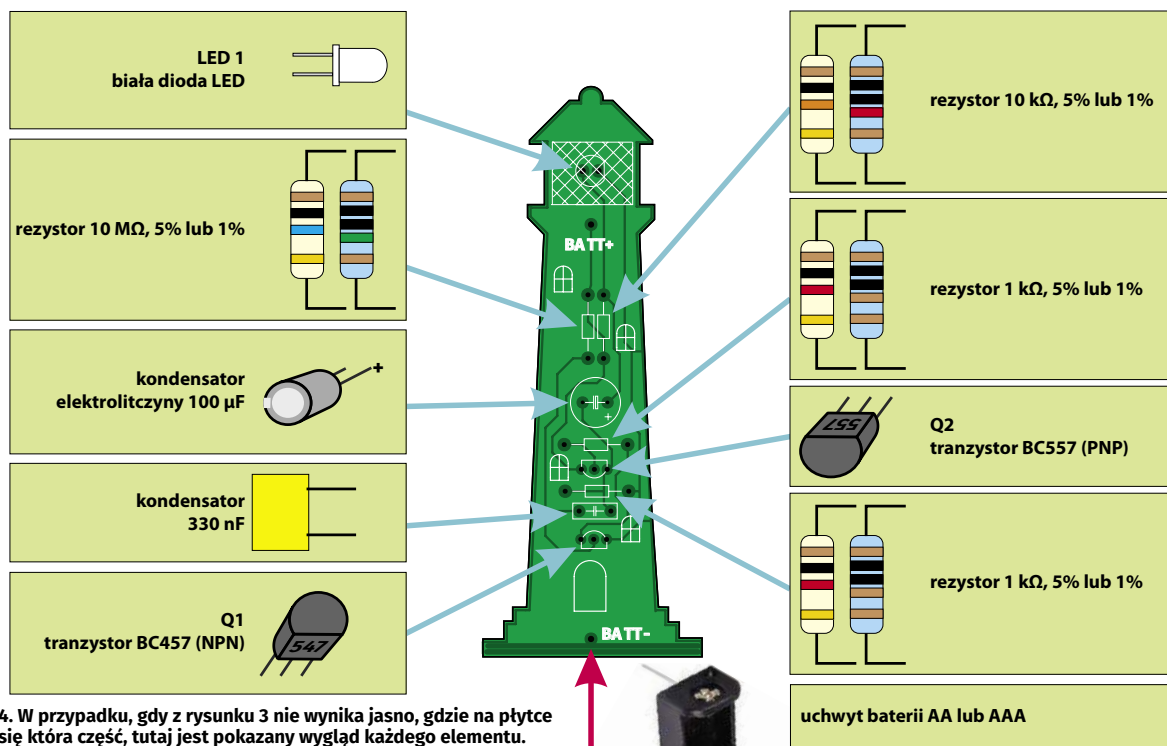
Odpowiednio zamontowany w pobliżu drzwi, włącznika światła lub umieszczony na półce, może pomóc wskazać drogę do miejsca lub wokół mebli w najciemniejszych nocach.

Zupełnie jak prawdziwa latarnia morska! ■

Andrew Woodfield

Adaptacja do wydania polskiego – Andrzej Nowicki

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



Rysunek 4. W przypadku, gdy z rysunku 3 nie wynika jasno, gdzie na płycie znajduje się która część, tutaj jest pokazany wygląd każdego elementu. Wystarczy podążać za strzałką, aby zobaczyć, gdzie się znajduje. Możesz również dopasować orientację części do rysunków; pięć elementów, których orientacja ma znaczenie to LED1, Q1, Q2, kondensator elektrolityczny (w kształcie puszeki) oraz uchwyt baterii. Reszta nie ma znaczenia, w którą stronę pójść.



Materiały dodatkowe dostępne są na stronie
Silicon Chip: <http://bit.ly/3GGmccS>
Materiały dodatkowe są również dostępne
na stronie edw.elportal.pl: <https://bit.ly/3GOzrlN>



W zeszłym miesiącu pokazaliśmy, jak działa i co potrafi nasze urządzenie. Teraz składamy wszystkie elementy razem i zaczynamy!

Programowalny termoregulator, część 2

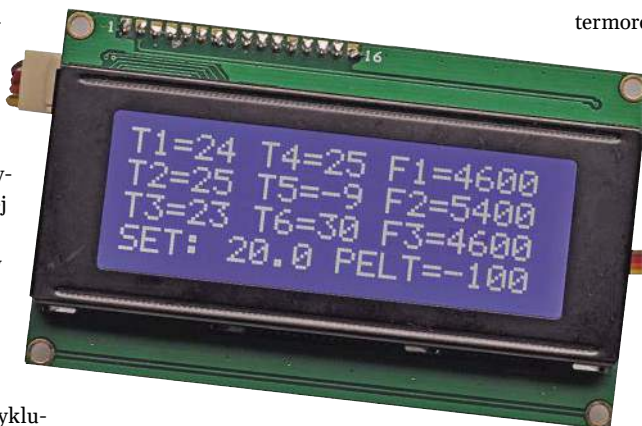
W numerze EdW z lutego 2023 r. przedstawiliśmy wszechstronne, sterowane modułem Arduino UNO urządzenie grzewcze/chłodzące. Wykorzystuje ono ogniwa Peltiera do podgrzewania lub schładzania wody w jednym lub więcej obiegów, a jego budowa jest w miarę łatwa (choć nieco pracochłonna). Można je używać do wielu zadań, w tym (ale z pewnością nie tylko!) do warzenia piwa, robienia sera i gotowania... a nawet wylęgania kurczaków (czy krokodyli)! Ten artykuł zawiera wszystkie instrukcje opisujące jak zbudować obie nakładki modułu Arduino, zaprogramować Arduino, zbudować obiegi wodne i dostosować je do swoich potrzeb.

Aby udowodnić, jak wiele możliwych zastosowań ma ten projekt, oto kolejne, które przyszło nam do głowy w zeszłym miesiącu: można by je użyć do zbudowania domowego inkubatora jaj, aby utrzymać jaja ptaków lub gadów w stałej temperaturze, tak aby się wykluły.

Często robi się to za pomocą lampy IR, ale to marnotrawstwo energii, a na dodatek ta metoda nie uwzględnia zmiennych warunków otoczenia.

Jajka kurze najlepiej trzymać w temperaturze 37,5°C do czasu wykucia; dotyczy to również jaj większości innych ptaków i gadów.

Poprzez ułożenie węży z rurki z wodą pod jajkami (najlepiej wykonać rurkę z dobrego przewodnika ciepła, np. miedzi) i umieszczenie wśród nich czujnika, można ustawić



Alfanumeryczny wyświetlacz LCD I²C umożliwia wyświetlanie wielu parametrów. Dostępne są temperatury odczytane ze wszystkich sześciu termistorów, a także prędkości wentylatorów, wartość zadana temperatury, tryb i poziomysterowania ogniwa Peltiera (temperatura zadana: 20,0°C; temperatura aktualna: 24°C; chłodzenie – pełna moc).

termoregulator tak, aby utrzymywał tę idealną temperaturę.

Termostat będzie zużywał tylko tyle energii ile potrzeba do utrzymania tej temperatury, a w upalny dzień (który może zabić embriony), może on nawet zapewnić trochę ochłody!

Jesteśmy pewni, że Czytelnicy pomyślą o innych zastosowaniach, które nie przyszły nam do głowy.

Ale wystarczy już tego wstępu; czas opisać jak poskładać i uruchomić kompletny termostat.

Budowa

Zacniemy od zbudowania obu nakładek, gdyż jest to warunek wstępny do uruchomienia całości. Chociaż, jeśli chcesz, możesz

przeprowadzić kilka podstawowych testów „obvodu wodnego” bez elementów sterujących.

Przed kontynuacją budowy możesz przystąpić do pracy wentylatory, pompy i ogniwa Peltiera zasilane bezpośrednio ze źródła napięcia 12 V, aby sprawdzić, czy wszystko działa..

Nakładka sterownika Peltiera

Nakładka sterownika Peltiera zbudowana jest zarówno z elementów do montażu przewlekane, jak i powierzchniowego; jej schemat montażowy pokazano na rysunku 7.

Żadna z części montowanych powierzchniowo nie jest zbyt trudna do przylutowania; najmniejsze elementy to kondensatory o symbolu 3216/1206, które, jak sama nazwa wskazuje, są stosunkowo duże – 3,2×1,6 mm.

Będą przydatne – jeśli nie obowiązkowe – do pracy z tymi podzespołami: pęseta, topnik lutowniczy i plecionka lutownicza. Zaczynaj od wlutowania kondensatorów. Przylegają one do dużych powierzchni miedzianych PCB, więc mogą wymagać znacznego podgrzania do prawidłowego przylutowania.

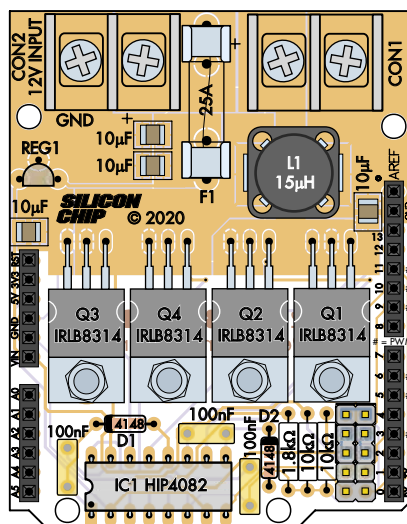
Nałóż niewielką ilość topnika na ich pola lutownicze, a następnie przylutuj jedno wyprowadzenie każdego kondensatora we właściwym miejscu. Jeśli kondensator przylega płasko do płytki, przylutuj drugie wyprowadzenie, w przeciwnym razie użyj pęsety i lutownicy, aby dopasować pierwszą końcówkę przed kontynuowaniem.

Drugą częścią do montażu powierzchniowego jest dławik. Oprócz tego, że łączy się z dużymi miedzianymi ścieżkami, ma również sporą pojemność cieplną; (jeśli możesz) w tym momencie możesz zwiększyć temperaturę lutownicy!

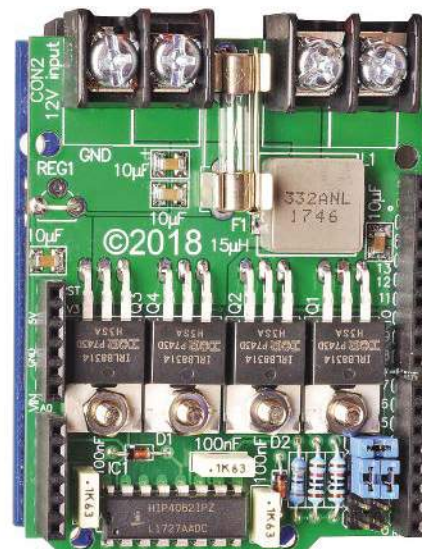
Podobnie jak w przypadku kondensatorów, nałóż topnik (tym razem sporo), a następnie przylutuj do płytki jedno wyprowadzenie. Kiedy element jest już we właściwym miejscu, przylutuj drugie wyprowadzenie. W tym momencie usuń nadmiar topnika przy użyciu dedykowanego środka do czyszczenia, np. alkoholu izopropylowego.

Następnie zamontuj oprawkę bezpiecznika, z założonym tymczasowo bezpiecznikiem. Zapewni to prawidłowe rozmieszczenie i orientację elementów. Po wlutowaniu oprawek bezpiecznik może pozostać na swoim miejscu.

Temperatura lutownicy może być teraz obniżona przy lutowaniu pozostałych części. Kontynuuj, montując diody D1 i D2 z paskami katodowymi zorientowanymi jak na rysunku, a następnie zamontuj trzy rezystory. Jeśli nie jesteś pewien, który z nich ma 1,8 kΩ, zmierz go za pomocą DMM. Następnie zamontuj IC1, upewniając się, że kropka/wcięcie przy końcówce 1 jest skierowana w lewo. Zalecamy przylutować IC1 bezpośrednio do płytki, bez używania podstawki.



Rysunek 7. Ten schemat montażowy i zdjęcie gotowej nakładki pokazują, gdzie należy zamontować elementy nakładki sterownika Peltiera. Widać pięć elementów SMD (cztery kondensatory i jeden dławik), ale wszystkie są dość duże. Pasta flux pomoże Ci je przylutować. REG1 nie jest potrzebny, jeśli do CON2 jest doprowadzone napięcie 12 V. W tym przypadku możesz zainstalować widoczną na fotografii zworkę dwóch dolnych pól lutowniczych.



Teraz wygnij wyprowadzenia MOSFET-ów Q1–Q4 tak, aby pasowały do obrysu na PCB i przymocuj każdy z nich do płytki za pomocą śruby M3 i nakrętki, przed przylutowaniem i przycięciem wyprowadzeń. Pamiętaj, aby dokręcić śrubę przed lutowaniem, ponieważ dokręcanie jej po lutowaniu może uszkodzić luty.

Wlutowaj przelotowo kondensatory, które są tego samego typu i nie są spolaryzowane. Upewnij się jednak, że wcisnąłeś je do końca przed lutowaniem, ponieważ nad tą płytką będzie ułożona kolejna.

Również REG1 wcisnij tak głęboko, jak tylko możesz, zanim go przylutujesz. Jak wspomniałem w zeszłym miesiącu, w zależności od tego, jak będziesz zasilal nakładki, możesz nie lutować REG1 lub połączyć zworką jego skrajne lewe i prawe pole lutownicze. Ale w większości przypadków, bezpieczniej jest zamontować REG1.

(Zdjęcie u góry po prawej pokazuje naszą nakładkę ze zworką w miejscu REG1).

Teraz można przylutować prostą listwę kołkową 5×2. Możesz użyć dwóch 5-kołkowych odcinków listwy jednorzędowej lutowanych obok siebie.

Następnie zamontuj CON1 i CON2. Ponieważ CON1 znajduje się nad gniazdem USB w UNO, a CON2 nad gniazdem DC, pamiętaj o tym, aby po przylutowaniu przyciąć ich wyprowadzenia jak najkrócej. Są to elementy o dużych stykach, umieszczone na miedzianych podstawkach, więc mogą wymagać lekkiego zwiększenia temperatury lutownicy.

Pozostają jeszcze najtrudniejsze do wlutowania elementy: cztery gniazda jednorzędowe różnej długości, żeńsko-męskie. Zalecamy,

po włożeniu gniazd do otworów, męskimi kołkami w dół, utworzenie „kanapki” z montowaną nakładką pomiędzy modułem Arduino na dole, a inną nakładką na górze, jeśli taką posiadasz. Pomoże to w wyrównaniu gniazd. Upewnij się, że wszystkie cztery gniazda płasko przylegają do PCB na każdym końcu i przyklej końce każdego gniazda na miejscu (np. odrobina kleju momentalnego). Może to być kłopotliwe, ponieważ przesunięcie jednego gniazda może spowodować przesunięcie pozostałych.

Zdejmij delikatnie moduł Arduino UNO z dołu i przylutuj od dołu skrajne kołki każdego gniazda. Jeszcze raz sprawdź, czy kołki pasują do gniazd modułu UNO. Jeśli tak, przylutuj dokładnie pozostałe szpilki i ew. odśwież lutowanie skrajnych kołków każdego gniazda.

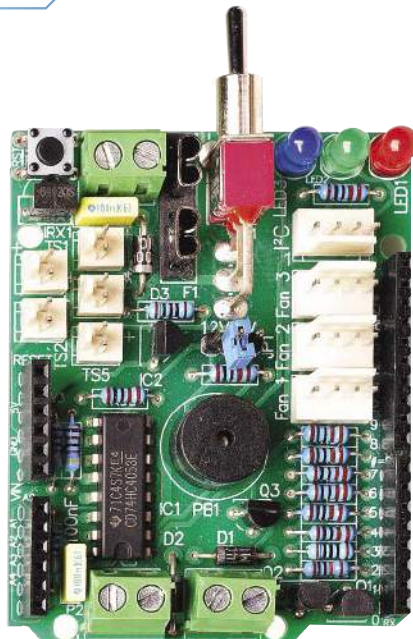
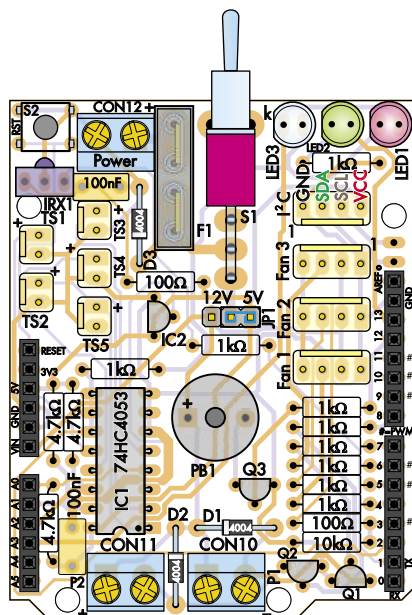
Od Red. EdW: z bardzo podobnym „kanapkowym” złożeniem modułu Arduino i nakładki możesz zapoznać się w styczniowym nr EdW z 2023 r., (w artykule o zdalnej stacji monitoringu), gdzie widać nieco dokładniej sposób połączenia płytek).

Zworki

Załącz trzy zworki, jak pokazano na rysunku 7. Nie powinieneś zmieniać ich położenia, chyba że radykalnie zmieniasz oprogramowanie na własne potrzeby. Zworka JP1 ustawia używanie końcówki D10 Arduino, zworka JP2 używanie końcówki D9, styki: LK3 zwarte a LK4 otwarte.

Budowa nakładki interfejsu

Użyj rysunku 8 jako schematu montażowego. Zaczynaj od rezystorów. Jak wspomniano wcześniej, najlepiej jest sprawdzić każdy element za pomocą DMM, aby zweryfikować



Rysunek 8. Budowa nakładki interfejsu jest prosta. Zalecamy zorientowanie spolaryzowanych wtyczek jak pokazano na schemacie i fotografii, ale tylko wtyczki interfejsu I²C i wentylatorów są krytyczne. S1, F1 i JP1 mogą być pominiete, jeśli napięcie 12 V będzie doprowadzane z nakładki sterownika Peltiera, a nie przez CON12. Możesz użyć wzduż krawędzi jednorzędowych gniazd żeńsko-męskich (do składania w stos), jak pokazano tutaj, lub zwykłych jednorzędowych listew kołkowych montowanych od spodu i lutowanych na górze.

ich wartość przed zamontowaniem. Jest to szczególnie ważne, ponieważ typy 100 Ω , 1 k Ω i 10 k Ω mają podobne oznaczenia kolorowymi paskami. Postępuj tak samo z trzema diodami, (są tego samego typu), ale upewnij się, że są zorientowane tak, jak pokazano na rysunku 8.

Następnie zamontuj zwirny przycisk chwilowy (S2). Wciśnij go aż do usłyszenia kliknięcia i płaskiego osadzenia na płytce drukowanej.

Są tylko dwa kondensatory, oba po 100 nF; typu MKT lub ceramiczne, po jednym na końcach płytki w pobliżu każdego układu scalonego. Przylutuj je w następnej kolejności. Potem zamontuj IC1; ponownie nie zalecamy używania podstawki. Upewnij się, że jest on zamontowany końcówką 1 w kierunku CON11. Przylutuj dwa wyprowadzenia i sprawdź czy IC1 przylega płasko do PCB; jeśli nie, podgrzej ponownie jedno z połączeń lutowanych i dopasuj je. Następnie przylutuj pozostałe wyprowadzenia.

Następnie zamontuj tranzystory Q1–Q3 oraz czujnik temperatury IC2, wszystkie w obudowach TO-92. Q3 jest innego typu niż Q1 i Q2, więc ich nie pomył. Dopasuj korpusy tranzystorów wg sitodruku konturów na płytce. Może być konieczne wygięcie ich wyprowadzeń, aby pasowały do otworów w płytce drukowanej.

Następnie zamontuj zaciski śrubowe CON10–CON12 i wszystkie kierunkowe wtyki 403-4 i 403-2. Jedynie orientacje wtyków złącza I²C i wentylatorów są krytyczne; upewnij się, że są one obrócone tak, jak pokazano

na rysunku 8, a także upewnij się, że bloki zacisków śrubowych są zamontowane tak, aby ich otwory na przewody były skierowane na zewnątrz płytki.

Konieczne sprawdź czy wszystkie te elementy przylegają płasko podstawami do PCB przed przylutowaniem wszystkich końcówek.

Zauważ, że pokazaliśmy wtyk interfejsu I²C wyświetlacza LCD obrócony względem wtyków wentylatorów; utrudnia to ich pomylenie, ponieważ uszkodzisz wyświetlacz, jeśli przypadkowo podłączysz go do wtyku wentylatora i włączysz zasilanie. Wszystkie wtyki 403-2 dwukołkowe powinny być tak samo zorientowane, aby później łatwiej można było zmienić kolejność podłączania czujników temperatury.

Następnie można wlutować trzy diody LED. Czerwona dioda znajduje się najbliżej krawędzi płytki, zielona w środku, a niebieska najbliżej przełącznika S1. Katody wszystkich trzech diod są skierowane w stronę tego przełącznika. W zależności od tego, jak zamierzasz używać gotowego projektu, możesz je podłączyć za pomocą swobodnych przewodów lub nawet zamontować w ich miejsce złącza typu wtyk-gniazdo i zamontować diody na panelu ew. obudowy.

Podobna uwaga dotyczy IRD1; ten również może być zamontowany poza płytką PCB, choć jeśli to zrobisz, najlepiej użyć jak najkrótszych przewodów, aby zapewnić niezawodnie działanie. Jeśli instalujesz odbiornik na płytce, upewnij się, że jego półkula soczewka jest skierowana w kierunku pokazanym na sitodruku PCB. Możesz wygiąć końcówki tak,

aby soczewka była skierowana do góry, ale musisz uważać, aby nie kolidowała z pobliskim dwukołkowym wtykiem.

Brzęczyk piezoelektryczny PB1 znajduje się w pobliżu środka płytki drukowanej. Przed zamontowaniem sprawdź jego polaryzację.

Jeśli planujesz zasilać gotowy zespół poprzez nakładkę sterownika Peltiera, możesz pominąć przełącznik S1, bezpiecznik F1 i zworę LK1/JP1. Nie zaszkodzi jednak zamontować je mimo wszystko. Jeśli je zamontujesz, postaraj się, aby wszystkie były płasko osadzone na PCB. Przełącznik i oprawka bezpiecznika są dość masywne i mogą wymagać więcej ciepła do lutowania niż mniejsze elementy.

Aby zakończyć budowę interfejsu, wystarczy zamontować złącza szyn Arduino. W większości przypadków będą wystarczające standardowe jednorzędowe męskie listwy kołkowe, chociaż w naszym prototypie zamontowaliśmy na wszelki wypadek gniazda żeńsko-męskie (do składania w stos, jak w nakładce sterownika Peltiera), jak widać na zdjęciach. Podobnie jak w przypadku gniazd nakładki sterownika Peltiera, powinieneś użyć innej płytki Arduino jako szablonu, aby upewnić się, że kołki są wlutowane równo i prosto.

Montaż pakietu

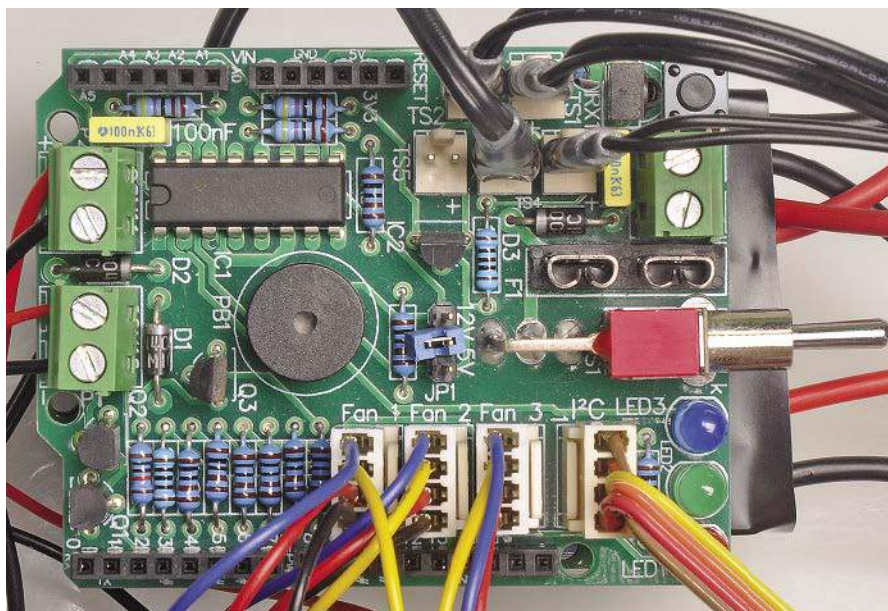
Nakładki są zaprojektowane tak, że nakładka sterownika Peltiera mieści się pomiędzy Arduino UNO na dole a nakładką interfejsu na górze. Nakładka interfejsu musi być na górze, aby można było uzyskać dostęp do jej licznych pionowych wtyków.

Najprostszym sposobem zasilania jest doprowadzenie go przez nakładkę sterownika Peltiera. Dostarczy ona napięcie 12 V do płytek powyżej i poniżej.

Należy jednak pamiętać, że jeśli zasilasz napięciem wyższym niż 1.5 V podłączonym do nakładki sterownika Peltiera, REG1 (który jest dość mały) nie może dostarczyć na tyle dużo prądu, aby uruchomić pompy lub wentylatory podłączone do nakładki interfejsu. W takim przypadku lepiej jest pominąć REG1 i dostarczyć napięcie 12 V bezpośrednio do CON12 na nakładce interfejsu.

Prąd doprowadzony do CON12 na nakładce interfejsu będzie również zasilał IC1 na nakładce sterownika Peltiera, ale nie pobiera on zbyt dużo prądu.

Podczas montażu pakietu, możesz znaleźć kilka miejsc, gdzie przewody lub końcówki dotykają elementów na płytce poniżej. Przytnij je tak krótko, jak tylko to możliwe; w przeciwnym razie oklej je samoprzylepną taśmą izolacyjną. Gniazdo USB modułu UNO powinno być też oklejone taśmą, aby chronić je przed zwarciami ze stykami zasilania nakładki sterownika Peltiera.



Nakładka interfejsu znajduje się na górze pakietu, ponieważ do jej pionowych wtyków są podłączone kable. Tak więc wysokość elementów na tej płytce nie jest krytyczna. Zauważ, że uchwyt bezpieczników jest pusty, ponieważ 12 V jest dostarczane przez VIN. Mogliśmy więc pominąć S1, F1 i LK1.

Jeśli to konieczne i musisz dołączyć kable zasilające do nakładki sterownika Peltiera, tymczasowo zdemonuj stos.

Przygotowanie wyświetlacza LCD

Możesz zakupić ekran LCD w SILICON CHIP ONLINE SHOP lub kupić części oddzielnie od Jaycar. Tak czy inaczej, będziesz musiał podłączyć adapter I²C do LCD. Ustaw odpowiednio styki 1 na module adaptera I²C i na płytce LCD i przylutuj jeden styk na miejscu. Przed przylutowaniem pozostałych styków sprawdź, czy obie płytki są równoległe, ale nie dotykają się. **Od Red. EdW: Odradzamy ten pomysł – bez problemu zakupisz alfanumeryczny wyświetlacz LCD 4×40 z podłączonym interfejsem I²C.**

Będziesz musiał także wykonać przewód łączący pomiędzy wtykiem I²C wyświetlacza LCD a wtykiem 403-4 I²C nakładki interfejsu. Do testowania naszego prototypu użyliśmy żeńsko-żeńskich przewodów zworkowych, ale były one dość krótkie. Być może znajdziesz dłuższe u lokalnego sprzedawcy.

Najlepszą opcją dla stałej konfiguracji jest wykonanie kabla z czterostykowymi kierunkowymi gniazdami 402-4 na każdym końcu. Zapoznaj się z Rysunkiem 8 co do wymaganych połączeń i sprawdź opisy kołków na płytce adaptera LCD I²C. Ponieważ kołki są w innej kolejności (GND, SDA, SCL, VCC na naszej płytce i GND, VCC, SDA, SCL na LCD), niektóre przewody będą musiały się krzyżować.

Wtyk na nakładce interfejsu jest kierunkowy, natomiast wtyk dostarczony z adapterem LCD nie jest. Warto zastąpić wtyk na LCD

typem 403-4, aby nie można było wykonać odwrotnego połączenia.

Zacznij składać wszystko razem

Na tym etapie, jeśli jeszcze tego nie zrobiłeś, musisz zdecydować o dokładnej konfiguracji wymaganej dla Twojego zastosowania (zastosowań). Najprawdopodobniej będziesz chciał zbudować coś, co wygląda jak na jednym z rysunków 3...6 w artykule w poprzednim miesiącu.

Krytyczne są połączenia wodne. W idealnym przypadku powinny być one jak najkrótsze, chociaż jeśli chcesz zaoszczędzić na kolanach,

rurki mogą być prowadzone po łagodnych łukach zamiast pod kątem prostym.

Pamiętaj, że możesz wybrać bloki wodne z przyłączami na tych samych lub przeciwnych końcach. Nie testowaliśmy, które z nich są sprawniejsze; podejrzewamy, że różnica jest nieznaczna.

Kolejną kwestią, którą należy wziąć pod uwagę podczas projektowania systemu jest to, że powietrze z chłodnicy lub radiatora nie powinno być wydychywane na inne części zespołu, ponieważ zmniejszy to jego ogólną efektywność.

W naszym przypadku upewniliśmy się również, że dwie chłodnice (jedna istniejąca na wycinarkę laserową i jedna na naszym nowym układzie wspomagania) wydychują powietrze w różnych kierunkach. Można to osiągnąć poprzez umieszczenie ich naprzeciw siebie, tak aby wciągały świeże powietrze z tego samego kierunku i wydychwały je równoległe w przeciwną stronę.

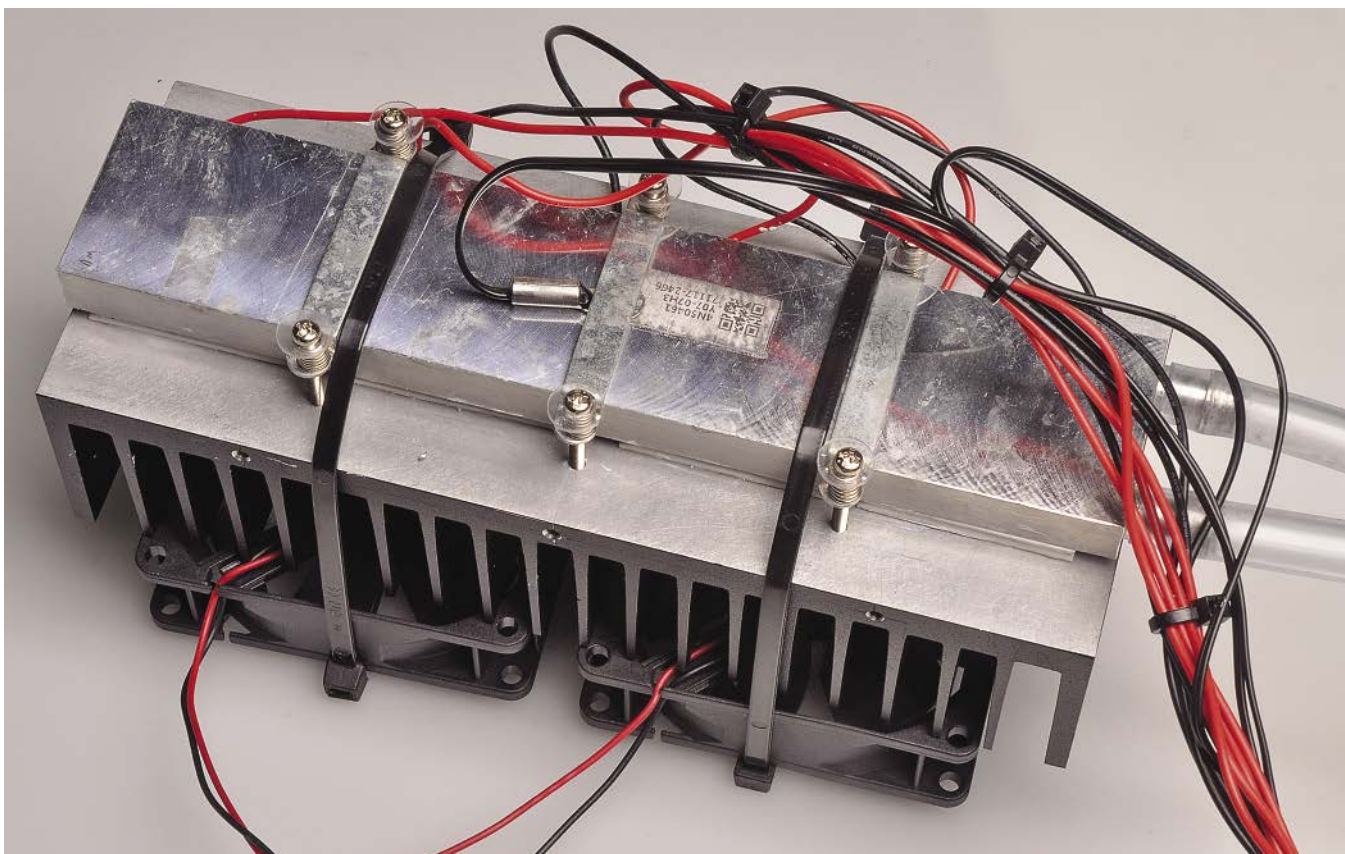
Zwróć też uwagę na nasze komentarze z zeszłego miesiąca dotyczące izolacji. W przypadku produkcji sera lub piwa w kąpieli wodnej o temperaturze zbliżonej do otoczenia, zapotrzebowanie na moc ogniw Peltiera nie będzie zbyt wysokie, ale gotowanie „sous-vide” w temperaturze około 60°C lub wyższej będzie wymagało porządnej izolacji, aby móc osiągnąć bardziej ekstremalne poziomy temperatury. Jeśli masz problemy z osiągnięciem docelowej temperatury, może pomóc lepsza izolacja.

Montaż urządzenia Peltiera

Nasz zestaw został wyposażony w sprzęt do montażu bloków wodnych po obu stronach ogniw Peltiera. Zawiera on kilka obejm, które są zaciskane przez śruby M4. Małe sprężyny



Do podłączenia zasilacza ATX użyliśmy pary złącz Molex (w tym przypadku Jaycar Cat PP0744). Złącza te mają obciążalność około 10 A każde, więc do naszego zastosowania potrzebne są dwa.



Minimalny układ hydrauliczny (odpowiadający Rysunkowi 5 z pierwszej części) wykorzystuje radiator żeberkowy uzupełniony o wentylatory do odprowadzania ciepła z ogniw Peltiera i bloku wodnego. Jest to ten sam układ, który jest używany w wielu wzmacniaczach i układach zasilania.

i podkładki zapewniają równomierną i nie za dużą siłę nacisku.

Obejmy te są przeznaczone do mocowania dwóch bloków wodnych, po jednym z każdej strony rzędu ogniw Peltiera. Jeśli używasz jednego bloku wodnego i radiatora, sposób montażu znajdziesz poniżej.

Zacznij od zmontowania bloków wodnych i ogniw Peltiera. Może to być trudne, ponieważ kilka elementów musi się połączyć w tym samym czasie i wszystkie będą pokryte dość śliską pastą termoprzewodzącą.

Wyczyść bloki wodne i ogniwa Peltiera alkoholem izopropylowym lub podobnym, aby usunąć wszelkie zanieczyszczenia i pozostałości po montażu. Pozostaw do wyschnięcia.

Ułóż rząd obejm na stole warsztatowym, ze śrubami i podkładkami zamontowanymi w otworach; główki powinny być skierowane w dół. Połóż jeden blok wodny na obejmach i nałóż minimalną ilość pasty termicznej na jedną stronę każdego ogniw Peltiera, rozprowadzając ją jak najcieńiej.

Optymalnie warstwa pasty termicznej powinna być jak najcieńsza, ale musi pokrywać całą powierzchnię stykających się części.

Upewnij się, że ogniwa Peltiera są zorientowane w ten sam sposób, przyciśnij je do bloku wodnego, rozprowadzając pastę termiczną pomiędzy ogniwem a blokiem wodnym.

Jeśli (na przykład) wszystkie czerwone przewody znajdują się po lewej stronie, a wszystkie czarne po prawej, to ogniwa są prawidłowo zorientowane.

Rozprowadź teraz pastę termiczną na górnej części ogniw Peltiera, a następnie połóż na nich drugi blok wodny, upewniając się, że złączki do rurek są zorientowane zgodnie z wymaganiami.

Umieść pozostałe obejmy na miejscu, a następnie załóż sprężyny, podkładki i nakrętki. Dokręć nakrętki, aż sprężyny zaczną się ścisnąć.

Upewnij się, że ogniwa Peltiera są równomiernie rozmieszczone; przynajmniej nie powinny wystawać poza bloki wodne. Następnie możesz dokręcić nakrętki, upewniając się, że sprężyny nie są ściśnięte do samego końca.

Użycie radiatora zamiast bloku wodnego

Aby sprawdzić, czy możemy obejść się bez chłodnicy/bloku wodnego, użyliśmy radiatora znacznie szerszego niż ogniwa Peltiera (40 mm). Dlatego nie mogliśmy użyć po obu stronach obejm, aby skompletować cały zespół razem. Jeśli posiadasz radiator o szerokości 40 mm, może to być możliwe, ale prawdopodobnie będziesz musiał przyciąć większy radiator, aby uzyskać odpowiedni rozmiar.

Zalecamy użycie większego radiatora, ponieważ pozwoli to na użycie większych wentylatorów, co zapewni bardziej efektywny transfer ciepła do powietrza.

Zakładając, że radiator ma znacznie więcej niż 40 mm szerokości, będziesz musiał wytrasować i wywiercić otwory na powierzchni radiatora, aby zamontować ogniwa Peltiera.

Rozłóż ogniwa Peltiera i blok wodny na radiatorze, aby określić, gdzie powinny znajdować się otwory i zaznacz je, w miarę możliwości w miejscach pomiędzy żebrami radiatora (pozwoli to na przewiercenie otworów na wylot).

Jeśli nie posiadasz gwintownika, a możesz wywiercić otwory w przestrzeniach między żebrami, zamiast gwintować, możesz przewiercić się na wylot i użyć długich śrub mocowanych przez nakrętki włożone pomiędzy żebrami radiatora. Z doświadczenia wiemy, że to działa, ale wykonanie tego jest bardzo kłopotliwe.

W przypadku gwintowania należy wywiercić otwory o średnicy określonej dla danego gwintu. Wymagane otwory są zazwyczaj nieco mniejsze niż rozmiar gwintu. Wiele gwintowników jest dostarczanych z wiertłami o odpowiednim rozmiarze.

Po wywierceniu otworów należy je starannie nagwintować. Nie spiesz się z tym i wykręć lekko gwintownik, jeśli się zacina; zwykle wystarcza to do usunięcia opiłków. Należy użyć

środka smarującego, który również pomoże; my użyliśmy z powodzeniem WD-40 lub oleju 3 w 1, chociaż mówi się, że spirytus lub nafta jest również idealna do gwintowania aluminium.

Oczyszczyć radiator z pozostałości i zeszlifuj ranty wokół gwintowanych otworów. Ponieważ obejmy mają spory prześwit od ogniw Peltiera, nie jest krytyczne, aby miejsce to było idealnie płaskie.

Wyczyścić blok wodny i ogniwa Peltiera alkoholem izopropylowym lub podobnym, aby usunąć wszelkie opilki i zabrudzenia i pozostaw do wyschnięcia.

Nałóż bardzo cienką warstwę pasty termicznej na obie strony każdego ogniwa Peltiera i umieść je na radiatorze w odpowiednim miejscu. Dopasowanie ich nie jest problemem, ale może wyglądać nieładnie, jeśli pasta termiczna dostanie się w niepożądane miejsca.

Upewnij się, że wszystkie ogniwa Peltiera są skierowane w tę samą stronę. Oprócz kolorowych przewodów, wiele z nich ma oznaczenia identyfikacyjne tylko po jednej stronie.

Położ blok wodny na górze na ogniwach, a następnie oprzyj na nim obejmy. W każdym otworze umieść najpierw sprężynę, a następnie podkładkę, i wkręć śrubę przez podkładkę, sprężynę i obejmę w radiator.

Po wkręceniu wszystkich śrub sprawdź, czy wszystko jest na miejscu, i dokręć śruby tak, aby sprężyny zostały ściśnięte, ale nie do końca.

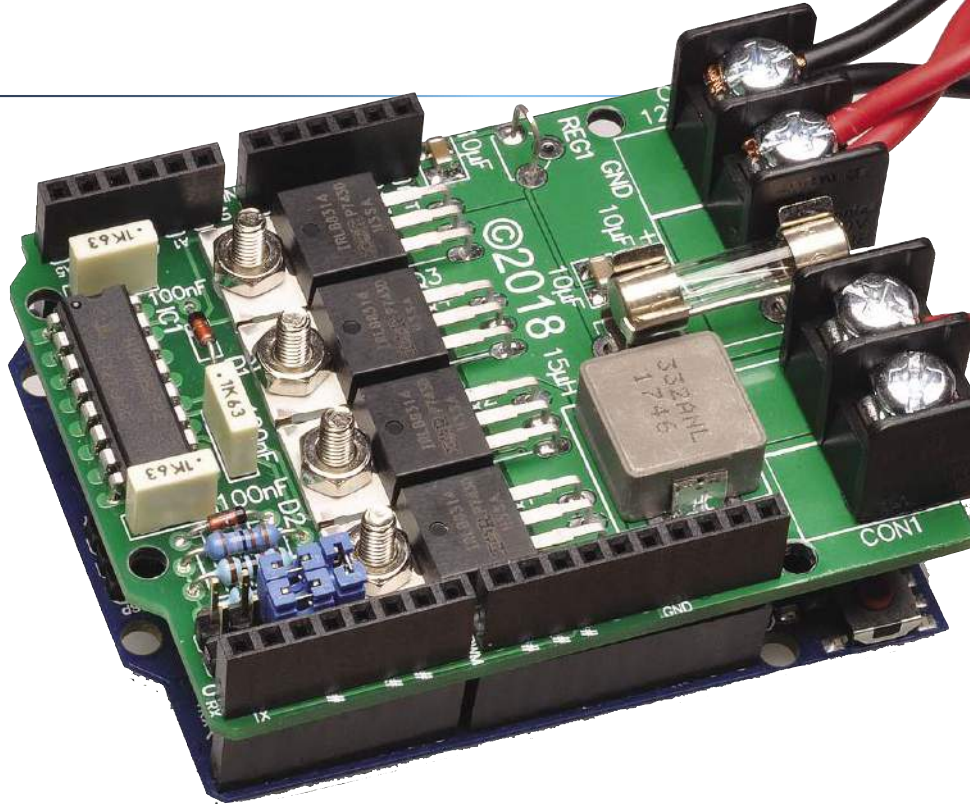
Do naszych testów zamontowaliśmy wentylatory za pomocą opasek kablowych przeciągniętych wokół całego zespołu. Twój radiator może być zaprojektowany tak, aby śruby były wkręcone bezpośrednio pomiędzy żeberka, w takim przypadku będzie to całkiem dobre rozwiązanie.

Inną opcją jest wywiercenie małych otworów przez żeberka w pobliżu ich końców. Następnie można przewlec opaski kablowe przez te otwory i otwory montażowe wentylatora. W każdym przypadku należy upewnić się, że strumień powietrza z wentylatora jest kierowany w stronę radiatora.

Pompy

Strona wejściowa (ssąca) pomp zatapialnych, które podaliśmy w spisie elementów, musi znajdować się całkowicie pod powierzchnią wody, ponieważ nie są one samozasysające. Zastosowanie pomp zatapialnych oznacza, że nie trzeba wycinać otworu w boku zbiornika wodnego, co pozwala uniknąć wycieków.

W przypadku naszej wycinarki laserowej, umieściliśmy pompę w pobliżu górnej części zbiornika; intencją było to, że w przypadku wycieku w obwodzie chłodzenia Peltiera, tylko niewielka część wody chłodzącej laser zostanie utracona. W zbiorniku pozostanie dosyć wody,



To zbliżenie na nakładkę sterownika Peltiera daje lepszy widok na zworki konfiguracyjne, a także pokazuje, jak nisko umieszczone są wszystkie części, aby uniknąć zwarcia z nakładką interfejsu zamontowaną powyżej.

aby pompa lasera zapewniła chłodzenie dzięki własnej chłodnicy z wentylatorami.

Pompa ogniwa Peltiera może pracować na sucho, ale to lepsze niż awaria lasera.

Udało nam się to osiągnąć poprzez wycięcie otworu w pokrywie, który jest mocnym połączeniem ciernym dla rurki wodnej. Jeśli wąż jest luźny, można użyć kilku opasek kablowych, aby ograniczyć jego ruch w pionie.

Stwierdziliśmy, że jeśli umieścimy pompę zbyt blisko powierzchni, powstanie wir, który pozwoli na zassanie powietrza. Rozwiązaniem jest obniżenie wlotu, co sprawi, że prawdopodobieństwo powstania wiru będzie mniejsze (ale w przypadku wycieku pompa ogniwa Peltiera wypompuje do ścieku więcej wody).

Ponieważ w tym naczyniu nasza pompa spoczywała na pompie lasera, nie mogliśmy jej obniżyć, więc przymocowaliśmy mały kawałek węża i kolanko skierowane w dół, aby obniżyć punkt ssania.

Inną opcją jest po prostu zwiększenie poziomu wody, jeśli jest na to miejsce. Może się okazać, że po uruchomieniu pomp poziom spada z powodu przemieszczania się wody do wężyków i trzeba będzie dolać wody.

Ponieważ woda przechodzi przez urządzenia takie jak blok wodny i chłodnica, powinna wchodzić od dołu i wychodzić od góry.

Ma to na celu zapewnienie, że wszelkie pęcherzyki powietrza mogą się podnieść i być usunięte. Wszelkie puste przestrzenie, w których zebrało się powietrze, nie będą przyczyniać się do wymiany ciepła, więc należy je zminimalizować.

Obieg wody powinien wracać do tego samego naczynia, aby zamknąć pętlę. Wycięliśmy drugi otwór w pokrywie, aby przymocować rurę powrotną na miejscu. Można ją również zablokować za pomocą opasek kablowych (lub uszczelnacza silikonowego).

Umieść powrót nieco powyżej poziomu wody. Pozwoli to na dostrzeżenie przepływu zwrotnego przy jednoczesnym zminimalizowaniu ilości powietrza. Powietrze nie jest dobrym przewodnikiem ciepła i należy unikać powietrza w przewodach wodnych.

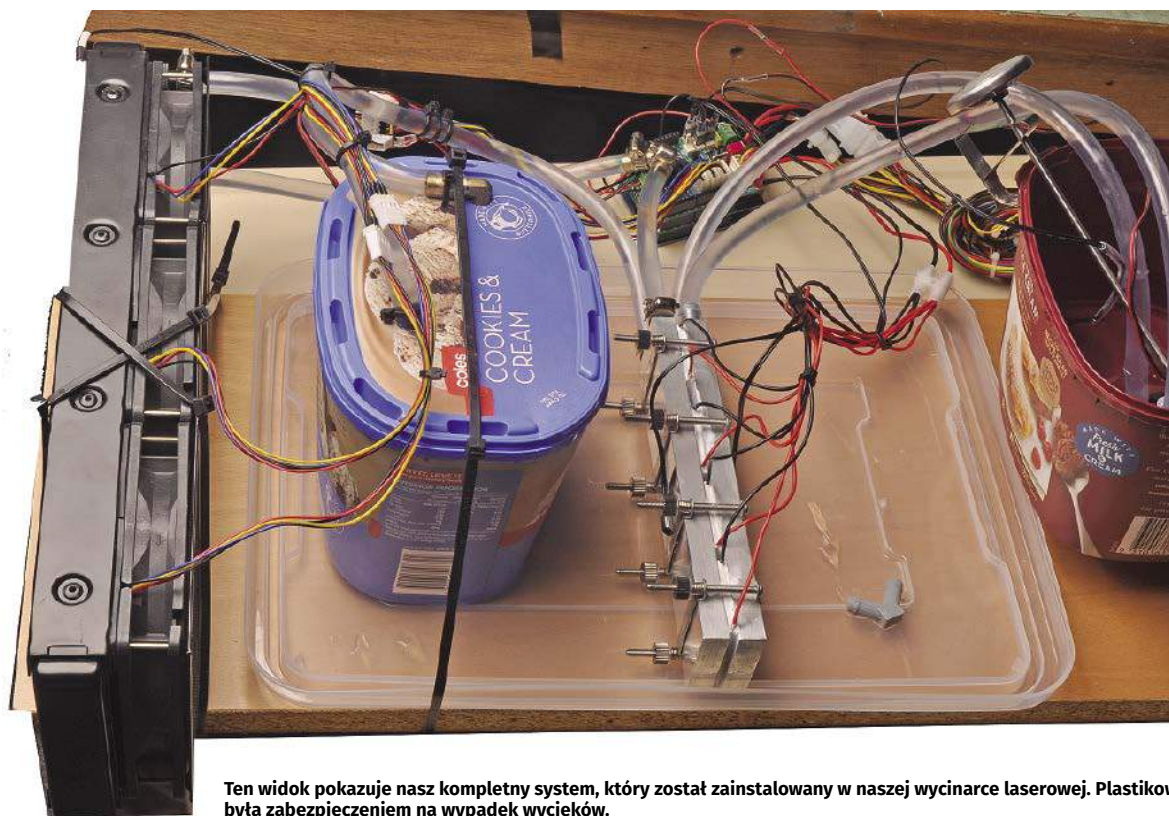
Jeśli to możliwe, należy umieścić powrót jak najdalej od wlotu pompy. Dzięki temu woda może się swobodnie mieszać w zbiorniku i nabierać jednolitej temperatury.

Po wykonaniu obiegu wody można przetestować pompę podłączając ją do zasilania 12 V. Powrotem powinien być stały, ciągły strumień wody, wskazujący, że występuje prawidłowy przepływ.

Sprawdź, czy nie ma wycieków i czy nie ma powietrza uwięzionego w rurkach. W razie potrzeby uzupełnij wodę. Jeśli nie ma przepływu, należy sprawdzić biegunowość zasilania pompy i kierunek przepływu. Używane przez nas pompy są ciche, ale słyszalne.

Przy działających pompach można również spróbować zasilić wentylatory i ogniwa Peltiera, aby zobaczyć, jaką wydajność może osiągnąć system. Należy pamiętać, że bez żadnej regulacji woda może się dość mocno nagrzewać.

Kiedy wszystko działa, zamontuj elementy na miejscu, aby nie przemieszczały się.



Ten widok pokazuje nasz kompletny system, który został zainstalowany w naszej wycinarce laserowej. Plastikowa tacka była zabezpieczeniem na wypadek wycieków.

Znaleźliśmy zapasową płytę półki meblowej, na której można wszystko zamontować.

Termistory

Termistory 10 kΩ, których używamy, zostały zahermetyzowane w małych pierścieniowych końcówkach do montażu śrub.

Miały również dołączony rozsądnej długości odcinek kabla, więc wszystko co musieliśmy zrobić to zakończyć każdy przewód termistora kierunkowym gniazdem 402-2 pasującym do nakładki interfejsu.

Termistory nie są spolaryzowane, więc nie ma znaczenia, który przewód idzie do którego styku.

Jeśli jednak chcesz umieścić czujnik w brzeczce piwnej (jak na naszym schemacie), nie zalecamy ich używania.

Zamiast tego należy użyć wersji, która posiada hermetyczną gilzę ze stali nierdzewnej dopuszczanej do kontaktu z żywnością. Są one dostępne, ale kosztują nieco więcej. Można mieszać i dopasowywać typy termistorów, o ile wszystkie mają tę samą wartość nominalną i podobne krzywe rezystancji (sprawdź określoną wartość Beta).

Nie byliśmy pewni, czy termistory, które dostaliśmy, są wodoodporne, więc na te, które miały być zanurzone w wodzie, nasunęliśmy sporej długości rurkę termokurczliwą, sięgającą poza termistor.

Następnie mocno zacisnęliśmy podgrzany wolny koniec w szczypcach, uszczelniając go,

choć wstrzyknięcie odrobiny silikonu do otwartego końca przed zaciśnięciem spowodowałoby bardziej niezawodne uszczelnienie.

Inną opcją jest zmontowanie ich od podstaw, używając zalewanych termistorów bez obudowy, drutu i gniazdek.

Nasze oprogramowanie zostało napisane do pracy z termistorami 10 kΩ lub 100 kΩ; pamiętaj, aby sprawdzić kod przed kompilacją, aby upewnić się, że oczekuje on wartości, których użyłeś.

Preferujemy typy 10 kΩ, ponieważ są one mniej podatne na wpływ EMI lub innych pól błędzących.

Montaż termistorów

Mała końcówka pierścieniowa na termistorach, których użyliśmy, sprawiła, że ich montaż był prosty.

Mimo, że nie skorzystaliśmy z opcji pomiaru temperatury chłodziarki, do zamocowania na niej termistorów wystarczyłyby zwykłe otwory gwintowane i wkręt.

W przypadku radiatorów, wykorzystano istniejącą śrubę montażową, aby przewlec ją przez otwór montażowy termistora i w ten sposób go przymocować.

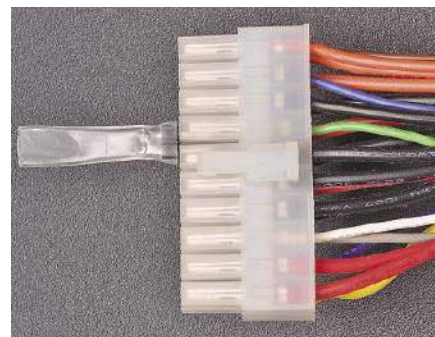
Jak wspomniano powyżej, termistor używany w wodzie obiegowej musi być dokładnie zabezpieczony przed wodą. Należy go również zamontować tak, aby nie wpadał do środka ponad uszczelnioną część, jeśli nie jest ona w pełni hermetyczna.

Jeśli nie trzeba go wyjmować, można wykonać parę małych otworów w boku zbiornika (powyżej linii wody!), aby przewlec opaskę zaciskową wokół przewodu termistora.

Przymocowanie termistorów do bloków wodnych (a więc w pobliżu ogniw Peltiera) było dość proste. Po prostu poluzowaliśmy jedną z obejm mocujących i wsunęliśmy płaski koniec termistora pod obejmę przed dokręceniem. Można użyć małej ilości pasty termoprzewodzącej.

Zasilanie

Do zasilania naszego termoregulatora użyliśmy zapasowego zasilacza ATX od komputera osobistego.



W zasilaczach ATX w celu uruchomienia trzeba zewrzeć zielony przewód do masy (dowolny czarny przewód). Zrobiliśmy prostą zworę ze złącza 2-szpilkowego i odcinka rurki termokurczliwej; teraz zasilacz włącza się, gdy otrzymuje zasilanie 230 V.

Jest to atrakcyjna opcja, jeśli masz dostępny wolny zasilacz. Dosłownie za grosze kupisz sprawny używany. Nawet jeśli musisz kupić nowy, są one stosunkowo niedrogie i mogą być całkiem wydajne.

Alternatywą jest jeden z wielu istniejących zasilaczy blokowych (typu open-frame) do montażu w wieżach/panelach. Altronics M8692 jest takim właśnie urządzeniem. Aby skorzystać z tego urządzenia, trzeba będzie wykonać pewne okablowanie sieciowe; przewody sieciowe są odsonowane, ale chronione pokrywką.

W założeniu tego typu zasilacz powinien być zainstalowany w obudowie i uważamy, że jest to rozsądne, niezależnie od rodzaju zasilacza, ponieważ pomoże umieścić elektronikę z dala od wody. Jeśli obudowa jest metalowa, należy pamiętać o jej odpowiednim uziemieniu.

Okablowanie 12 V potrzebne do tego rodzaju zasilania jest proste i wymaga jedynie podwójnego kabla 30 A (najlepiej czerwono-czarnego) zakończonych przewodami do zaciśnięcia na każdym końcu.

Zasilacze ATX wymagają nieco więcej pracy po stronie 12 V, ale wymagają jedynie podłączenia przewodu typu IEC do zasilania sieciowego.

Zazwyczaj istnieje wiele przewodów 12 V (żółty) i GND (czarny); będziesz musiał użyć kilku najgrubszych z nich, aby być pewnym, że możesz pobrać prąd o wystarczającym natężeniu.

Zasilacze ATX mają również przewód sygnału zasilania, który musi być zwarty do masy, aby zasilacz wystartował. Ten przewód jest zwykle koloru zielonego; my po prostu użyliśmy zworki, aby zewrzeć go z sąsiednim przewodem masy. Zobacz zdjęcia, które pokazują jak podłączyliśmy nasz zasilacz.

Jeśli jesteś pewien, że w przyszłości nie będziesz potrzebował zasilacza do pracy w komputerze, to kilka żółtych przewodów (dodatknie 12 V) i czarnych (masa) można połączyć razem i spiąć w jedną parę przewodów wysokoprądowych.

Niezależnie od źródła zasilania, podłącz je do zacisków wejściowych 12 V na nakładce sterownika Peltiera. Zacisk dodatni to ten, który znajduje się najbliżej bezpiecznika.

Podłączenie przewodów

Być może będziesz musiał rozebrać pakiet Arduino, aby podłączyć ogniwa Peltiera do ich nakładki sterownika. Orientacja, z jaką ogniwa Peltiera są podłączone, określi polaryzację napięcia potrzebną do ogrzewania lub chłodzenia, ale łatwo jest zmienić oprogramowanie, jeśli polaryzacja jest odwrócona, więc nie przejmuj się tym zbyt. Upewnij się tylko, że wszystkie ogniwa są podłączone z tą samą polaryzacją.

Dla ułatwienia użyliśmy małego kawałka listwy zaciskowej (8 zacisków) do rozdzielania połączeń; pozwala nam to również na poprowadzenie krótkich przewodów ogniwa Peltiera daleko od nakładki sterownika.

Zamontuj UNO poniżej, a osłonę nakładki interfejsu powyżej. Podłącz wentylatory, kabel I²C wyświetlacza LCD i termistory. Zobacz w tabeli 1, który termistor powinien być podłączony do którego wtyku. W razie potrzeby przypisanie czujników może być również zmienione w oprogramowaniu.

Pompę (pompy) podłącz do dwóch zacisków śrubowych w pobliżu IC2. Sprawdź, czy biegunowość jest prawidłowa, ponieważ pompy nie będą działać prawidłowo, jeśli będą obracać się do tyłu.

Jeśli masz osobne zasilanie 12 V dla nakładki interfejsu, podłącz je teraz. Potrzebny jest tylko dość mały bezpiecznik (wystarczy 3 A, typu samochodowego), chyba że masz bardzo duże wentylatory i pompy.

Oprogramowanie sterujące

Oprogramowanie, które napisaliśmy jest w zasadzie podstawowe, ale zapewnia większość lub wszystkie funkcje niezbędne do różnych zadań. Mierzy temperaturę ze wszystkich sześciu czujników, ale wykorzystuje dane tylko z trzech do podejmowania decyzji. Pozostałe temperatury są wyświetlane, ale nie są wykorzystywane przez oprogramowanie sterujące.

Aby zaprogramować płytke UNO, musisz zainstalować Arduino Integrated Development Environment (IDE), które zawiera również wszystko, co jest potrzebne do modyfikacji oprogramowania, jeśli zdecydujesz się to zrobić.

Użyliśmy IDE w wersji 1.8.5 i sugerujemy, abyś zrobił to samo, aby uniknąć problemów, które mogą pojawić się z powodu zmian pomiędzy wersjami.

Jak w przypadku wielu zaawansowanych projektów Arduino, potrzebne są pewne zewnętrzne biblioteki. Mogą one wydawać się skomplikowane, ale korzystanie z nich jest łatwiejsze niż konieczność pisania własnych funkcji interfejsu. Wszystkie one znajdują się w pakiecie do pobrania, wraz z samym „szkieletem” (kodem programu) Arduino.

Biblioteka I2CLCD jest zaadaptowana z innej biblioteki typu „open-source”. Dodaliśmy możliwość automatycznego wykrywania adresu I²C wyświetlacza LCD. Najprostszym sposobem dodania tej biblioteki jest skopiowanie folderu „I2CLCD” z archiwum .ZIP do folderu „Libraries” (w systemie Windows jest to folder „Documents”, w podkatalogu o nazwie „Arduino”).

Możesz również skopiować pozostałe trzy dostarczone biblioteki, ponieważ wiadomo, że wersje, które dołączyliśmy, działają.

Tabela 1. Podłączenia termistorów. Pokazane są tutaj podłączenia wykonane w naszym prototypie. Tylko pierwsze trzy są krytyczne dla pracy oprogramowania, aby móc sterować ogniwami Peltiera.

	Lokalizacja termistorów
TS1	Regulowana temperatura
TS2	Na bloku wodnym Peltiera, pętla TS1
TS3	Na bloku wodnym Peltiera, pętla przeciwna do TS1 i TS2
TS4	Na chłodnicy/radiatorze, ta sama pętla co TS3
TS5	Zapasowy (obecnie nieużywany)
TS6	Czujnik temperatury płytki interfejsu

Te trzy biblioteki można również zainstalować, znajdując je po nazwie w „Library Manager”. Aby to zrobić, wyszukaj „OneWire”, „DallasTemperature” i „Irrremote” i zainstaluj każdą po kolei. Jeśli masz już foldery o jednej z tych nazw, możesz mieć już zainstalowaną tę bibliotekę, więc prawdopodobnie nie chcesz jej nadpisywać, chyba że stwierdzisz, że nasz szkic nie działa.

Jeśli instalujesz biblioteki poprzez skopiowanie plików, być może będziesz musiał zamknąć i ponownie otworzyć Arduino IDE, aby je wykryło.

Przygotowanie pliku źródłowego

Nie będziemy się tutaj zbyt zagłębiać w szczegóły działania szkicu, ponieważ możesz łatwo zbadać kod źródłowy.

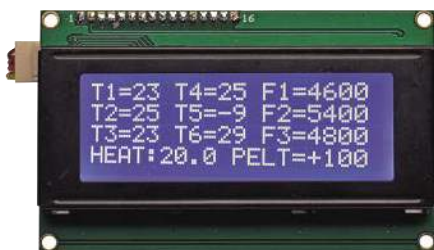
Jego działanie polega na skanowaniu termistorów raz na sekundę, wraz z sygnałami tachometru wentylatora. W tym samym czasie przetwarzane są wszelkie odebrane komendy IR. W zależności od nich oprogramowanie wybiera odpowiedni tryb (grzanie, chłodzenie lub wyłączenie), a następnie aktualizuje sygnały sterujące wentylatorami, pompami i ogniwami Peltiera.

Szkic jest dobrze udokumentowany komentarzami „in-line”, więc są one dobrym miejscem do rozpoczęcia analizy kodu, jeśli chcesz go zrozumieć i zmienić.

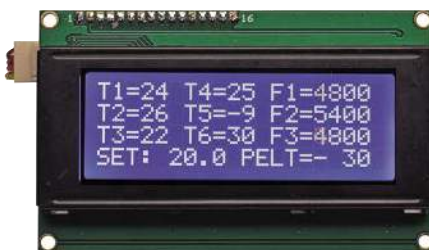
Szkic nazywa się „Peltier_Controller_V10”, choć może się to zmienić, jeśli będziemy go dalej aktualizować.

Na etapie programowania warto odłączyć UNO od stosu płytek i podłączyć go (samodzielnie) do portu USB komputera. Pozwoli to uniknąć problemów, które mogłyby się pojawić w związku z tym, że sygnał odbiornika IR jest współdzielony z jednym ze styków używanych do programowania.

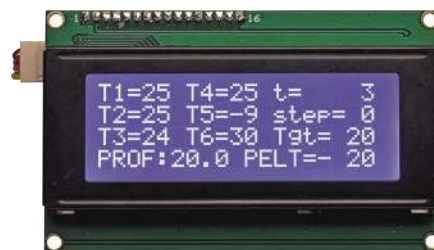
Jeśli Twój „panel” Peltiera znajduje się dalej od komputera, to również ułatwi Ci życie.



W większości trybów wyświetlana jest temperatura i prędkość wentylatorów. Tu pokazujemy tryb ogrzewania przy pełnej mocy, PELT=+100%; tryb pełnego chłodzenia wykorzystuje PELT=-100%.



W trybie Set, sterownik ogniw Peltiera moduluje PWM, aby doprowadzić temperaturę T1 (górną po lewej) do wartości zadanej (dół po lewej). W tym przypadku potrzebne jest umiarkowane chłodzenie na poziomie 30%.



W trybie profilu wartość zadana jest zmieniana zgodnie z serią punktów temperatury z jej interpolacją pomiędzy nimi. Zamiast prędkości obrotowej wentylatora, po prawej stronie wyświetlany jest czas, numer kroku i docelowa temperatura.

Otwórz plik szkicu, wybierz UNO z menu „Tools>Board” i upewnij się, że wybrany jest właściwy port szeregowy. Prześlij szkic (CTRL+U), a jeśli to się uda, odłącz kabel USB i zamieść UNO w stosie płytek.

Wyświetlacz powinien ożyć, pokazując tablicę temperatur. Nic więcej nie powinno się wydarzyć.

Domyślnie szkic akceptuje polecenia z pilota Jaycar XC3718 lub uniwersalnego pilota Altronics A1012 ustawionego na kod TV 089. Mogą działać inne piloty zaprogramowane protokołem TV Philips.

Podstawowe funkcje pracy

Istnieją cztery podstawowe tryby pracy: pełne grzanie, pełne chłodzenie, sterowanie proporcjonalne z ustaloną temperaturą docelową lub sterowanie proporcjonalne według zdefiniowanego w szkicu profilu temperatury.

Dla dwóch pierwszych trybów ogniw Peltiera są zasilane pełnym napięciem z jedną lub drugą polaryzacją. W każdym trybie na wyświetlaczu LCD pojawiają się różne informacje o stanie urządzenia, co widać na załączonych zdjęciach.

W dwóch ostatnich trybach urządzenie stara się utrzymać temperaturę głównego termistora (T1) na żądanym poziomie poprzez ogrzewanie lub chłodzenie w stopniu zależnym od potrzeb.

Do sterowania można użyć następujących przycisków na pilocie:

- „CH+” i „CH-” (na obu typach pilota) włączają odpowiednio pełne grzanie i pełne chłodzenie. Drugie naciśnięcie któregośkolwiek z tych samych przycisków wyłączy termoregulator.
- Aby zaprogramować wartość zadaną dla trzeciego trybu (stałej temperatury), należy wprowadzić trzy cyfry na klawiaturze numerycznej; wprowadzona liczba jest dzielona przez dziesięć, aby uzyskać temperaturę docelową. Na przykład wprowadzenie 1, 2, 3 spowoduje ustawienie wartości docelowej na 12,3°C. Można to zrobić tylko w czasie bezczynności urządzenia, ponieważ w przeciwnym razie mogłoby to spowodować gwałtowne przełączenie między grzaniem a chłodzeniem.

- Naciśnięcie przycisku „Power” (na pilocie Altronics) lub „Play” (na pilocie Jaycar) spowoduje rozpoczęcie lub zakończenie pracy w trybie wartości zadanej. W tym trybie można zmieniać wartość zadaną za pomocą przycisków zwiększania i zmniejszania głośności. Można to zrobić podczas pracy urządzenia, ponieważ w tym przypadku niewielkie zmiany są dopuszczalne.
- Tryb profilu temperatury jest aktywowany przez naciśnięcie przycisku „EQ” na pilocie Jaycar lub „-/-” na pilocie Altronics.

Zamiast pokazywać prędkości wentylatorów, wyświetlacz LCD wskazuje czas, numer kroku i następny cel czasowy. Urządzenie przechodzi przez tablicę kroków jednostkowych: temperatury/czasu; ustawionych w szkicu, interpolując temperaturę pomiędzy każdym punktem.

Można to wykorzystać do implementacji kuchenki „sous-vide” opartej na sterowniku czasowym, o której wspomnieliśmy wcześniej, lub profilu warzenia lub serowarstwa określonego przez dokładny produkt, który próbujesz zrobić. Zazwyczaj można uzyskać pojęcie o profilu, który będzie potrzebny, z przepisu, ale mogą być wymagane eksperymenty i zmiany w celu uzyskania najlepszego wyniku.

Rozwiązywanie problemów

Możesz sprawdzić, czy ogniwa Peltiera są podłączone z oczekiwaną polaryzacją, ustawiając urządzenie w trybie pełnego chłodzenia, a następnie sprawdzając, czy temperatura głównego czujnika (T1) spada, a nie rośnie. Jeśli wzrasta, to dezaktywuj tę linię w kodzie dodając „//” na początku:

```
setBipolar((pDrive*PWM_TOP)/100);
// wyjście skalowane
tj,
// setBipolar((pDrive*PWM_TOP)/100);
//scaled output
i usunąć „//” z początku tego:
// setBipolar(-(pDrive*PWM_TOP)/100);
//scaled output,
tj,
setBipolar(-(pDrive*PWM_TOP)/100);
//scaled output,
```

Jeśli twój LCD nie świeci się lub nic nie wyświetla, sprawdź czy czerwona dioda LED szybko miga. Jeśli tak, to program nie wykrył modułu I²C, więc nie mógł zainicjować i sterować wyświetlaczem.

Nasz szkic zawiera kod do automatycznego wykrywania adresu I²C wyświetlacza, więc powinien działać, jeśli wyświetlacz LCD jest podłączony prawidłowo. Sprawdź swoje okablowanie i zresetuj Arduino, naciśkając przycisk RST na nakładce interfejsu. Jeśli to nie rozwiąże problemu, może to być problem z Twoim modułem LCD.

Co teraz?

Przedstawiliśmy sporą liczbę opcji i zastosowań, w których można wykorzystać ten układ, ale nie mamy miejsca, aby szczegółowo omówić wszystkie możliwości.

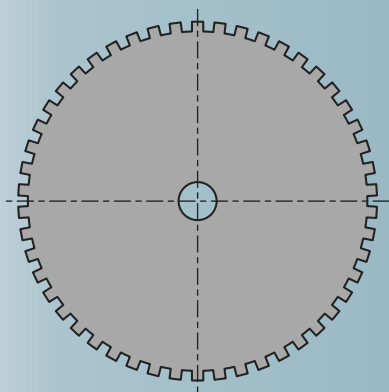
Możesz zmodyfikować na wiele sposobów kod źródłowy, aby dostosować go do swojej aplikacji. Na przykład, mógłbyś dodać do swojego Arduino moduł zegara czasu rzeczywistego oparty na DS3231, podłączając go do portu I²C (sprzedajemy je za kilka dolarów w Silicon Chip online shop). Pozwoliłoby to na skonfigurowanie kodu tak, aby automatycznie uruchamiał i zatrzymywał urządzenie o zadanych godzinach.

Możesz też zmodyfikować kod tak, aby można było ustawić wiele profili temperatury, aby dostosować je do różnych procesów, z możliwością wyboru między nimi (np. naciśkając różne przyciski na pilocie).

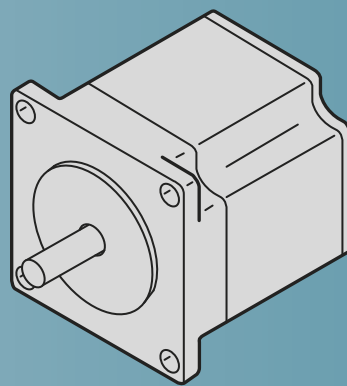
Jest tak wiele możliwości wykorzystania tego projektu; chcielibyśmy usłyszeć od naszych Czytelników, jakie zastosowania wymyślił dla termoregulatora! ■

Tim Blythman i Nicholas Vinen

Adaptacja do wydania polskiego – Andrzej Nowicki



Silniki krokowe w praktyce



Sterowniki bipolarnych silników krokowych, część 4

W tym miesiącu w ramach nauki obsługi silników krokowych omawiamy sterowniki bipolarnych silników krokowych. Dla tych samych rozmiarów bipolarnie silniki krokowe oferują na ogół wyższy moment obrotowy i lepszą wydajność w porównaniu z opisanymi wcześniej silnikami unipolarnymi. Jednak te korzyści mają swoją cenę, ponieważ sterowniki silników bipolarnych w porównaniu z unipolarnymi są konstrukcyjnie bardziej złożone. Na szczęście, na rynku jest obfitość tanich układów sterowników i modułów.

Podstawowym wymogiem dla sterownika silnika bipolarnego jest możliwość odwrócenia kierunku przepływu prądu w każdym z dwóch uzwojeń. Osiąga się to za pomocą pary obwodów „mostka H”.

Mostek H

W zasadzie zaprojektowanie mostka typu H jest stosunkowo proste, to znaczy do momentu, kiedy się tego nie spróbuje i nie zniszczy wielu MOSFET-ów. W najprostszej formie, mostek H to po prostu cztery półprzewodnikowe przełączniki ułożone w konfiguracji „H”, która jest zapewne znana wielu czytelnikom. **Rysunek 27** pokazuje podstawowy układ mostka H – topologię obwodu, którą można znaleźć nie tylko w sterownikach silników krokowych, ale również jest szeroko stosowana do odwracania kierunku obrotów w sterownikach silników prądu stałego.

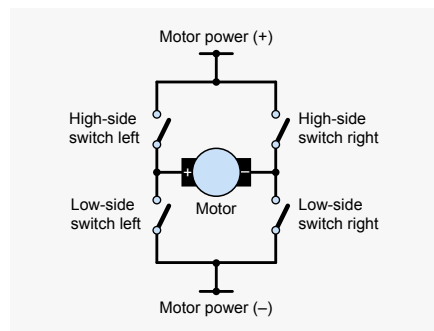
Odniesienia na rysunku 27 do przełączników „high-side”/„low-side” to terminologia powszechnie stosowana w przypadku sterowników mostków H. Przełączniki high-side łączą się z dodatnią szyną zasilania, a przełączniki low-side łączą się z szyną ujemną lub masą, więcej na ten temat później. Aby silnik prądu stałego obracał się w jednym kierunku, należy zamknąć lewy górny i prawy dolny przełącznik, jak pokazano na **rysunku 28**. Aby odwrócić kierunek, zamknij prawy górny i lewy dolny przełącznik (**rysunek 29**). Ta sama metodologia jest stosowana do odwrócenia prądu w uzwojeniu silnika krokowego.

Zamieńmy przełączniki na tranzystory MOSFET, typu p dla strony wysokiej i typu n dla strony niskiej (**rysunek 30**). MOSFETY typu N są preferowane, ponieważ mają zwykle mniejszą rezystancję

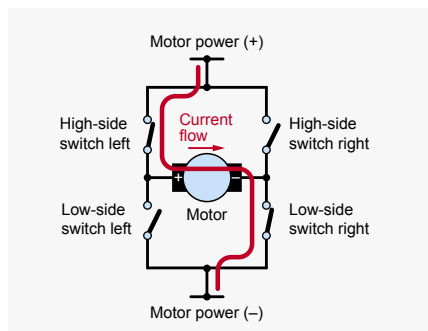
włączania, co oznacza mniejsze straty mocy i niższą temperaturę pracy. Niestety, zastosowanie typu n dla strony wysokiej nie jest proste, ponieważ napięcie bramki musi przekraczać napięcie zasilania mostka. Można to jednak osiągnąć za pomocą obwodu pompy ładunkowej DC-DC, aby podnieść napięcie powyżej dodatniej szyny zasilania.

Istnieje tu potencjalny problem. Gdybyś zamknął *oba* prawe lub *oba* lewe przełączniki, to doszłoby do zwarcia w zasilaniu silnika. Oczywiście można zastosować środki, które – teoretycznie – zapobiegają temu. Środki te mogłyby być zaimplementowane w logice dyskretnej lub oprogramowaniu, ale zwarcie nadal mogłoby się zdarzyć.

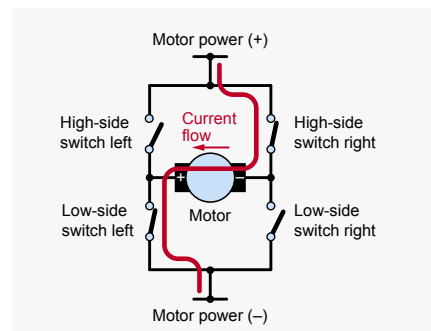
W energoelektronice używamy terminu „shoot-through” dla sytuacji, gdy *oba* MOSFET-y po tej samej stronie mostka są chwilowo zamknięte w tym samym



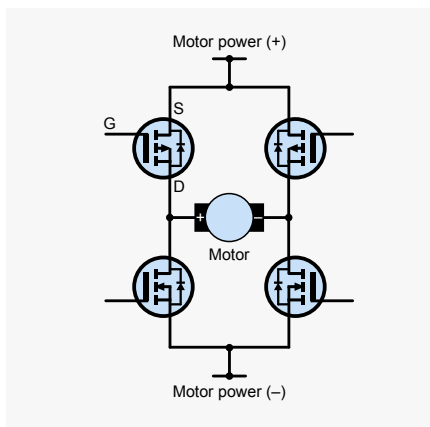
Rysunek 27. Cztery przełączniki i źródło zasilania – podstawa sterownika H-Bridge



Rysunek 28. Przepływ prądu w mostku H dla określonego kierunku obrotów silnika



Rysunek 29. Kierunek prądu w mostku H dla odwrócenia kierunku obrotów silnika z rysunku 28



Rysunek 30. Mostek H z tranzystorami MOSFET – typu p dla strony wysokiej i typu n dla strony niskiej. Diody są samoistne, stanowią część tranzystora MOSFET

czasie, nawet tylko przez kilka mikrosekund stanu przejściowego podczas przełączania tranzystorów. Otwarcie i zamknięcie przełączników MOSFET może nakładać się na siebie z powodu opóźnień spowodowanych stałymi czasowymi wynikającymi z impedancji bramki, pojemności bramki, prądu sterowania bramki, a nawet związanymi z projektem płytki drukowanej. Aby temu zapobiec, wprowadza się czas martwy, dzięki któremu załączenie tranzystora jest opóźnione o okres dłuższy niż czas wyłączenia. Oznacza to, że podczas tych martwych czasów wszystkie cztery MOSFET-y będą wyłączone. Jednak rozwiązanie problemu „shoot-through” prowadzi do nowego problemu – przy otwartych wszystkich przełącznikach nie ma którędy płynąć prąd silnika (prąd wynikający z zaniku pola magnetycznego w uzwojeniu). To spowoduje gwałtowny wzrost napięcia w cewce silnika, który potencjalnie może przekroczyć napięcie znamionowe MOSFET-ów i spowodować ich awarię – często dramatyczną awarię z dymem, a nawet płomieniami na MOSFET-ach dużej mocy. (Mówiłem, że projektowanie i działanie mostków H jest zdradliwe!)

Musimy zapewnić, aby przez cały czas istniała droga dla przepływu prądu w cewce i jest to osiągane za pomocą „wewnętrznych” diod MOSFET-ów; lub za pomocą zewnętrznych diod, jeśli używamy tranzystorów bipolarnych. Warto zauważyć, że diody wewnętrzne MOSFET mogą nie być wystarczająco szybkie lub zdolne do odprowadzenia wystarczającej ilości energii prądów uzwojenia podczas okresów martwych, dlatego w konstrukcjach MOSFET mogą być wymagane również diody zewnętrzne.

Istnieją jeszcze inne problemy projektowe dla mostka H, ale wykraczają one poza ramy tego artykułu. Celem rozwodzenia się nad

pułapkami w projektowaniu mostków H jest pokazanie, że pomimo ich powierzchownej prostoty nie są one łatwe do zbudowania i niezawodnego działania. Odpowiedzią jest korzystanie z gotowych sterowników silników krokowych – ich kluczową zaletą jest to, że problemy złożonych zależności czasowych są już załatwione wewnątrz sterownika. Możesz skupić się na aplikacji sterownika silnika i nie martwić się o projekt mostków.

Główne atrybuty sterownika krokowego

Jakie cechy powinien mieć nasz sterownik silnika krokowego i dlaczego? Przeanalizujemy trzy ważne wymagania.

Mikrokroki

Nasza aplikacja może wymagać bardzo małych kroków silnika do ustalania pozycji. Rozważmy drukarkę 3D, która porusza swoją głowicą wytłaczającą w krokach ułamków mm, być może 20 μm lub mniej. Zakładając, że silnik krokowy ma 200 kroków na obrót i napędza wytłaczarkę głowicy przez 18-zębowy pasek GT2 (skok 2 mm) to układ ten daje długości kroku:

$$\begin{aligned} 1 \text{ obrót} &= 18 \text{ zębów} \times 2 \text{ mm} \\ &= 36 \text{ mm przesuwu liniowego} \end{aligned}$$

$$\begin{aligned} 1 \text{ pełny krok} &= 36/200 \\ &= 0,18 \text{ mm przesuwu liniowego} \end{aligned}$$

Gdyby istniał sposób na zmniejszenie wielkości kroku, moglibyśmy poprawić rozdzielczość. W tym miejscu wykorzystywana jest ważna technika silnika krokowego zwana „mikrokrokiem”. Zamiast pełnego załączenia lub wyłączenia uzwojenia, stosowane są pośrednie wartości prądu. Dzięki temu silnik może zajmować pozycje pomiędzy odstępami między zębami na wirniku. Pół, ćwierć, 1/8, 1/16 kroku są powszechne; niektóre sterowniki mogą nawet zejść do 1/256 kroku, ale kosztem innych parametrów wydajności. Używając powyższego przykładu z drukarki 3D, mikrokroki 1/16 dają teraz rozdzielczość 0,18 mm/16, czyli 0,01125 mm lub 11,25 μm .

Aby ująć to w kategoriach kąta kroku; pełny obrót daje $360^\circ/200=1,8^\circ$ na jeden krok, ale 1/16 kroku daje $360^\circ/(200 \times 16)=3200$ kroków na obrót, z kątem kroku $1,8^\circ/16=0,1125^\circ$. Wiele sterowników używa dwóch pseudo sinusoid do wysterowania uzwojeń techniką PWM (modulacja szerokości impulsu) z przesunięciem fazy o 90° . Gdy prąd w jednym uzwojeniu wzrasta, prąd w drugim uzwojeniu maleje. Zapewnia to bardziej płynny ruch w porównaniu z grubszymi krokami dyskretnymi przy stosowaniu obrotów o pół lub nawet ćwiartkę kroku. Ważne jest, aby zrozumieć, że mniejszy

rozmiar kroku nie zawsze jest pożądany lub korzystny; na przykład:

- Może zostać osiągnięty punkt, w którym silnik nie porusza się, ponieważ wielkość kroku jest tak mała, że moment obrotowy podczas obrotu nie pokonuje strat tarcia w silniku. Na przykład, 1/16 mikrokroku ma około 10% dostępnego momentu trzymającego w porównaniu do pełnego kroku, 1/256 mikrokroku ma około tylko 0,6%. Może to powodować utratę mikrokroków, co doprowadzi do błędów pozycyjnego i zwiększonej histerezy.
- Mniejsza wielkość kroku wymaga wyższej częstotliwości sygnału zegarowego, który jeśli pochodzi z mikrokontrolera może być poza jego możliwościami. Na przykład, napędzanie powyższego silnika krokowego z prędkością 600 RPM oznacza prędkość zegara $3200 \text{ kroków} \times 600/60 = 32 \text{ kHz}$.

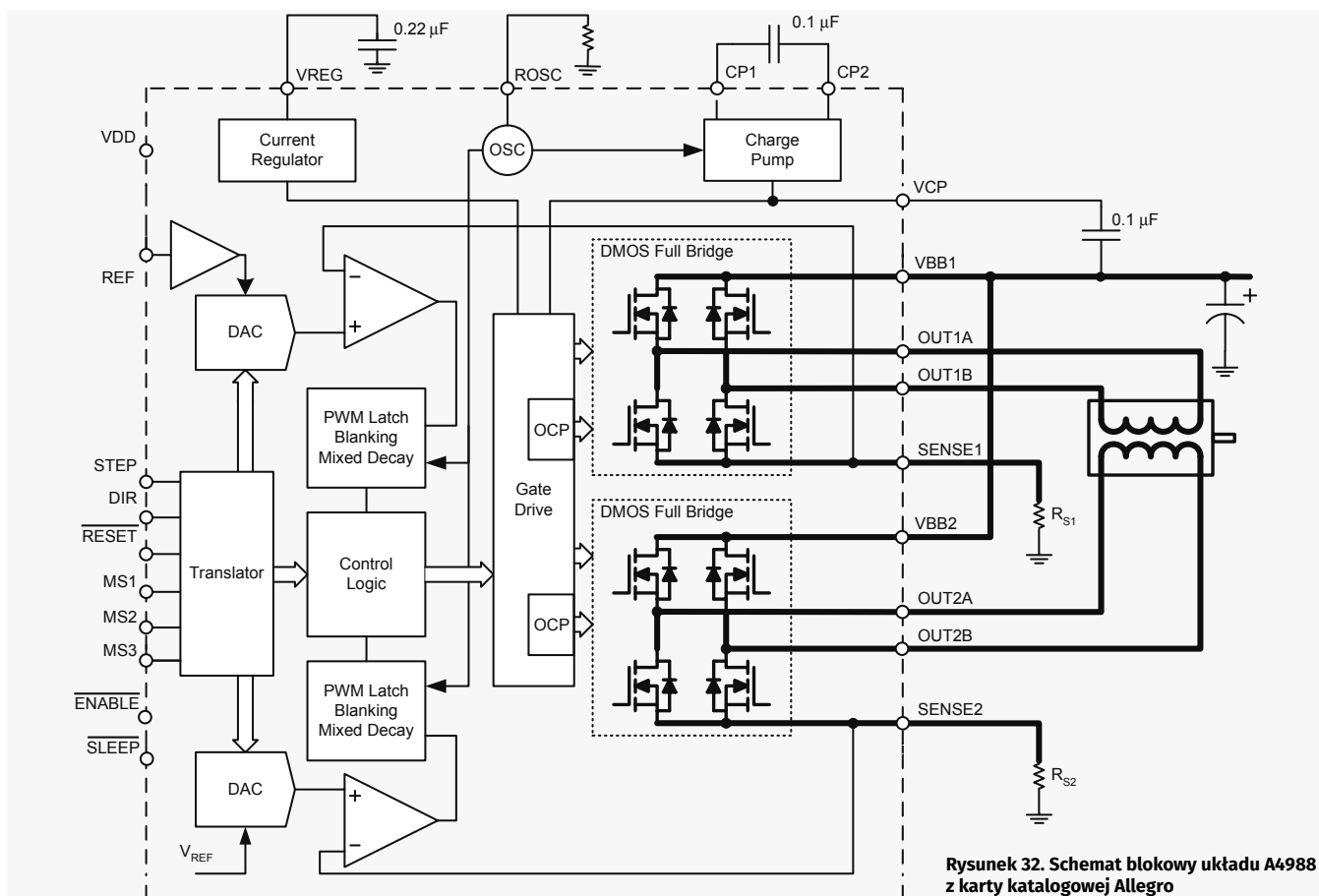
Ograniczenie prądu

Jest to być może największe źródło zamieszania dla nowych użytkowników bipolarnych silników krokowych. W zeszłym miesiącu, opisaliśmy ograniczenie prędkości silnika krokowego wynikające z szybkości zmian prądu uzwojenia spowodowanego jego indukcyjnością. Zwiększenie napięcia napędzającego uzwojenie pozwala na uzyskanie większych prędkości obrotowych silnika i jego przyspieszenie, ale działanie ograniczające prąd sterownika zapewnia, że nie zostanie przekroczony prąd znamionowy silnika.

To jest ważne, więc warto to rozwinąć. Hybrydowy silnik krokowy może mieć prąd fazowy 1 A przy napięciu 3,7 V, co jak można rozsądnie sądzić oznacza 3,7 V maksimum. Cóż, to prawda, ale dopiero po osiągnięciu przez prąd uzwojenia stanu ustalonego. Po skokowej zmianie prądu w uzwojeniu, indukcyjność uzwojenia przeciwstawia się zmianie prądu. Opozycja ta zmniejsza się aż do momentu, gdy szybkość zmian staje się zerowa. Jeżeli sterownik krokowy pracuje przy napięciu powiedzmy 12 V lub nawet 24 V, to zastosuje on tak wysokie napięcie, jak tylko



Rysunek 31. Płytki breakout Pololu A4988



Rysunek 32. Schemat blokowy układu A4988 z karty katalogowej Allegro

może bez przekraczania maksymalnego prądu. Oznacza to krótsze tempo zmian prądu uzwojenia, czyli możliwe są szybsze przyspieszenia i wyższe prędkości obrotowe. Dlatego drukarki 3D, które potrzebują dużych przyspieszeń, używają napięć zasilania znacznie większych niż znamionowe (w stanie ustalonym) silnika krokowego.

Minimalne/maksymalne napięcie i prąd roboczy

Dopasowanie sterownika silnika krokowego do silnika jest istotne. Sterownik musi być w stanie dostarczyć pełny maksymalny prąd fazowy w sposób ciągły, jeśli chcemy osiągnąć optymalną wydajność momentu obrotowego silnika. Minimalne napięcie znamionowe sterownika musi być poniżej wybranego zasilania i maksymalne napięcie znamionowe powyżej zasilania.

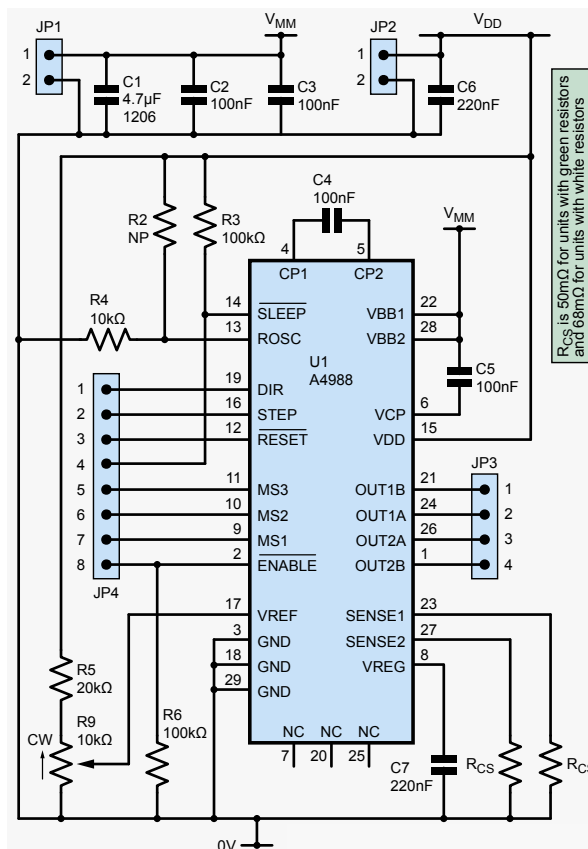
Układy scalone sterownika silnika krokowego A4988 i DRV8825

W społeczności majsterkowiczów zajmujących się drukarkami 3D bardzo dobrze znane są układy sterownika krokowego A4988 i DRV8825, odpowiednio z Allegro i Texasu, zwłaszcza A4988. Są one dostępne tylko w obudowach do montażu powierzchniowego,

ale nie jest to ograniczenie, ponieważ są one łatwo dostępne na kompaktowych, niedrogich płytkach uruchomieniowych typu breakout; na przykład rysunek 31. Płytki te „dopasowują” rozstaw pinów układu scalonego montowanego powierzchniowo, takiego jak A4988, do rozstawu pinów, który bardziej odpowiada komponentom z otworami przelotowymi i często nadają się do bezpośredniego użycia jako płytki prototypowe. Dodatkowe komponenty wymagane przez układ mogą również pojawić się na płytce breakout.

Schemat blokowy płytki A4988 (rysunek 32) warto omówić nieco dokładniej, ponieważ pokazuje, jak zwiększona złożoność sterowników bipolarnych jest załatwiona w jednym, łatwym w użyciu urządzeniu.

Na schemacie blokowym widać dwa ‘mostki H’,



Rysunek 33. Schemat układu Pololu dla ich płytki breakout opartej na A4988

Tabela 7. Opcje wyboru mikro-kroków A4988 z wykorzystaniem pinów sterujących MS1, MS2 i MS3

MS1	MS2	MS3	Rozdzielczość mikro-kroków
niski	niski	niski	pełny krok
wysoki	niski	niski	pół kroku
niski	wysoki	niski	ćwierć kroku
wysoki	wysoki	niski	jedna ósma krok
wysoki	wysoki	wysoki	jedna szesnaśta krok

po jednym dla każdego uzwojenia. Zauważ, że MOSFETy typu high-side są również typu n, więc A4988 zawiera pompę ładunkową dla napięć bramek wyższych niż napięcie zasilania. Dla każdego mostka wymagany jest zewnętrzny rezystor zabezpieczający, aby można było kontrolować prąd uzwojenia jako część obwodu ograniczającego prąd A4988. Układ logiczny sterujący implementuje różne tryby mikrokroków, interfejsy i ochronę przed „przestrzeleniem” (shoot-through) poprzez zapewnienie stałego czasu wyłączenia. Ponieważ układ A4988 implementuje wszystkie funkcje, z którymi do tej pory się zetknęliśmy, to spodziewaj się, że będzie potrzebował kilku zewnętrznych komponentów, które są na płytce breakout.

Rysunek 33 pokazuje schemat płytki breakout A4988 amerykańskiego producenta, firmy Pololu – **pololu.com**.

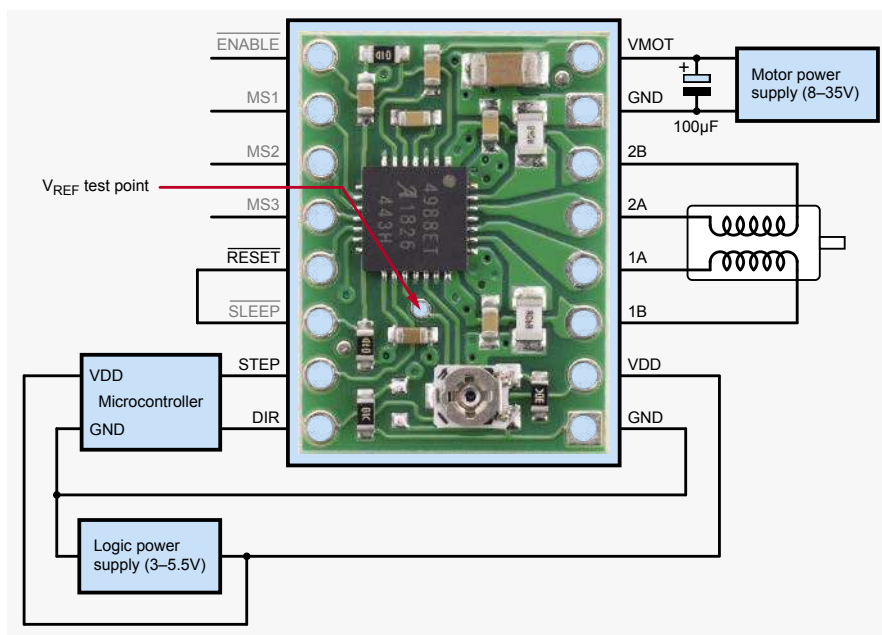
Płytką ta jest dostępna w Wielkiej Brytanii w **technobotsonline.com** za około 5 funtów. Oprócz zewnętrznych elementów wyszczególnionych w karcie katalogowej A4988, zawiera ona również kondensatory odsprężające zasilanie oraz potencjometr (R9) do ustawiania limitu prądu.

Płytką breakout A4988 jest bardzo łatwa w podłączeniu i użytkowaniu, minimalny schemat połączeń pokazany na **rysunku 34** jest w dużej mierze samoobjaśniający, ale krótko:

Potrzebne są dwa zasilacze, jeden dla silnika w zakresie 8–35 V i drugi 5 V dla zasilanie logiki, które może pochodzić z mikrokontrolera lub innego stabilizowanego zasilania. VMOT, zasilanie silnika musi mieć kondensator 100 µF odsprężający jego linie zasilania; jest to jedyny element, o który musisz się zatroszczyć, wszystkie inne są na płytce.

Wejście Enable może pozostać niepodłączone, ale jeśli zostanie podłączone do +5 V, wyłączy wyjścia do silnika, natomiast cała pozostała logika sterowania pozostanie aktywna.

MS1, MS2 i MS3 – jeśli pozostawimy je odłączone sterownik przejdzie w tryb pełnokrokowy. Różne tryby mikrokrokowania



Rysunek 34. Schemat połączeń dla płytki breakout A4988 firmy Pololu

można wybrać poprzez wykonanie odpowiednich połączeń do +5 V, jak pokazano w **tabeli 7**. Ponieważ A4988 posiada wewnętrzne rezystory podciągające, to piny mogą być pozostawione otwarte, jeśli wymagany jest stan niski.

Końcówki reset i sleep są podobne do wejścia enable w tym, że wyłączają wyjścia do silnika, ale dodatkowo posiadają dalsze funkcje oszczędzania energii. Zaleca się korzystanie z arkusza danych Allegro, jeśli musisz używać tych funkcji, ale jeśli nie, to po połączeniu tych końcówek razem, będą one miały ten sam rezystor podciągający na płytce breakout.

Krok i kierunek poznałeś już zapewne z poprzednich artykułów tej serii.

W końcu dochodzimy do połączeń silnika. Tak długo jak jedno uzwojenie jest podłączone do 2A i 2B, a drugie do 1A i 1B, twój silnik krokowy będzie się obracał. Możesz zmienić kierunek za pomocą pinu sterującego kierunkiem, lub po prostu zamienić 1A i 1B.

Ustawienie ograniczenia prądu

Ostatnim tematem, który omówimy w tym odcinku, jest ograniczenie prądu. Ważne jest, aby było ono ustawione prawidłowo, ponieważ w przeciwnym razie silnik może się przegrzewać lub może brakować mu momentu obrotowego. Istnieje kilka sposobów ustawienia limitu prądu, ale metoda, której ja używam okazała się skuteczna.

Z rysunku 33 widać, że trymer ustawia napięcie na wejściu V_{REF} układu A4988. Najpierw obliczasz napięcie V_{REF} , a następnie ustawiasz je za pomocą trymera *przed* podłączeniem silnika.

Z karty katalogowej silnika będzie podany prąd fazowy. Załóżmy, że prąd wynosi 1 A i chcemy ograniczenia prądu przy tej wartości, aby osiągnąć pełną wydajność. Aby to zrobić musimy znać wartość rezystorów current-sense (do pomiaru prądu) – wszystkie płytki wysyłane z Pololu i Technobotów mają rezystory current-sense o wartości $RCS=0,068 \Omega$.

Korzystamy więc z następującego wzoru:

$$\begin{aligned} V_{REF} &= I_{MAX} \times R_{CS} \\ &= 1 \times 0,068 \\ &= 544 \text{ mV} \end{aligned}$$

Uzbrojony w woltomierz zmierz napięcie pomiędzy masą logiczną a przelotką znajdującą się pomiędzy trymerem a układem A4988 – patrz rysunek 34. Reguluj trymer tak długo, aż zmierzysz obliczone napięcie. (Wziąłem również mały śrubokręt zaciskowy, przymocowałem dodatni przewód woltomierza do trzonu śrubokręta za pomocą klipsa, a czarny przewód do masy. Możesz wtedy zmierzyć napięcie bezpośrednio na trymerze podczas jego regulacji, ponieważ suwak trymera jest podłączony do V_{REF} pin).

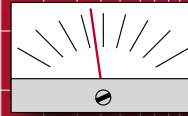
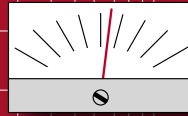
W następnym miesiącu

W części 5, będziemy badać cechy wielu innych bipolarnych sterowników silnika krokowego, bo chociaż A4988 jest świetny, poczekaj aż zobaczysz co możesz zrobić z bardziej zaawansowanymi sterownikami! ■

Paul Cooper

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

AUDIO OUT



PE Mini-monitorowa zwrotnica dla przetworników Wavector, część 2

Spędziliśmy trochę czasu na analizowaniu parametrów głośników i konstrukcji zwrotnic i jesteśmy prawie w punkcie, w którym możemy zbudować kompletny głośnik, ale najpierw musimy ukończyć naszą nową zwrotnicę. (Red. W tym artykule terminy głośnik i kolumna głośnikowa są synonimami.)

Układ podstawowy

W artykule z zeszłego miesiąca doszliśmy do układu pokazanego na **rysunku 16** (powtórzony poniżej). Nie jest to ostateczny projekt zwrotnicy, ale jest to zadowalający fundament, na którym można budować.

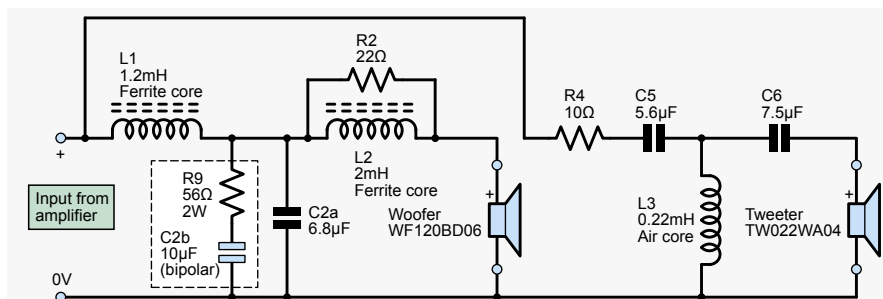
Zauważyłem, że jest trochę garbów na charakterystyce częstotliwościowej imoedancji i co bardziej niepokojące, dołek (przy 2 kHz) poniżej 8 Ω (**rysunek 17**). Ważne jest, aby podczas projektowania zwrotnic kontrolować krzywą impedancji. Projektant kolumn nie powinien utrudniać życia projektantowi wzmacniacza poprzez niskie spadki i bardzo szybkie (reaktywne) zmiany impedancji. Ten problem został rozwiązany przez dodanie rezystora tłumiącego 56 Ω (R9) równolegle do kondensatora C2a filtra dolnoprzepustowego. Oczywiście nie chcemy, aby 14% naszej cennej mocy basu było marnowane w tym rezystorze, który jest w rzeczywistości równoległy do jednostki basowej.



Para prototypów PE Mini-monitora w akcji – zwróć uwagę na piankę antyrefleksyjną na biurku!

Ponieważ działanie rezystora tłumiącego R9 jest potrzebne tylko w okolicach 2 kHz, to szeregowo z tym rezystorem włączono kondensator

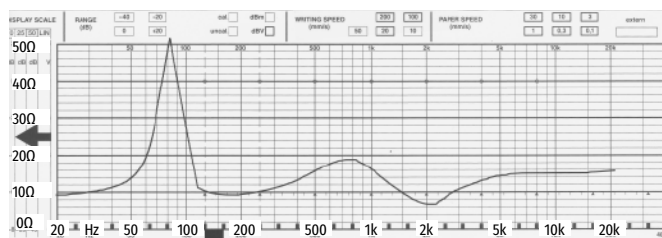
bas-bloker 10 μF (C2b). Kondensator ten nie musi mieć niskiego ESR, więc można zastosować tani elektrolit bipolarny, jeśli ktoś jest oszczędny. Otrzymana krzywa impedancji z obwodem tłumiącym jest pokazana na **rysunku 18**.



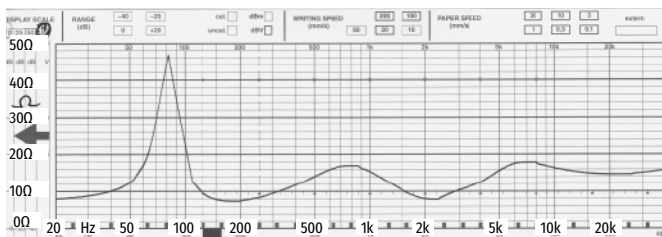
Rysunek 16. Kompletna prowizoryczna zwrotnica z ubiegłego miesiąca, plus obwód tłumiący (R9, C2b)

Trochę szczęścia z fazą

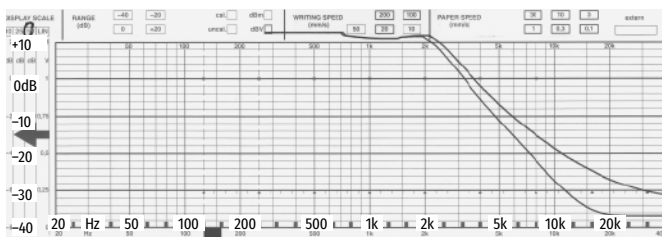
Miałem szczęście, właściwa faza jest często trudna do uzyskania w zwrotnicach, wymaga filtrów wszechprzepustowych lub przesuwania pozycji przetworników, z głębokimi stopniami w poziomie lub niepożądanymi dużymi odstępami w pionie. W tym przypadku różnica w poziomej płaszczyźnie



Rysunek 17. Krzywa impedancji pierwszej zwrotnicy z Rysunek 16 bez obwodu tłumiącego



Rysunek 18. Krzywa impedancji z obwodem tłumiącym (R9, C2b). Krzywa impedancji dla obwodu końcowego na rysunku 21 jest taka sama



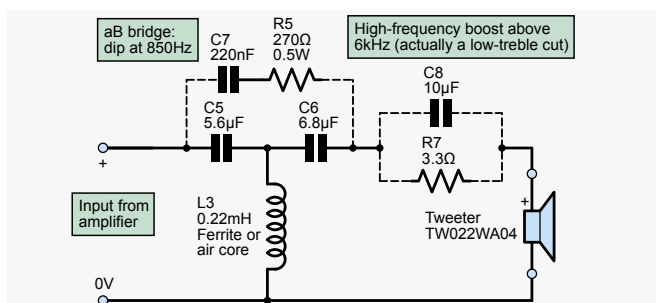
Rysunek 19. Poprawa tłumienia dolnoprzepustowego przy 10 kHz poprzez dodanie C3a i C4

promieniowania pomiędzy dwoma przetwornikami wynosiła tylko około 15 mm, przy czym położenie dużej kopułki głośnika niskotonowego i wpuszczonej kopułki głośnika wysokotonowego pomaga tę różnicę zminimalizować.

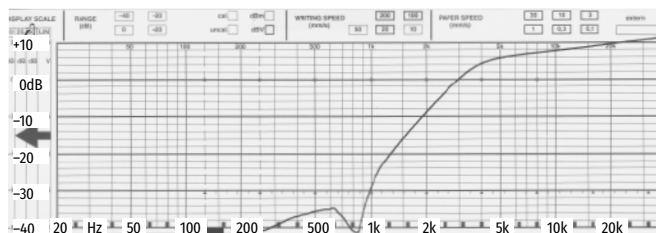
Ponadto, niska częstotliwość zwrotnicy sprawiła, że przesunięcie fazowe w sensie falowym było niewielkie; tak małe, że zniwelowało elektryczne przesunięcie fazowe z filtrów. Nie musiałem nawet odwracać fazy głośnika wysokotonowego, co rzadko zdarza się w przypadku zwrotnicy o tej topologii. Słuchając białego szumu, można zidentyfikować wcięcia fazowe wokół punktu zwrotnicy, poruszając się w górę i w dół przed głośnikiem. (Róbcie to, gdy nikt nie patrzy!) Biały szum powinien „żelować się” na osi obudowy w punkcie pomiędzy dwoma jednostkami napędowymi. Obiektywnym testem każdej zwrotnicy jest delikatne podłączenie głośnika wysokotonowego poza fazą. Powinno powstać głębokie symetryczne zero (jak pokazano na rysunku 32).

Małe poprawki

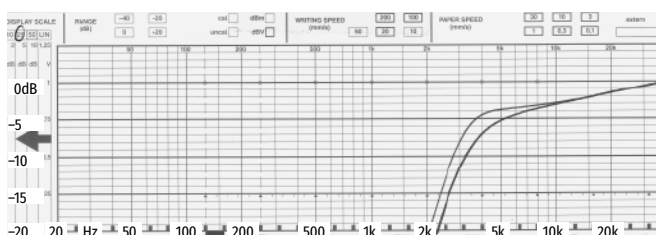
Kiedy podstawowa zwrotnica jest już uporządkowana, można dokonać pewnych udoskoaleń, często obejmujących testy odsłuchowe, zwane „voicingiem”. Drogimi elementami są cewki (zwykle ok. £4 do £6 za sztukę), więc unikamy dodawania kolejnych. Większość poprawek polega na zastosowaniu względnie tanich rezystorów i kondensatorów, generalnie dodając dodatkowe udoskoaleń w obszarze górnych i średnich tonów. Niektóre poprawki nie są widoczne na krzywej odpowiedzi akustycznej, ponieważ ich wpływ demonstruje się nisko w obszarze pasma zaporowego krzywych zwrotnicy, ale nadal są słyszalne. Z tego powodu, niektóre krzywe elektryczne (rysunek 19 i rysunek 20b) są wykreślone w skali 50 dB.



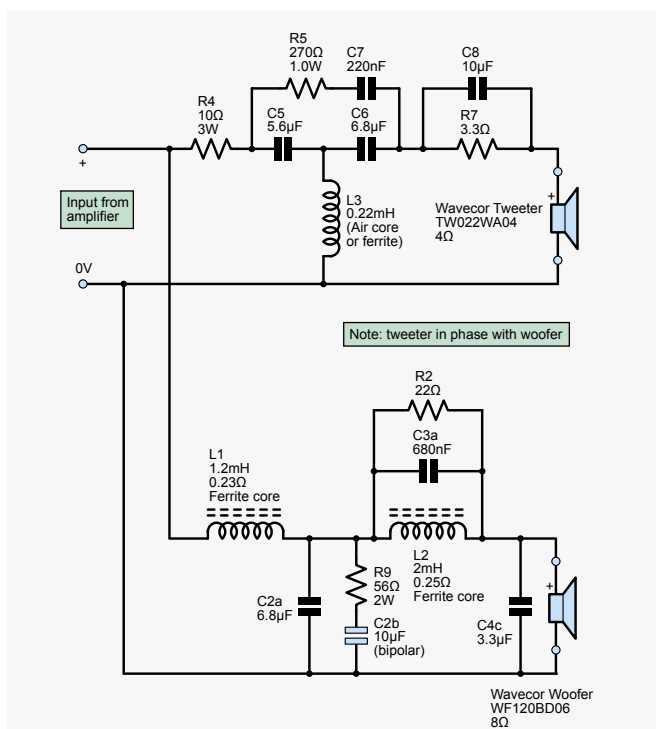
Rysunek 20a. Dodanie obwodu mostka Butterwortha (C7, R5) dla tłumienia rezonansu głośnika wysokotonowego. Uwaga dodano również obwód podbijający tony wysokie (C8, R7)



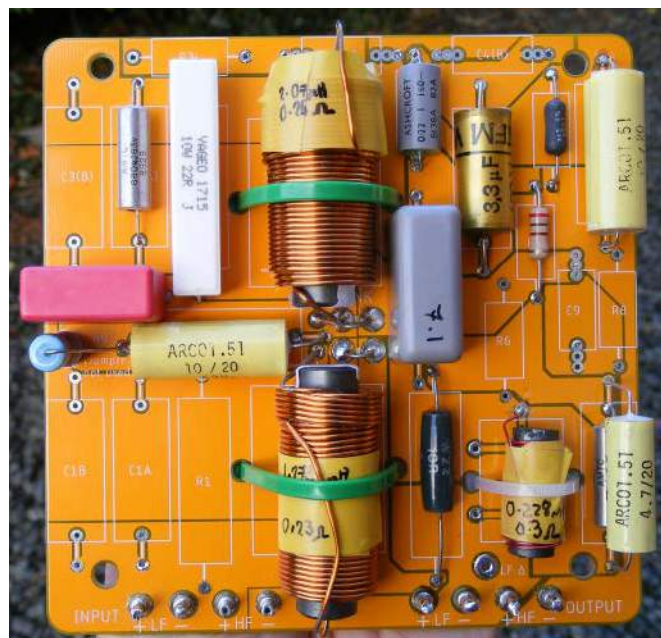
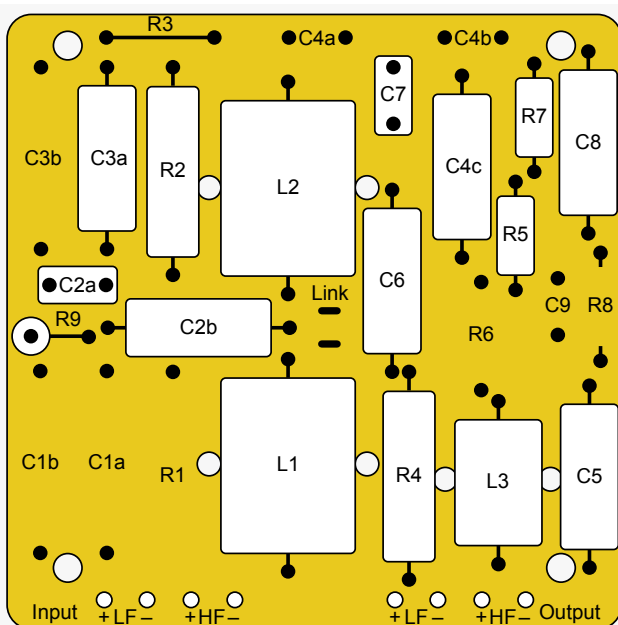
Rysunek 20b. Wynikowy dip antyrezonansowy głośnika wysokotonowego z wykorzystaniem obwodu mostka Butterwortha (C7, R5)



Rysunek 20c. Dodanie obwodu podbijającego tony wysokie (C8, R7) obcina nieco sygnał niskich tonów



Rysunek 21. Ostateczny układ zwrotnicy ze wszystkimi poprawkami. Nie zmieniaj one krzywej impedancji



Rysunek 22. (powyżej po lewej) Rozkład komponentów dla zwrotnicy PE Mini-Monitor na płycie drukowanej uniwersalnej zwrotnicy pasywnej. (Uwaga: oznaczenie fazy głośnika wysokotonowego na płycie jest nieprawidłowe i powinno być odwrócone – tzn. na dole po prawej stronie zamienić '+' i '-' po obu stronach wyjścia 'HF'. (Jest to poprawne dla wersji LS3/5A, która ma odwrócone połączenie głośnika wysokotonowego, ale głośnik wysokotonowy dla PE Mini-Monitor jest podłączony w fazie). Zwróć też uwagę, że dwie centralne zwory („Link” na rysunku) trzeba przylutować do płytki. Rysunek 23a. (powyżej po prawej) pokazuje ukończoną zwrotnicę PE Mini-Monitor (uwzględnij połączenie drutem R3!). Rysunek 23b. (poniżej) Zwróć uwagę na pionowy montaż R9

Tłumienie szczytu

Szczyt 10 kHz jest nadal lekko słyszalny na wyjściu głośnika niskotonowego na osi, przy odłączonym głośniku wysokotonowym. Dalsze tłumienie można uzyskać przez dostrojenie cewki końcowej z równoległym kondensatorem C3a. Pomaga również kondensator C4 równoległe z głośnikiem niskotonowym. Różnica dla krzywej elektrycznej dolno- i wysokoprzepustowej jest pokazana na **rysunku 19**. Charakterystyka obniża się o około 6 dB dla 10 kHz. Ponieważ nie ma teraz wysokich tonów promieniowanych z głośnika niskotonowego, ogólna odpowiedź wysokich tonów brzmi nieco płynniej, nie ma wcięt anulujących wysokie częstotliwości.

Bardziej miękkie tony wysokie

Zwrotnica w obecnym kształcie może brzmieć nieco „twardo” lub „średnio-tonowo” dla starszych słuchaczy (50+) z powodu związanego z wiekiem pewnego zaniku wysokich tonów lub starczego przytępienia słuchu, które ja mam. Można to poprawić poprzez lekkie podniesienie wysokich tonów we wzmacniaczu lub poprzez włączenie obwodu składającego się z C8 i R7 w szereg z głośnikiem wysokotonowym za filtrem górno- i wysokoprzepustowym. Zmniejsza to nieco moc wyjściową w rejonie częstotliwości podziału i podbija ją o kilka dB przy 10 kHz. Ten delikatny wzrost kompensuje również zmniejszenie wyjścia wysokich częstotliwości poza osi z głośnika wysokotonowego. Z powodu zwiększonej impedancji C6 jest

również zmniejszony do bardziej powszechnej wartości 6,8 μ F.

Głośnik wysokotonowy – sekcja akustycznego mostka Butterwortha

Oto kolejny subtelny zabieg, rozważany przez KEF-a i innych w celu wygładzenia krzywej górno- i wysokoprzepustowej w dolnym zakresie. Chodzi o tłumienie szczytu rezonansowego głośnika wysokotonowego o częstotliwości 850 Hz za pomocą akustycznego mostka Butterwortha, składającego się z C7 i R5. Mogłem go usłyszeć tylko na białym szumie z odłączonym głośnikiem niskotonowym, ale kosztuje mniej niż 1 funt, więc warto. Zmodyfikowany układ i otrzymana krzywa są pokazane na **rysunkach 20a i 20b**. Ta cecha jest ważniejsza dla filtrów górno- i wysokoprzepustowych trzeciego rzędu, w których głośnik wysokotonowy nie ma podłączonej cewki ani innych środków tłumiących rezonans. Istnieje opcja, aby głośnik wysokotonowy miał w szczeliny magnetycznej ferrofluid, który zapewnia duże tłumienie i moc. Nie podoba mi się to, bo olej po kilku latach wysycha/gęstnieje i źle brzmi przy niskich poziomach głośności.

Końcowy obwód zwrotnicy

Jest pokazany na **rysunku 21**. Jest tam wiele elementów 'R i C', ale minimum cewek. Poprosiłem Volt Loudspeakers w Devon o ich nawinięcie i z zadowoleniem stwierdziłem, że ich rezystancja DC była niższa niż pierwotnie

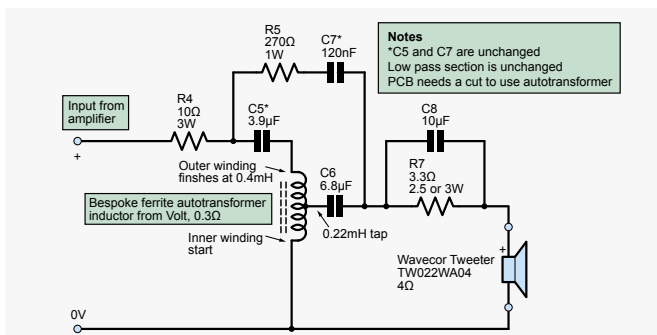


określona, a ich tolerancja mieściła się w granicach kilku procent.

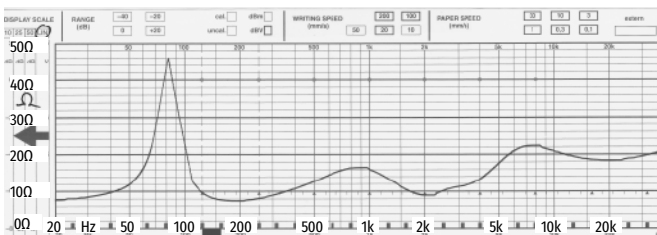
Końcowa odpowiedź częstotliwościowa

Zbudowanie zwrotnicy na płycie drukowanej Universal stanowi odświeżającą zmianę w stosunku do zwykłych, dzisiejszych mikro- i minikomponentów. **Rysunek 22** pokazuje rozkład elementów, a **rysunek 23** zdjęcie zmontowanej płytki. Metalizowane otwory przelotowe i elementy z końcówkami osiowymi dają maksymalną wytrzymałość, idealną dla zwrotnicy wstrząsanej w głośniku. Dobrym pomysłem jest owinięcie kawałka taśmy izolacyjnej wokół cewki w miejscu, gdzie koniec drutu nawojowego jest zaciśnięty pod opaską zaciskową.

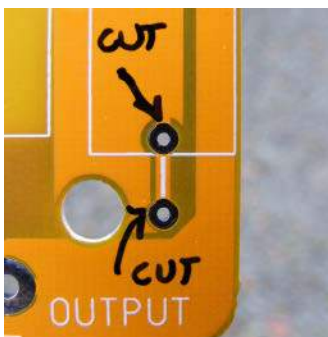
Zapobiega to powstawaniu zwarcia z powodu ścierania się emalii powstającej w wyniku wibrowania. Klej na gorąco, pianka i silikon elektryczny (bezkwasowy) mogą zredukować ewentualne brzęczenia.



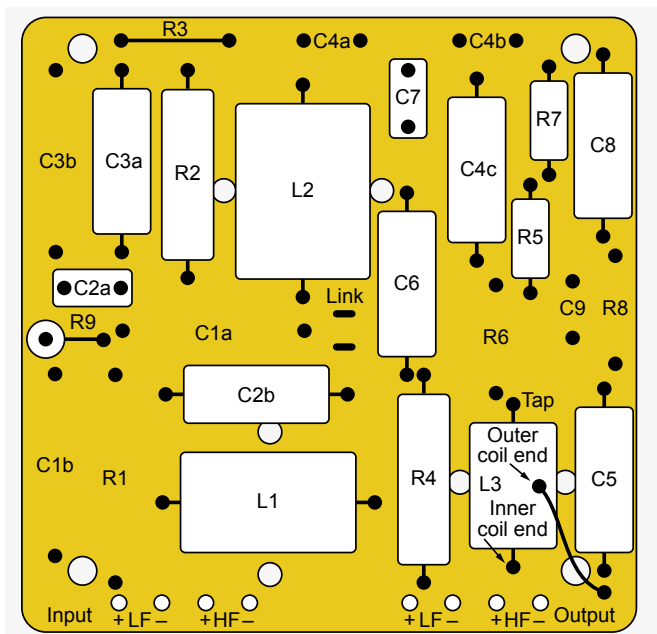
Rysunek 24. Nowa sekcja górnoprzepustowa z wykorzystaniem autotransformatora



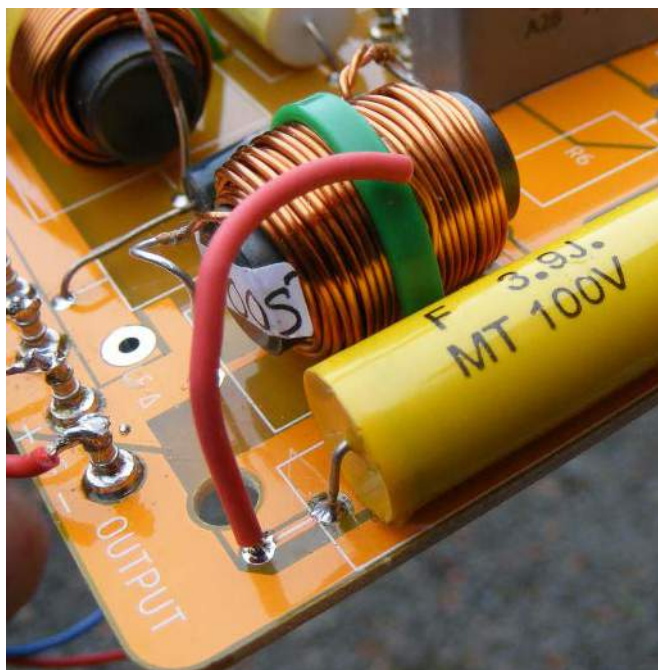
Rysunek 25. Impedancja wersji z autotransformatorem, znacznie wyższa w górnej części pasma



Rysunek 26. Aby użyć autotransformatora, płytka drukowana wymaga odizolowania ścieżki poprzez jej przecięcie, jak pokazano na rysunku. Zostanie to w pełni omówione w przyszłym artykule



Rysunek 28. Rozkład elementów nowej płytki z cewkami niskiej częstotliwości (L1 i L2) ustawionymi pod kątem prostym. Zauważ też, że zmienię się pozycje kondensatorów C1a i C2b. Uwaga (rysunek 22) dotycząca wyjścia głośnika wysokotonowego ma zastosowanie również tutaj.

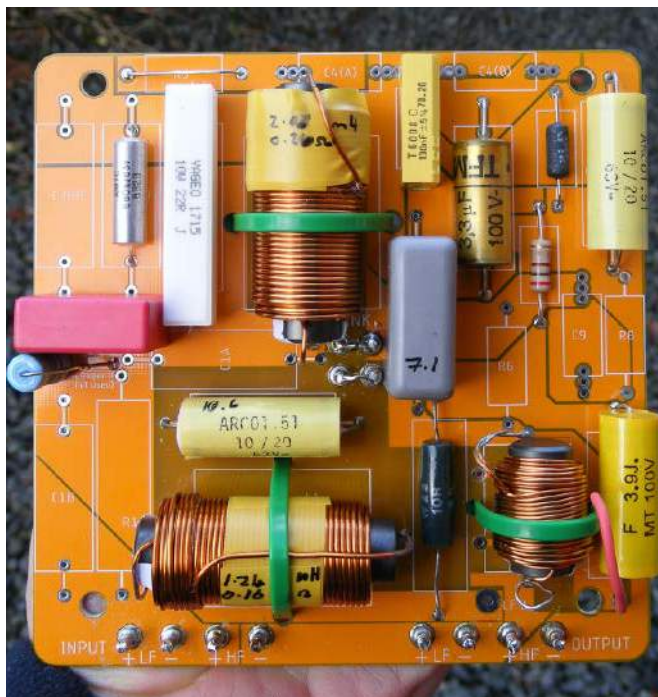


Rysunek 27. Przewód z zewnętrznego uzwojenia cewki odgałęzionej wprowadzamy do koszulki izolacyjnej i doprowadzamy do zapasowego otworu montażowego kondensatora

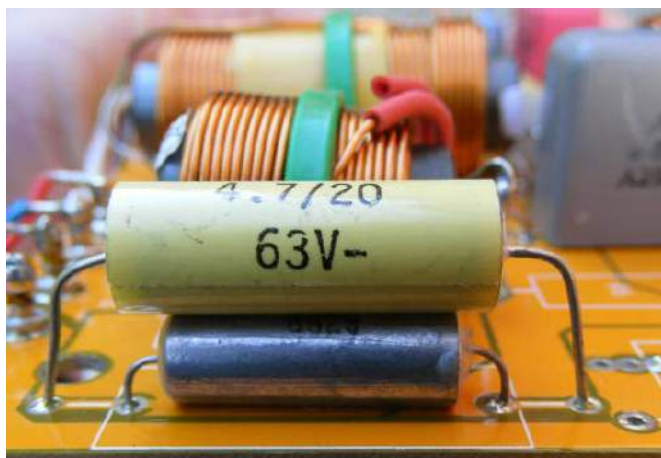
Lista części

Uwagi:

- Poniższa lista dotyczy jednej zwrotnicy – do pary potrzebne są dwa zestawy!
- Wszystkie komponenty mają tolerancję $\pm 5\%$ lub lepszą. Najważniejsze jest, aby komponenty na lewych i prawych zwrotnicach były dopasowane dla dobrego obrazowania stereo. Jeśli więc np. masz dwa kondensatory $10 \mu\text{F}$, które mają wartości o 10% za wysokie, umieść je w tym samym miejscu na każdej płytce.
- 'WW' oznacza wire-wound (drutowy).



Rysunek 29. Zdjęcie nowej płytki z autotransformatorem



Rysunek 30. Połączenia wielopunktowe mogą być wykorzystane do równoległego łączenia elementów w celu wyrównania tolerancji lub uzupełnienia wartości. W tym przypadku C5 (5,6 μF) został złożony z „wysokiego” 4,7 μF i 680 nF

Rezystory:

R2 22 Ω 6 W WW

R4 5,6 do 12 Ω 2,5 W WW, wybierz wartość według gustu, aby ustawić poziom głośnika wysokotonowego

R5 270 Ω 0,5 W węglowy R7 3,3 Ω 2,5 W WW

R9 56 Ω 2 W metalizowany lub WW

Kondensatory:

Należy pamiętać, że płytka drukowana uniwersalnej zwrotnicy pasywnej ma wiele otworów, więc konstruktorzy mogą stosować typy radialne lub osiowe. Dla kondensatorów foliowych należy stosować polipropylen (MKP), poliwęglan (MKC) lub poliestr (MKT) – preferencje w tej kolejności, dopuszczalne napięcie minimum 63 V.

C2a, C6 6,8 μF C2b, C8 10 μF C3a 680 nF C4c 3,3 μF

C5 5,6 μF od Blue Aran – nr części.

CVMPC25560 (lub połączenie równoległe: 4,7 μF i 820 nF) 220 nF

Cewki:

L1 1,2 mH 0,23 Ω 1 mm WW na 12,5×

Rdzeń 50 mm (Neosid F6 manganese ferrite rod part number 36-702-26). Numer części Volt F002.

L2 2 mH 0,25 Ω taka sama konstrukcja jak L1. Numer części Volt F020

L3 0,22 mH 0,3 Ω Ferryt o długości 9 mm lub 12 mm × 25 mm. Alternatywnie można użyć typu rdzeń powietrzny, 0,21 mH 0,4 Ω (Volt typ 72).

Cewki L1, L2, L3 (rdzeń ferrytowy) wykonane dla PE przez Volt Loudspeakers dostępne są w komplecie z rezystorami i kondensatorami ze sklepu PE – część WAVXO.

Różne:

Opaski kablowe 3 mm na 100 mm, 3 szt.

Universal Passive Crossover PCB ze sklepu PE – Part UPC0320

Przywieszki wieżowe 2 mm, 12 szt.

Drut miedziany 22swg ocynowany do linek, 100 mm

Taśma izolacyjna PCV o szerokości 19 mm, 150 mm

Autotransformator – przyszła aktualizacja

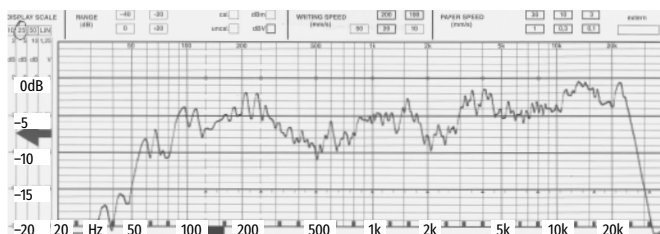
Zleciłem firmie Volt nawinięcie kilku cewek indukcyjnych, aby wdrożyć metodę wysokiej impedancji zastosowaną w zwrotnicy LS3/5A. Pozwala to uniknąć utraty efektywności przy dopasowaniu głośnika wysokotonowego 4 Ω do głośnika niskotonowego 8 Ω . Są to 0,4 mH odhaczone na 0,22 mH. Nowa sekcja górnoprzepustowa jest

pokazana na **rysunku 24**, a ponieważ impedancja wejściowa jest wyższa, pierwszy kondensator C5 musi być zmniejszony do 3,9 μF (**mouseer.co.uk** część ECW-FD2W395K) Wartość kondensatora C7 również musi być zmieniona na 120 nF (dostępny w firmie Tayda). Krzywa odpowiedzi jest taka sama, jednak krzywa impedancji pokazana na **rysunku 25** jest teraz o 4 Ω wyższa w zakresie wysokich tonów, co zmniejsza zniekształcenia wzmacniacza. Jest jeszcze jeden ważny szczegół konstrukcyjny, na który należy zwrócić uwagę, gdy ta część będzie dostępna: płytka drukowana będzie musiała być wycięta, patrz **rysunek 26**. Podłączenie autotransformatora pokazano na **rysunku 27**: Końcówki cewki wykonano z wyprowadzeń jednego z zapasowych kondensatorów (C5). Zostanie to w pełni omówione w późniejszym artykule.

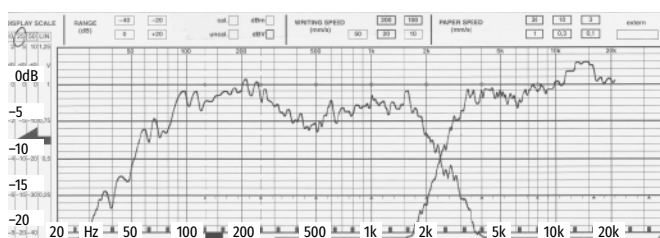
Edycja końcowa

Jednym z „niebezpieczeństw” projektowania CAD PCB jest to, że łatwo jest zmieniać i modyfikować.

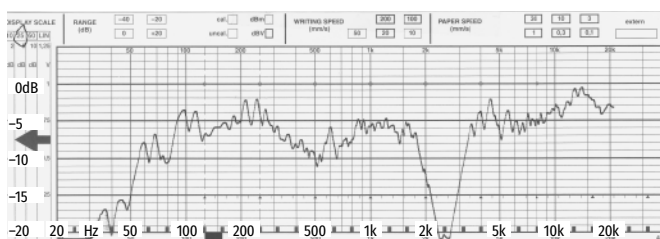
Nie czułem się dobrze z tym, że dwie cewki (L1 i L2) w sekcji niskich częstotliwości są w jednej linii, ponieważ w każdym podręczniku jest napisane „nie rób tego, umieść je pod kątem prostym”. Ta rada dotyczy głównie wysokich częstotliwości, a nie tych niskich, o których tutaj mowa. Ponadto, całkowita ścieżka obwodu magnetycznego była dość długa. W rezultacie nie było widać różnicy, ale i tak zmieniłem płytke. Jedyne co musiało się zmienić to położenie C1a i C2b. **Rysunek 28**



Rysunek 31. Końcowa charakterystyka częstotliwościowa z wykorzystaniem poprawionej zwrotnicy na dwustopowej otwartej podstawie w pół-antechoywalonie



Rysunek 32. Krzywa zwrotnicy pokazująca osobno wyjścia głośnika niskotonowego i wysokotonowego. Zbocze głośnika niskotonowego wynosi około -15 dB na oktawę, a wysokotonowego bliżej -24 dB/okt. Zbocza są monotoniczne, to znaczy, że schodzą w dół w sposób ciągły bez ponownego skoku w górę. Różne tempo tłumienia jest pożądane z powodu różnicy w kierunkowości każdego z przetworników w punkcie podziału zwrotnicy. (Głośnik niskotonowy zaczyna „promieniować”, a wysokotonowy jest prawie wszechkierunkowy)



Rysunek 33. Odpowiedź z głośnikiem wysokotonowym podłączonym antyfazowo w stosunku do głośnika niskotonowego. Jest głęboka i symetryczna, wskazuje na dobrą kontrolę fazy w rejonie podziału zwrotnicy. Jeśli tego nie słyszysz, proponuję znaleźć audiologa

Quiz: Półprzewodnikowe elementy mocy. Historia

Pierwszymi elementami prostującymi prąd przemienny, stosowanymi w sieciach wysokonapięciowych prądu stałego, były:

- ☐ diody kuprytowe
- ☐ diody selenowe
- ☐ lampy rtęciowe

Pierwsze diody germanowe, blokujące napięcie do 200 V i przewodzące prąd do 35 A, pojawiły się w roku:

- ☐ 1933
- ☐ 1948
- ☐ 1952

Tranzystory germanowe o mocy 100 W pojawiły się w roku:

- ☐ 1948
- ☐ 1952
- ☐ 1957

Bipolarne tranzystory krzemowe pojawiły się na rynku w roku:

- ☐ 1952
- ☐ 1957
- ☐ 1961

W 1957 roku firma General Electric wdrożyła do produkcji element pod nazwą SCR. Był to:

- ☐ tyrystor triodowy
- ☐ diak
- ☐ dynistor

W 1960 roku opracowano element, który nazwano GTU. Był to rodzaj:

- ☐ tyrystora
- ☐ tranzystora bipolarnego
- ☐ tranzystora z izolowaną bramką

W roku 1978 firma International Rectifier wprowadziła na rynek MOSFET mocy o parametrach:

- ☐ 100 V/5 A
- ☐ 200 V/10 A
- ☐ 400 V/25 A

Elementy o nazwie IGBT są obecne na rynku od:

- ☐ lat sześćdziesiątych
- ☐ połowy lat osiemdziesiątych
- ☐ od roku 2005

Pierwsze tranzystory mocy z SiC pojawiły się na rynku:

- ☐ ok. 1980 roku
- ☐ ok. 2000 roku
- ☐ ok. 2010 roku

Pierwsze tranzystory mocy z GaN pojawiły się na rynku:

- ☐ ok. 1980 roku
- ☐ ok. 2000 roku
- ☐ ok. 2010 roku

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 17.03.2023.

przedstawia rozmieszczenie elementów na poprawionej płytce, a rysunek 29 zdjęcie gotowej płytki z zamontowanym autotransformatorem L3.

Zastosowano odrobinę „voicing” w stylu BBC, aby skompensować mały rozmiar. Zwróćcie uwagę na lekkie podbicie basu, aby skompensować brak głębokiego basu.

Jest lekkie obniżenie w punkcie podziału częstotliwości, aby skompensować zbyt „wysunięcie” dźwięku do bliskiego monitorowania. Ponadto, występuje delikatny wzrost wysokich częstotliwości na osi, aby skompensować zmniejszenie dyspersji poza osią. Ta odpowiedź nie jest tak gładka jak oryginalna krzywa (rysunek 16b, patrz poprzednia część artykułu), ale w rzeczywistości „brzmi” bardziej gładko.

Innym ciekawym rozwiązaniem jest optymalizacja wartości kondensatorów poprzez ich łączenie równolegle: montowanie ich jeden nad drugim, jak pokazano na rysunku 30. Jest to szczególnie przydatne, gdy trzeba dopasować parę wartości – na przykład na dwóch płytkach zwrotnicy, lub gdy chcemy stworzyć nieosiągalny element, który leży

między dwoma wartościami szeregu E; na przykład 4 μ F jest w granicach 0,5% od 3,3 μ F + 680 nF.

Odstuchanie końcowe

Ostateczna charakterystyka częstotliwościowa, uwzględniająca wprowadzone poprawki, pokazana jest na rysunku 31. Osobne krzywe dla obu przetworników pokazano na rysunku 32. Połączenie antyfazowe pokazane jest na rysunku 33. Może jestem stronniczy, ale miałem właśnie dobrą sesję z *Mini-Monitorami PE* i uważam, że dają one znakomitą przejrzystość i obraz przestrzenny jak na swoją cenę (około 300 funtów). Słuchałem Goldfrapp ze łzami w oczach, prawdziwy wskaźnik efektywnego projektowania audio. (Tu jestem jeszcze bardziej nieobiektywny – Alison Goldfrapp kupiła moje thereminy)! ■

Jake Rothman

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

Quiz: Praktyczny kurs op-ampów

Wzmacniacz operacyjny ma:

- ☐ jedno wejście
- ☐ dwa wejścia
- ☐ cztery wejścia

Współczynnik wzmocnienia op-ampa przy otwartej pętli sprzężenia zwrotnego wynosi:

- ☐ 10×
- ☐ 1000×
- ☐ kilkaset tysięcy razy

Typowy wzmacniacz operacyjny ma impedancję wejściową w zakresie:

- ☐ omów
- ☐ kiloomów
- ☐ megaomów

Typowy wzmacniacz operacyjny ma impedancję wyjściową w zakresie:

- ☐ omów
- ☐ kiloomów
- ☐ megaomów

Op-amp jest zwykle zasilany symetrycznie. Typowe wartości napięcia zasilającego wynoszą:

- ☐ ± 5 V
- ☐ ± 12 V lub ± 15 V
- ☐ ± 30 V

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 10.03.2023.

Maksymalne napięcie wyjściowe (napięcie nasycenia) op-amp wynosi zwykle:

- ☐ kilkanaście miliwoltów poniżej napięcia zasilania
- ☐ kilka woltów poniżej napięcia zasilania
- ☐ jest dokładnie równe napięciu zasilania

Wzmacniacz buforowy (wtórnik napięciowy) ma:

- ☐ impedancję wejściową dużą, a wyjściową małą
- ☐ impedancję wejściową równą wyjściowej
- ☐ impedancję wejściową małą, a wyjściową dużą

Wzmocnienie wzmacniacza buforowego wynosi:

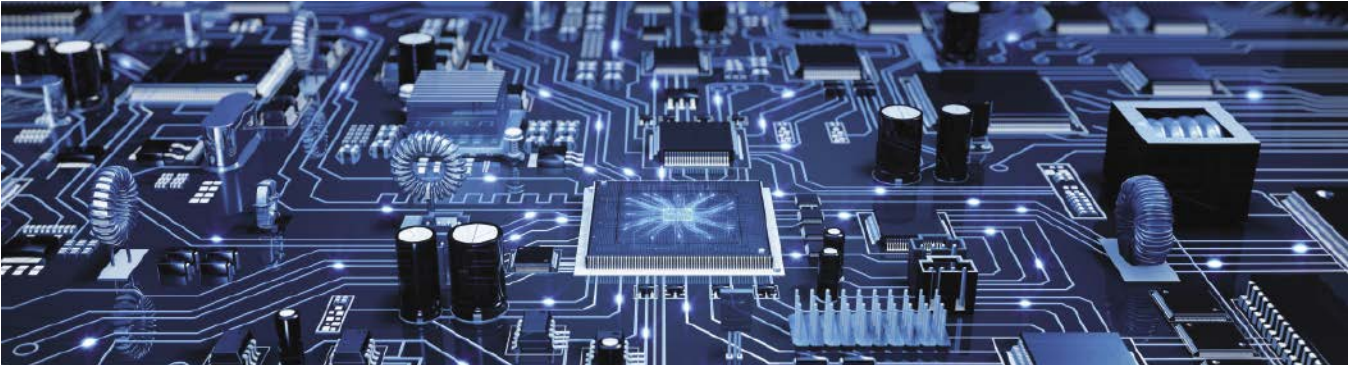
- ☐ dokładnie 1
- ☐ minimalnie poniżej 1, np. 0,99999
- ☐ minimalnie powyżej 1, np. 1,00001

Na wyjściu op-amp o wzmocnieniu A = 200 000 razy jest napięcie 1 V. Ile wynosi różnica napięć na jego wejściach?

- ☐ 5 μ V
- ☐ 0,5 mV
- ☐ 0

Na wejście wzmacniacza buforowego podano napięcie +1 V. Jakie jest napięcie na jego wyjściu?

- ☐ -1 V
- ☐ +1 V
- ☐ -1 V lub +1 V w zależności od sposobu zasilania op-ampa



Powrót do układów tensometrycznych

W numerze grudniowym EdW przyjrzelśmy się tensometrom i sygnałom różnicowym w artykule zainspirowanym przez post na forum EEWeb od Scotta Silera, który napisał: „Jestem inżynierem lotnictwa i pracuję nad pobocznym projektem, polegającym na zbieraniu danych z tensometru za pomocą NI DAQ (urządzenie do akwizycji danych firmy National Instruments). Sygnał wyjściowy wynosi 0...36 mV DC, więc potrzebuję trzykrotnego wzmocnienia, aby lepiej wykorzystać możliwości DAQ. Moje doświadczenie elektroniczne jest ograniczone do kilku zajęć, które odbyłem podczas moich studiów licencjackich. Mam wzmacniacz operacyjny AD8628 i zbudowałem podstawowy, nieodwracający układ z ujemnym sprzężeniem zwrotnym, jak widać na załączonym rysunku (patrz rysunek 1). Napięcie zasilające wynosi 0 do 3 V DC. Wydaje się, że układ nie ma żadnego wzmocnienia, a zamiast tego napięcie wyjściowe jest w rzeczywistości niższe niż wejściowe”.

Otrzymaliśmy e-mail od Johna Ellisa, który uważał, że nie zdiagnozowałem problemów z układem Scotta wystarczająco szczegółowo. John napisał: „Z zainteresowaniem przeczytałem artykuł Iana Bella na temat tensometrów i wzmacniaczy różnicowych.

Pomyślałem jednak, że artykuł mógł dostarczyć nieco więcej informacji na temat problemu, z którym borykał się Scott Siler, co mogło być bardziej pomocne dla Waszych czytelników przy projektowaniu ich własnych układów sensorycznych”.

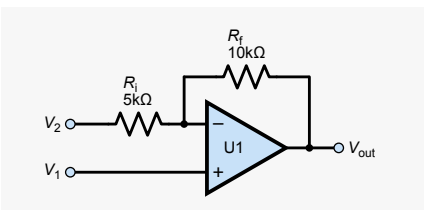
John przedstawił również własną analizę układu, wyprowadzając wzór na niesymetryczne wzmocnienie dla układu na rysunku 1 (co pokazuje, że nie jest to dobry wzmacniacz różnicowy), a także wskazując na możliwe problemy z napięciem wejściowym wspólnym (co może być głównym powodem, że układ

nie zadziałał). Przyjrzymy się tym aspektom projektowania układów ze wzmacniaczami operacyjnymi.

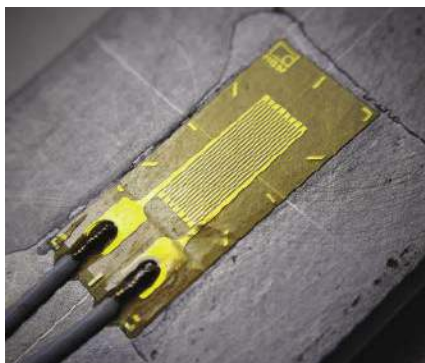
W odpowiedzi na krytykę Johna – jest kilka powodów, dla których niekoniecznie przyglądamy się szczegółom problemów zamieszczanych na forum. Po pierwsze, pełne szczegóły nie zawsze są zamieszczane w wątku dyskusyjnym – w pewnym stopniu tak było w tym przypadku – nie mamy pełnego schematu układu Scotta (choć możemy się domyślać) i nie znamy wszystkich szczegółów użytego ogniwa obciążnikowego. Po drugie, konkretne problemy są często rozwiązywane przez

forumowiczów na długo przed tym, jak artykuł może zostać opublikowany, więc warto spojrzeć na temat szerzej. Po trzecie, często istnieje kilka tematów, które można rozwinąć z zagadnień omawianych w wątku na forum i nie ma miejsca, aby zająć się nimi wszystkimi w jednym artykule. Dlatego mamy wieloczęściowe artykuły i czasami wracamy do tematu na życzenie – co właśnie robimy tutaj.

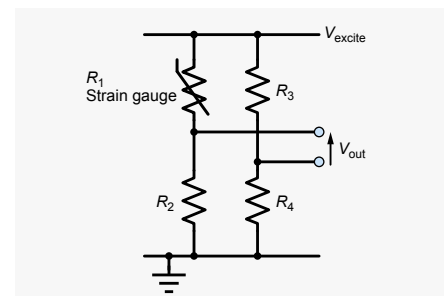
Dla dobra Czytelników, którzy nie mają pod ręką artykułu z EdW 12/2022., krótko zrekapitułujemy wiedzę na temat tensometrów, ogniwa obciążnikowego oraz sygnałów różnicowych i wspólnych, aby zapewnić tło i kontekst dla dyskusji o układach wzmacniacza operacyjnego.



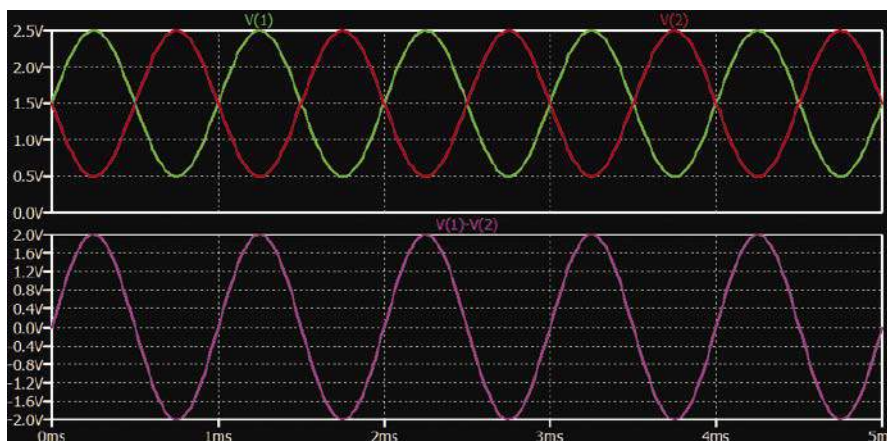
Rysunek 1. Obwód Scotta, omówiony w tym artykule



Rysunek 2. Typowy tensometr zacementowany do mierzonego podłoża



Rysunek 3. Popularny obwód mostka tensometrycznego – w tym przykładzie R1 jest tensometrem, ale stosuje się też inne układy



Rysunek 4. Przykładowy sygnał różnicowy: szczyt 2 V, sygnał sinusoidalny 1 kHz z napięciem wspólnym 1,5 V DC. Górne przebiegi (czerwony, zielony) to poszczególne napięcia. Dolny przebieg (magenta) to sygnał różnicowy

Czujniki tensometryczne

Odształcenie jest miarą deformacji (zmiany rozmiaru lub kształtu) obiektu. Tensometr jest czujnikiem, którego opór elektryczny zmienia się wraz z odkształceniem. Zazwyczaj składa się on z cienkiej, elastycznej folii izolacyjnej, na której umieszczony jest długi pasek przewodzący, zwykle w kształcie zygzaka (patrz **rysunek 2**) z końcówkami do podłączenia do obwodu pomiarowego. Tensometry są często wbudowane w większe urządzenia, zwane „ogniwnami obciążnikowymi”, w których tensometry są przymocowane do specjalnie zaprojektowanego metalowego korpusu, który odkształca się, gdy na urządzenie działa siła. Ogniwa obciążnikowe mają wiele zastosowań przemysłowych w pomiarach siły i masy.

Tensometry są zwykle stosowane w obwodach mostka Wheatstone’a, jak pokazano na **rysunku 3**. Mostek składa się z dwóch równoległych dzielników potencjału, a napięcie wyjściowe jest różnicą napięć obu dzielników. Napięcie różnicowe z mostka nie ma składowej stałej występującej w prostym dzielniku potencjału, więc może być wzmacniane bez składowej stałej powodującej nasycenie wzmacniacza. Układ na **rysunku 3** ma pojedynczy tensometr z trzema rezystorami stałymi, istnieje wiele możliwych scenariuszy pomiarów fizycznych z użyciem wielu tensometrów. Mostek może zawierać jeden, dwa, trzy lub cztery tensometry w zależności od zastosowanej konfiguracji. W niektórych sytuacjach można zastosować dodatkowe rezystory (np. pomiędzy napięciem zasilającym a mostkiem), aby uzyskać pożądane właściwości obwodu.

Sygnały różnicowe i sygnały wspólne

Sygnał różnicowy jest przenoszony na dwóch przewodach innych niż masa, dzięki czemu możemy obserwować napięcie na każdym

przewodzie z osobna (np. V_1 i V_2). Sygnał rzeczywisty jest równy różnicy napięć pomiędzy dwoma przewodami, mierzonych w odniesieniu do masy. Jeśli więc dwa napięcia na przewodach oznaczmy jako V_1 i V_2 , to sygnał różnicowy wynosi $(V_1 - V_2)$. Jednakże, aby w pełni opisać sygnał różnicowy musimy podać zarówno sygnał różnicowy jak i napięcie wspólne, czyli średnią z napięć na dwóch przewodach $(V_1 + V_2)/2$.

Na **rysunku 4** przedstawiono przykładowy sygnał różnicowy – jest to sinusoida o częstotliwości 1 kHz z napięciem szczytowym 2 V i składową napięcia wspólnego, która wynosi 1,5 V DC. Górny wykres pokazuje poszczególne sygnały (V_1 i V_2) i pozwala zaobserwować sygnał wspólny 1,5 V. Dolny wykres pokazuje sygnał różnicowy jako pojedynczy przebieg.

Wzmacniacze różnicowe

Różnicowy wzmacniacz napięcia to układ, który ma dwa wejścia i wzmacnia różnicę napięć między nimi – wzmacnia więc sygnał różnicowy na swoim wyjściu. Jeśli napięcia wejściowe wynoszą V_1 i V_2 , a wzmacnienie napięcia różnicowego wynosi A_d , to napięcie na jego wyjściu wynosi $A_d(V_1 - V_2)$. Jego wyjście może być różnicowe (wtedy określa się go jako wzmacniacz w pełni różnicowy) lub niesymetryczne. Standardowy wzmacniacz operacyjny jest wzmacniaczem różnicowym

z wyjściem niesymetrycznym. Jednak wzmacniacze operacyjne są bardzo często wykorzystywane do budowy układów wzmacniających, które mają zarówno wejście jak i wyjście niesymetryczne – są to dobrze znane konfiguracje odwracające i nieodwracające (patrz **rysunek 5** i dalsza dyskusja). Idealny wzmacniacz różnicowy wzmacnia tylko różnicę pomiędzy V_1 (wejście nieodwracające) i V_2 (wejście odwracające).

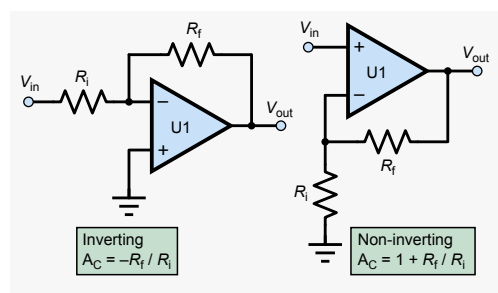
Jednym ze sposobów analizy pracy wzmacniacza różnicowego jest potraktowanie każdego napięcia niesymetrycznego na wejściu jako sygnału odrębnie wzmacnianego. Napięcie wyjściowe jest więc tworzone z nieodwracającego sygnału wejściowego razy jego wzmacnienie (A_1) minus sygnał wejściowy odwracający razy jego wzmacnienie (A_2). Możemy to zapisać w postaci równania w następujący sposób:

$$V_{out} = A_1 V_1 - A_2 V_2$$

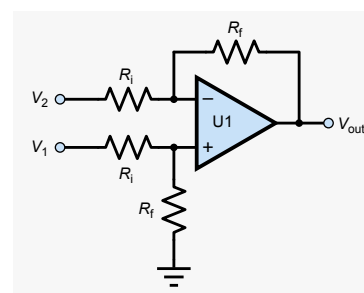
Jeśli $A_1 = A_2 = A_d$ to w takim idealnym przypadku $V_{out} = A_d(V_1 - V_2)$. W idealnym przypadku wzmacnienia dla obu wejść są równe; jednak nie jest tak w rzeczywistych wzmacniaczach różnicowych. Przy odrobinie algebraicznej manipulacji możemy przekształcić powyższe równanie tak, aby zawierało ono człon $(V_1 - V_2)$ plus kilka innych członów. Warto to zrobić, ponieważ możemy wtedy zobaczyć „część idealną” – wzmacnienie różnicowe pomnożone przez $(V_1 - V_2)$, oraz „część błędną” – resztę równania. Otrzymujemy równanie:

$$V_{out} = \frac{A_1 + A_2}{2} (V_1 - V_2) + (A_2 - A_1) \frac{V_1 + V_2}{2}$$

w którym widzimy, że $(A_1 + A_2)/2$ odpowiada A_d w równaniu idealnym i jest średnią z dwóch wzmacnień. Drugi człon zawiera $(V_1 + V_2)/2$, czyli sygnał wejściowy w trybie wspólnym (średnia z napięć wejściowych), który jest mnożony przez wzmacnienie, które określamy jako wzmacnienie sygnału wspólnego A_{cm} . Wzmacnienie sygnału wspólnego jest równe różnicy między dwoma odrębnymi wzmacnieniami, więc w idealnym przypadku, gdy oba wzmacnienia są równe, wartość $(A_2 - A_1)$ wynosi zero i napięcie wejściowe sygnału wspólnego nie ma wpływu na napięcie sygnału wyjściowego.



Rysunek 5. Podstawowe wzmacniacze operacyjne: odwracający (po lewej) i nieodwracający (po prawej)



Rysunek 6. Standardowy wzmacniacz różnicowy operacyjny

Powyższe równanie możemy przepisać jako:

$$V_{out} = A_d (V_1 - V_2) + A_{cm} \frac{V_1 + V_2}{2}$$

Gdzie A_d jest wzmocnieniem napięcia różnicowego, a A_{cm} jest wzmocnieniem napięcia wspólnego.

Im mniejszy jest wpływ sygnału wspólnego na sygnał wyjściowy wzmacniacza różnicowego, tym lepszy jest wzmacniacz. Zdolność wzmacniacza operacyjnego do tłumienia sygnału wspólnego jest wyrażana jako stosunek wzmocnienia sygnału różnicowego do wzmocnienia sygnału wspólnego; nazywa się to „współczynnikiem tłumienia sygnału wspólnego” (CMRR), który często wyraża się w decybelach w następujący sposób:

$$CMRR = 20 \log_{10} \left[\frac{A_d}{A_{cm}} \right] \text{ dB}$$

CMRR układu AD8628 użytego w układzie Scotta jest bardzo wysoki i wynosi 130 dB (wzmocnienie sygnału wspólnego jest ponad trzy miliony razy mniejsze od wzmocnienia sygnału różnicowego).

Wzmacniacze operacyjne

Przydatne jest rozpatrywanie układu Scotta w kategoriach omówionych powyżej ogólnych właściwości wzmacniacza różnicowego. Najpierw jednak krótko opiszemy najbardziej podstawowe i powszechnie stosowane układy wzmacniacza operacyjnego, tj. – konfiguracje wzmacniacza odwracającego lub nieodwracającego, jak pokazano na rysunku 5. W obu przypadkach ujemne sprzężenie zwrotne tworzy para rezystorów, które działają jak dzielnik potencjału, podając ułamek napięcia wyjściowego z powrotem do wejścia odwracającego wzmacniacza. Wzmocnienie tego układu określa stosunek wartości rezystorów, tworzących dzielnik potencjału. Wzmacniacz jako całość jest albo odwracający (wzmocnienie ujemne) albo nieodwracający (wzmocnienie dodatnie) w zależności od tego czy sygnał wejściowy jest kierowany do wejścia odwracającego czy nieodwracającego. Układ Scotta (rysunek 1) jest podobny do układu z rysunku 5 – układ sprzężenia zwrotnego jest taki sam, ale ma dwa wejścia zamiast jednego. W rzeczywistości układ Scotta jest jakby połączeniem dwóch podstawowych konfiguracji wzmacniacza. Jeśli uziemy wejście V_1 w układzie Scotta otrzymamy wzmacniacz odwracający z V_2 jako jego wejściem. Jeśli uziemy wejście V_2 otrzymamy wzmacniacz nieodwracający z V_1 jako jego wejściem.

Superpozycja

Po zidentyfikowaniu tego związku pomiędzy obwodem Scotta i podstawowymi

konfiguracjami wzmacniaczy operacyjnych, możemy użyć pewnej teorii obwodów zwanej „twierdzeniem superpozycji”, aby uzyskać wzór na jego wzmocnienie. Twierdzenie superpozycji mówi, że dla obwodu liniowego możemy znaleźć interesujące nas napięcie lub prąd biorąc każde źródło po kolei i ustawiając wszystkie inne źródła na zero. Obliczając odpowiednie wartości dla obwodu z aktywnym tylko tym jednym źródłem i powtarzając to kolejno dla wszystkich źródeł, a następnie dodając wyniki z poszczególnych źródeł, wyznaczamy poszukiwane wartości.

W przypadku układu Scotta mamy dwa źródła – wejścia V_1 i V_2 . Nie musimy się przejmować zasilaczami – są to napięcia stałe, które nie przenoszą sygnałów i (przynajmniej w idealnym przypadku) nie zmieniają napięcia wyjściowego ani wzmocnienia, dopóki wzmacniacz operacyjny działa normalnie. Napięcie wyjściowe ze wzmocnienia V_2 przy uziemionym V_1 wynosi:

$$\left(-\frac{R_f}{R_i} \right) \cdot V_2$$

Napięcie wyjściowe ze wzmocnienia V_1 , przy uziemionym V_2 wynosi:

$$\left(1 + \frac{R_f}{R_i} \right) \cdot V_1$$

Dodając je do siebie otrzymujemy:

$$V_{out} = \left(1 + \frac{R_f}{R_i} \right) \cdot V_1 - \left(\frac{R_f}{R_i} \right) \cdot V_2$$

Możemy porównać to równanie z nieidealnym wzmacniaczem różnicowym z oddzielnymi wzmocnieniami (A_1 i A_2), jak opisano powyżej; znajdując wzmocnienie różnicowe jako średnią z A_1 i A_2 , otrzymujemy:

$$Ad = \left(\frac{R_i + 2R_f}{2R_i} \right)$$

Zły układ?

Przy wartościach rezystorów Scotta wzmocnienie różnicowe wynosi: $[5 + (2 \times 10)] / (2 \times 5) = 2,5$. Podobnie możemy znaleźć wzmocnienie sygnału wspólnego z $A_1 - A_2$, które jest po prostu jednością. Zatem CMRR wynosi zaledwie 2,5 (ok. 8 dB) i jest ponad milion razy gorsza niż CMRR samego wzmacniacza operacyjnego! Użycie wysokiej klasy wzmacniacza operacyjnego nie gwarantuje wysokiej klasy układu. We wcześniejszych dyskusjach odrzucono układ Scotta jako nie będący wzmacniaczem różnicowym, ale to nie jest do końca prawda – to jest wzmacniacz różnicowy, tylko bardzo kiepski. Wzmocnienie różnicowe nie ma wartości wymaganej przez Scotta (3), ale można to naprawić za pomocą innych rezystorów. Jest mało prawdopodobne,

aby ta nieprawidłowa wartość wzmocnienia różnicowego była powodem, że obwód nie działał tak jak Scott miał nadzieję.

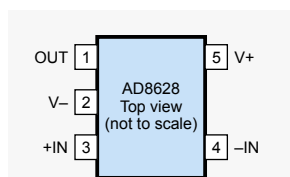
Układ na rysunku 1 może być złym wzmacniaczem różnicowym, ale nie jest to układ bezużyteczny w innych kontekstach. Częstym zastosowaniem tego układu jest jednoczesne wzmocnienie sygnału i przesunięcie jego poziomu DC. Na przykład, rozważmy sygnał o zakresie $\pm 0,5$ V, który musi być wprowadzony do ADC z wejściem 0...5 V. Wzmocnienie o wartości 5 jest wymagane do odwzorowania zakresu $\pm 0,5$ V na 0...5 V, ale konieczne jest również przesunięcie poziomu DC z 0 V do 2,5 V. Konfiguracja układu na rysunku 1, zasilana z napięcia ± 5 V lub większego, może obsłużyć wymagane sygnały. Jeżeli do wejścia nieodwracającego (V_1 na rysunku 1) podłączymy sygnał o stosunku rezystorów spełniającym wymagania $1 + R_f/R_i = 5$, czyli $R_f/R_i = 4$, otrzymamy wymagane wzmocnienie 5 w stosunku do tego wejścia. Wzmocnienie w stosunku do wejścia odwracającego wyniesie wtedy $-4(-R_f/R_i)$. Napięcie wejściowe DC na wejściu V_2 na rysunku 1 zostanie wzmocnione -4 razy, więc jeśli podłączymy to wejście do poziomu DC o wartości $2,5 / -4 = -0,625$ V (na przykład uzyskane z ujemnej linii zasilania) to da wymagane $+2,5$ V DC na wyjściu, gdy sygnał jest 0 V.

Układy alternatywne

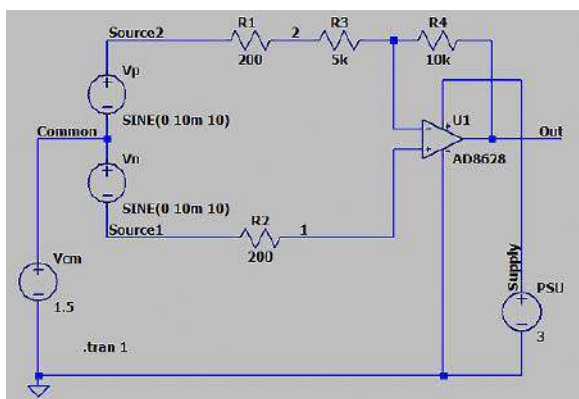
Jak omawialiśmy w poprzednim artykule, istnieje standardowy układ wzmacniacza operacyjnego dla „właściwego” wzmacniacza różnicowego (patrz **rysunek 6**). Porównując go z rysunkiem 1 widzimy, że jest to w zasadzie ten sam układ z dzielnikiem potencjału na wejściu V_1 . Dzielnik tłumy wejście V_1 , aby skompensować większe wzmocnienie przez wejście odwracające z rysunku 1. Wzmocnienie różnicowe układu z rysunku 6 to po prostu R_f/R_i – zauważ, że dwie pary rezystorów mają takie same wartości. Układ ten osiąga wysoki współczynnik CMRR (ograniczony przez wzmacniacz operacyjny) tylko wtedy, gdy pary rezystorów R_1 i R_i mają bardzo dokładnie dopasowane wartości – trudno jest więc osiągnąć bardzo wysoki współczynnik CMRR przy użyciu dyskretnych rezystorów. Innym problemem tego układu jest jego stosunkowo niska impedancja wejściowa, co oznacza, że może on obciążać źródła sygnału, takie jak czujniki, potencjalnie prowadząc do błędów pomiarowych.

Wejście o wspólnym trybie pracy

Jak wskazano powyżej, słabe działanie układu Scotta jako wzmacniacza różnicowego prawdopodobnie nie jest powodem,



Rysunek 7. Układ AD8628 w obudowie 5-końcówkowej TSOT i SOT-23



Rysunek 8. Schemat LTSpice do symulacji układu z rysunku 1

że nie zadziałał. Wzmocnienie 2,5 nadal zapewniłoby rozsądne wyjście (nie „niższe od wejścia”). Bardziej prawdopodobne jest, że powód jest związany z wejściem trybu wspólnego (common-mode) zastosowanym w układzie Scotta. Zakładając, że użyty przez Scotta czujnik wagowy jest podobny do tego z rysunku 3 i że wszystkie cztery rezystory mostka mają w przybliżeniu taką samą wartość (co jest typowe), to sygnał z czujnika wagowego będzie niewielkim napięciem różnicowym (Scott podaje od 0 do 36 mV) z napięciem wspólnym równym połowie napięcia zasilania czujnika wagowego (z powodu prawie równych rezystancji w dzielnikach potencjału). Scott stwierdził, że jego ogniwo obciążnikowe pracuje na 12 V, ale nie podał dalszych szczegółów. Jeśli założymy najprostszy przypadek – ogniwo obciążnikowe jest podłączone do zasilania 12 V względem masy – to jego napięcie wyjściowe w trybie wspólnym wynosiłoby 6 V. Jeśli weźmiemy pod uwagę tylko wyliczone wcześniej wzmocnienia, to możemy się spodziewać, że napięcie wyjściowe układu Scotta będzie 2,5 razy większe od sygnału różnicowego (około 0 do 90 mV), a napięcie wyjściowe w trybie wspólnym będzie jednokrotnie większe od napięcia wejściowego 6 V – czyli 6 V. Tu pojawia się problem – Scott podaje, że układ wzmacniacza operacyjnego AD8628 (rysunek 1) pracuje z zasilania 3 V (a jego maksymalne zasilanie to 5 V). Jeśli nasze założenia dotyczące mocy ogniwa obciążnikowego są poprawne, to napięcie trybu wspólnego w układzie z rysunku 1 byłoby znacznie wyższe niż napięcie zasilania. Typowo wzmacniacze operacyjne nie pracują w takich warunkach.

Wzmacniacze operacyjne mają charakterystykę zwaną „common-mode input range”, która określa skrajne wartości sygnałów wejściowych common-mode, przy których będą

one działać poprawnie. Niektóre wzmacniacze operacyjne (w tym AD8628) radzą sobie z sygnałami trybu wspólnego równymi napięciu zasilania – często określa się to jako zdolność „rail-to-rail input”. Musimy sprawdzić, co arkusz danych mówi o zakresie wejściowym common-mode AD8628 i jego zachowaniu przy wejściach przepięciowych.

AD8628

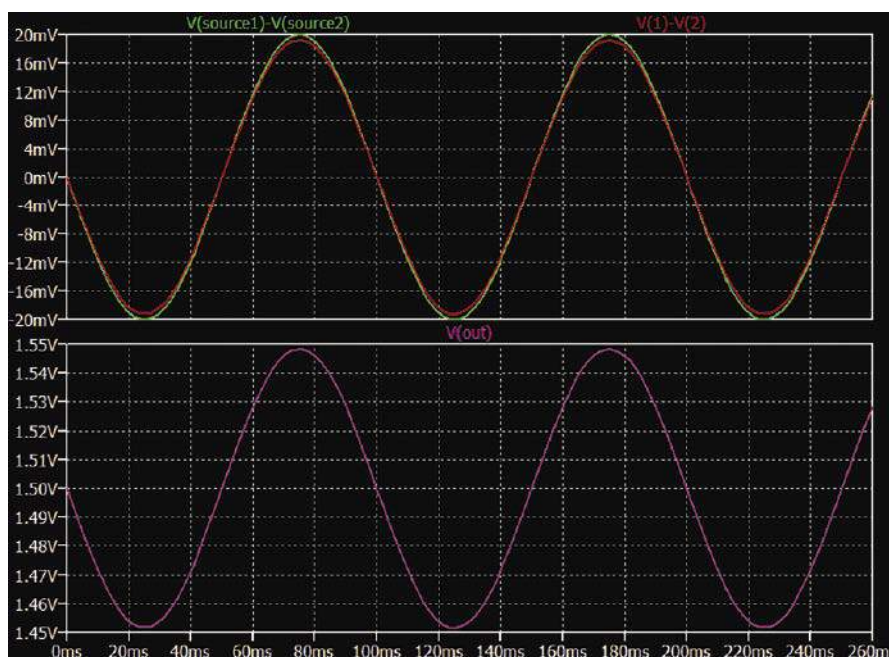
AD8628 jest produkowany przez Analog Devices i opisany jako zero-drift, pojedyncze zasilanie, wejście/wyjście typu rail-to-rail. Nie jest to zwykły wzmacniacz operacyjny, ponieważ swoją wysoką precyzję osiąga dzięki połączeniu technik automatycznego zerowania i choppingu (opatentowanych przez Analog Devices). Analog Devices twierdzi, że pozwala to wzmacniaczowi operacyjnemu utrzymać niskie napięcie offsetowe w szerokim zakresie temperatur i przez cały okres eksploatacji oraz charakteryzuje się on niskim poziomem szumów w porównaniu z innymi wzmacniaczami z automatycznym zerowaniem. Dostępny jest

w różnych obudowach – popularnych 8-wyprowadzeniowych SOIC i MSOP oraz mniej typowych (dla wzmacniaczy operacyjnych) 5-wyprowadzeniowych TSOT i SOT-23 (stosowanych przez Scotta). Patrz rysunek 7, gdzie przedstawiono rozkład wyprowadzeń dla obudowy 5-końcówkowej. Wzmacniacze tensometryczne są wymienione wśród sugerowanych zastosowań tego układu.

Wzmacniacze z auto zerowaniem i choppingiem zostały zaprojektowane, aby poradzić sobie z trudnościami związanymi z precyzyjnym wzmocnieniem sygnałów o małej częstotliwości (typowo kilka herców i poniżej). W zależności od aplikacji, wzmacniacze te mogą być również wymagane do obsługi znacznie wyższych częstotliwości, co zwiększa trudności w tworzeniu projektów odpowiednich układów. Sygnały o niskiej częstotliwości występują zazwyczaj w pewnych implikacjach czujnikowych, takich jak tensometry, gdzie mierzony sygnał zmienia się powoli. Problemem są szумы i przesunięcia o niskiej częstotliwości, które są nieodłącznym elementem wzmacniaczy elektronicznych i które sprawiają, że zwykle wzmacniacze operacyjne nie nadają się do tych zadań. Rozwiązanie wykorzystuje układy przełączające i, jak wskazuje podwójne podejście zastosowane w AD8628, istnieją dwa podstawowe sposoby implementacji tych układów – auto-zero i chopping.

Przebieg

Część opisu AD8628 zatytułowana „rail-to-rail input” wskazuje, że może on obsługiwać sygnały wejściowe bliskie zasilaniu, ale karta katalogowa zawiera więcej szczegółów. Stwierdza, że jeśli ktoś z wejść przekroczy



Rysunek 9. Wyniki symulacji układu z rysunku 8

Pliki LTSpice omawiane w tym artykule są dostępne do pobrania ze strony PE (www.epemag3.com).

napięcie którejś z szyn zasilania o więcej niż 0,3 V, to duże prądy mogą płynąć przez diody układu ESD (electrostatic discharge), zabezpieczające przed wyładowaniami we wzmacniaczu. Diody te są podłączone pomiędzy wejściami a każdą z szyn zasilających i normalnie są spolaryzowane wstecznie. Jednakże mogą one zostać spolaryzowane w kierunku przewodzenia, jeśli napięcie wejściowe przekroczy napięcie zasilania. Jeśli w wyniku tego płynie nadmierny prąd, układ może ulec trwałemu uszkodzeniu. Dla układów, w których może wystąpić przepięcie, karta katalogowa zaleca zastosowanie szeregowych rezystorów na wejściach, aby ograniczyć prąd do 5 mA. Nie wiemy wystarczająco dużo o układzie Scotta, aby mieć pewność co do poziomu prądu, który mógł być dostarczony przez tensometr – chociaż Scott nie wspominał nic o zniszczeniu jego wzmacniacza operacyjnego.

Oprócz uszkodzeń, gdy prądy nie są ograniczone, przepięcia w trybie wspólnym mogą powodować dziwne zachowania w niektórych wzmacniaczach różnicowych, w szczególności napięcie wyjściowe może nagle przeskoczyć w kierunku przeciwnym do szyny zasilania – zjawisko znane jako odwrócenie fazy napięcia wyjściowego. W przypadku AD8628 ograniczenie wejściowego prądu przepięcia do 5 mA powinno temu zapobiec. Ponownie, nie mamy wystarczających informacji, aby stwierdzić, czy wystąpiło to w układzie Scotta. Innym

potencjalnym problemem związanym z przepięciem w trybie wspólnym, który dotyczy sytuacji, gdy występuje ono chwilowo, jest to, że niektóre wzmacniacze auto-zero potrzebują długiego czasu, aby odzyskać sprawność po usunięciu przeciążenia. Karta katalogowa układu AD8628 podaje, że ma on znacznie krótszy czas powrotu do stanu wyjściowego niż typowe wzmacniacze auto-zero. Ten problem nie jest istotny w przypadku Scotta, ponieważ domniemane przeciążenie jest trwałe.

Symulacja

LTspice miał model dla AD8628, więc możemy spróbować symulacji. Schemat pokazano na rysunku 8. Tutaj ustawiliśmy wejście z określonymi sygnałami wspólnym i różnicowymi – 1,5 V napięcie wspólne, i 40 mV peak-to-peak sygnał różnicowy otrzymujemy dla wartości podanych na schemacie, ale można je łatwo zmienić. Takie podejście jest stosowane, ponieważ zapewnia bezpośrednią kontrolę sygnału wejściowego oraz dlatego, że nie znamy dokładnego układu czujnika wagowego Scotta. Zaczynamy od napięcia wspólnego na poziomie połowy napięcia zasilania wzmacniacza, tak aby układ działał poprawnie i abyśmy mogli potwierdzić obliczone wcześniej wzmocnienia. Źródła mają pewną impedancję wyjściową (R_1 i R_2 po 200 Ω), ponieważ czujnik wagowy miałby niezerową impedancję wyjściową, ale użyte wartości są wyborem

arbitralnym, gdyż nie znamy szczegółów dotyczących ogniwa obciążnikowego.

Wyniki symulacji przedstawiono na rysunku 9. Widzimy, że sygnał wejściowy ($v(1)-v(2)$) jest nieco niższy od sygnału źródłowego ($v(\text{źródło } 1)-v(\text{źródło } 2)$) z powodu obciążenia źródła przez układ wzmacniacza. Amplituda wejściowa sygnału różnicowego wynosi około 38,3 mV peak-to-peak, a wyjściowa około 96,7 mV peak-to-peak, co potwierdza obliczone powyżej oczekiwane wzmocnienie różnicowe 2,5. Wyjście jest wyśrodkowane na poziomie 1,5 V – równym wejściowemu napięciu wspólnemu, co potwierdza wzmocnienie sygnału wspólnego równe 1. Eksperymenty z napięciem wspólnym wykazały prawidłowe działanie przy wejściach trybu wspólnego 0,05 V i 2,95 V, potwierdzając pracę w trybie rail-to-rail. Przy 6 V napięcie wyjściowe jest w zasadzie nasycone przy zasilaniu 3 V, a sygnał pokazuje się na poziomie kilkudziesięciu mikrowoltów. Nie możemy być pewni, że ta symulacja dokładnie odwzorowuje sytuację w układzie Scotta, ale jeśli wyjście trybu wspólnego jego ogniwa obciążnikowego wynosiło 6 V, to jest to najbardziej prawdopodobny powód, dla którego jego układ nie działał. ■

Ian Bell

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

Quiz: Miks wiedzy z różnych artykułów w EdW 03/2023

Aby polakierować obudowę kolumn głośnikowych, trzeba oszlifować ścianki ze sklejk. Szlifierki tarczowe są do tego bezużyteczne. Dlaczego?

- ☐ szlifierka ma za duże obroty
- ☐ należy szlifować w poprzek włókien
- ☐ należy szlifować wzdłuż włókien

Z czego składa się przetwornik ultradźwiękowy Langevina?

- ☐ z kapsułek mylarowych
- ☐ z płytek krzemowych
- ☐ z płytek piezoelektrycznych

Czyszczenie ultradźwiękami odbywa się wskutek generowania pęcherzyków:

- ☐ próżniowych
- ☐ gazowych
- ☐ transportujących

Aby uniknąć powstawania fali stojących w myjce ultradźwiękowej stosuje się:

- ☐ tłumienie amplitudy drgań płytki piezoelektrycznej
- ☐ okresowe zmiany częstotliwości drgań
- ☐ damping interferencyjny

Moc przetwornika w myjce ultradźwiękowej wynosi?

- ☐ 30 do 50 W
- ☐ 1 do 2 W
- ☐ 200 do 300 W

Ogniwo Peltiera w programowalnym termoregulatorze wymaga zasilania prądem:

- ☐ 5 do 40 A
- ☐ 1 do 3 A
- ☐ 0,1 do 0,5 A

Do programowalnego termoregulatora wybrano termistor 10 k Ω , a nie 100 k Ω , ze względu na:

- ☐ większą stabilność temperaturową
- ☐ lepsze dopasowanie do obwodu włączenia termistora
- ☐ mniejszą podatność na wpływy pól elektromagnetycznych



Mostek typu H stosowany do sterowania kierunkiem obrotów silnika krokowego składa się z:

- ☐ czterech kondensatorów
- ☐ czterech kluczy
- ☐ czterech rezystorów

Jeśli silnik krokowy ma 200 kroków na jeden obrót, a jeden krok składa się z 16 mikrokroków, to kąt obrotu przypadający na jeden mikrokrok wynosi:

- ☐ 0,1125°
- ☐ 0,011125°
- ☐ 1,125°

Dlaczego drukarki 3D używają napięć zasilania znacznie większych niż znamionowe silnika krokowego?

- ☐ dla poprawy dokładności ruchu
- ☐ dla osiągnięcia dużych przyspieszeń
- ☐ dla zwiększenia szybkości obrotów silnika

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 03.03.2023.

Praktyczny kurs op-ampów

Tak jest, chodzi o wzmacniacz operacyjny, prawdziwy koń pociągowy elektroniki. Jeśli już potrafisz zaprzęgnąć Arduino do zapalenia lampki, a może nawet do generowania sygnałów, to zdziwisz się, że można to zrobić 100 razy taniej i prościej przy użyciu op-ampów. Pozostaniemy przy nazwie op-amp, stosowanej w literaturze angielskiej, bo to dźwięczna nazwa i takiej nazwy używa autor tego kursu Jos Verstraten, guru holenderskich pasjonatów elektroniki, twórca niezliczonych projektów i setek opublikowanych artykułów. Ten kurs to jego prawdziwy popis belferski. Wykonasz serię ćwiczeń pod jego nadzorem i op-ampy będziesz miał w jednym palcu. Do ćwiczeń wystarczy mieć płytkę stykową, parę op-ampów 741 oraz garść rezystorów i kondensatorów. Jeszcze dwa źródła zasilania i multimetr. Wszystko znajdziesz łatwo i tanio kupisz w sklep.avt.pl. W AVT możemy przygotować kompletny zestaw do tego kursu, jeśli będzie takie zapotrzebowanie. Zaczynamy!

Redakcja EdW

1. Wprowadzenie do wzmacniacza operacyjnego – 741

Używając prostych eksperymentów, nauczymy Cię jak projektować wszelkiego rodzaju układy z op-ampem 741. Układy stanowiące podstawę elektroniki analogowej, czyli tej części elektroniki, w której pracuje się z napięciami o wszystkich możliwych wartościach. Mimo rewolucji cyfrowej, analogowa elektronika nadal odgrywa ważną rolę. Nie można pomyśleć o urządzeniu elektronicznym bez elektroniki analogowej w nim. Czy jest to starożytny 741, czy bardziej współczesny typ nie ma znaczenia. Wszystkie op-ampy działają identycznie i jeśli po przebrnięciu przez ten kurs potrafisz zaprojektować dość skomplikowany układ z 741-kami, to możesz zrobić to samo z każdym innym op-ampem.

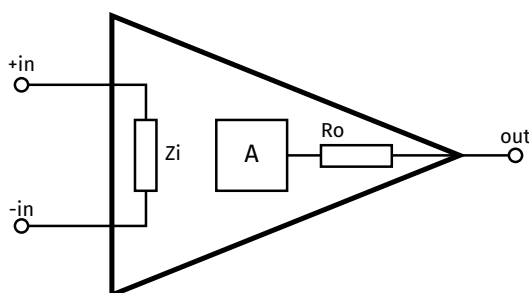
Uczymy się przez eksperymentowanie

W przypadku niektórych eksperymentów kończymy rozdział pustym wykresem, w którym możesz wpisać zmierzone wartości. W ten prosty i nieco zabawny sposób nauczysz się rozumieć korelację między napięciami wejściowymi i wyjściowymi.

Oczywiście można też bez zbędnych analiz budować opisane układy z użyciem innych op-ampów we własnych projektach różnych urządzeń. Ponieważ op-ampy, na tym poziomie rozważań, mają prawie identyczne właściwości, twoje układy na pewno będą działać.

Co to jest op-amp?

Pytanie może zbędne, ale ten kurs powinien zacząć się od odpowiedzi na to pytanie. Wzmacniacz operacyjny to układ scalony złożony z rezystorów, diod, tranzystorów i czasami kondensatora, którego wewnętrzna struktura sprawia, że idealnie nadaje się do wzmacniania napięć.



Rysunek 1. Trójkątny symbol op-ampa, z rysowanymi w niego głównymi wielkościami charakterystycznymi

Oprócz niezbędnych połączeń zasilających i niektórych połączeń pomocniczych, op-amp ma zawsze dwa wejścia i jedno wyjście. Jedno wejście nazywane jest **wejściem nieodwracającym** lub **dodatnim**, drugie zaś **wejściem odwracającym** lub **ujemnym**. Op-amp jest w zasadzie **wzmacniaczem różnicowym**: wzmacnia różnicę napięcia, która występuje między jego dwoma wejściami.

Właściwości

Wzmacniacz operacyjny charakteryzuje się szeregiem właściwości, które schematycznie przedstawiono na **rysunku 1**.

Współczynnik wzmocnienia

Po pierwsze, masz do czynienia ze współczynnikiem wzmocnienia A, który wskazuje jak bardzo op-amp wzmacnia różnicę napięć pomiędzy dwoma wejściami. Wartość współczynnika wzmocnienia A samego op-ampa (który często nazywany jest współczynnikiem wzmocnienia w otwartej pętli, czyli wzmocnienia bez wpływu obwodów zewnętrznych układu scalonego) jest bardzo duża. Im większa tym lepiej, dlatego na rynku są op-ampy, które wzmacniają aż pół miliona razy! Normalny op-amp, taki jak 741, jest nieco skromniejszy: jego wzmocnienie w otwartej pętli wynosi około 200 000 razy. Jest to niewyobrażalnie duża wartość: różnica napięcia 1 mV między wejściami teoretycznie doprowadziłaby do napięcia wyjściowego 200 V, gdyby op-amp potrafił generować takie napięcia! W praktyce należy więc zwykle podejmować działania ograniczające wzmocnienie układu.

Impedancja wejściowa

Drugą ważną wielkością jest impedancja wejściowa, reprezentowana przez Zi na rysunku 1. Jest to opór pozorny pomiędzy dwoma wejściami. Podłączanie omomierza nie ma sensu, wspomniana wielkość jest obecna, ale nie można jej zmierzyć taką metodą. Zi nie jest więc obecny w układzie scalonym w postaci rezystora, ale wynika z działania stopnia wejściowego układu. Impedancja wejściowa określa jednak realne obciążenie, jakiego doświadczają obwody poprzedzające układ scalony.

Jeśli Zi jest niskie, to IC będzie pobierał duży prąd z poprzedzającego go obwodu, co oczywiście nie jest takie idealne. Stąd ta wielkość w op-ampach powinna być również jak największa. Obecnie znane są wzmacniacze operacyjne o Zi 15.000 MΩ! Są to tzw.

wzmacniacze BiMOS, zbudowane z FET-ów na wejściu. Tu również 741 jest dość skromny: jego impedancja wejściowa wynosi tylko 2 MΩ.

Rezystancja wyjściowa

Ostatnim ważnym parametrem jest rezystancja wyjściowa R_o . To również jest rezystor pozorny, nie występujący jako taki, ale wyrażający się w momencie obciążenia wyjścia układu scalonego. To obciążenie (np. kolejny stopień) żąda prądu od op-ampa i w efekcie na R_o powstaje spadek napięcia. Napięcie wyjściowe op-ampu spada. Im mniejsza rezystancja wyjściowa, tym mniejszy spadek napięcia i tym mniej problemów może się pojawić. Stąd wymóg, aby R_o było jak najmniejsze. W przypadku 741 obowiązuje skromna wartość 75 Ω.

Podsumowanie

Podsumowując, można stwierdzić, że wzmacniacz operacyjny jest określony przez:

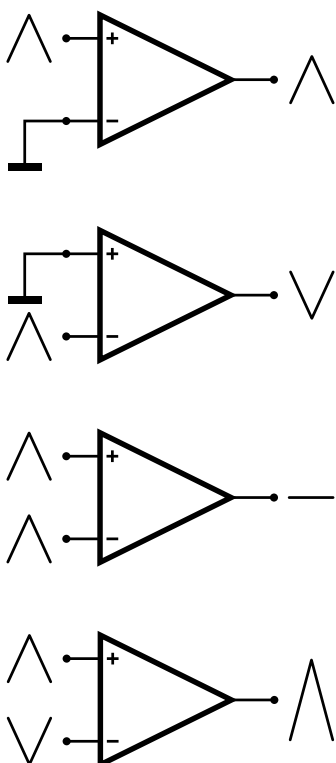
- jego współczynnik wzmocnienia A , który powinien być jak największy;
- jego impedancję wejściową Z_i , która również powinna być jak największa;
- jego rezystancję wyjściową R_o , która powinna być jak najmniejsza.

Zasilanie op-ampa

Op-amp jest zwykle zasilany symetrycznie. Oznacza to, że do dodatniego pinu zasilania podłączasz napięcie dodatnie względem masy $+V_s$, a do ujemnego pinu zasilania napięcie ujemne względem masy $-V_s$, które są równe co do wartości bezwzględnej. Typowe wartości to: ± 12 V lub ± 15 V. Pod wpływem napięć na wejściach, napięcie na wyjściu op-ampa może zmieniać się pomiędzy dwoma wartościami napięć zasilających. Mówi się wtedy o symetrycznym wyjściu układu scalonego.

Cztery warianty sterowania

Z faktu, że op-amp ma dwa wejścia wynika, że możnaysterować układ na cztery różne sposoby. Zostały one narysowane na rysunku 2.



Rysunek 2. Cztery różne możliwości sterowania wzmacniaczem operacyjnym

- Jeśli wejście ujemne jest podłączone do masy, a sygnał, który ma być wzmocniony jest podawany na wejście dodatnie, to wyjście będzie zmieniać się w fazie z napięciem wejściowym. Jeśli napięcie wejściowe wzrośnie, to napięcie wyjściowe również wzrośnie.
- Jeśli połączysz wejście dodatnie z masą i podasz sygnał wejściowy na wejście ujemne, to napięcie wyjściowe będzie w przeciwfazie do wejściowego. Jeśli napięcie wejściowe wzrośnie, napięcie wyjściowe spadnie.
- Trzeci rodzaj sterowania zakłada równe napięcia na obu wejściach. Ponieważ op-amp jest wzmacniaczem różnicowym i w tym przypadku nie ma różnicy napięć pomiędzy obu wejściami, to wyjście nie dostarczy żadnego sygnału.
- Czwarty rodzaj sterowania zakłada przeciwstawne sygnały na obu wejściach. Jeśli napięcie na wejściu dodatnim wzrośnie, to w tym samym czasie napięcie na wejściu ujemnym spadnie. Różnica napięć między wejściami będzie wtedy zawsze maksymalna, a więc i napięcie wyjściowe będzie maksymalne.

Oczywiście, narysowane obwody są bardzo zgrubnymi uproszczeniami. Jeśli użyjesz op-ampa w ten sposób, to duży współczynnik wzmocnienia układu scalonego spowoduje, że przy najmniejszej różnicy napięć między wejściami wyjście „skoczy” do poziomu jednego z napięć zasilających. (Dla większości op-ampów napięcie wyjściowe może osiągać wartości maksymalne o kilka woltów poniżej napięć zasilających, tylko dla specjalnych op-ampów, nazywanych „rail-to-rail” tzw. swing napięcia wyjściowego osiąga poziom zaledwie 100 mV poniżej napięć zasilających.)

Praktyczne układy wzmacniające ze wzmacniaczami operacyjnymi są możliwe tylko dzięki dołączeniu obwodów sprzężenia zwrotnego. Stosując rezystory między wyjściem a wejściem ujemnym, można zmniejszyć A op-ampa do wartości użytkowych.

Pierwszy eksperyment

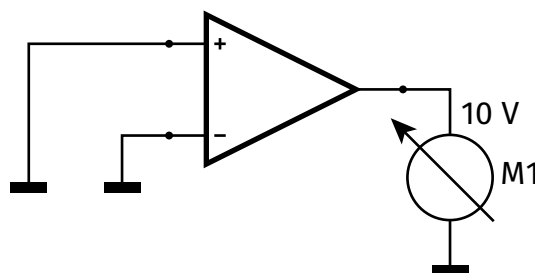
W poprzednim punkcie stwierdziliśmy, że op-amp nie dostarcza żadnego sygnału na wyjściu, jeśli na oba wejścia są podane identyczne napięcia. Można to wypróbować podłączając oba wejścia do masy. Schemat tego układu jest narysowany na rysunku 3. Podłącz miernik do wyjścia 5.

Napięcie wyjściowe układu nie jest zerowe, jak powinno być, ale na przykład ok. +10 V lub -8 V. Jak to możliwe?

Przesunięcie

W tym pierwszym eksperymencie figle płyta zła cecha wzmacniaczy operacyjnych: offset, czyli przesunięcie.

Przyjrzyj się rysunkowi 4, prostemu przedstawieniu stopnia wejściowego op-ampa. Dwa tranzystory T_1 i T_2 są połączone w układzie wzmacniacza różnicowego, czyli ze wspólnymi złączami emiterowymi. Baza jednego tworzy wejście dodatnie op-ampa, baza drugiego wejście ujemne. Oba tranzystory przewodzą, więc między bazą a emiterem jest normalne napięcie przewodzenia U_{be} wynoszące około 0,7 V.



Rysunek 3. Pierwszy eksperyment, który pozwala zbadać, czy układ rzeczywiście dostarcza 0 V na wyjściu, gdy nie ma napięcia różnicowego między wejściami

W teorii U_{be1} powinno być równe U_{be2} . W praktyce nigdy tak nie jest, bo ten parametr zależy od wielu czynników, takich jak budowa układu scalonego (IC), temperatura otoczenia itp. Ta różnica w U_{be} jest nazywana offsetem op-ampa.

Dwa U_{be} są interpretowane przez IC jako część sygnału wejściowego. Jeśli uziemisz oba wejścia, to IC „myśli”, że na wejściu dodatnim jest napięcie U_{be1} , a na ujemnym U_{be2} . Pierścieniowa różnica między tymi dwoma wielkościami jest wzmacniana 200 000 razy i prowadzi do pełnego dodatniego lub ujemnego nasycenia wyjścia układu. Stąd +10 V lub -8 V na wyjściu układu scalonego.

Kompensacja przesunięcia

Przesunięcie jest bardzo niepożądane, dlatego prawie wszystkie op-ampy mają dwa zaciski, do których można podłączyć potencjometr regulacyjny i skompensować to przesunięcie. – Podłącz potencjometr 10 kΩ do pinów 1, 5 (Offset Null), suwak potencjometru do $-V_s$. Obracaj ostrożnie suwak potencjometru. W pewnym momencie napięcie wyjściowe przeskoczy z jednej polaryzacji do drugiej. Ten punkt to prawidłowe ustawienie kompensacji przesunięcia.

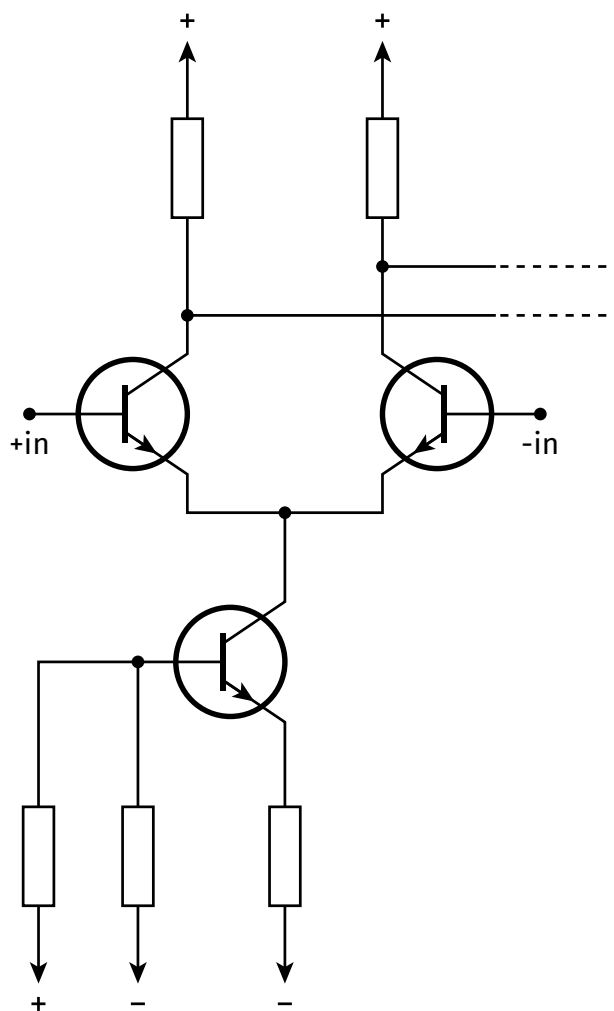
Jest to punkt przejściowy od przesunięcia dodatniego do ujemnego. To, że nie można wyregulować wyjścia IC na zero jest logiczne. Nawet szczątkowe przesunięcie o ułamek mV jest wzmacniane 200 000 razy i wysyła wyjście układu do jednego z napięć zasilających.

Jeśli w kolejnych eksperymentach nauczysz się jak okiełznać to duże wzmocnienie w otwartej pętli, to w praktyce okaże się, że kompensacja offsetu jest bardzo płynna.

Wnioski

Wzmacniacz operacyjny jest bardzo dobrym wzmacniaczem, bo ma bardzo duże wzmocnienie. Ta właściwość ma jednak tę wadę, że z gołym układem nie można zrobić zbyt wiele. W prawie wszystkich przypadkach, włączając obwody rezystorów pomiędzy wejściami i wyjściami, zmusisz op-amp do zrobienia tego, czego oczekujesz.

I to właśnie zamierzamy zrobić w kolejnych ćwiczeniach, abyś stał się znakomitym konstruktorem układów stosujących op-ampy!



Rysunek 4. Wyjaśnienie zjawiska offsetu wynika z tego schematu stopnia wejściowego op-ampa

2. Op-amp jako wzmacniacz buforowy

Wstęp

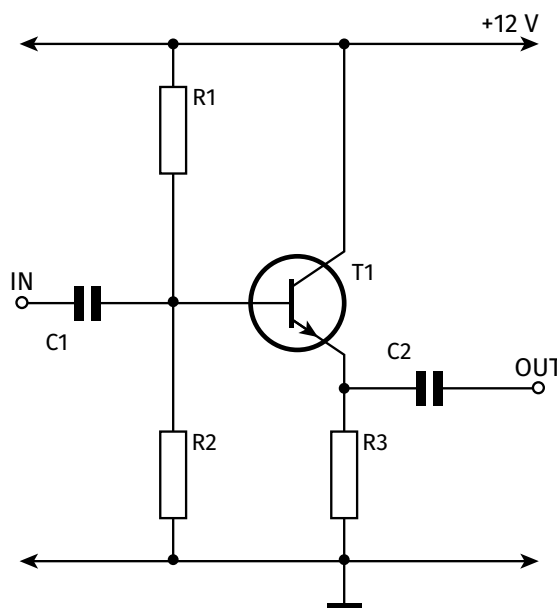
Wzmacniacze buforowe to układy, które w jak najmniejszym stopniu wpływają na sygnał podawany na wejście, ale pełnią rolę swojego transformatora impedancji. Impedancja wejściowa bufora jest bardzo duża, impedancja wyjściowa bardzo zaś mała. Wzmocnienie napięciowe jest równe jeden.

Wzmacniacze buforowe są stosowane tam, gdzie bezpośrednie sprzężenie jednego układu z drugim może powodować problemy. Najprostszym przykładem bufora jest wtórnik emiterowy, patrz rysunek 5, często używany do przesyłania przez długi kabel małych sygnałów wrażliwych na zakłócenia, takich jak napięcie wyjściowe mikrofonu.

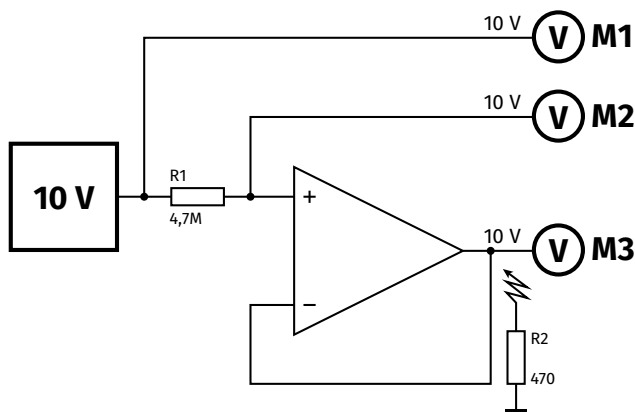
Układ bufora op-ampa jest prosty: sygnał, który ma być buforowany jest podawany na wejście dodatnie, wejście ujemne jest bezpośrednio połączone z wyjściem. Rysunek 6 przedstawia schemat testowy, który można wykorzystać do badania właściwości takiego bufora.

Układ z rysunku 5 wymaga sześciu elementów. Z op-ampem znacznie łatwiej jest zrobić wzmacniacz buforowy, co jest dość logiczne, jeśli weźmiemy pod uwagę dość wysoką impedancję wejściową i raczej niską wyjściową układu scalonego.

Do wejścia podłącz napięcie stałe ze źródła – zasilacza na zakres ± 10 V. Trzy obwody pomiarowe są również na zakresie ± 10 V. Rezystor



Rysunek 5. Najprostszym przykładem wzmacniacza buforowego jest wtórnik emiterowy



Rysunek 6. Układ wykorzystany do badania właściwości wzmacniacza buforowego op-amp

R1 o wartości 4,7 M Ω jest włączony w obwód wejściowy, rezystor R2 o wartości 470 Ω jest włączony w taki sposób, że w razie potrzeby można go na krótko doprowadzić do kontaktu z wyjściem układu scalonego. Te dwa rezystory nie są potrzebne do pracy samego bufora, ale są pomocniczymi elementami do pomiaru trzech ważnych parametrów układu: współczynnika wzmocnienia, impedancji wejściowej i impedancji wyjściowej.

Nasz drugi eksperyment

Po włączeniu zasilania widać, że we wszystkich trzech punktach pomiarowych jest to samo napięcie.

Co można z tego wynioskować?

• Wzmocnienie napięcia

Po pierwsze, że wzmocnienie napięciowe układu jest równe jeden. 2 V na wejściu daje 2 V na wyjściu.

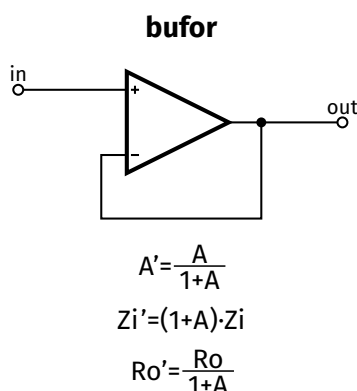
• Impedancja wejściowa

Po drugie, że impedancja wejściowa układu jest najwyraźniej bardzo wysoka. Wejście żąda tak małego prądu od źródła napięcia wejściowego, że na bardzo dużym rezystorze R1 nie występuje mierzalny spadek napięcia. Gdyby tak było, M2 wskazywałby niższe napięcie niż M1. Z tej różnicy napięć mogliśmy obliczyć prąd przez R1 i w konsekwencji impedancję wejściową bufora.

• Impedancja wyjściowa

Ustaw napięcie wejściowe układu na +5 V, podłącz na krótko wolną końcówkę rezystora R2 do wyjścia i obserwuj miernik M3. Brak wskazań jakiegokolwiek zmiany.

Z tego możemy wynioskować, że rezystancja wyjściowa bufora jest bardzo mała. Przecież obciążenie o rezystancji 470 Ω powoduje przepływ prądu o wartości około 10 mA. Ten prąd pochodzi z wyjścia op-ampa, a więc płynie również przez opór wewnętrzny wyjścia op-ampa. Gdyby ten opór był powiedzmy 10 Ω , to podłączenie rezystora obciążenia spowodowałoby spadek napięcia na wyjściu o 0,1 V, co zauważylibyśmy na M3. Teraz miernik nie wskazuje żadnej zmiany, z czego wynika, że opór wewnętrzny bufora jest bardzo niski.



Rysunek 7. Podsumowanie specyfikacji wzmacniacza buforowego

Wnioski

W skrócie, bufor ma wzmocnienie napięciowe równe 1, bardzo dużą impedancję wejściową i bardzo małą impedancję wyjściową. Teoretycznie sprowadza się to do pomnożenia wartości impedancji wejściowej „gołego” op-ampa przez współczynnik wzmocnienia op-ampa i podzielenia wartości rezystancji wyjściowej „gołego” op-ampa przez ten sam współczynnik wzmocnienia.

Objaśnienie działania

Jak można wytłumaczyć takie zachowanie? Jak to się dzieje, że ten niewiarygodnie wysoki współczynnik wzmocnienia 200 000 został zredukowany do prostego przeniesienia bez zmian napięcia z wejścia do wyjścia? Najprostszym sposobem fizycznego zademonstrowania tego jest założenie na chwilę, że op-amp działa bardzo wolno, lub innymi słowy, że napięcie na wejściu niekoniecznie powoduje natychmiast napięcie na wyjściu. Jeśli przy takim założeniu nagle podłączysz do wejścia dodatniego napięcie o wartości, powiedzmy, 1 V, to wyjście najpierw przez chwilę pozostanie zerowe. Jest to wtedy również napięcie na wejściu ujemnym. Czyli między dwoma wejściami jest różnica napięć nie mniejsza niż 1 V, a op-amp zadziała na pełną moc.

Dla tej różnicy napięć zadziała wzmocnienie 200 000 razy. Napięcie wyjściowe bardzo szybko wzrośnie. Jednak natychmiast ten wzrost napięcia znajdzie się na wejściu ujemnym. W związku z tym, różnica napięć między dwoma wejściami staje się coraz mniejsza.

Układ z op-ampem zawsze dąży do uzyskania minimalnej różnicy napięć między dwoma wejściami op-ampa.

W pewnym momencie napięcie wyjściowe wzrosło do prawie +1 V. Wtedy między wejściem dodatnim a ujemnym jest różnica napięcia rzędu znikomo małego ułamka wolta. Po wzmocnieniu przez 200 000 to niezmiernie małe napięcie wytwarza napięcie wyjściowe, które jest prawie równe napięciu na wejściu.

Nie do końca równa się.

Z omówienia działania bufora można wynioskować, że wzmocnienie nie jest tak naprawdę dokładnie równe 1, ale o włos mniejsze. Przecież to bardzo małe napięcie różnicowe zapewnia stan ustalony układu. Dlaczego więc zawsze mówi się, że bufor ma wzmocnienie równe 1? Bo odchylenie jest naprawdę znikome. Ściśle biorąc, wzmocnienie napięciowe bufora jest równe:

$$A' = A / (1 + A)$$

gdzie A' to wzmocnienie bufora, a A to wzmocnienie samego op-ampa. Wystarczy wypełnić ten wzór dla 741 z jego A wynoszącym 200 000:

$$A' = 200.000 / (1 + 200.000)$$

$$A' = 200.000 / 200.001$$

$$A' = 0,99999999...$$

Brak różnicy napięć między wejściami

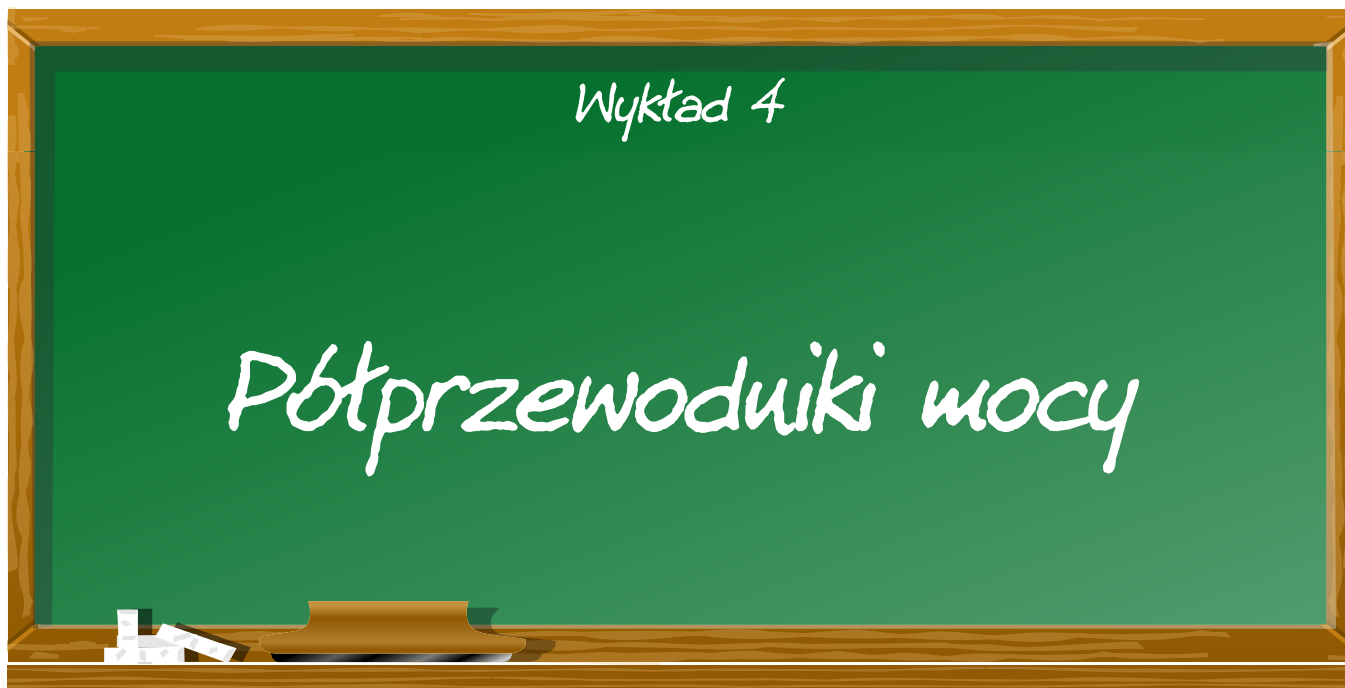
Z tej dyskusji wynika bardzo ważna cecha op-ampów: ponieważ różnica napięć między wejściami nie jest nawet tak naprawdę mierzalna, mówi się, że op-amp zawsze zapewnia, że jego dwa wejścia mają to samo napięcie. Działanie wszystkich układów budowanych z op-ampów można wytłumaczyć za pomocą tej prostej zasady, wystarczy poszukać jej w kolejnych eksperymentach.

Podsumowanie

Wzory na rysunku 7 wyraźnie grupują wszystkie właściwości wzmacniacza buforowego z op-ampem. Symbole z primem oznaczają właściwości całego układu, symbole bez prima oznaczają właściwości samego op-ampa. ■

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 listów od uczniów, nauczycieli i studentów. Dla nich jest ta rubryka.



Napisałem tytuł „półprzewodniki mocy” i natychmiast chciałem przekreślić i poprawić na „półprzewodnikowe elementy mocy”, ale zaraz potem pomyślałem sobie – przecież po angielsku używa się nazwy „power semiconductors” i każdy rozumie, że chodzi o elementy, a nie o materiały półprzewodnikowe. To dlaczego po polsku tak nie można? Można.

Temat wykładu dotyczy elementów półprzewodnikowych przeznaczonych do pracy z dużymi mocami. Z reguły chodzi o pracę przełącznikową jako niesterowane lub sterowane klucze, czyli elementy o dwóch stanach pracy: przewodzenie, nieprzewodzenie. Określenie duże moce jest bardzo szerokie, oznacza prądy w stanie przewodzenia od kilku amperów do kiloamperów i napięcia w stanie blokowania od kilkuset woltów do kilowoltów. A moce w obwodach sterowanych tymi kluczami mogą się mieścić w przedziale od watów do setek megawatów. Celem tego wykładu jest uporządkowanie podstawowej wiedzy o rodzajach półprzewodnikowych elementów mocy, w kontekście określonych segmentów ich zastosowań w energoelektronice.

Współczesna elektronika kojarzy nam się przede wszystkim z układami scalonymi zawierającymi miliony, a nawet miliardy tranzystorów. Przyzwyczajeni też jesteśmy do myśli, że postęp w elektronice wyznacza prawo Moore’a, czyli coraz większa skala integracji, coraz mniejszych wymiarowo tranzystorów, pobierających coraz mniejsze moce i zasilanych coraz mniejszym napięciem. W tym kontekście może się wydawać, że temat tego wykładu dotyczy peryferyjnej dziedziny elektroniki, a może nawet już nie elektroniki, ale elektryki. Zgoda, jest to część elektroniki (nazywana energoelektroniką), która operuje dużymi mocami, prądami i napięciami, czyli wkracza na teren elektryki. Nie ma jednak zgody na twierdzenie, że jest to peryferyjna dziedzina elektroniki. Przeciwnie, w świetle wyzwań, jakie stoją przed światem walczącym o zeroemisyjną cywilizację, energoelektronika ma do odegrania kluczową rolę. Znaczenie elektroniki mocy, opartej na aplikacjach półprzewodnikowych elementów mocy, będzie z roku na rok coraz większe. Ponieważ ten cykl wykładów jest adresowany głównie do studentów i uczniów, przyszłych profesjonalnych elektroników, to temat tego wykładu nie jest dla nich peryferyjny, gdyż dotyczy strategicznie kluczowego sektora gospodarki, w którym wielu z nich znajdzie zatrudnienie.

Zachowamy tradycyjną strukturę wykładu, tj. podział na dwie części: historyczną i merytoryczną. Ze względu na dużą różnorodność elementów ograniczymy się tylko do ich krótkich prezentacji. Jest to wykład przeglądowy, prezentujący półprzewodniki mocy w ujęciu panoramicznym. Na opisy pasjonującego świata elektronów i dziur, na wyjaśnienie jak działają poszczególne elementy mocy, przyjdzie czas innym razem.

A. Historia, czyli lektura lekka na rozgrzewkę

Pierwszymi elementami prostującymi prąd przemienny, stosowanymi w sieciach wysokonapięciowych prądu stałego, były lampy rtęciowe. Wynalezione w 1914 roku, dominowały w latach 1920–1940. Koniec ich zastosowań nastąpił dopiero w połowie lat siedemdziesiątych, gdy ostatecznie zostały wyparte przez elementy półprzewodnikowe, głównie tyrystory. Pierwszymi półprzewodnikowymi elementami prostowniczymi znacznej mocy były diody kuprytowe (z tlenku miedzi – 1927 r.) i selenowe (1933 r.). To jednak

zamierzchną prehistoria, o diodach kuprytowych i selenowych już dawno nikt nie pamięta. Za punkt startu współczesnej historii elementów półprzewodnikowych mocy możemy uznać rok 1952, gdy pojawiły się na rynku diody germanowe blokujące napięcie do 200 V i przewodzące prądy do 35 A. Od tego momentu zaczynamy kreślić naszą historyczną oś czasu (rysunek 1). Daty przypisane poszczególnym wdrożeniom na tym wykresie należy traktować jako przybliżone, gdyż w wielu przypadkach wprowadzenie na rynek nowego produktu to proces rozciągnięty w czasie i wybór konkretnej daty zależy od przyjętego kryterium. Już druga pozycja na osi czasu może być dyskusyjna.

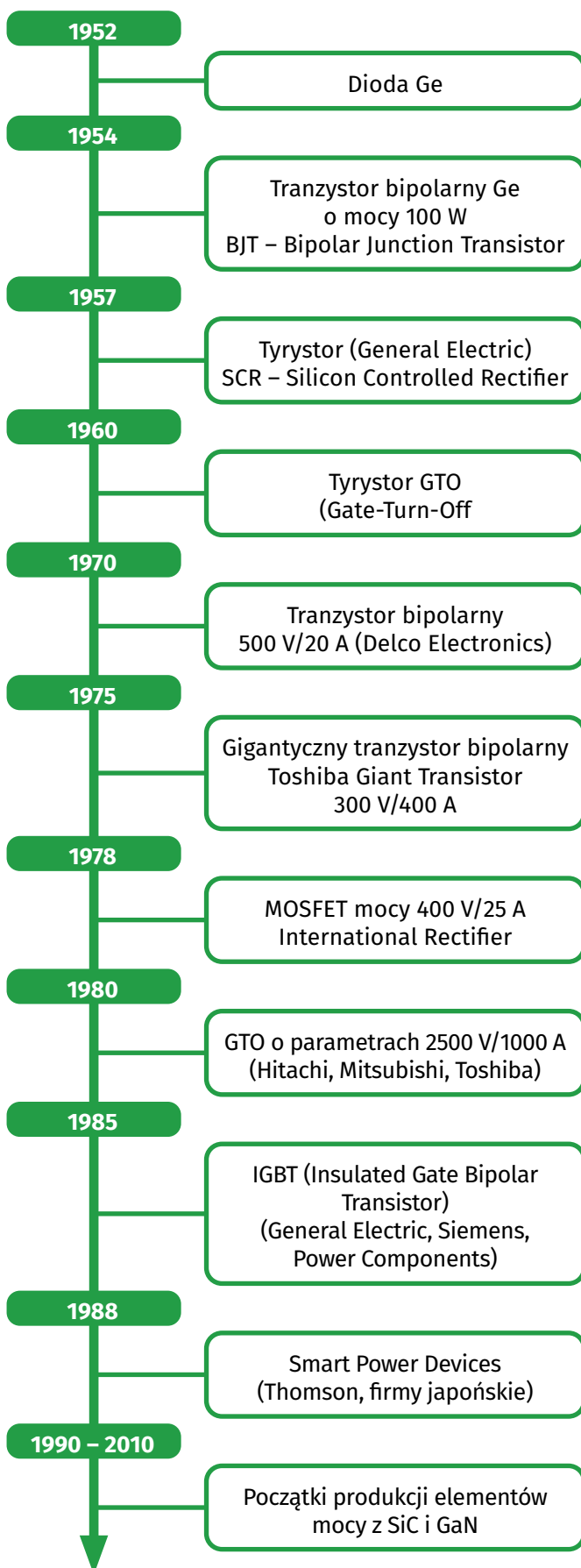
Już w roku 1952 produkowano tranzystory germanowe o prądzie kolektora do 100 mA i taki tranzystor w niektórych zastosowaniach (na przykład we wzmacniaczu słuchawkowym) był nazywany tranzystorem mocy. W naszej opowieści nie o takie moce chodzi, więc przesuwamy datę na osi czasu do roku 1954, gdy pojawiły się tranzystory Ge o mocy 100 W.

Już wtedy wiadomo było, że lepsze pod względem osiąganych mocy będą tranzystory krzemowe, jak tylko zostanie opanowana technologia ich wytwarzania. Wynika to z modelu pasmowego półprzewodników, a dokładniej z szerokości pasma zabronionego W_g . Dla Ge $W_g = 0,7$ eV i maksymalna dopuszczalna temperatura złącza p-n wynosi 85°C, natomiast dla Si $W_g = 1,12$ eV, co przekłada się na znacznie wyższą temperaturę pracy złącza p-n, wynoszącą aż 155°C. Jak już jesteśmy przy tym temacie, to spójrzmy na **tablicę 1**. Zobaczymy, że są materiały o jeszcze szerszym paśmie zabronionym (tzw. półprzewodniki szerokopasmowe – GaN, SiC), które mogą pracować przy jeszcze wyższych temperaturach złącza p-n. Stąd z tymi materiałami od dziesiątków lat łączono duże nadzieje na postęp w elektronice mocy. W ostatnich latach te nadzieje zaczynają się spełniać. Ale zajmijmy się naszą osią czasu, a do tematu elementów z półprzewodników innych niż krzem jeszcze wrócimy.

Bipolarne tranzystory krzemowe pojawiły się na rynku w roku 1957 i już w latach pięćdziesiątych osiągały większe moce niż tranzystory germanowe. W rozwoju bipolarnych tranzystorów krzemowych mocy odnotujemy na osi czasu dwa wydarzenia: rok 1970 – tranzystor o parametrach 500 V/20 A (firma Delco Electronics) i rok 1975 – Toshiba Giant Transistor o parametrach 300 V/400 A. To prawdziwy gigant.

Wróćmy jeszcze do lat pięćdziesiątych. Pamiętamy zapewne z ostatniego wykładu, że W. Shockley od początku lat pięćdziesiątych intensywnie pracował nad strukturą czterowarstwową n-p-n-p. Wreszcie w roku 1957 firma General Electric wdrożyła do produkcji ten element pod nazwą SCR – Silicon Controlled Rectifier, znany powszechnie jako tyrystor, również w odmianach o nazwach diak, triak. Wadą tyrystora jest „zatrząskanie się” (latch-on), tj. po przejściu w stan przewodzenia nie może być wyłączony sygnałem podanym na bramkę, wyłącza się dopiero po wyłączeniu napięcia zasilania. W roku 1960 opracowano tyrystor przełączany sterowaniem bramki ze stanu przewodzenia do stanu blokowania. Ten rodzaj tyrystora nazwano GTO (Gate Turn-Off).

Dotychczas opisywaliśmy rozwój elementów mocy wywodzących się ze struktury tranzystora bipolarnego. W latach sześćdziesiątych pojawiła się druga ścieżka rozwoju elementów mocy,



Rysunek 1. Historyczna oś czasu w rozwoju półprzewodnikowych elementów mocy

wywodząca się ze struktury tranzystora MOS. I to był przełom. Uproszczo się sterowanie bramką, wzrosły częstotliwości pracy i otworzyły się możliwości zwiększania zakresu natężenia prądu przez łączenie równoległe elementów. Pierwsze MOSFETy mocy pojawiły się już w latach siedemdziesiątych. Zaznaczmy na osi czasu rok 1978 – International Rectifier wprowadza na rynek MOSFET mocy o parametrach 400 V/25 A.

Z czasem MOSFETy mocy zdominowały rynek. Obecnie są najczęściej stosowanymi elementami mocy. Wielką rolę odgrywa też struktura hybrydowa tranzystora bipolarnego ze sterowaniem bramką odizolowaną od półprzewodnika – IGBT (Insulated Gate Bipolar Transistor), obecny na rynku od połowy lat osiemdziesiątych. Dalszy rozwój półprzewodników mocy przebiegał głównie w kierunku opracowywania modułów integrujących w jednej obudowie kilka elementów mocy i układ sterujący – smart modules (moduły inteligentne) oraz w kierunku zastąpienia Si przez półprzewodniki szerokopasmowe SiC, GaN. Pierwsze eksperymentalne tranzystory polowe z GaN zademonstrowano w 1993 r., a na rynku są dostępne od 2010 roku. W tym samym czasie obserwujemy początki elementów mocy z SiC. W roku 1993 opracowano laboratoryjnie diody Schottky z SiC, w 2008 roku pojawiły się na rynku JFET o napięciu 1200 V, a w 2011 były dostępne pierwsze MOSFETy z SiC o napięciu 1200 V.

Tablica 1. Porównanie niektórych właściwości trzech materiałów półprzewodnikowych stosowanych do wytwarzania elementów mocy

Właściwości	Si	SiC	GaN
Szerokość przerwy energetycznej (pasma zabronionego) E_g (eV)	1,12	3,26	3,5
Ruchliwość elektronów μ_n (cm^2/Vs)	1400	900	1500...2000
Ruchliwość dziur μ_p (cm^2/Vs)	600	100	200
Natężenie pola elektrycznego przebicia E_b (V/cm) $\times 10^6$	0,3	3	3
Przewodność cieplna (W/cm)	1,5	4,9	1,3
Nasylenie szybkości unoszenia elektronów V_s (cm/s) $\times 10^7$	1	2,7	2,7
Względna przenikalność elektryczna	11,8	9,7	9,5

B. Meritum, czyli sedno sprawy

Zacznijmy od **klasyfikacji**. Musimy wprowadzić jakiś ład w tej mnogości różnych elementów mocy. Każda klasyfikacja powinna opierać się na określonych kryteriach podziału. Najbliżej realiów katalogowych będzie podział na diody, tranzystory i tyrystory z uwzględnieniem trzech rodzajów materiałów: Si, GaN, SiC. Dalsze kryteria klasyfikacji dotyczą różnic fizyko-technologicznych. Ostatecznie otrzymujemy schemat klasyfikacji pokazany na **rysunku 2**. Za istotne kryterium podziału przyjęliśmy **rodzaj materiału**. Z dwóch powodów:

- po pierwsze, warto zaakcentować ten fakt, że inaczej niż w mikroelektronice bazującej wyłącznie na krzemie, w dziedzinie elementów mocy do głosu doszły również inne materiały półprzewodnikowe. Wprawdzie ponad 90% elementów mocy wytwarza się nadal z krzemu, ale znaczenie azotku galu (GaN) i węgla krzemu (SiC) z roku na rok rośnie.
- po drugie, różnice w parametrach materiałowych (tablica 1), szczególnie wartości przerwy energetycznej (szerokości pasma zabronionego) istotnie wpływają na parametry elementów.

Przyjrzyjmy się dokładnie tablicy 1 i skomentujmy, jaki wpływ mogą mieć różnice właściwości materiałowych Si, SiC, GaN na parametry osiągane w elementach mocy wytwarzanych z tych materiałów.

Si vs SiC vs GaN

Zacznijmy od porównania płytek, w których wytwarza się elementy. Płytki Si do produkcji elementów mocy to nie te same płytki, które stosuje się do produkcji układów scalonych. Pewnie pamiętamy, że w produkcji układów scalonych używa się ogromne talerze o prawie półmetrowej średnicy, uzyskane z pocięcia monokryształu krzemu hodowanego metodą Czochralskiego. Krzem Czochralskiego nie nadaje się do produkcji elementów wysokonapięciowych, bo jest wyciągany z tygla, a więc ma sporo zanieczyszczeń pochodzących z tygla. Krzem dla elementów wysokonapięciowych musi być znacznie czystszy. Monokryształy Si o wymaganej czystości wytwarza się metodą beztyglową, nazywaną metodą strefy topionej (FZ-float zone). Taki krzem jest droższy niż otrzymywany metodą Czochralskiego i płytki mają mniejsze średnice (do 300 mm).

Nie wdając się w szczegóły technologii wytwarzania kryształów SiC i GaN, powiedzmy tylko że jest to technologia jeszcze trudniejsza i droższa niż dla Si beztyglowego (FZ). I średnice płytek są mniejsze, dla SiC do 200 mm, a dla GaN do 100 mm. Te różnice wpływają istotnie na poziom cen wytwarzanych elementów.

Przechodząc do porównania właściwości fizycznych Si, GaN i SiC, w kontekście ich wpływu na parametry elementów mocy, posłużmy się tablicą 1, lub dla lepszej ekspresji **rysunkiem 3**. Najważniejsza różnica to szerokość pasma zabronionego. SiC i GaN są nazywane półprzewodnikami szeroko-przerwowymi, bo mają 3× szersze pasmo zabronione (przerwę energetyczną). Stąd wynika, że dla wywołania przeskoku elektronu z pasma walencyjnego do pasma przewodnictwa potrzebny jest 3 razy większy kwant energii. Dlatego elementy z SiC i GaN mogą pracować w zakresie temperatury złącza nawet powyżej 200°C, podczas gdy dla krzemu graniczna temperatura wynosi około 150°C. Wyższa dopuszczalna temperatura pracy to możliwość rozpraszania większej mocy. Innym parametrem dającym istotną przewagę SiC i GaN nad Si jest 10× większe natężenie pola elektrycznego wywołującego przebicie. Trzecią przewagą obu materiałów szeroko-przerwowych nad krzemem jest możliwość regulowania w znacznie szerszym zakresie poziomu domieszkowania warstw typu p i typu n. Dodajmy do tego 3-krotnie większą przewodność cieplną SiC w porównaniu z Si, co ma kapitalne znaczenie dla elementów mocy, w których trzeba odprowadzać ze złącza p-n duże ilości ciepła. Z kolei przewagą specyficzną dla GaN jest większa ruchliwość elektronów i możliwość osiągnięcia znacznie większych szybkości unoszenia elektronów. Ta właściwość predysponuje GaN szczególnie do zastosowań w elementach mocy pracujących w zakresie większych częstotliwości. Główne obszary zastosowań elementów wytwarzanych z Si, SiC i GaN przedstawia **rysunek 4**. Warto na koniec zwrócić uwagę, że ostatnie 10 lat stało pod znakiem rozwoju technologii SiC i GaN, ale przyszłość być może należy do diamentu, który ma szerokość przerwy energetycznej

5,5 eV. Przejdźmy teraz do panoramicznego przeglądu poszczególnych elementów według klasyfikacji na rysunku 1. Będą to krótkie prezentacje, w których postaram się zwrócić uwagę na cechy charakterystyczne, różniące właściwości określonych rodzajów elementów od innych. Zaczynamy od diod, po czym przyjdzie kolej na tranzystory i wreszcie różne rodzaje tyrystorów.

Diody

Najistotniejsze jest zrozumienie z czego wynikają różnice aplikacyjne diod PIN i diod Schottky’ego.

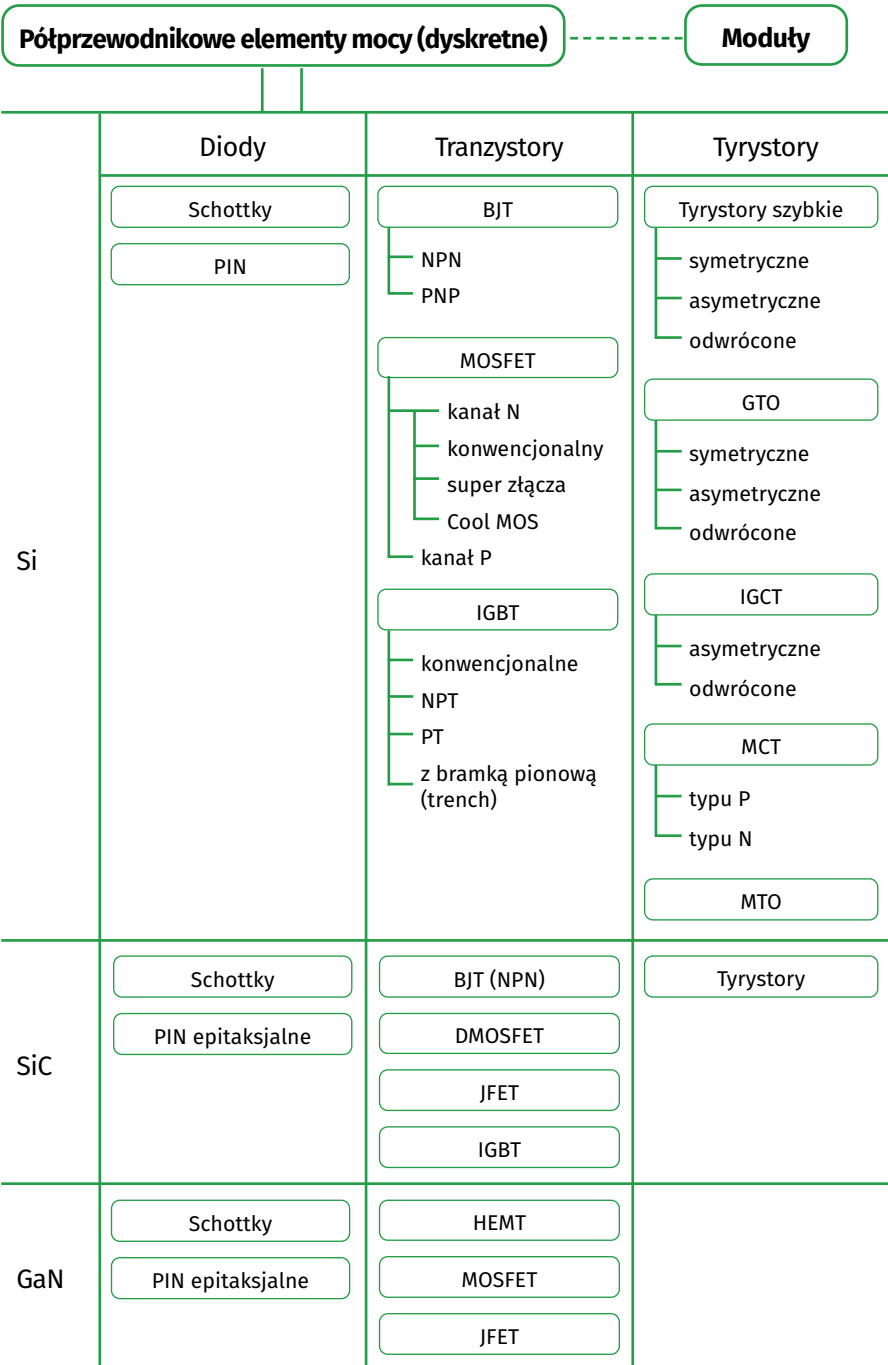
W latach sześćdziesiątych diody krzemowe zdominowały rynek diod prostowniczych, wypierając diody germanowe. Mimo rozwoju w ostatnich 20 latach technologii SiC i GaN, nadal diody są produkowane głównie z krzemu. Kierunek rozwoju diod wyznacza dążenie do poprawy trzech parametrów:

- zwiększanie napięcia w kierunku wstecznym
- zwiększanie prądu przewodzenia
- zwiększanie częstotliwości przełączania lub zmniejszanie czasu przełączania przy pracy impulsowej.

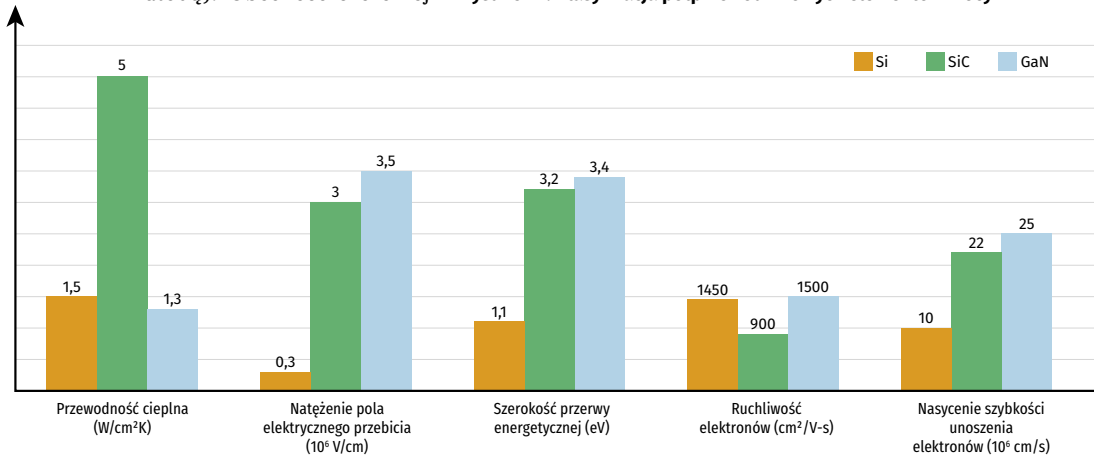
Wysokonapięciowe diody mocy są fizycznie złączami p – n (diody PIN) lub m – s (diody Schottky’ego).

Diody PIN

Cechą charakterystyczną diody mocy ze złączem p-n jest obecność co najmniej jednej słabo domieszkowanej warstwy wysokorezystywnej (warstwa n⁻ na rysunku 5a) pomiędzy złączem p-n a powierzchniami kontaktowymi struktury (anodą i katodą). Obecność szerokiej



Rysunek 2. Klasyfikacja półprzewodnikowych elementów mocy



Rysunek 3. Graficzne porównanie właściwości Si, SiC i GaN

warstwy słabo domieszkowanej (stąd jej nazwa I od intrinsic, czyli półprzewodnik samoistny) jest niezbędna dla osiągnięcia wysokich napięć pracy w kierunku zaporowym.

Dioda Schottky'ego

W diodzie Schottky'ego funkcję prostowniczą pełni złącze metal – półprzewodnik (m – s) przy kontakcie metalowym. Nie ma złącza p-n, a więc nie ma w ogóle warstwy p (rysunek 5b), ale podobnie jak w diodzie PIN, dla osiągnięcia wysokich napięć w kierunku zaporowym, niezbędna jest szeroka warstwa wysokooporowa. Jest to bardzo słabo domieszkowana warstwa n⁻, której kontakt z elektrodą katody odbywa się poprzez silnie domieszkowaną warstwę n⁺, zapewniającą styk omowy o bardzo małej rezystancji. Dioda Schottky'ego, w przeciwieństwie do diody PIN jest elementem unipolarnym, tj. działa na nośnikach prądu jednego rodzaju. Ma to swoje zalety i wady. Podczas działania diody Schottky'ego nie dochodzi do gromadzenia się w warstwie n⁻ nadmiarowych par elektron – dziura, dlatego rezystancja tej warstwy nie maleje wraz ze wzrostem prądu stałego. Wobec tego uzyskanie akceptowalnie małego spadku napięcia na diodzie w stanie przewodzenia jest możliwe tylko przy odpowiednio małej grubości warstwy n⁻, co nakłada znaczne ograniczenia na wielkość dopuszczalnego napięcia wstecznego. Stąd dla diod Schottky'ego maksymalne napięcie wsteczne nie przekracza 1000 V, podczas gdy dopuszczalna gęstość prądu przewodzenia jest o rząd wielkości mniejsza niż w przypadku diody PIN zaprojektowanej dla tego samego napięcia wstecznego. Brak akumulacji nadmiarowych par elektron – dziura w warstwie n⁻ ma też pozytywne skutki – niskie straty energii i krótki czas przełączania podczas przejścia ze stanu przewodzenia do stanu zaporowego, bo brak ładunku nadmiarowego zmagazynowanego w warstwie n⁻ skutkuje niewielkim i krótkim impulsem prądu w stanie nieustalonym. Zatem dioda Schottky'ego jest znacznie „szybsza” niż dioda PIN. Diody Schottky'ego wytworzone z SiC lub GaN zachowują zalety krzemowych diod Schottky'ego pod względem wysokich częstotliwości pracy.

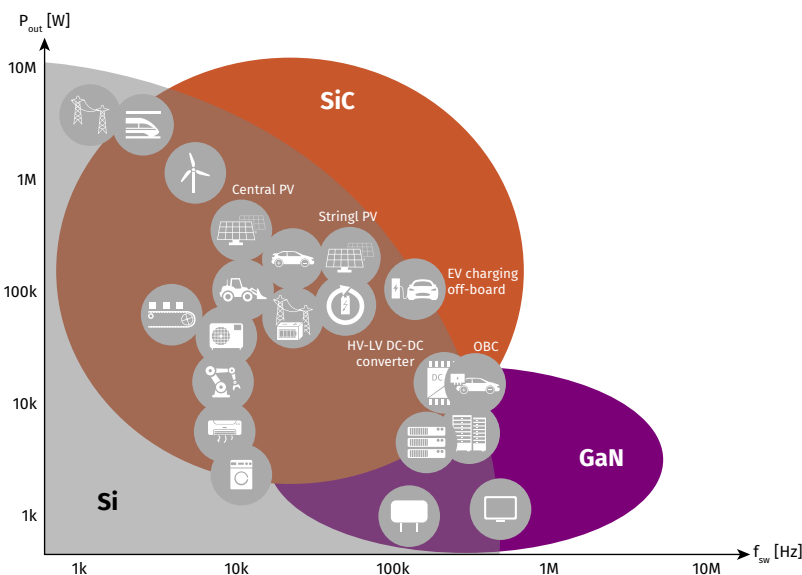
W katalogach rozróżnia się cztery rodzaje diod opartych na złączu p-n:

- diody prostownicze (rectifier diodes)
- diody szybkiego odzyskiwania (fast recovery diodes)
- diody lawinowe (avalanche diodes)
- diody tłumiące stany przejściowe (transient-voltage-suppression diodes)

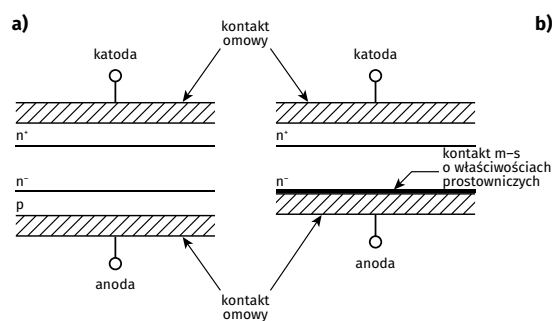
Diody prostownicze są stosowane głównie w układach prostowniczych pracujących z niską częstotliwością, zwykle jest to sieć 50 Hz. Dla tak niskiej częstotliwości parametry impulsowe diody nie mają istotnego znaczenia. W szczególności ładunek Q_{rr} (reverse recovery charge) może być duży. Liczą się tylko minimalne straty mocy w stanie przewodzenia i zapewnienie zadanego napięcia wstecznego. Natomiast dla **diod szybkiego odzyskiwania** najważniejszym parametrem jest jak najmniejszy ładunek Q_{rr}, tj. ładunek, który trzeba odprowadzić po przełączeniu ze stanu przewodzenia do wstecznego, by dioda odzyskała właściwości zaporowe. **Diody lawinowe** są przystosowane do pracy w zakresie przebiecia lawinowego złącza p-n. **Diody tłumiące stany przejściowe** (diody TVS) to te same diody lawinowe, ale z dodatkową cechą: określoną wartością napięcia przebiecia lawinowego, która jest nieprzekraczalną wartością napięcia bezpieczną dla chronionego tą diodą urządzenia. Zwykła dioda lawinowa nie ma takiej funkcji, ponieważ maksymalna wartość napięcia wstecznego, przy której zaczyna się proces przebiecia lawinowego nie jest wielkością znormalizowaną.

Tyrystor

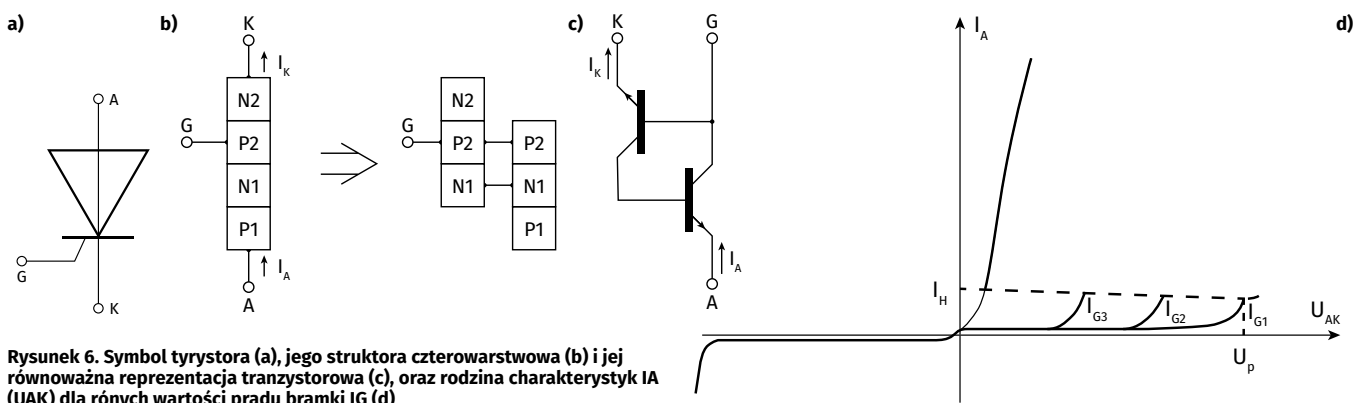
Nazwa tyrystor, wskazująca na analogię funkcjonalną tego elementu do tyratronu (przełącznikowa lampa gazowana), oznacza półprzewodnikowy element dwustanowy o trzech lub więcej złączach (czterech lub więcej warstwach), który może być przełączany ze stanu blokowania do przewodzenia i odwrotnie. Istnieje wiele rodzajów tyrystorów. Najbardziej rozpowszechniony jest tyrystor triodowy. W literaturze angielskiej nosi on nazwę półprzewodnikowego zaworu sterowanego (SCR – Semiconductor Controlled Rectifier), a w literaturze polskiej jest po prostu nazywany tyrystorem. Zatem nazwa tyrystor jest używana, zależnie od kontekstu, w znaczeniu wąskim dotyczącym tyrystora triodowego, albo w znaczeniu szerokim, obejmującym wiele rodzajów struktur wielowarstwowych, jak widać w klasyfikacji na rysunku 2. **Tyrystor triodowy** ma strukturę i charakterystyki pokazane



Rysunek 4. Zakresy zastosowań elementów Si, SiC, GaN na wykresie sterowanej mocy i częstotliwości pracy, na przykładzie produktów firmy Infineon



Rysunek 5. Uproszczony schemat struktury diody PIN (a) i diody Schottky (b)



Rysunek 6. Symbol tyrystora (a), jego struktura czterowarstwowa (b) i jej równoważna reprezentacja tranzystorowa (c), oraz rodzina charakterystyk I_A (U_{AK}) dla różnych wartości prądu bramki I_G (d)

na rysunku 6. Jego trzy końcówki noszą nazwy katoda, anoda i bramka. Charakterystyki prądowo-napięciowe (rysunek 6d) pokazują, że tyrystor działa jako jednokierunkowy przełącznik bistabilny. W kierunku zaporowym tyrystor działa podobnie jak dioda, tj. nie przewodzi prądu aż do wystąpienia przebicia. W kierunku przewodzenia też nie przewodzi prądu, dopóki nie zostanie wyzwolony zapłon przez odpowiedni sygnał prądowy w obwodzie bramki. Po wyzwoleniu sygnałem bramkowym tyrystor przechodzi w stan przewodzenia, czyli jego rezystancja katoda-anoda gwałtownie maleje do bardzo małych wartości. Istotną cechą tyrystora jest brak możliwości wyłączenia, tj. przejścia ze stanu przewodzenia do stanu blokowania, za pomocą sterowania bramką. Po zapłonie, czyli przejściu do stanu przewodzenia, tyrystor pozostaje w tym stanie dopóki napięcie anoda-katoda nie zmieni polaryzacji lub prąd anodowy nie zmaleje poniżej pewnej wartości krytycznej, nazywanej prądem trzymania – I_H . Struktura tyrystora pokazana na rysunku 6 jest obrazem uproszczonym, wystarczającym dla uproszczonej analizy sposobu działania tego elementu. Rzeczywisty przekrój struktury tyrystora jest bardziej złożony, szczególnie dla elementów bardzo dużej mocy, w których powierzchnia struktury półprzewodnikowej może być bardzo duża i pojawia się problem jak włączyć całą powierzchnię tyrystora w akceptowalnie krótkim czasie. Rozwiązaniem tego problemu jest zastosowanie bramki rozgałęzionej o odpowiedniej topologii.

Z kolei zwiększenie powierzchni bramki oznacza wzrost prądu zapłonowego bramki. Aby włączyć tyrystor z rozgałęzioną bramką, mogą być wymagane impulsy prądu bramki o amplitudzie dziesiątek, a nawet setek amperów, co mocno komplikuje konstrukcję układu sterującego zapłonem bramki. Aby uniknąć tych problemów stosuje się pomocniczą strukturę tyrystorową, czyli w jednej strukturze półprzewodnikowej mamy tyrystor główny (silnoprądowy) i „tyrystorek” pomocniczy, który jest wyzwolany niewielkim prądem bramki, a jego prąd anodowy wpływa do rozgałęzionej bramki tyrystora głównego. Trochę to przypomina bombę wodorową, w której zapalnikiem jest znacznie słabsza od niej bomba atomowa.

Swoją drogą, zapłon tyrystora to pasjonujące zjawisko fizyczne, swoją gwałtownością przypominające wybuch bomby. Chętnie to zjawisko bym opisał, ale wykład ma charakter panoramiczny, więc musimy powściągać ochotę do wycieczek w głąb poszczególnych tematów. Mogę do tego tematu wrócić w przyszłości, jeśli byłoby takie zapotrzebowanie.

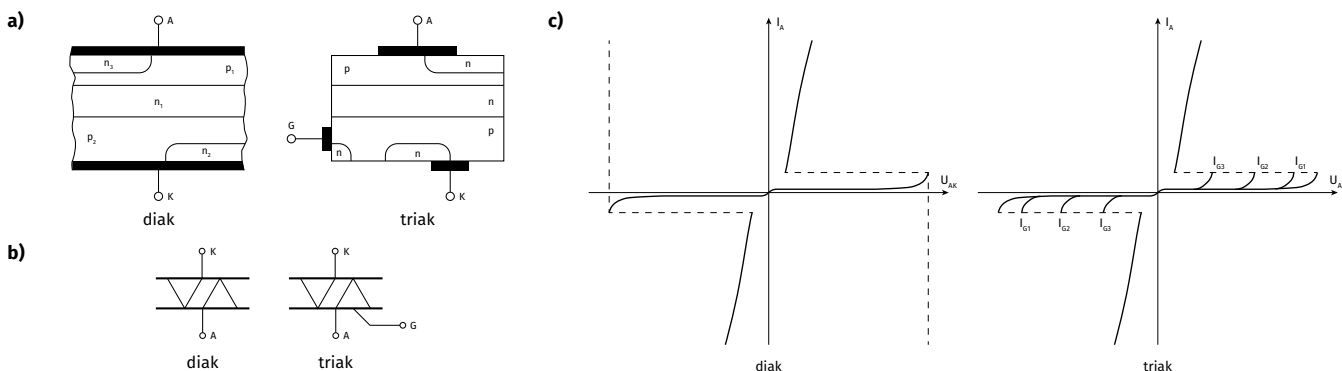
Tyrystor dwukierunkowy (diak, triak)

Na rysunku 7 przedstawiono podstawową strukturę pięciowarstwową n-p-n-p-n tyrystora dwukierunkowego dla wariantu dwukońcówkowego (diak) oraz trójkątkowego (triak) wraz z symbolami i charakterystykami tych elementów. Sposób działania tych tyrystorów można wyjaśnić na podstawie superpozycji dwóch struktur czterowarstwowych ($p_1-n_1-p_2-n_2$ oraz $p_2-n_1-p_1-n_3$). Taki tyrystor ma symetryczne właściwości dla obu polaryzacji anoda-katoda. W triaku za pomocą sterowania prądem w obwodzie bramki można regulować wartość napięcia przełączania. Przełączanie do stanu przewodzenia dodatniego lub ujemnego może odbywać się przezysterowanie obwodu bramki zarówno prądem dodatnim jak i ujemnym.

Przedstawione dotychczas rodzaje tyrystorów, znane już w latach sześćdziesiątych, nie spełniały wszystkich potrzeb konstruktorów układów.

Przede wszystkim, brakowało dwóch rzeczy:

- możliwości wyłączania (przechodzenia ze stanu przewodzenia do blokowania) sygnałem sterującym bramką.



Rysunek 7. Tyrystor dwukierunkowy (triak, diak): struktura pięciowarstwowa (a), symbole (b), charakterystyki (c)

- sterowania bramką bez poboru dużego prądu bramki.

Dążenie do spełnienia tych potrzeb określiło dalsze kierunki rozwoju technologii elementów mocy.

Tyrystory GTO (Gate Turn Off)

Jest to taki rodzaj tyrystora, w którym elektroda bramki służy zarówno do włączania (zapłonu) jak i do wyłączania (przejścia ze stanu przewodzenia do stanu blokowania). GTO nie-

wiele się różni od wcześniej omawianego tyrystora triodowego. Jest to struktura czterowarstwowa z bramką, jak w tyrystorze triodowym, ale w tym przypadku konstrukcja bramki umożliwia przepływ dużego prądu bramki przy polaryzacji ujemnej (wyłączanie z przewodzenia do blokowania). Aby wyłączyć tyrystor trzeba „ukraść” znaczną część (ok. 25%) prądu płynącego od anody do katody. Dzieje się to przez odprowadzenie tego prądu przez elektrodę bramki. Dlatego konstrukcja GTO musi pozwalać na przepływ dużych prądów bramkowych. Na rysunku 8 pokazano uproszczony przekrój struktury GTO. W rzeczywistości jest to struktura zawierająca równolegle połączone setki, a nawet tysiące elementarnych „tyrystorków”. W kierunku zaporowym (III-a ćwiartka wykresu $I_A [U_{AK}]$ na rysunku 8c) tyrystory GTO mogą mieć dwa rodzaje charakterystyk – symetryczne lub asymetryczne. **Symetryczny GTO** może przewodzić w kierunku zaporowym w zakresie dużych napięć (jak dioda), porównywalnych z napięciem blokowania w kierunku przewodzenia.

Natomiast **niesymetryczny GTO** już przy niewielkich (dziesiątki woltów) napięciach zaporowych zaczyna przewodzić prąd. W połowie lat 90. wielu producentów opracowało nową generację tyrystorów wyłączanych bramką, nazywanych GCT (Gate Commutated Thyristors), tj. **tyrystory z komutacją bramkową**. Ta grupa dzieli się z kolei na trzy rodzaje tyrystorów:

- GCT z twardym przełączaniem, bez stosowania sterownika
- IGCT (Integrated Gate Commutated Thyristors), tj. GCT zintegrowany ze sterownikiem
- ETO (Emitter Turn–Off Thyristor) lub ECT (Emitter Controlled Thyristor) to GCT zintegrowany ze sterownikiem i przełączany przez obwód kaskadowy.

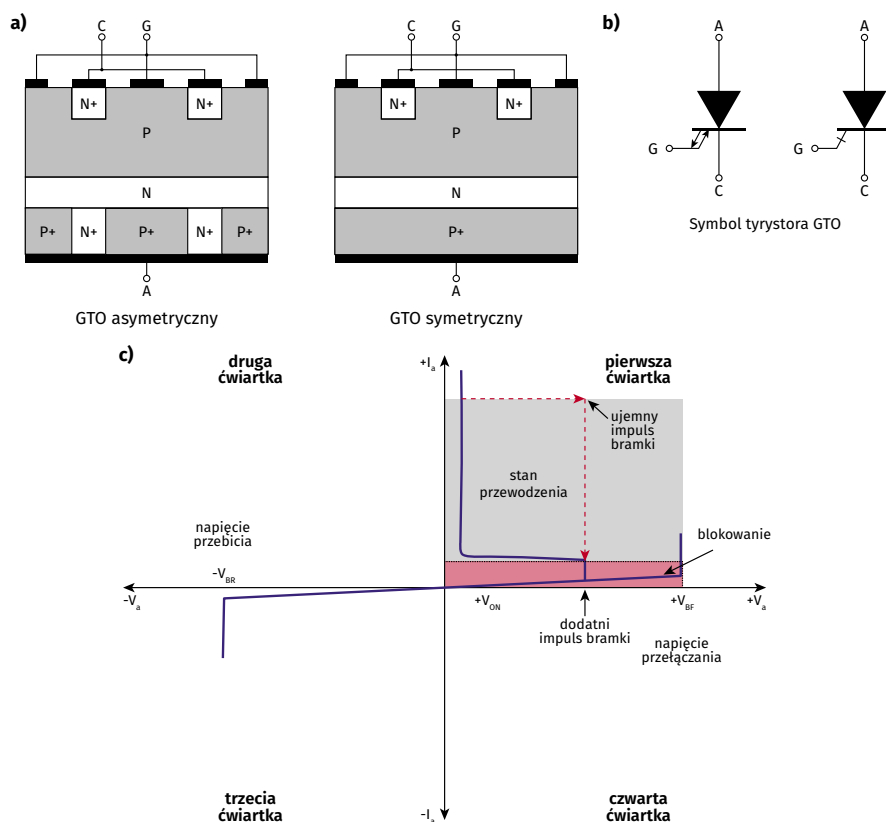
W terminologii rynkowej, tj. w katalogach producentów, wyróżnia się jeszcze kilka rodzajów tyrystorów, więc poświęćmy im parę zdań.

PCT (Phase Control Thyristor) – tyrystory do sterowania fazy, to tyrystory do zastosowań przemysłowych, pracujące przy względnie niskich częstotliwościach, zwykle jest to częstotliwość sieci 50 Hz. Mogą to być sterowane prostowniki, softstarty do silników elektrycznych, różne falowniki itp. W przypadku tego typu tyrystorów główny nacisk kładzie się na minimalizację strat mocy w stanie przewodzenia, przy jednoczesnym zapewnieniu zadanego napięcia blokującego dla obu polaryzacji – w przód i wstecz. Tyrystory PCT są często stosowane w przekształtnikach na napięcia 6...10 kV i wyższe, gdzie konieczne jest szeregowo łączenie poszczególnych tyrystorów. Dlatego istotne jest wysokie napięcie blokujące i zapewnienie synchronicznego włączania i wyłączania (powrotu do właściwości blokujących). Tyrystory PCT mają zwykle duże wartości ładunku Q_{rr} , bo dla częstotliwości 50 Hz ten parametr nie ma znaczenia, ale wyklucza to możliwość ich stosowania do pracy impulsowej przy wyższych częstotliwościach. Czas wyłączania wynosi z reguły ponad 50 μs i często nie jest podawany przez producentów.

Tyrystory przeznaczone do pracy przy wysokich częstotliwościach, które mają małe wartości ładunku Q_{rr} i czasów włączania/wyłączania, są niekiedy nazywane **PPT (Pulse Power Thyristors)** lub **FT (Fast Thyristors)**. Odrębną grupę stanowią **LTT (Light Triggered Thyristors)**, tj. **tyrystory wyzwalane przez światło**, a nie przez impuls prądu bramki. Nasze pierwsze skojarzenie

Tablica 2. Porównanie parametrów IGBT z MOSFET i tranzystorem bipolarnym

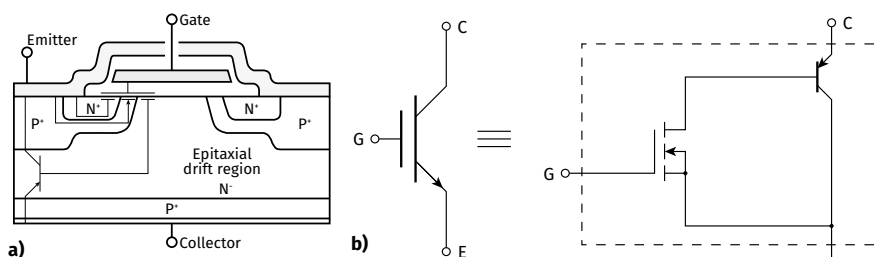
Parametr elementu	Bipolarny tranzystor mocy	MOSFET mocy	IGBT
Maksymalne napięcie	Wysokie < 1 kV	Wysokie < 1 kV	Bardzo wysokie > 1 kV
Maksymalny prąd	Duży < 500 A	Mały < 200 A	Duży > 500 A
Sterowanie	Prądowe $\beta=20\div200$	Napięciowe $V_{GS}=3\div10 V$	Napięciowe $V_{GE}=4\div8 V$
Impedancja wejściowa	Niska	Wysoka	Wysoka
Impedancja wyjściowa	Niska	Średnia	Niska
Szybkość przełączania	Mała (μs)	Duża (ns)	Średnia
Koszt	Niski	Średni	Wysoki



Rysunek 8. Struktura tyrystora GTO (a), jego symbol (b) i charakterystyka (c)

podpowiada nam zastosowanie LTT jako fotodetektora, ale w tym przypadku chodzi o coś innego. Tyrystory LTT są obecnie szeroko stosowane w wysokonapięciowych układach mocy, gdzie jest wymagane łączenie szeregowo tyrystorów. W LTT nie ma metalu w obszarze bramki. Zamiast bramki jest okienko fotodetektora. Impulsy świetlne są generowane przez laser i dostarczane do okienka fotodetektora za pomocą kabla światłowodowego. Wyzwalanie tyrystorów światłem zamiast prądem bramki ma szereg zalet, takich jak:

- precyzyjna kontrola czasu przełączania grupy tyrystorów,
- galwaniczna izolacja obwodu sterującego od obwodu silnoprądowego,
- brak wpływu zakłóceń elektromagnetycznych,
- kompaktowe ułożenie kabla światłowodowego w sprężenie
- wysoka sprawność i niezawodność.



Rysunek 9. Przekrój struktury tranzystora IGBT pokazujący wewnętrzne połączenie tranzystora bipolarnego i MOSFET (a), symbol i równoważny schemat zastępczy tranzystora IGBT

MOSFETy

Tranzystory MOS znamy jako niezwykle miniaturowe, wielkości kilku nanometrów, elementy układów scalonych pamięci lub mikroprocesorów. Okazuje się, że na tej samej zasadzie (sterowanie napięciem bramki – źródło przewodnością kanału łączącego dren ze źródłem) działają tranzystory MOS o gabarytach miliony razy większych, nazywane tranzystorami mocy. W tym przypadku chodzi o pracę dwustanową tranzystora używanego jako przełącznik zwarty lub rozarty. Najistotniejsze cechy MOOSFETa mocy to jak najmniejsza rezystancja w stanie włączenia i jak największe napięcie blokowane w stanie wyłączenia. W stanie włączenia napięcie podane na bramkę względem źródła powoduje inwersję typu przewodnictwa przy powierzchni krzemu, z typu p na n, i warstwa inwersyjna tworzy kanał przewodzący, który łączy źródło z drenem (zauważmy, że MOSFETy mocy są wytwarzane wyłącznie jako tranzystory ze wzbogaconym kanałem typu n). Przy zerowym napięciu na bramce nie ma kanału i napięcie dren – źródło odkłada się w całości na złączu p-n (źródło – podłoże) spolaryzowanym w kierunku zaporowym. Aby to złącze wytrzymało jak największe napięcie wsteczne obszar drenu (typ n) musi być słabo domieszkowany, co skutkuje dużą szerokością warstwy zaporowej. Z drugiej strony, słabo domieszkowany obszar drenu wnosi dużą rezystancję, włączoną szeregowo do przewodzącego kanału w stanie przewodzenia tranzystora, co stoi w sprzeczności z dążeniem do minimalizacji rezystancji przewodzącego MOSFETa. Ta sprzeczność stanowi podstawowy problem w konstrukcji MOSFETów mocy. Aby rezystancja przewodzącego tranzystora była jak najmniejsza, jego konstrukcja powinna charakteryzować się ekstremalnie małą długością i maksymalnie dużą szerokością kanału. Dlatego konstrukcyjnie MOSFET mocy składa się z równoległe połączonych tysięcy tranzystorów MOS o krótkim kanale. Efektywnie szerokość kanału osiąga w ten sposób kilkanaście milimetrów. W porównaniu z innymi elementami mocy zaletą MOSFETów jest duża szybkość przełączania i mała rezystancja w stanie przewodzenia, czyli małe straty energii na kluczu zwartym.

Dodajmy, że z istoty budowy i zasady działania MOSFETów, tj. sterowania bramką odizolowaną od podłoża krzemowego, zaletą tego elementu jest mała wartość prądu sterowania (praktycznie zerowa w stanie statycznym). Konstruktorzy wpadli na pomysł, żeby tę właściwość sterowania odizolowaną bramką połączyć hybrydowo ze strukturą tranzystora bipolarnego. Tak powstał bardzo popularny element mocy – tranzystor IGBT (Insulated Gate Bipolar Transistor).

IGBT

Tranzystory IGBT łączą cechy budowy i zalety tranzystorów bipolarnych i MOSFET dużej mocy, co widać na **rysunku 9**. Od strony wejściowej tranzystor IGBT zachowuje się jak MOSFET (zerowy prąd bramki w stanie statycznym), natomiast obwód obciążenia „widzi” tranzystor bipolarny. To hybrydowe połączenie dało świetne efekty – powstał element mocy przełączający obciążenia dużej mocy (do megawatów) z wysokimi częstotliwościami (dziesiątki kHz), przewodzący bardzo duże prądy (do wielu kiloamperów), blokujący wysokie napięcia (do wielu kilowoltów), sterowany napięciowo (przynajmniej w stanie statycznym i przy niewielkich częstotliwościach) i wprowadzający niewielkie tylko straty mocy w obwodzie kluczowanym (choć spadek napięcia w stanie przewodzenia jest znaczny – ok. 2,5 V).

Zakończenie

Ten wykład, jak każdy inny, trzeba w określonym momencie zakończyć, chociaż temat półprzewodników mocy został zaledwie zasygnalizowany. Zdaje sobie sprawę, że amatorski pasjonat elektroniki raczej omija to pogranicze elektroniki i elektryki. I dobrze. To nie jest pole dla amatorskich eksperymentów. Łatwo o nieszczęście. Natomiast Czytelnicy EdW kształcący się na zawodowych elektroników zapewne potrzebują znacznie bardziej pogłębionej wiedzy o poszczególnych elementach mocy. Chętnie wrócę do niektórych tematów, jeśli będzie takie zainteresowanie. Polecam też podręczniki publikowane na stronach niektórych producentów (lista najważniejszych producentów półprzewodnikowych elementów mocy – w ramce). ■

Wiesław Marciniak

Lista najważniejszych producentów elementów półprzewodnikowych mocy:

- Infineon
- Mitsubishi Electric
- Vishay
- ON Semiconductor
- ST Microelectronics
- ABB Semiconductors
- Littelfuse
- Dynex Semiconductor
- Semikron
- Fuji Electric
- Rohm
- Toshiba

Programowanie wizualne z XOD



Przedstawiamy wizualne programowanie XOD dla Arduino

Muszę się przyznać: masowe i rosnące zastosowanie programowalnymi mikrokontrolerami w elektronice hobbyistycznej w dużej mierze mnie ominęło. Dlaczego? Cóż, szczerze mówiąc, nie lubię listingów kodu – jeden znak w złym miejscu i program się nie uruchomi; a poza tym, kto chce się uczyć mnóstwa skomplikowanych instrukcji, zanim uda się zmusić mikrokontroler do zrobienia nawet najprostszej rzeczy? (OK, wielu z was! Ale inni, jak ja, raczej nie)!

Cóż, w ciągu ostatnich kilku miesięcy miałem objawienie – i to wszystko dzięki XOD (wymawiane 'zod'). XOD to darmowy pakiet oprogramowania do programowania wizualnego dla Arduino. Nie ma w nim

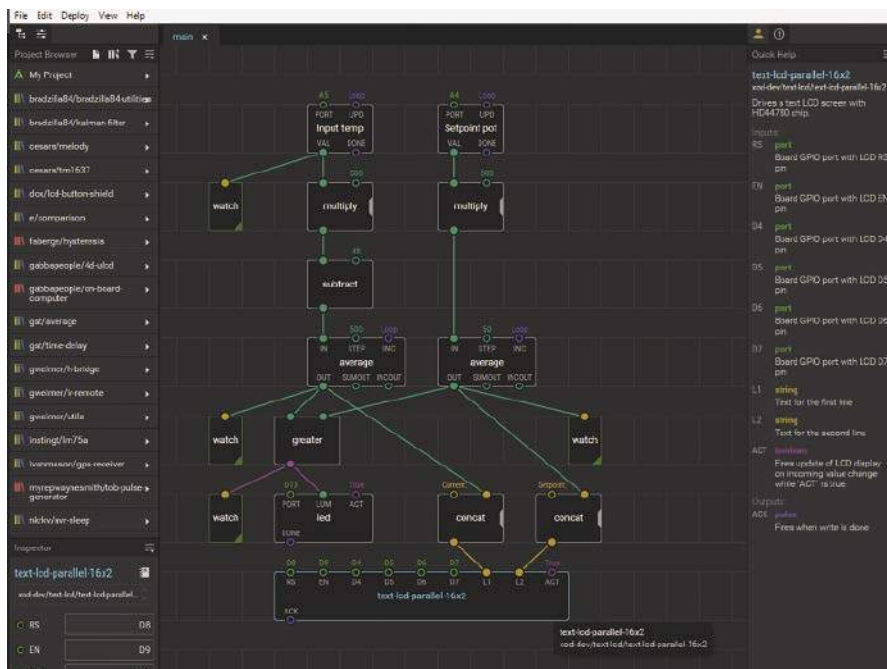
kodu – po prostu umieszczasz gotowe bloki na ekranie, a następnie „łączysz” je ze sobą. Jeśli potrafisz zrozumieć prosty schemat układu, możesz zaprogramować Arduino w XOD.

Nie mam wątpliwości, że ekspert piszący kod znalazłby wiele niedociągnięć w XOD, ale dla kogoś bez takiego doświadczenia, kto chce szybko sprawić, by Arduino robiło rzeczy z prawdziwego zdarzenia, jest to po prostu wspaniałe. Co więcej, poziom wsparcia jest dobry. Zamiast setek kiepskich filmów na YouTube, oprogramowanie ma wbudowany samouczek z 42 jasnymi i istotnymi lekcjami. Prowadzą one od zerowej wiedzy do posiadania wystarczającej wiedzy, aby zbudować całkiem złożone systemy sterowania. Istnieje również forum użytkowników, które uważam za pomocne i wspierające.

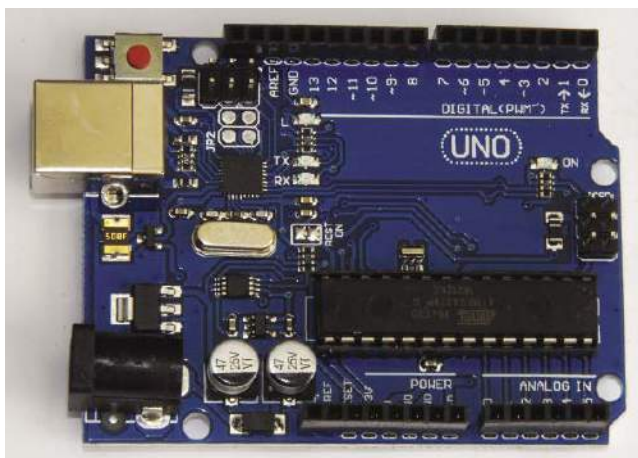
Wreszcie, XOD pozwala na emulację programu na ekranie – to znaczy, że możesz zobaczyć jak program będzie działał, nawet gdy go piszesz. To znacznie ułatwia debugowanie programu.

Pobieranie oprogramowania

Oprogramowanie można pobrać ze strony <https://xod.io/downloads/> i jest dostępne dla platform Windows, macOS i Linux. (Jeśli nie chcesz początkowo pobierać wersji desktopowej, możesz też trochę poznać oprogramowanie, korzystając z wersji przeglądarkowej). Przed pobraniem, lub skorzystaniem z przeglądarkowej wersji XOD, wymagane będzie utworzenie konta z użyciem adresu e-mail



Ekran programowania XOD. W głównym obszarze roboczym znajduje się program ('sketch'), który jest budowany. Tutaj podświetlone zostało pole programowania LCD – lewa dolna kolumna ekranu pokazuje konfigurowalne ustawienia dla LCD, a prawa kolumna wyświetla plik pomocy dla wyświetlacza LCD.



Rysunek 1. Płytkę Arduino Uno – jest niezwykle tania, a przy użyciu oprogramowania do programowania wizualnego XOD, teraz bardzo łatwa do zaprogramowania. Można mieć projekty gotowe i działające w dostępie kilkanaście minut.



Rysunek 3. Moduł DFRobot 1602 LCD Keypad Shield podłączony do płytki Arduino Uno. Uno jest zaprogramowane oprogramowaniem XOD omówionym w tym artykule, a na wyświetlaczu widoczna jest aktualna temperatura oraz wartość zadana, przy której zapala się dioda alarmowa.

i hasła. Przypominamy – XOD jest całkowicie darmowy, więc nie dostaniesz wersji niekompletnej lub ograniczonej w swoich możliwościach.

Instalacja i uruchomienie oprogramowania jest proste; jedynym problemem, jaki miałem, było określenie właściwej płytki Arduino podczas przesyłania do niej danych po raz pierwszy. Dla Arduino Uno lub klonu wybierz domyślną opcję „Arduino/ Genuino Uno”. Możesz również zauważyć, że musisz określić właściwy port komunikacyjny komputera.

Zauważ, że aby korzystać z oprogramowania, płytka Arduino musi być podłączona do portu USB komputera – połączenie

USB zasilą płytkę. Istnieje również funkcja symulacji online dostępna dla XOD, ale stwierdziłem, że jest ona nieco obciążona błędami.

Używanie oprogramowania

Pierwszym krokiem po ściągnięciu XOD jest przećwiczenie wbudowanego samouczka. Jest on dość szczegółowy i jeśli wykonasz wszystkie sugerowane czynności (samouczki są wykonywane w rzeczywistym środowisku programistycznym), zajmie to kilka godzin. Jednak już po pierwszych dziesięciu lekcjach będziesz miał wystarczająco dużo informacji, aby rozpocząć programowanie. W każdej chwili możesz wrócić do tutoriala, przechodząc do: Help >

Open Tutorial Project. Pomocna okazała się również ta strona: <https://dronebotworkshop.com/gettingstarted-with-xod/>

Sprzęt

W tym przykładzie użyję:

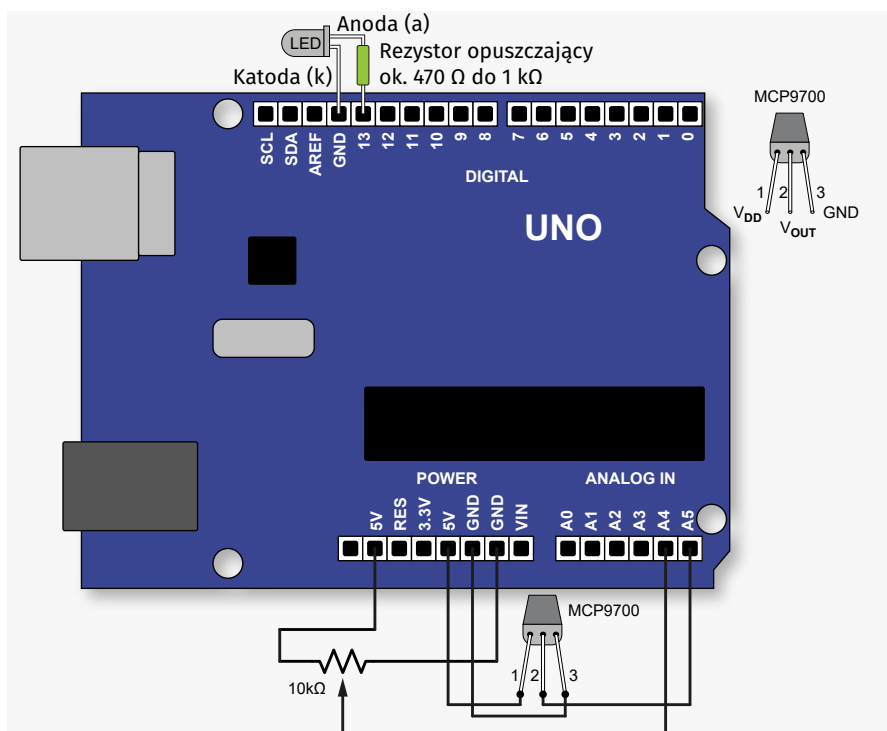
- Arduino Uno, na przykład płytka pokazana na **rysunku 1**.
- Scalony czujnik temperatury Microchip MCP9700.
- Dioda LED z odpowiednim rezystorem opuszczającym.
- Potencjometr 10 kΩ.
- Moduł DFRobot 1602 LCD Keypad Shield (**rysunek 3**).

Zacznijmy od czujnika temperatury. Podłącz odpowiednie piny (patrz **rysunek 2**) z czujnika do zasilania 5 V i masy na płytce, a wyjście do portu A5 (analog 5) (uwaga: wszystkie połączenia są oznaczone na płytce). Podłącz diodę LED pomiędzy piny D13 i masę. (Aby przyspieszyć pracę, kupuję diody LED wstępnie okablowane z rezystorem do użycia na 12 V. Są one nadal wystarczająco jasne, nawet przy pracy z 5 V). Upewnij się, że polaryzacja jest prawidłowa. Na koniec, podłącz potencjometr do zacisków 5 V i masy, z jego wyprowadzeniem ślizgacza podłączonym do portu A4. Powinieneś teraz mieć to, co pokazano na **rysunku 2**. Do tarczy (płytki) LCD przejdziemy później.

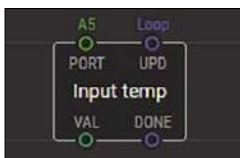
Pisanie programu

Wybierz File > New Project – zobaczysz w gruncie rzeczy pusty ekran. To jest Twój obszar roboczy, w którym będziesz budował program. Po lewej i prawej stronie znajdują się dwie wąskie kolumny – wyświetlają one informacje, które pomogą Ci w trakcie pracy.

Zróbmy alarm temperaturowy, a żeby było trudniej, uśrednijmy odczyt temperatury w ciągu około pięciu sekund i ustawmy wyjście



Rysunek 2. Wstępna konfiguracja alarmu temperatury na płytce Arduino z wykorzystaniem czujnika MCP9700 i diody LED z rezystorem opuszczającym. Tarcza LCD jeszcze nie dodana.



Rysunek 4. Dodanie pierwszego bloku XOD – bloku funkcjonalnego 'analog-sensor'.

czujnika tak, abyśmy pracowali w stopniach Celsjusza (°C). Następnie dodamy wejście potencjometryczne, aby ustawić temperaturę zadaną, a następnie wyświetlacz LCD, aby pokazać zarówno temperaturę zadaną, jak i temperaturę w czasie rzeczywistym. Brzmi ciężko? Spokojnie!

Funkcjonalne bloki konstrukcyjne

Z obrazków w tym artykule widać, że program jest zbudowany z bloków funkcjonalnych (zwanych „nodes” – węzły) – ale gdzie one są? Dwukrotne kliknięcie w głównym obszarze roboczym spowoduje pojawienie się okna 'search nodes'. Programowanie chcemy rozpocząć od bloku łączącego analogowe wejście płytki Arduino z czujnikiem temperatury.

Po wpisaniu 'analog' w pole wyszukiwania pojawia się xod/ common-hardware/analog-sensor. Wybierz go, klikając dwukrotnie. Do głównego obszaru roboczego zostanie dodany blok funkcjonalny o nazwie analog-sensor – patrz rysunek 4.

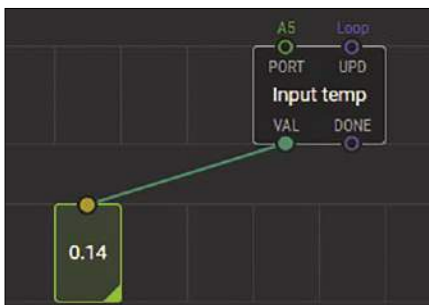
Kliknięcie na ten blok powoduje wyświetlenie informacji w dwóch kolumnach ekranu. Po prawej stronie wyjaśnione są poszczególne etykiety na bloku funkcjonalnym, natomiast po lewej stronie mogą one być konfigurowane przez użytkownika. Etykiety na tym pierwszym bloku to:

PORT – port analogowy, który będzie używany dla tego wejścia na płytce Arduino. Zmień go na A5.

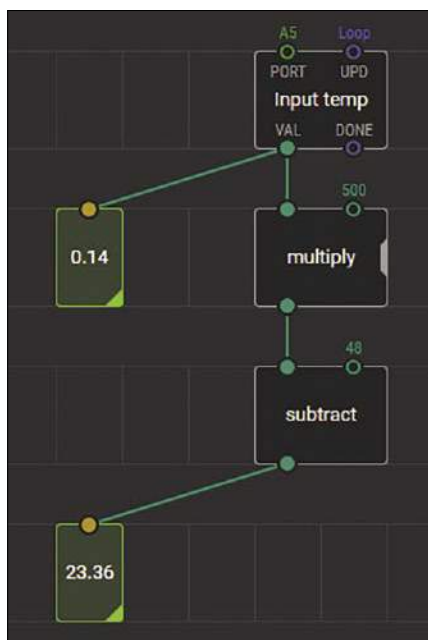
UPD – sposób w jaki wejście jest aktualizowane. Dostępne opcje to 'never', 'on boot' oraz 'continuously'. Ustaw ją na 'continuously' – pokazane na bloku jako 'loop'.

VALUE – wartość wyjścia, skonfigurowana automatycznie na zakres 0–1.

DONE – wyprowadza impuls po zakończeniu odczytu (nie używamy tego).



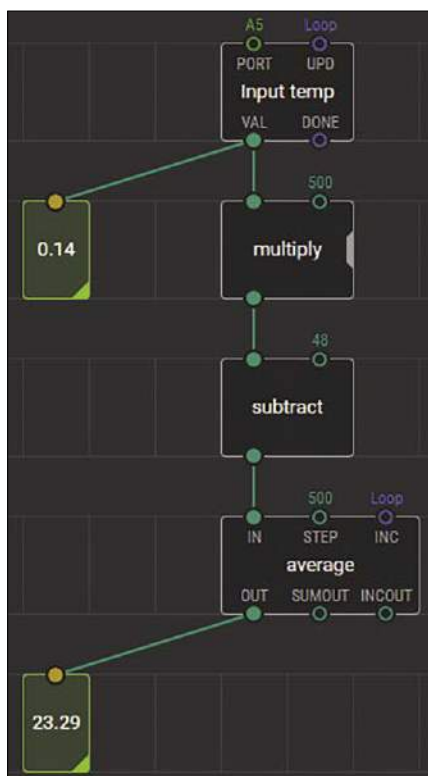
Rysunek 5. Dodanie bloku watch i podłączenie go do bloku Input temp.



Rysunek 6. Dodanie bloków do wyświetlania mierzonej temperatury w °C.

Możemy zmienić etykietę bloku, aby przeczytać coś innego niż analog-sensor – i teraz, korzystając z kolumny po lewej stronie, oznaczyć je jako Input temp.

Tak więc dodaliśmy teraz wejście analogowe na porcie A5, które jest stale aktualizowane i jest oznaczone w sposób, który możemy natychmiast zrozumieć.



Rysunek 7. Wykorzystanie generowanej przez użytkownika funkcji average do generowania pięciosekundowych średnich z odczytów temperatury.

Tworzenie wejścia

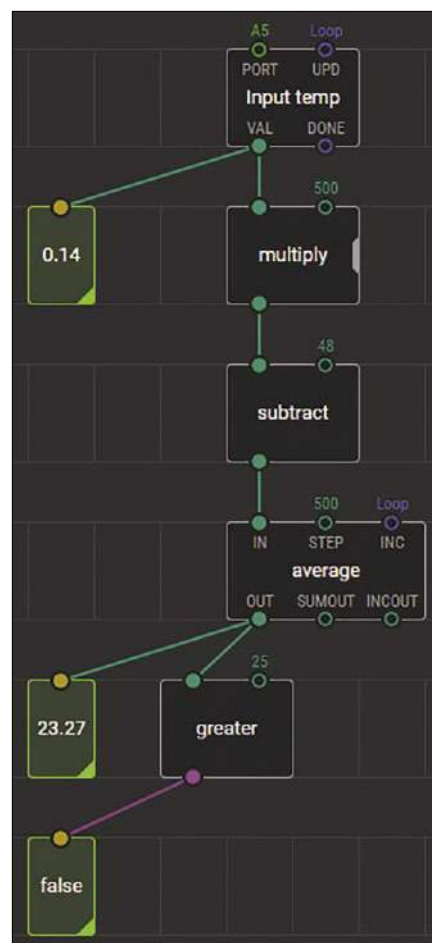
Ale czy nasze wejście rzeczywiście coś odczytuje? Kliknij dwukrotnie obszar roboczy i wpisz debug, a pojawi się seria opcji – wybierz xod/debug/ watch i podłącz ten blok do wyjścia (VAL) bloku Input temp. Załaduj program do płytki, naciskając symbol 'bug' w prawym dolnym rogu ekranu. W ten sposób program zostanie załadowany w formie debugowania – możesz wtedy zobaczyć działanie programu w czasie rzeczywistym, jak pokazano na rysunku 5.

Podgrzewając i schładzając czujnik temperatury powinniśmy teraz zobaczyć, jak zmienia się wartość w okienku watch – mamy wejście!

Hmm, ale jest też pewien problem – w temperaturze pokojowej wyjście czujnika pokazuje 0,14 – nie jest to „prawdziwy” odczyt temperatury. Musimy przeliczyć ten odczyt na stopnie Celsjusza, a w przypadku tego czujnika oznacza to pomnożenie tego odczytu przez 500 i odjęcie 50. Jak więc to zrobić?

Przetwarzanie wartości wejściowych

Dwukrotne kliknięcie w obszarze roboczym i wyszukanie pod 'multiply' przynosi xod/core/



Rysunek 8. Blok greater umożliwia dodanie alarmu systemowego. Blok true/false pozwala sprawdzić, czy alarm jest aktywny.

multiply, który następnie wybieramy. Łączymy IN1 bloku multiply z wyjściem VAL bloku input temp i korzystając z lewej kolumny na ekranie, zmieniamy wartość IN2 na 500.

Następnie dodajemy blok xod/core/subtract i odejmujemy 50. (Cóż, jak widać tutaj, faktycznie użyłem 48 – ten offset dawał lepszą dokładność z konkretnym czujnikiem temperatury, którego używałem).

Następnie możemy dodać kolejny watch box do wyjścia z subtract box i przesłać program w trybie debugowania na płytke. Powinno się wtedy zobaczyć odczyt w stopniach Celsjusza, jak pokazano na **rysunku 6**.

Przetwarzanie danych

Teraz chcemy uśrednić odczyt temperatury; np. w ciągu pięciu sekund. Jeśli jednak klikniesz dwukrotnie na obszar roboczy i zaczniesz wpisywać 'average', to okaże się, że nie uzyskasz żadnego wyniku – w standardzie nie ma bloku 'averaging'.

Jednak jedną z zalet systemu XOD jest to, że użytkownicy mogą wgrywać swoje własne, użyteczne bloki, aby inni mogli z nich korzystać. Można je znaleźć na stronie XOD w zakładce „libraries”. Wchodząc w File > Add

Library i wpisując nazwę biblioteki, tak jak jest to pokazane na stronie internetowej, można dodać te dodatkowe bloki. Average jest jednym z nich – znajduje się pod 'gst/average/average'. Ten blok ma sześć etykiet, ale nas interesują tylko trzy – in, out i step.

Kliknij i przeciągnij linię od wartości wyjściowej bloku subtract do IN na bloku average. Zmień STEP na 500 i INC na Continuously (pokazane na bloku jako LOOP). Teraz uśredniamy wartość wyjściową naszego czujnika temperatury w ciągu około pięciu sekund. Ale czy to zadziała? Dodaj blok watch do wyjścia bloku average, załaduj program w trybie debugowania na płytke, a następnie sprawdź, czy odczyt temperatury jest teraz uśredniony w jakimś okresie (**rysunek 7**).

Dodanie alarmu

Teraz dodamy do programu część alarmową. Wybierz blok xod/core/greater i połącz IN1 z wyjściem bloku average. Zmień IN2 na 25 (jako wejście, 25°C). Dodaj blok watch do wyjścia bloku greater (**rysunek 8**).

Załaduj program w formie debugowania i sprawdź, czy blok watch zmienia się z false na true, gdy temperatura czujnika

przekracza 25°C. (Zauważ, że blok watch może pokazywać zarówno wartości [liczby], jak i stany logiczne [false/true].)

Zapalenie diody LED

Następnym krokiem jest zapalenie diody LED, gdy alarm zostanie uruchomiony. Dodaj do ekranu xod/common-hardware/led – wyszukiwanie pod 'led' spowoduje pojawienie się tego bloku. Zmień jego port na D13 i połącz jego pin LUM z wyjściem bloku greater (patrz lewy dolny róg, **rysunek 9**).

Załaduj program w trybie debugowania. Teraz dioda LED świeci, gdy zmierzona temperatura jest powyżej 25°C, uśredniona w ciągu około pięciu sekund. (Zauważysz, że świeci się również dioda LED na płycie – jest ona również podłączona do D13).

Wejście zewnętrzne potencjometru

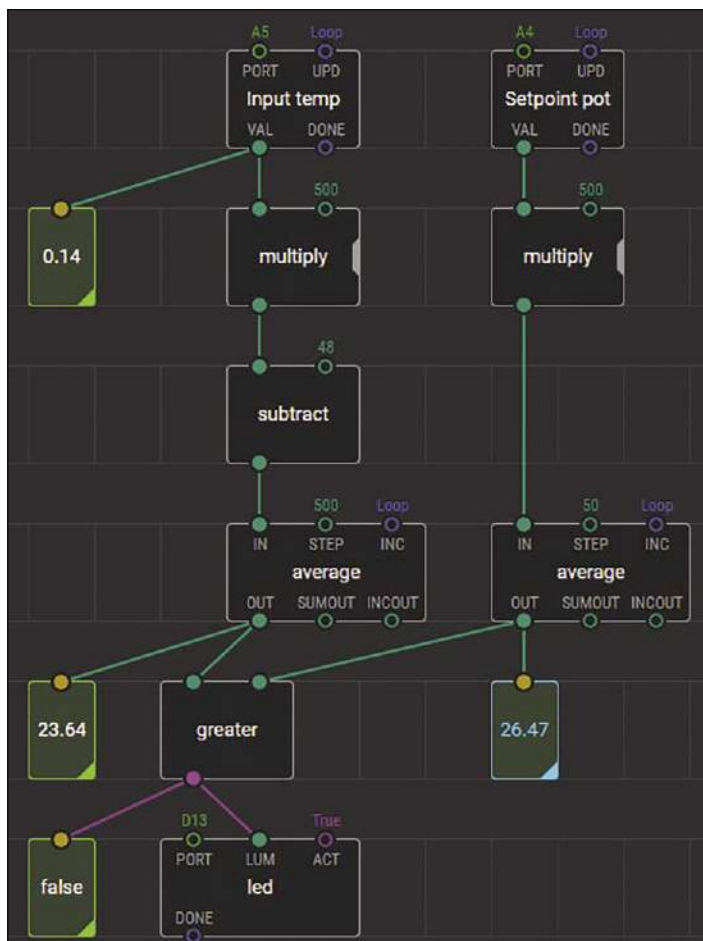
W obecnym formacie programu, temperatura zadana alarmu jest wpisywana do programu. Ale co by było, gdybyśmy chcieli zmienić wartość zadaną za pomocą zewnętrznego potencjometru? Pokażemy, że zrobienie tego zajmuje tylko chwilę. Zauważ, że używamy portu A4 dla wejścia potencjometru i ponownie mnożymy jego wyjście przez 500 (nie trzeba odejmować), a następnie uśredniamy wyjście. Tutaj wartość zadana wynosi 26,47°C. (Czy chcesz zaokrąglić tę liczbę? Jeśli tak, znajdź blok xod/math/round i dodaj jego). Patrz prawa strona **rysunku 9**.

Podłączenie wyświetlacza LCD

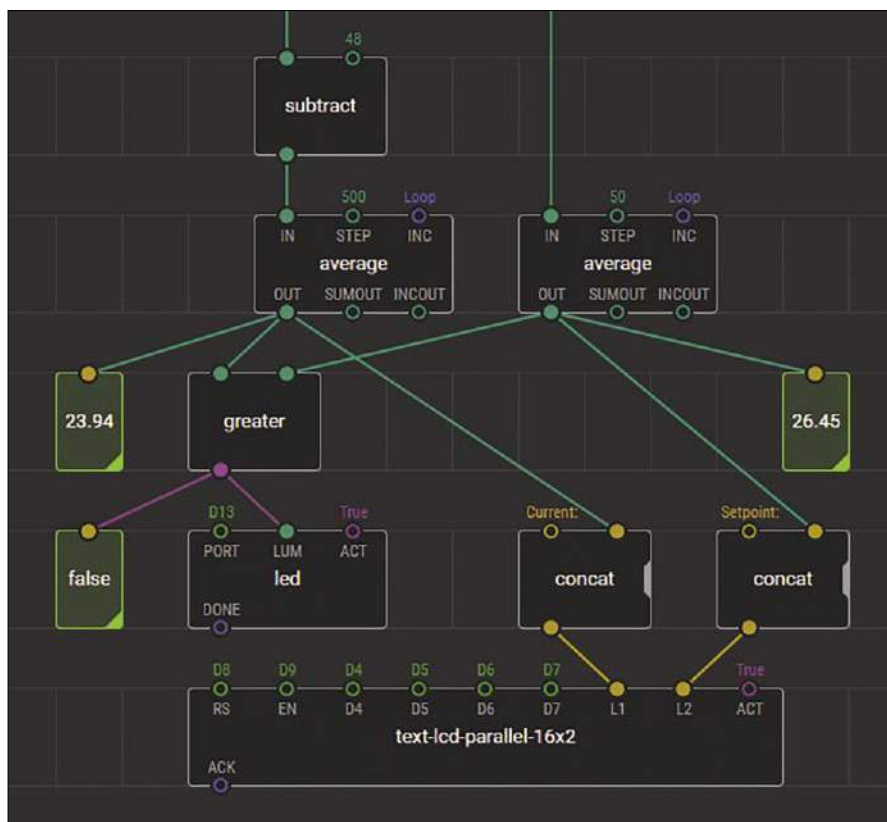
A może by tak teraz pokazać na ekranie LCD zarówno temperaturę zadaną jak i monitorowaną? Muszę być jednak szczerzy. Istnieje tu potencjalna przeszkoda sprzętowa – a jest nią to, że moduły LCD mają wiele różnych pinów. Domyślny blok programowania LCD 16×2 w oprogramowaniu XOD nie pasował do modułu DFRobot 1602 LCD Keypad Shield, którego używałem – to znaczy, nie był „kompatybilny z wyprowadzeniami”, więc musiałem dokonać pewnych zmian w portach. Oznaczało to sprawdzenie karty katalogowej modułu LCD i zmianę wyznaczonych portów na bloku XOD LCD (o czym za chwilę), aby dopasować funkcje poszczególnych pinów. W przypadku DFRobota LCD wyglądało to następująco:

pin	LCD port Arduino
RS	D8
PL	D9
D4	D4
D5	D5
D6	D6
D7	D7

Aby dodać tarczę (shield) LCD, należy najpierw usunąć istniejące połączenia czujnika temperatury, diody LED i potencjometru.



Rysunek 9. Tutaj dodajemy bloki, aby: a) umożliwić Arduino zapalenie diody LED, gdy zostanie wywołany alarm; oraz b) użyć potencjometru do regulacji wartości zadanej.



Rysunek 10. Dodanie modułu DFRobot 1602 LCD Keypad Shield kończy projekt alarmu. (Uwaga: ten obraz pokazuje tylko dolną część bloków; górna część jest identyczna jak dwa górne rzędy na rysunku 9). Problemy z kompatybilnością pinów oznaczały, że był to jedyny skomplikowany etap projektu – ale XOD podjął się tego i tarcza LCD działa dobrze – patrz rysunek 3.

Następnie podłącz tarczę (shield) LCD, upewniając się, że jest umieszczona jak najdalej od płytki Arduino. Po wykonaniu tych czynności należy ponownie podłączyć wejścia potencjometru i czujnika temperatury za pomocą pinów wejść analogowych z czoła shielda. Wyjście LED również musi zostać przeniesione. Numery portów pozostają takie same, więc zajmuje to tylko chwilę – chociaż będziesz musiał zmienić złącza z męskich na żeńskie.

Następnie dodajemy blok text-lcd-parallel-16x2. L1 odnosi się do linii 1, a L2 do linii 2 na wyświetlaczu; każda z nich jest osobno adresowana (patrz rysunek 10). Linia 1 jest następnie połączona z uśrednioną temperaturą wejściową poprzez blok xod/core/concat,

które jest używany do wyświetlania słowa „Current:” oraz liczbowego odczytu temperatury. Linia 2 wykorzystuje podobnie blok concat do wyświetlania słowa Setpoint: i uśrednionego wyjścia potencjometru. Zmieniony program jest następnie wgrany w trybie debugowania i powinieneś zobaczyć, że cały system działa. Jeśli LCD nic nie pokazuje, spróbuj wyregulować wbudowany potencjometr kontrastu.

Podsumowanie

No i mamy to! Zbudowaliśmy cyfrowy system ostrzegania przed nadmierną temperaturą, który zapala diodę LED, gdy temperatura zadana zostanie przekroczona i wyświetla

Pliki XOD omawiane w tym artykule są dostępne do pobrania ze strony EPE (www.epemag3.com).

na LCD aktualną temperaturę oraz temperaturę zadaną. Tak, jest to podstawowy układ, ale nie wymagał on żadnego kodu i był szybki i łatwy do złożenia.

Teraz zastanów się, jak łatwo można go rozszerzyć lub zmienić – na przykład:

- Zmiana czasu uśredniania temperatury wymaga zmiany tylko jednej liczby w oprogramowaniu.
- Dokładne dostrojenie dokładności czujnika temperatury to po prostu zmiana liczby w offsecie.
- Dodanie histerezy wymaga tylko jednego dodatkowego pola programowania (blok histerezy jest dostępny w bibliotekach).
- Zmiana tego, co jest wyświetlane na wyświetlaczu LCD (np. wyświetlanie czasu uśredniania temperatury zamiast wartości zadanej) jest tak prosta, jak narysowanie nowej linii łączącej i zmiana etykiety tekstowej.
- Porównywanie dwóch temperatur wymaga jedynie dodania kolejnego czujnika temperatury i kilku bloków programowania.

A oprogramowanie XOD? Dowód na to, że pudding jest w jedzeniu – w tej historii nie tylko przedstawiliśmy oprogramowanie do programowania, ale także udało nam się zbudować kompletny projekt!

Co dalej?

Mam nadzieję, że ten artykuł zainspirował Cię do wypróbowania XOD z Arduino – w kolejnych numerach pojawi się wiele innych projektów opartych na XOD, które skuszą nowicjuszy do odkrycia mocy cyfrowego sterowania. ■

Julian Edgar

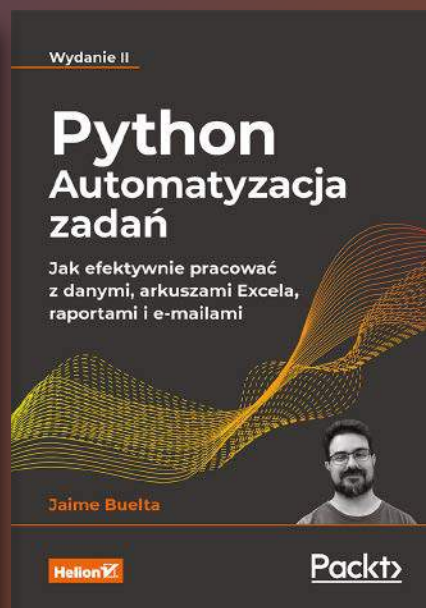
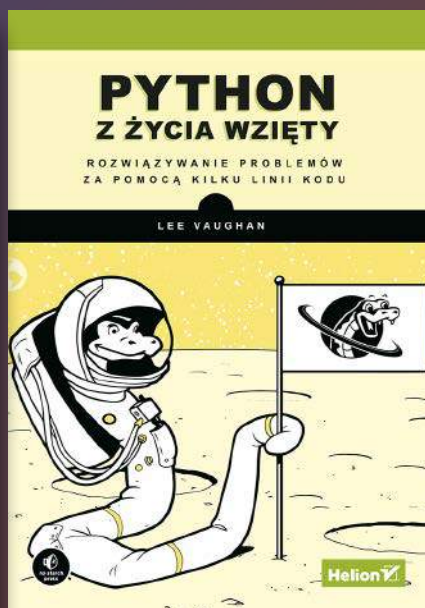
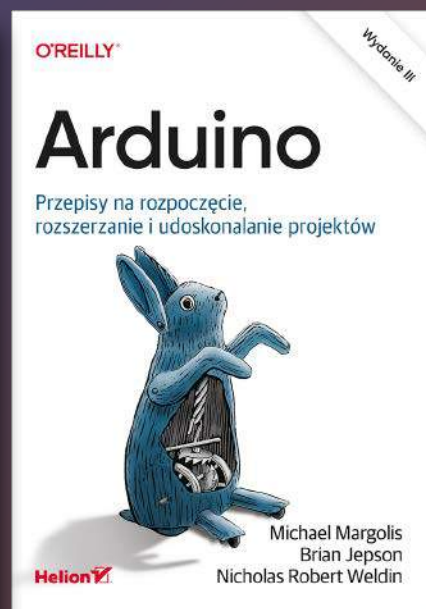
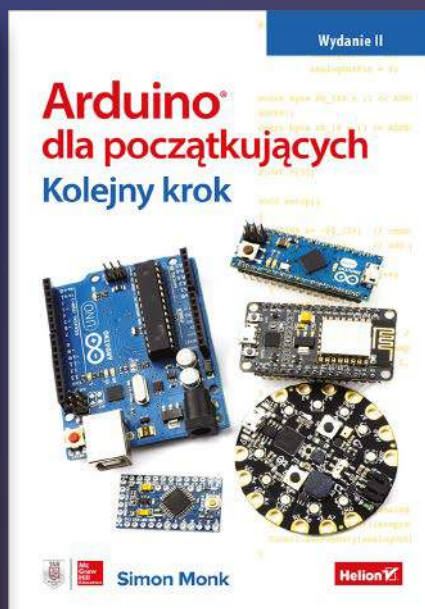
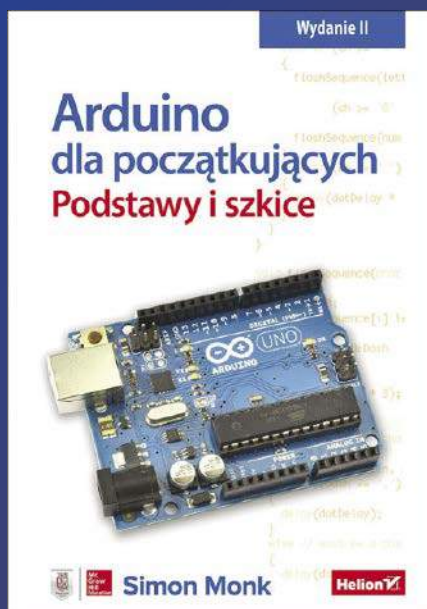
Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, marzec 2020 (www.epemag3.com)

świat radio
Magazyn wszystkich użytkowników eteru
KROTKOFALARSTWO CB · RADIOTECHNIKA

przejrzyj i kup na
www.ulubionykiosk.pl



KSIĄŻKI W ULUBIONYM KIOSKU Z RABATEM DO 30%



Zobacz pełną ofertę – ponad 500 tytułów!
Zamów wygodnie na UlubionyKiosk.pl

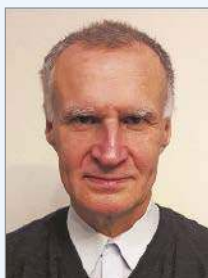
Uczmy się na cudzych błędach

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 natrafiliście 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.

Alarm do jednośladu

Ten prosty układ zabezpieczy przed kradzieżą twój rower lub skuter. Przy próbie „odprowadzenia” twojego jednośladu, rozlegnie się alarm informujący o próbie kradzieży. Wszystko co musisz zrobić, to przeciągnąć cienki drucik przez szprychy koła (najlepiej przedniego), tak aby zerwał się gdy nastąpi obrót koła. Sytuacja taka wystąpi gdy ktoś spróbuje ukraść rower. Wówczas z głośniczka wydobydzie się dźwięk przypominający wystrzały z broni maszynowej (lub podobny; w układzie scalonym jest możliwość wyboru trzech trybów-dźwięków), co zaalarmuje otoczenie.

Rysunek 1 pokazuje zmontowany na płytce uniwersalnej prototyp poddany testom w laboratorium „Electronics For You”. Schemat ideowy alarmu pokazano na rysunku 2.

Opis układu i jego działanie

Alarm wykorzystuje układ scalony UM3562, który jest generatorem dźwięków przypominających wystrzały z broni maszynowej. Dodatkowymi elementami są: tranzystor BC547 (T1), SL100 (T2), dioda Zenera 3,1 V (ZD1), dioda prostownicza 1N4007 (D1) i głośniczek (LS1).

Sercem alarmu jest układ scalony UM3562. To małej mocy CMOS dużej skali integracji LSI zamknięty w obudowie ośmio-nóżkowej. Układ scalony zawiera generator przestrzajany napięciem VCO, generator-timer, klucz rozpoznający trzy stany wejścia, obwód modulatora obwiedni generatora oraz mieszacz. Nóżka 2 tego IC przewidziana jest do wyboru trybu dźwięku generowanego na wyjście. Rozpoznaje on trzy stany. Zwarcie do masy wybiera specyficzny tryb przerywanych dźwięków. Zwarcie nóżki 2 do zasilania +3 V wybiera tryb – strzelby-karabinu. Pozostawienie nóżki 2 nie podłączonej skutkuje dźwiękiem przypominającym strzały broni maszynowej. Nóżka 4 to Trigger, czyli wyzwolenie alarmu-dźwięku. W wykorzystanej aplikacji połączono ją na stałe z masą. Noga 2 jest natomiast połączona z masą drucikiem, który ma ulec zerwaniu. Wykorzystano tu jeden tryb „broni maszynowej”, aczkolwiek potencjometrem VR1 można ustawić częstotliwość wystrzałów.

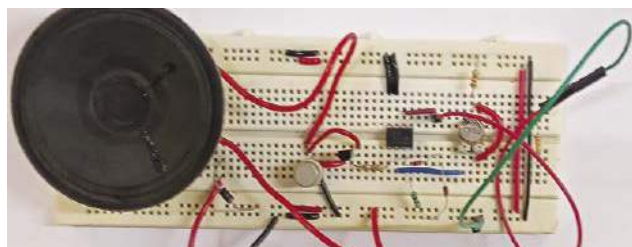
Układ scalony UM3562 jest przewidziany do zasilania z baterii 3-woltowej i może to być przedział od 2,4 V do 3,6 V. Tutaj zastosowano baterijkę 9 V, aby wydawany z głośniczka dźwięk był wystarczająco głośny. Zatem, w celu zasilania układu scalonego napięcie baterii jest redukowane do 3,1 V przy pomocy rezystora R2 i diody Zenera ZD1.

Wyjściem sygnału audio jest nóżka 5. Dwa tranzystory pracujące w połączeniu Darlingtona służą do wzmocnienia sygnału.

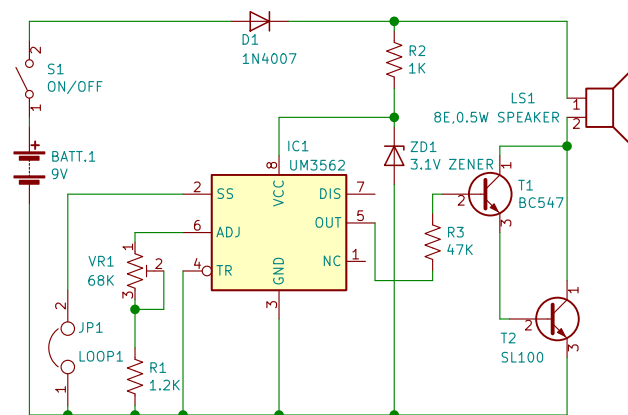
Praca układu jest zatem prosta. Po zamontowaniu, przeciągnięciu drucika przez koło roweru, należy przełączyć przełącznik S1 w pozycję ON. To uzbroi alarm. Jeśli teraz ktoś będzie chciał rower „podprowadzić”, drucik ulegnie zerwaniu i alarm poinformuje o próbie kradzieży.

Konstrukcja mechaniczna i przetestowanie działania alarmu

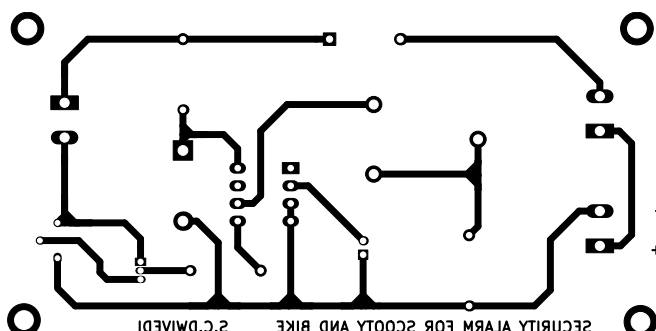
Układ zmontowano na jednostronnej płytce PCB. Można wprost skorzystać z rysunku 3, gdyż jest to projekt druku w naturalnych wymiarach



Rysunek 1.



Rysunek 2.



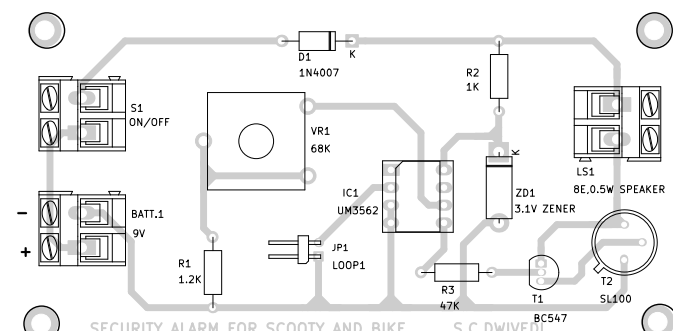
Rysunek 3.

(w papierowej wersji EdW). Rozmieszczenie elementów na płytce pokazuje rysunek 4. Do umieszczenia płytki PCB wraz z baterią i głośnikiem należy przygotować odpowiedniej wielkości pudełko-obudowę. Należy przygotować wygodne złącze dla podłączenia drucika (od Red. EdW: tę czynność trzeba wykonać nie tylko przy każdorazowej próbie kradzieży).

Po podłączeniu baterijki, przełącz S1 w pozycję ON i schowaj gdzieś układ, może pod koło zapasowe (w rowerze może być trudno, a bateria może raczej być obecna cały czas!). Teraz drucik należy przeprowadzić przez koło starannie, tak aby mieć pewność, że ulegnie zerwaniu gdy nastąpi obrót koła.

Na rysunku 5 pokazano wykonany przez autora prototyp alarmu zmontowany na płytce uniwersalnej.

Ten artykuł był wcześniej opublikowany na łamach „EFY”, lipiec 2022 (efymag.com)



Rysunek 4.



Rysunek 5.

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

Półprzewodniki:

IC1 – UM3562 machine-gun sound generator
T1 – BC547 tranzystor NPN
T2 – SL100 tranzystor NPN
D1 – 1N4007 dioda prostownicza
ZD1 – dioda Zenera 3,1 V

Rezystory: (wszystkie 0,25 W, ±5%; jeśli nie zaznaczono inaczej)

R1 – 1,2 kΩ
R2 – 1 kΩ
R3 – 47 kΩ
VR1 – potencjometr 68 kΩ

Inne:

S1 – przełącznik on/off
LS1 – głośniczek 8 Ω/0,5 W
BATT.1 – bateria 9 V
ponadto: cienki elastyczny drucik

Uwagi i poprawki

Brak noty aplikacyjnej układu scalonego UM3562. Natomiast budzi wątpliwość trwałe połączenie wyprowadzenia 4 z masą oraz połączenie nóżki 2 także z masą przez drucik który ma ulec zerwaniu. Skoro n.4 to Trigger (prawdopodobnie aktywny stanem wysokim, choć „kółeczko” przy tej nóżce sugeruje co innego – stan niski), a n.2 SS to „mode” rozpoznający stan niski, wysoki i „floating”, to prawdopodobnie należy zewrzeć nogi 2 i 4, i drucikiem przeciągnąć do masy. Wtedy po zerwaniu drutu rozlegnie się dźwięk strzałów. Przy połączeniu jak na schemacie (rysunek 2), układ powinien być cały czas cicho, jeśli Trigger jest aktywny stanem wysokim. Lub powinien cały czas wydawać dźwięki wg trybu odpowiadającego zwarcia SS z masą (jakieś „special kind of pulsating sound” – z oryginalnego tekstu) jeśliby Trigger był aktywny stanem niskim.

Wydaje się, że w tej „zabawce” zbyteczny jest potencjometr dla ustawienia częstotliwości wystrzałów alarmu.

To proste rozwiązanie jest być może ciekawym, jednym z wielu pomysłów, jak odstraszyć złodzieja twojego roweru. Można jednak wątpić

w jego skuteczność. Układ scalony UM3562 jest ciekawym układem wydającym odgłosy strzałów broni maszynowej lub strzelby. To IC przeznaczony do zabawek i jako taki zapewne spełni swoje zadanie. UM3562 zasilany jest napięciem 3 V, ale nawet przy baterijnym zasilaniu 9-cio woltowym głośnika, wątpliwe aby skutecznie złodzieja odstraszył (aczkolwiek możliwym jest, aby dźwięk był „dość głośny”).

Tranzystory T1 i T2 trudno traktować jako wzmacniacz, jak można wnioskować z oryginalnego tekstu. To jedynie klucz, aczkolwiek skutecznie przeniesie modulację częstotliwości, a raczej fazy wyjściowego sygnału audio. W takim razie, za cenę jakości generowanego dźwięku można się spodziewać całkiem przyzwoitej amplitudy i głośności. Ograniczenie w tym zakresie leży raczej po stronie baterii.

Zabezpieczenie roweru wymaga każdorazowego podłączenia drucika, który przy próbie kradzieży ma się zerwać. Pomijając niewygodę takiego sposobu, który jest pewnie do przezwyciężenia montując (przy kole) jakiś kontaktron, pozostaje problem trudniejszy. Gdzie w rowerze (a nawet w skuterze czy motocyklu) schować układ a szczególnie głośniczek, aby złodziej nie mógł go natychmiast urwać? Pewnie jakimś rozwiązaniem może być szczelna obudowa, aby alarm był nawet gdy złodziej schowa go do kieszeni. No właśnie, czy będzie wy? Tzn. strzelał wystarczająco głośno? A przecież wyłącznik S1 musi być na zewnątrz obudowy. Trudno liczyć na złodzieja, który nie zorientuje się w sprawie. Osobiście, przypnę jednak mój rower do słupa lub płotu tradycyjną linką. ■

REKLAMA



Elementy do Twojego projektu możesz kupić na: sklep.avt.pl

Miliomierz w zakresie od 0,1 do 1 Ohma

W obwodach elektronicznych mamy do czynienia z bardzo szerokim zakresem stosowanych rezystancji. W szczególności tzw. „shunt rezystory” (boczniki) mają zwykle wartości poniżej jednego ohma. Pomiar standardowym miernikiem tak małych wartości jest problematyczny z uwagi na zakres i rozdzielczość najniższego zakresu pomiarowego typowych multimetrów. W tym zakresie dobrym rozwiązaniem są mostki pomiarowe, aczkolwiek w amatorskiej praktyce dość niewygodne. Niżej prezentowany projekt pozwala na pomiar rezystorów w zakresie 0,1 do 1 Ω . Prototyp „on breadboard” wykonany w laboratorium „Electronics For You” widzimy na zdjęciu – rysunek 1.

Opis układu i jego działanie

Schemat ideowy miliomierza pokazuje rysunek 2. Układ wykorzystuje płytkę Arduino w wersji UNO, stabilizator liniowy LDO typu MIC5219 oraz wyświetlacz LCD, dwa rzędy po 16 znaków.

Gdy przez rezystancję R płynie prąd I , napięcie na rezystorze jest proporcjonalne do obu tych wielkości. To prawo Ohma i proponowany układ wykorzystuje je do pomiaru nieznanego rezystancji R . Mierzony rezystor należy podłączyć między punktami A i B. Układ omiara mierzy w istocie napięcie podane na wejście analogowe A0 płytki Arduino. Układ wykorzystuje przetwarzanie analogowo-cyfrowe, a jako napięcie odniesienia włączone jest referencyjne 1,1 V. Zakres pomiarowy zdefiniowany jest tą wartością, wydajnością źródła prądowego oraz rozdzielczością przetwornika w mikrokontrolerze. Mózgiem układu jest Arduino UNO, który po stronie cyfrowej dokonuje kalkulacji wartości mierzonej rezystancji i obsługuje wyświetlacz LCD1 w celu wskazania tej wartości. W celu wykonania pomiaru, należy „wystartować” naciskając niestabilny przycisk S1. Arduino włączy źródło prądowe 100 mA i dokona pomiaru napięcia na mierzonej rezystancji. Wartość przeliczy zgodnie z obowiązującym prawem Ohma i wyświetli ją na „displayu”. Mierzona rezystancja powinna mieścić się w zakresie do jednego ohma. Jeśli będzie wyższa, mikroprocesor wyśle na wyświetlacz komunikat OL lub NC. W tym przypadku zostanie także wyłączone źródło prądowe będące podstawą działania proponowanego miernika. Ponowne wystartowanie przyciskiem S1 będzie także nieskuteczne.

Arduino Uno to system „open source” wykorzystujący mikrokontroler ATmega 328P. Zawiera on czternaście linii – cyfrowych wejść/wyjść. Sześć linii może być skonfigurowane jako wyjścia PWM, także sześć jako analogowe wejścia. Oprócz mikrokontrolera, płytkę Arduino zawiera rezonator kwarcowy 16 MHz, złącze USB, złącze zasilania typu jack, a także „ICSP header”, które można wykorzystać w procesie programowania procesora. Przepisanie programu mającego obsługiwać aplikację jest bardzo proste. W chronionym obszarze pamięci ATmega328 zainstalowany jest bootloader i dla procesu programowania nie potrzeba żadnego dodatkowego hardware-u. Naszą aplikację obsługuje software mR_Meter. Ino zwany popularnie szkicem. Program napisano w bliskim C języku Arduino, a kompilacją i przepisaniem software-u zajmie się darmowe oprogramowanie Arduino IDE.

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

Półprzewodniki:

płytkę Arduino UNO
IC1 – MIC5219 – stabilizator LDO

Rezystory: (wszystkie 0,25 W, $\pm 5\%$)

R1 – 33 Ω

R2 – 150 Ω

R3 – 10 k Ω

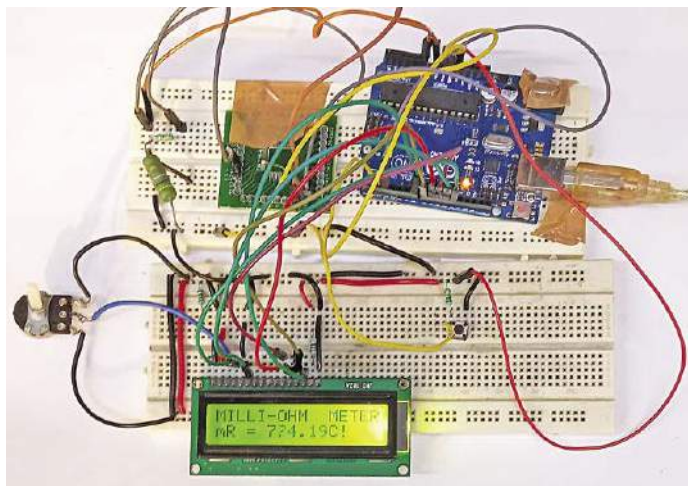
VR1 – potencjometr 10 k Ω

Inne:

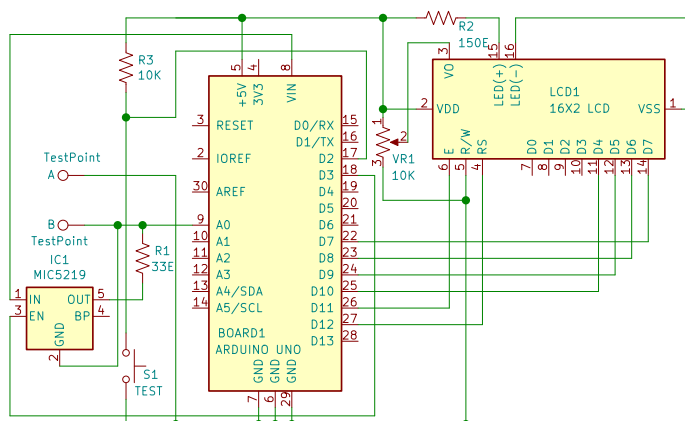
LCD1 – wyświetlacz LCD (2 rzędy 16 znaków)

S1 – przycisk niestabilny

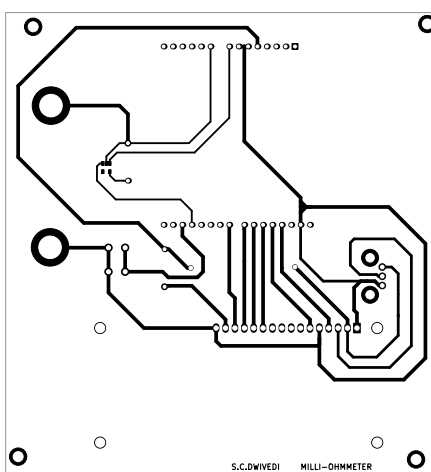
końcówki pomiarowe, zworki „jumpers”, bateria 9 V lub adapter 9 V/500 mA, rezystor o znanej wartości (np. 680 m Ω) w celu skalibrowania i przetestowania miernika



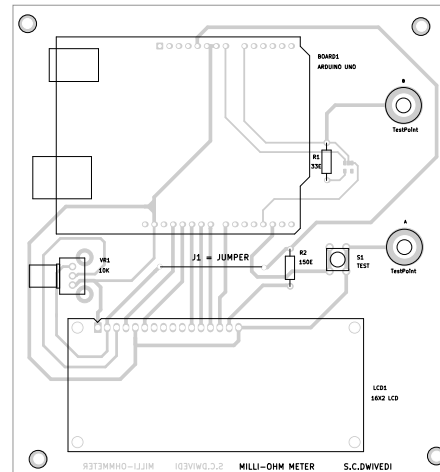
Rysunek 1. Prototyp autora zmontowany na płytce uniwersalnej



Rysunek 2. Schemat ideowy mili-ohmo-mierza



Rysunek 3. Jednostronna płytka PCB w skali 1:1



Rysunek 4. Rozmieszczenie elementów na PCB

Konstrukcja miernika i jego testowanie

Na **rysunku 3** jest proponowany projekt druku płytki PCB. To jednostronny druk i rysunek w papierowej wersji EdW powinien odpowiadać rzeczywistej wielkości płytki. Rozmieszczenie elementów na PCB pokazuje **rysunek 4**.

Przed montażem mechanicznym należy „wypalić” software mR_Meter.ino w mikrokontrolerze Arduino. Prace mechaniczne sprowadzą się do umieszczenia uzbrojonej płytki PCB w odpowiednio przygotowanej obudowie. Przycisk S1 należy umieścić w wygodnym miejscu, gdyż każdy pomiar wymaga „startowania”. Jako źródło zasilania można wykorzystać 9-cio voltową baterijkę lub adapter 9 V/500 mA, który należy podłączyć do wejścia zasilania obecnego na płytce Arduino. Dla wygody użytkownika, należy także zadbać o odpowiednie złącza A i B, do których będzie podłączana mierzona rezystancja.

Uwagi „Testera” Redakcji EFY

Obwód mierzy niskie rezystancje w zakresie 0,1 do 1 Ω . Należy się liczyć ze sporym błędem pomiarowym jeśli porównywalnej wielkości są pasożytnicze rezystancje styków, przewodów i/lub ścieżek w części analogowej miernika.

W miejsce stabilizatora LDO MIC5219 – IC1 można zastosować zamiennik – podobny stabilizator typu Low Drop Out.

Można zaproponować zastosowanie dodatkowej diody Zenera (o napięciu ok. 1,5 V) między punktami pomiarowymi A i B, w celu by-passu dla prądu stabilizatora, gdy zostanie on włączony, a nie podłączono rezystora pomiarowego z prze-

Ten artykuł był wcześniej opublikowany na łamach „EFY”, czerwiec 2022 (efymag.com)

A. Samiuddhin

Uwagi i poprawki

Układ zapewne działa, ale zawiera kardynalne błędy projektowe.

Pomiar rezystancji jak autor pisze, bazuje na prawie Ohma, które wszyscy znamy. $U=RI$. To wzór, którego uczą już w szkole podstawowej. Ale co dalej? Pomiar bazuje na przetworniku analogowo-cyfrowym wbudowanym w mikrokontroler ATmega328. To przetwornik 10-cio bitowy, a więc o rozdzielczości tysiąca „schodków”. Mikroprocesor ma możliwość włączenia napięcia referencyjnego 1,1 V i takie wykorzystano. Oznaczałoby to, że jest możliwy pomiar napięcia z rozdzielczością jednego miliwolta. No tak, ale my mamy mierzyć rezystancję. Autor zaprzął do roboty prawo Ohma i przepuszcza przez nieznaną rezystancję znany prąd. Źródło prądowe jest ustawione na wartość 100 mA. A rezystancję chcemy mierzyć w zakresie do 1 Ω . 0,1 A przemnożone przez 1 Ω to 0,1 V. I to ma odpowiadać pełnemu zakresowi pomiarowemu! Wykorzystane jest więc mniej niż 10% zakresu pomiarowego przetwornika A/C. Układ zapewne działa. Przetwornik w μP ATmega328 gwarantuje dokładność pomiaru na poziomie 2 LSB, a więc całkiem nieźle. Ale to przykład, jak nie powinno się projektować układu wykorzystującego przetwarzanie analogowo-cyfrowe.

To nie jedyny błąd. Omomierz niskich wartości rezystancji powinien koniecznie wykorzystywać technikę pomiaru zwaną pomiarem Kelwina. Trzeba dołożyć dwa przewody, co jednak uniezależni pomiar nieznaną rezystancję od porównywalnych z nią „pasożytniczych” rezystancji obwodu pomiarowego.

Na tym lista błędów się nie kończy. Autor wykonał źródło prądowe (którego wymaga prawo Ohma) w oparciu o stabilizator LDO MIC5219. Jest to zapewne wersja „nieregulowana” 3,3-voltowa. Ze schematu odczytujemy $R1=33 \Omega$, a więc źródło prądowe powinno pompować prąd równy 100 miliamperów. Ale MIC5219 nie jest tu najszcześliwszym wyborem. Wydaje się, że lepszym byłby sędziwy „pływający” LM317 lub jakiś jego odpowiednik. Co prawda byłby problem z jego włączaniem/wyłączaniem sygnałem logicznym z mikrokontrolera. Problem niezbyt trudny. MIC5219 ma wejście Enable, i wydaje się, że cecha ta problem eliminuje całkowicie. Ale to nieprawda! Wykonując źródło prądowe na liniowym stabilizatorze napięcia autor wykorzystał technikę pływającego potencjału masy (stabilizatora LDO). Wejście Enable „rozumie” sygnał cyfrowy o standardowych poziomach logicznych CMOS lub TTL. Stan niski nie wyżej niż 0,4 V i jedynka logiczna nie gorzej niż 2 V. Wyjście

cyfrowe Arduino spełnia tę tolerancję, a więc w czym problem? Problem jest i to poważny, bo wejście Enable w MIC5219 odniesione jest względem jego GND, a tutaj masa „pływa”. Zasilanie stabilizatora liniowego pobrane jest z wejścia VIN płytki Arduino, a to poziom 9 V. Potencjały pływające są zależne od rezystancji pomiarowej między punktami A i B. Przy braku lub „większej” rezystancji między tymi punktami, potencjały pływające wyjścia i masy MIC5219 pójdą w kierunku napięcia VIN. Jest więc realna sytuacja, gdy zakresy wyjścia cyfrowego Arduino i wejścia EN stabilizatora LDO będą zupełnie niekompatybilne. Poprawka proponowana przez „Testera EFY” w punkcie 3-cim problemu tego nie rozwiązuje. ■

REKLAMA

KEY PRODUCENT AUTOMATYKI GRZEWCZEJ
11-200 Bartoszyce ul. Bohaterów Warszawy 67 **pwkey@onet.pl**
tel. (89)7635050 fax (89)7635051

TANIE REGULATORY

DO KOTŁÓW WĘGLOWYCH I NA DREWNO

z wbudowanym termostatem pokojowym
zapewniającym komfort i oszczędność



REGULATORY DO KOTŁÓW Z PODAJNIKIEM

REGULATORY POGODOWE

- Prosta obsługa, bogate możliwości programowania
- Możliwość dopasowania do każdego kotła i rodzaju paliwa
- Wysoka jakość
- Gwarancja 24 miesiące

www.pwkey.pl

Wyświetlacz na matrycy diod LED

Pełny tytuł tego projektu, to wyświetlacz wykonany w oparciu o matrycę punktów (diod LED) z możliwością przewijania obrazu i wykorzystujący sterownik z mikrokontrolerem serii STM32. Wyświetlacze tego typu zyskują dużą popularność w branży reklamowej. Projekt, który tu prezentujemy potrafi wyświetlać informację wprost z twojego smartfona wykorzystując łącznie bezprzewodowe.

W tym celu wykorzystano moduł Bluetooth HC-05. Cały system koordynuje mikrokontroler rodziny STM32. Software w nim zawarty musi zaprogramować współpracujące komponenty. Musi także odczytać informację z modułu Bluetooth i przesłać ją do drivera, który obsłuży wyświetlacz graficzny złożony z kilku paneli matrycy, złożonych z kolorowych diod LED. Funkcja przewijania obrazu pozwala na wyświetlenie większej ilości informacji aniżeli statycznie mieści się w ramach rozdzielczości wyświetlacza graficznego. „Scrolling” obrazu realizowany jest też programowo. Płytką z mikrokontrolerem STM32 nosi nazwę Blue Pill i zawiera mikrokontroler STM32F103C8T6. To wydajny procesor w architekturze trzydziestodwu bitowej z szybkim zegarem, z rdzeniem ARM Cortex M3. To procesor firmy STMicroelectronics predysponowany do aplikacji wymagających dużych mocy obliczeniowych, a równocześnie oszczędnych w zakresie mocy zasilania. Matryce LED-owe są powszechnie spotykane na dworcach kolejowych czy lotniskach. Wyświetlają bieżący czas oraz godziny przyjazdu i odjazdu pociągów. To samo na lotniskach, oczekiwania i potrzeby w zakresie szeroko rozumianych wyświetlaczy są te same.

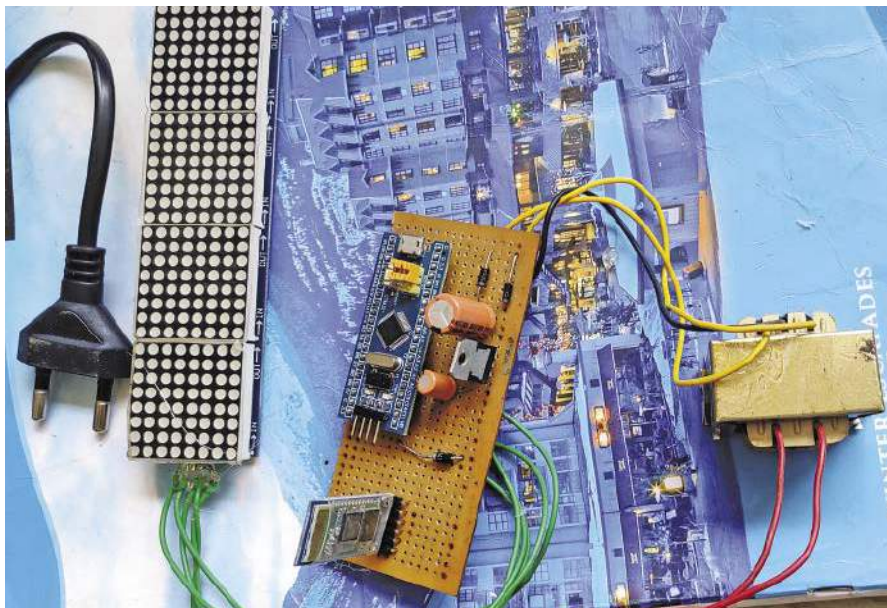
Bieżący projekt to wyświetlacz o mniejszych rozmiarach, aczkolwiek w zakresie obsługiwanej go elektroniki nie ustępuje przytoczonym wyżej aplikacjom. Prototyp wykonany przez autora pokazuje zdjęcie na **rysunku 1**, a spis

potrzebnych podzespołów zawiera poniższe zestawienie.

Na **rysunku 2** pokazano matrycę wyświetlacza LED-owego, gdzie widać sterownik MAX7219. Pojedynczy panel składa się z 64-ech diod LED w matrycy 8×8. LED-y są kolorowe i obsługiwane są w trybie wyświetlacza graficznego. Dlatego mimo skromnej rozdzielczości, można tworzyć stosunkowo bogate obrazy. Komunikacja między mikrokontrolerem a driverem wyświetlacza jest szeregową,

dlatego wystarczy połączyć te moduły za pomocą dwóch przewodów. Tu wykorzystano mikrokontroler STM32, aczkolwiek może być np. ESP32, NodeMCU czy popularne Arduino.

Moduły 8×8 LED można łączyć kaskadowo zwiększając powierzchnię i rozdzielczość wyświetlacza. Autor wykorzystał cztery takie moduły. Hardware w zakresie łączenia szeregowego modułów jest bardzo prosty i nie wymaga zwielokrotnienia ilości połączeń z mikrokontrolerem. Matryce 8×8



Rysunek 1. Prototyp wyświetlacza wykonany przez autora

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

MAX7219 – matryca wyświetlacza LED-owego (1088AS)
STM32 – płyta Blue Pill z mikrokontrolerem
HC-05 – moduł Bluetooth
zasilacz stabilizowany 5 V
uniwersalna płytka stykowa „breadboard”
zworki dla „breadboard”

Elementy zasilacza

Półprzewodniki:

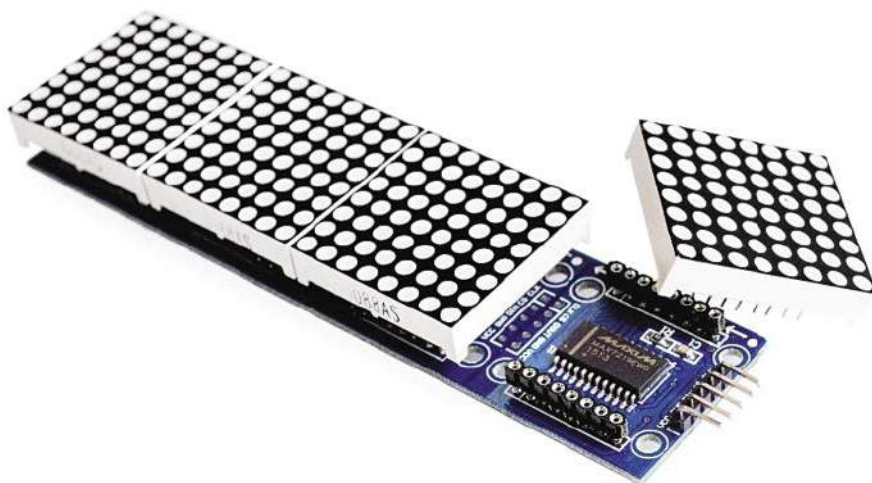
IC1: 7805 – scalony stabilizator 5 V
D1, D2: 1N4007 – diody prostownicze

Kondensatory:

C1: 1000 µF/25 V – elektrolityczny
C2: 100 µF/25 V – elektrolityczny

Inne:

CON1: złącze 2-pinowe
X1: transformator sieciowy – 230 VAC – uzwojenie pierwotne; wtórne: 12-0-12 VAC 500 mA



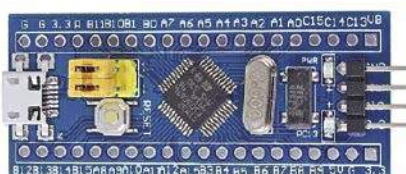
Rysunek 2. Wyświetlacz czterech paneli LED-owych ze sterownikiem MAX7219

LED wyposażone są w piny wejścia i wyjścia. Wystarczy połączyć wyjście wcześniejszej kaskady z wejściem następnej, tworząc w ten sposób wyświetlacz o dowolnej długości (zachowując wysokość ośmiu diod).

Driver MAX7219 obsługuje matryce LED w konfiguracji wspólnej katody. Szeregowa transmisja upraszcza hardware, komplikując nieco stronę programową. Aczkolwiek każdy z popularnych mikrokontrolerów można zaprogramować do takiej obsługi. Tak utworzona magistrala przesyła dane i potrafi zaadresować każdy segment wyświetlacza. Przesłanie danych do dalszych segmentów nie wymaga odświeżania zawartości całego wyświetlacza. Nie



Rysunek 3. Wygląd modułu Bluetooth HC-05



Rysunek 4. Płytki Blue Pill z mikrokontrolerem STM32F103C8T6

ma jedynie możliwości programowego sterowania jasnością wyświetlanej informacji. Na każdym module znajduje się rezystor, który ustala prąd diod w każdym segmencie matrycy.

Do komunikacji Bluetooth wykorzystano popularny moduł o oznaczeniu HC-05. To niewielkich rozmiarów płytka o wszechstronnych możliwościach. Można ją zastosować np. w bezprzewodowych słuchawkach, kontrolerach gier, bezprzewodowej klawiatury lub myszy i wielu innych podobnych aplikacjach. Zasięg łącza Bluetooth sięga do 100 metrów w wolnej przestrzeni, co jest także uzależnione czynnikami atmosferycznymi i ukształtowaniem terenu. W praktyce jest on ograniczony czynnikami urbanistycznymi, czyli przeszkodami na drodze łącza bezprzewodowego. Moduł HC-05 wykorzystuje protokół zgodny ze standardem IEEE802.15.1, który pozwala na tworzenie prywatnych sieci PAN (Personal Area Network). Radiowa łączność bezprzewodowa odbywa się w paśmie rozproszonym z wykorzystaniem techniki „frequency hopping”. Łączność między modulem Bluetooth i mikrokontrolerem wykorzystuje standard transmisji szeregową z typową szybkością 9600 bitów na sekundę. Dzięki transmisji szeregową wielkość modułu i ilość wyprowadzeń ograniczona jest do minimum, co widać na zdjęciu – **rysunek 3**. Oprócz dwóch wyprowadzeń koniecznych do zasilania, komunikacja jest dwustronna po liniach RXD i TXD. Dodatkowe dwie linie stanowią sygnały EN i STATE. HC-05 można zaprogramować do pracy w konfiguracji master lub slave na magistrali łączącej go z mikrokontrolerem.

Na **rysunku 4** widać zdjęcie modułu Blue Pill z mikroprocesorem STM32F103C8T6. Hardware tego modułu jest „open source” i jest to rozwiązanie bardziej ekonomiczne aniżeli wykorzystanie np. płytki Arduino. Mimo dużej ilości pinów mikrokontrolera, płytka Blue Pill jest niewielkich rozmiarów i zawiera rezonator kwarcowy 8 MHz i 32 kHz. Pierwszy z nich stanowi zegar systemu, 32 kHz przewidziany jest dla zegara czasu rzeczywistego RTC (Real Time Clock). W oznaczeniu mikrokontrolera

STM32F103C8T6 zawarte są następujące informacje:

- STM oznacza rodzinę μP firmy STMicroelectronics
- 32 – oznacza trzydziestodwubitową architekturę ARM
- F103 – oznacza architekturę ARM Cortex M3
- C – obudowa 48-mio pinowa
- 8 – oznacza wielkość pamięci Flash 64 kB
- T – obudowa LQFP
- 6 – zakres temperatury -40°C do $+85^{\circ}\text{C}$

Mikrokontroler ten, mimo rozbudowanej architektury, zadowala się niewielką mocą zasilania. Ma wbudowane mechanizmy płytkiego i głębokiego uśpienia, dzięki czemu nadaje się do aplikacji o zasilaniu baterijnym. W poniższej tabeli wyszczególniono zestawienie cech mikrokontrolera STM32 na module Blue Pill.

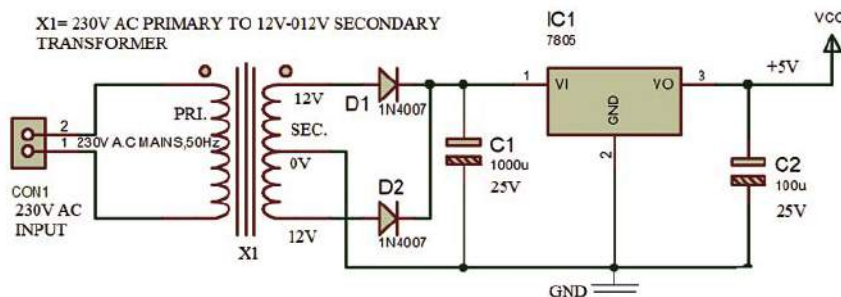
Schemat i działanie układu

Schemat wyświetlacza wykorzystującego matrycę diod LED pokazano na dwóch rysunkach. Osobno zasilacz i zasadniczą część komunikacji między zastosowanymi modułami i matrycą wyświetlacza. Zasilacz jest na **rysunku 5**. Wykorzystano transformator sieciowy o dwu uzwojeniach wtórnych, dający symetryczne napięcie 12 VAC (X1 230 VAC/12 V-0-12 V 500 mA). Dzięki temu prostownik dwupołkowy zawiera jedynie dwie diody 1N4007 (D1 i D2). Z uwagi na niewielką moc, zasilacz ten wykorzystuje stabilizator liniowy – IC1 typu 7805. Dodatkowymi elementami są jedynie dwa kondensatory filtrujące napięcie na wejściu i na wyjściu stabilizatora: C1=1000 μF /25 V i C2=100 μF /25 V.

Zasilanie pięciowoltowe jest adekwatne dla wszystkich zastosowanych modułów. Bluetooth HC-05 może być zasilany napięciem z przedziału 3,3 V do 6 V, driver matrycy LED-ów przyjmie zasilanie od 3,3 V do 5 V, a dla płytki Blue Pill nominalne zasilanie to +5 V.

Specyfikacja STM32 (BLUE PILL):

Architektura – 32-bitowa ARM Cortex M3
Napięcie zasilania: 2,7 V do 3,6 V
Częstotliwość zegara CPU – 72 MHz
Ilość wyprowadzeń GPIO – 37
Ilość pinów PWM – 12
Wejścia analogowe – 10 (dwunasto-bitowe)
Liczba USART – 3
Magistrale I²C – 2
Magistrale SPI – 2
Liczba CAN 2.0 – 1
Timery – 3 (16-to bitowe), (1 PWM)
Wielkość Flash Memory – 64 kB
Pamięć RAM – 20 kB



Rysunek 5. Schemat ideowy zasilacza

Schemat ideowy i blokowy, wg którego łączone są podzespoły pokazano na rysunkach 6 i 7. Nie ma dużej różnicy między ujęciem schematu jako blokowy i ideowy. Łączone ze sobą podzespoły traktowane są bowiem jako elementy. Na rysunku 6 widać, iż STM32 łączony jest z jednej strony z driverem matrycy wyświetlacza MAX7219, a z drugiej strony z modulem HC-05. Czwartym niezbędnym podzespołem jest telefon bądź inne urządzenie z systemem Android. Tu połączenie narysowano „symbolicznie”, gdyż jest to łącze bezprzewodowe.

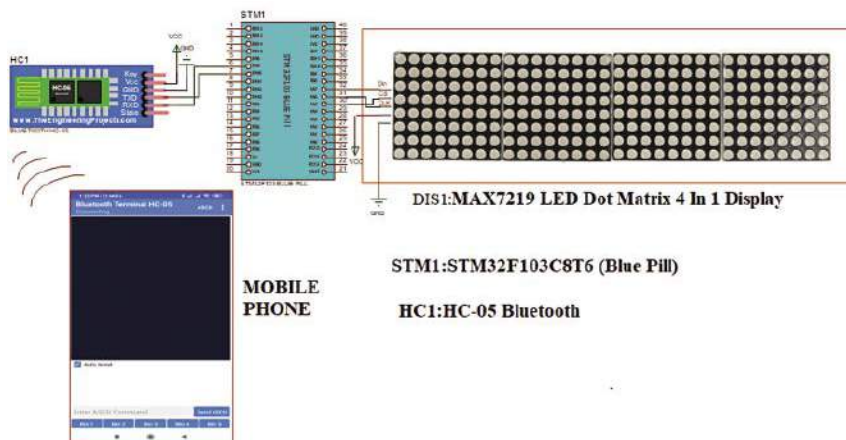
Komunikacja między STM32 i MAX7219 wykorzystuje standard SPI i jest to transmisja jednokierunkowa. Angażuje ona trzy wyprowadzenia mikrokontrolera: Data, Clock i Chip Select. Od strony MAX7219 wykorzystane są piny o oznaczeniu: CS, CLK i DIN. W STM32 do tej transmisji oddelegowano (oprogramowano) 3 piny portu PA: PA4, PA5 i PA7. Łącze między HC-05 i STM32 jest również „oszczędne”. Poza masą i zasilaniem, wystarczą dwa przewody: TX do PA9 i RX do PA10 mikrokontrolera.

Sterowanie wyświetlacza LED-owego z aplikacji urządzenia Android

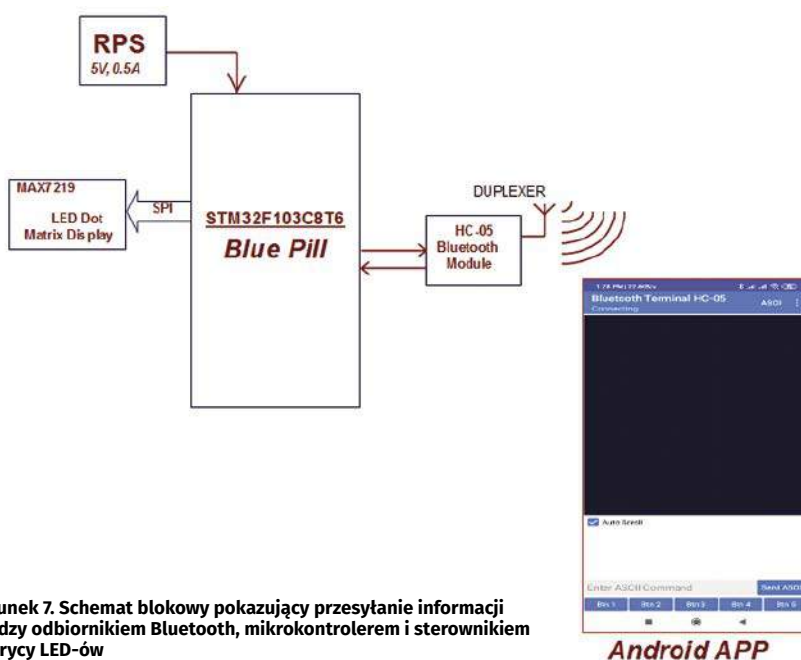
Aby połączyć telefon (lub inne urządzenie z systemem operacyjnym Android) z odbiornikiem HC-05 konieczna jest dedykowana aplikacja na telefonie. Jest ona dostępna w sklepie Google Play Store. To aplikacja darmowa i można ją ściągnąć pod adresem: https://play.google.com/store/apps/details?id=project.bluetoothterminal&hl=en_IN.

Aplikacja ta powinna zgłosić się jako „Bluetooth Terminal HC-05” ze stroną główną jak na rysunku 8.

W ustawieniach są do wyboru dwa formaty przesyłanych znaków: HEX lub ASCII. Wybraliśmy format ASCII, w którym można bezpośrednio przysyłać znaki tekstowe, cyfry i znaki specjalne. Można teraz przetestować, czy łącze jest aktywne i czy cały system działa. Oszczędność hardware-u i „przewodów” łączących płytki, okupiona jest komplikacją po stronie programowej. Nie wnikając w software ściągniętej aplikacji, rozbudowany protokół transmisji Bluetooth oraz oprogramowanie STM32 tworzące kanały łączności między źródłem (HC-05) i portem docelowym, którym są rejestry drivera MAX7219, znak „wystukany” na telefonie powinien się pojawić na wyświetlaczu matrycy LED-owej. Poprawne działanie wymaga stosownych ustawień w aplikacji na telefonie. Rysunek 9 pokazuje zrzut ekranu, w którym oprócz wyboru formatu znaków ASCII, należy zaznaczyć symbole końca transmisji: \r i \n.



Rysunek 6. Schemat ideowy łączonych podzespołów



Rysunek 7. Schemat blokowy pokazujący przesyłanie informacji między odbiornikiem Bluetooth, mikrokontrolerem i sterownikiem matrycy LED-ów

Uwaga od Redakcji EFY

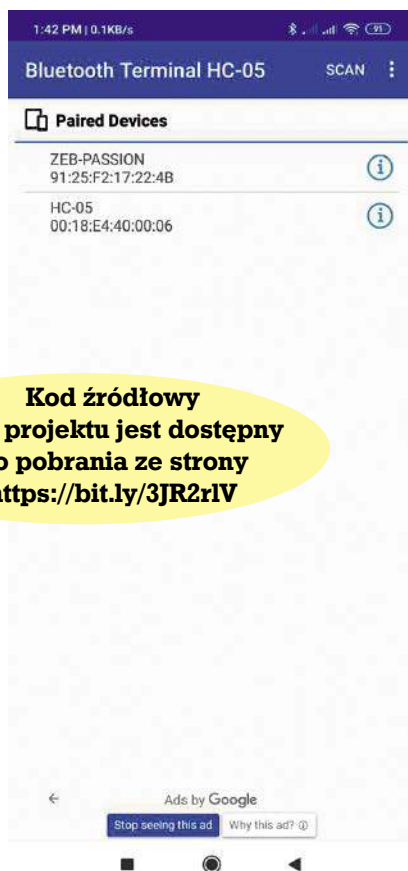
Aplikacja wykorzystuje typowe cechy łączności Bluetooth. Konieczne jest zatem sparowanie telefonu z odbiornikiem HC-05. Aplikacja pewnie zapyta o hasło konieczne do identyfikacji parowanych urządzeń. Domyślnym hasłem jest tu kod 1234.

Oprogramowanie

STM32F103C8T6 jest kolejnym mikrokontrolerem w rodzinie STMicroelectronics. Działanie zgodne z oczekiwaniami wymaga zaprogramowania – wpisania do pamięci Flash odpowiedniego, dedykowanego programu. W tym celu można użyć wszystkich dostępnych narzędzi programistycznych przygotowanych dla procesorów rodziny STM. Można użyć oprogramowania IDE Keil ARM MDK lub np. IAR Workbench, Atollic TrueStudio, MicroC Pro ARM, Crossworks ARM, Ride 7, czy PlatformIO+STM32.

W związku z popularnością Arduino i łatwością programowania na tej platformie, wiele osób jest obeznanych z Arduino i z zadowoleniem przyjmie możliwość wykorzystania Arduino do zaprogramowania mikrokontrolera STM32. Stworzenie takiej możliwości, być może w największym stopniu przyczynia się do popularyzacji takiego, nie innego systemu, mikrokontrolera. Do popularyzacji Arduino przyczynia się też bogata biblioteka programów (szkiców) na tę platformę. Wykorzystamy tę możliwość i zaprogramujemy STM32 z użyciem środowiska Arduino IDE.

Co prawda STM32 na płytce Blue Pill nie ma zaszytego bootloadera czyniącego go kompatybilnym w środowisku Arduino IDE. Blue Pill wyposażony jest w złącze mikro-USB i stosowny bootloader można ściągnąć i do STM32 wpisać. Niżej zamieszczony link kieruje do filmu YouTube, gdzie pokazano krok po kroku, jak należy postępować w celu



Kod źródłowy
tego projektu jest dostępny
do pobrania ze strony
<https://bit.ly/3JR2rIV>

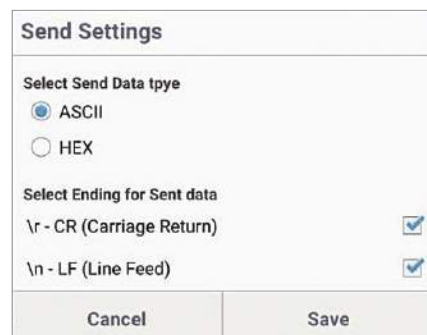
Rysunek 8. Wygląd strony głównej aplikacji – Bluetooth Terminal HC-05

wgrania odpowiedniego bootloadera: <https://www.youtube.com/watch?v=Tm7IWQLrKYs>.

Poniżej w kilku krokach wyszczególniono tok postępowania, który ułatwi przygotowanie Arduino IDE dla programowania mikrokontrolera z rodziny STM32.

Krok 1: Jeśli jeszcze nie zainstalowałeś Arduino IDE na swoim komputerze, w pierwszej kolejności uczynić to. Oprogramowanie to ściągniesz z Internetu; pamiętaj o poprawnym wyborze systemu operacyjnego zgodnego z zainstalowanym na twoim komputerze.

Krok 2: Po zainstalowaniu Arduino IDE, musisz jeszcze ściągnąć bootloader dla STM32. Znajdziesz go w zakładce File → Właściwości



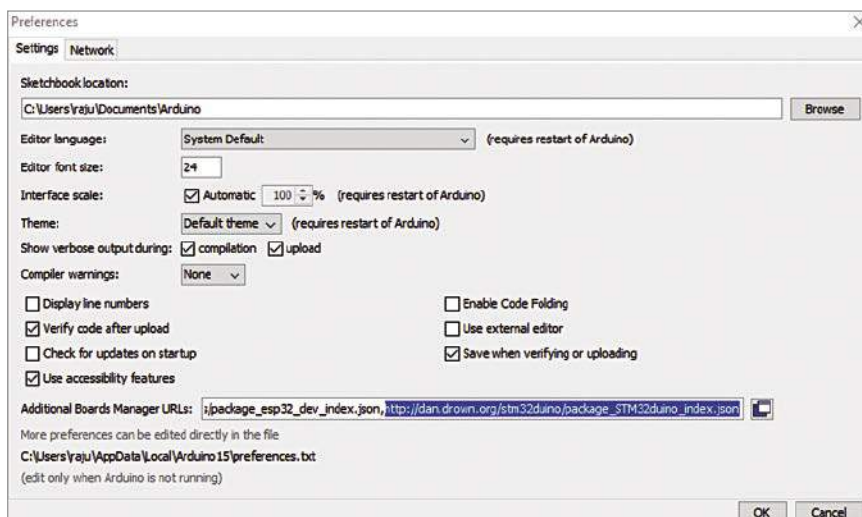
Rysunek 9. Konieczne ustawienia skojarzone z wyborem formatu znaków ASCII



Krok 3: Klikając zakładkę „Właściwości” powinno pokazać się okno dialogowe jak na **rysunku 10**. W linii Additional Boards Manager URL wybierz link: http://dan.drown.org/stm32duino/package_STM32duino_index.json.

Zatwierdź wybrane ustawienia klikając na OK.

Krok 4: Przejdź do narzędzi Tool → Boards → Board Manager. Powinno otworzyć się



Rysunek 10. Wygląd zakładki, gdzie należy wybrać właściwy link z menadżera „dodatkowych płytek” – „Additional Boards”

następne okno dialogowe menadżera „Boards Manager”. Wybierz STM32F1 i zainstaluj to oprogramowanie.

Krok 5: Po instalacji tego pakietu przejdź do „narzędzi”, gdzie powinno pokazać się okno zgodne z tym co pokazuje **rysunek 11**. Teraz zaznacz linię: Generic STM32F103C series. Zaznacz też prawidłowo pola: 64 kB pamięci Flash i zegar CPU – 72 MHz. Upewnij się, że „variant” jest zgodny z wersją mikrokontrolera, który chcesz zaprogramować. Upewnij się też, że zaznaczyłeś właściwy bootloader i ustaw „upload method” na – Serial.

Krok 6: Połącz płytke z STM32 z komputerem wykorzystując obecne na Blue Pill złącze mikro-USB. W menadżerze urządzeń sprawdź, który port COM został przydzielony dla tego łącza. Ten sam numer portu ustaw w narzędziach: Tools → Port.

Krok 7: Po wykonaniu powyższych czynności, zwróć uwagę na pozycję w prawym dolnym rogu zakładki programu Arduino IDE. Powinna tam być linia, jak pokazuje **rysunek 12**: Generic STM32F103C (Blue Pill) podłączony do COM9 (przykładowo). Jeśli poprawnie wykonałeś powyższe ustawienia, oprogramowanie Arduino IDE jest gotowe do zaprogramowania twojej płytki z mikrokontrolerem STMicroelectronics.

Wyżej wymienione kroki przygotowują sprzęt, aby można było programować STM32 z użyciem narzędzi platformy Arduino. Celem jest zaprogramowanie pamięci Flash tak, aby STM32 wykonywał zadanie jakie narzuca hardware zgodny ze schematem na rysunku 6. Program jest dostępny jako „szkielet” w bibliotece Arduino. Jest to software napisany w istocie w języku Arduino. Dla kompilacji i „załadowania” software-u użyto IDE w wersji 1.8.11. Dla poprawnego działania drivera matrycy

LED-ów, trzeba dołączyć pakiet programu o nazwie:

```
#include <MD_MAX72xx.h>
```

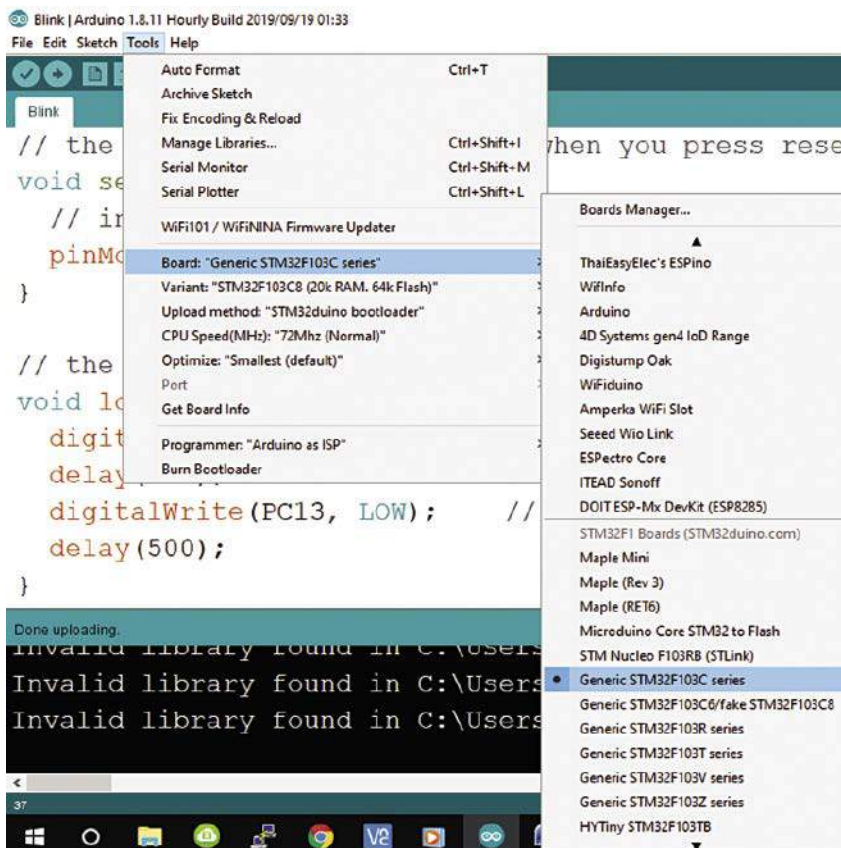
Tu jest zawarta biblioteka sterowników, aby MAX72xx obsługiwał matrycę 64-ech diod LED w sposób graficzny. Czyli, aby potrafił zaadresować i zaświecić lub zgasić każdy pojedynczy piksel, co jest zgodne z ideą wyświetlacza graficznego.

Konstrukcja i testowanie działania wyświetlacza

Pierwszą czynnością jest wpisanie kodu źródłowego *mains.ino* do STM32. Następnie należy połączyć moduły zgodnie ze schematem ideowym i podłączyć zasilanie 5 V z wcześniej przygotowanego zasilacza. Jeśli nie popełniłmy błędów, urządzenie powinno być gotowe i działać zgodnie z przewidywaniami.

Zatem, w celu przetestowania bieżącego projektu, należy rozpocząć od części software-owej projektu, którego celem jest wpisanie kodu źródłowego *mains.ino* do mikrokontrolera STM32. Wcześniej jednak musieliśmy zainstalować bootloader Arduino w mikrokontrolerze STM32. Można się też posłużyć modulem FTDI, który pozwoli na bezpośrednie połączenie interfejsu USB w komputerze z USART-em mikrokontrolera. Potrzebny jest odpowiednio przygotowany kabel USB-to-TTL i driver CP2102. Jedną stronę takiego konwertera łączymy z komputerem, drugi bezpośrednio z wyprowadzeniami mikrokontrolera skonfigurowanymi jako port szeregowy USART.

Łącząc moduł FTDI do STM32 potrzebne są cztery połączenia. RX i TX odpowiednio do pinów PA-9 i PA-10 STM32, oraz masę i zasilanie 3,3 V. Tryb „boot mode” w STM32 należy ustawić jako Boot-1. To powinno umożliwić transmisję szeregową między komputerem i docelowym modulem STM.



Rysunek 11. Dalsze ustawienia właściwe dla mikrokontrolera STM32F103C



Rysunek 12. Linia informująca o właściwym wyborze ustawień dla uP STM32F103C

Na etapie testowania można zaniedbać zasilacz 5-cio woltowy opisany wcześniej. Można moduł zasilic bezpośrednio ze złącza mikro-USB obecnego na płytce Blue Pill. Zasilaczem może być typowy adapter (ładowarka) z kablem zakończonym mikro-USB.

Dla uruchomienia łącza Bluetooth z modulem HC-05, należy zainstalować w telefonie wskazaną wcześniej aplikację. Testowanie urządzenia sprowadza się do sprawdzenia, czy znaki wysyłane z telefonu pojawiają się na wyświetlaczu złożonym z paneli matrycy diod LED. ■

Pamarthi Kanakaraja

REKLAMA

Kursy w Ulubionym Kiosku

IT i Hi-tech • Muzyka i Dźwięk

Pełna oferta na stronie

www.ulubionykiosk.pl

Już ponad rok publikujemy dla projektantów i programistów dla elektroniki. Odwiedź

ELPORTAL.pl

Podręczny monitor serca EKG

Każdy elektrokardiogram jest narzędziem monitorującym zdrowie twojego serca. Informacja „bicia serca” pozyskana jest elektrodami umieszczonymi na piersi i na ramionach pacjenta. Wzmocniony (i odfiltrowany) sygnał biopotencjałów pozyskany elektrodami jest zobrazowany przebiegiem wyświetlanym na monitorze bądź może być kreślony przy pomocy plotera na papierze. Przebieg EKG jest graficzną reprezentacją zmian potencjału w wybranych punktach ciała, wywołanych skurczami mięśnia sercowego.

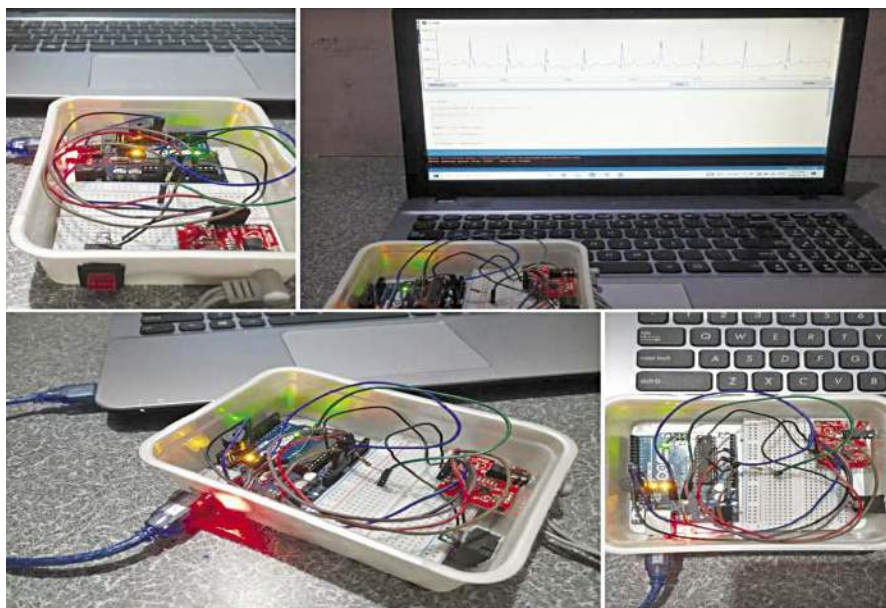
Prototyp podręcznego EKG wykonany przez autora pokazano na **rysunku 1**. Na **rysunku 2** zamieszczono schemat blokowy urządzenia. Wykorzystano tu tani moduł EKG oraz płytkę Arduino. To prosty system jednokanałowy nie wymagający wielu elektrod. Opcjonalnie jest też buzzer wydający sygnał dźwiękowy informujący o ew. nieprawidłowościach. Mimo, że układ kreśli przebieg

jednokanałowy, potrzebne są 3 sondy-elektrody. LA należy podłączyć do lewego ramienia. RA w okolice ramienia prawego i RL – prawej nogi.

Od Red. EdW: LA, RA i RL to standardowe oznaczenia sond, aczkolwiek inne konfiguracje podłączeń są stosowane. Generalnie LA i RA to wejście różnicowe – 1 kanał, a RL jest wyjściem stosowanym w celu poprawy, eliminacji

zakłóceń common-mode; sygnał wspólny (zakłócający nałożony na „wątki” użyteczny sygnał różnicowy) pojawiający się na elektrodach LA i RA jest wzmocniony, odwrócony w fazie i należy go podłączyć „zwrotnie” do prawej nogi.

Na **rysunku 3** mamy schemat ideowy, aczkolwiek niewiele odbiegający od blokowego z **rysunku 2**. Schemat połączeń modułu EKG z płytką Arduino widzimy także na **rysunku 5**. Cały



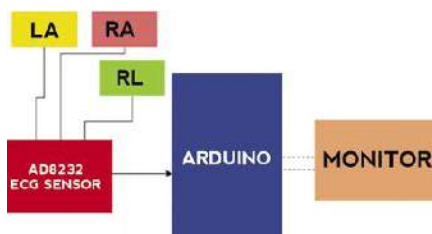
Rysunek 1.

```

Arduino Code:
int ekg = 0;
int buzzer = 11;
int Signal;
int Threshold = 590;
void setup()
{
  pinMode(buzzer, OUTPUT);
  Serial.begin(9600);
  pinMode(9, INPUT);
  pinMode(10, INPUT);
}
void loop()
{
  if ((digitalRead(9) == 1) || (digitalRead(10) == 1))
  {
    Serial.println("!");
  }
  Signal = analogRead(ekg);
  Serial.println(Signal);
  if (Signal > Threshold)
  {
    tone(buzzer, HIGH);
  }
  else
  {
    noTone(buzzer);
    delay(12);
  }
}

```

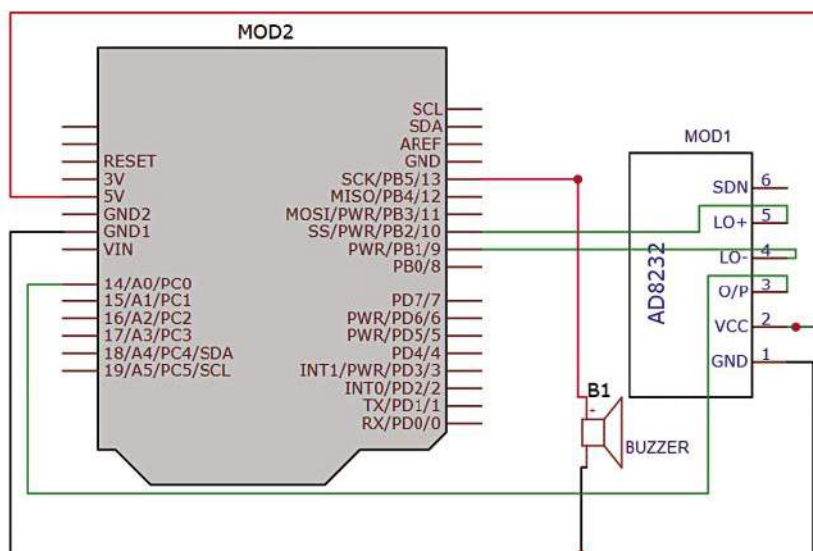
Rysunek 4.



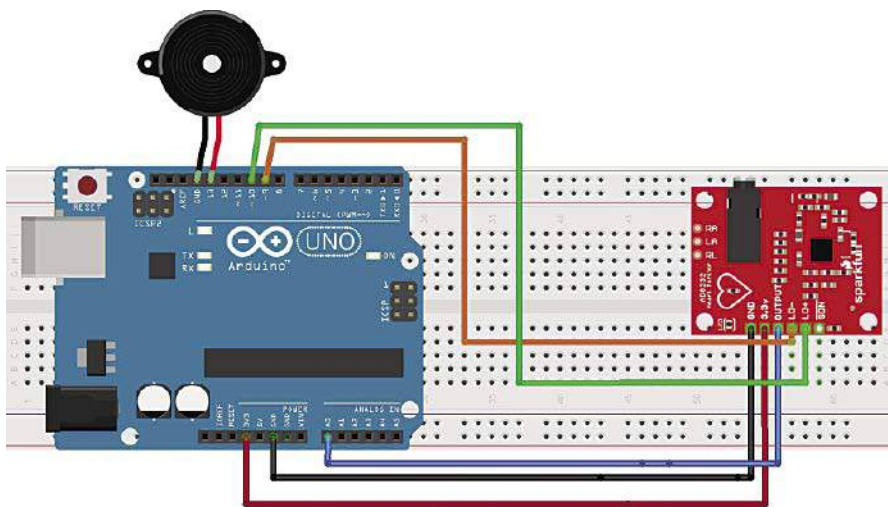
Rysunek 2.

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

Moduł AD8232
3-metrowe elektrody z końcówkami umożliwiającymi przyklejenie – 3 szt.
Płytkę Arduino UNO
Buzzer 3–5 V



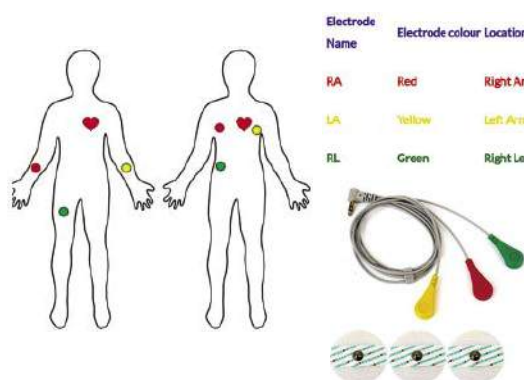
Rysunek 3.



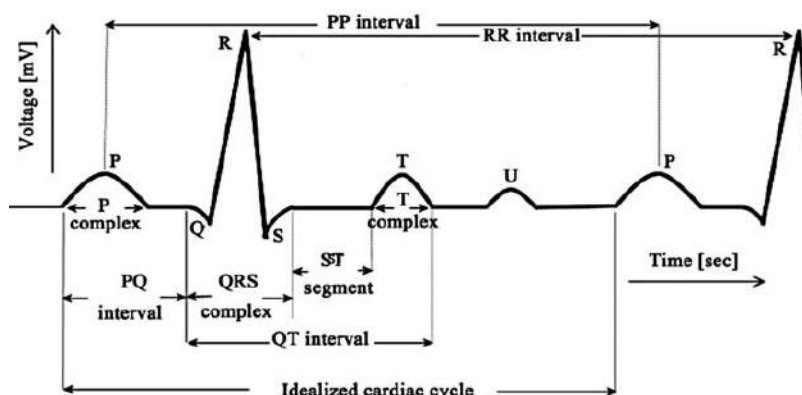
Rysunek 5.

hardware składa się z modułu AD8232 (MOD1) i Arduino UNO (MOD2). Dodatkowymi elementami są sondy i opcjonalnie piezoelektryczny

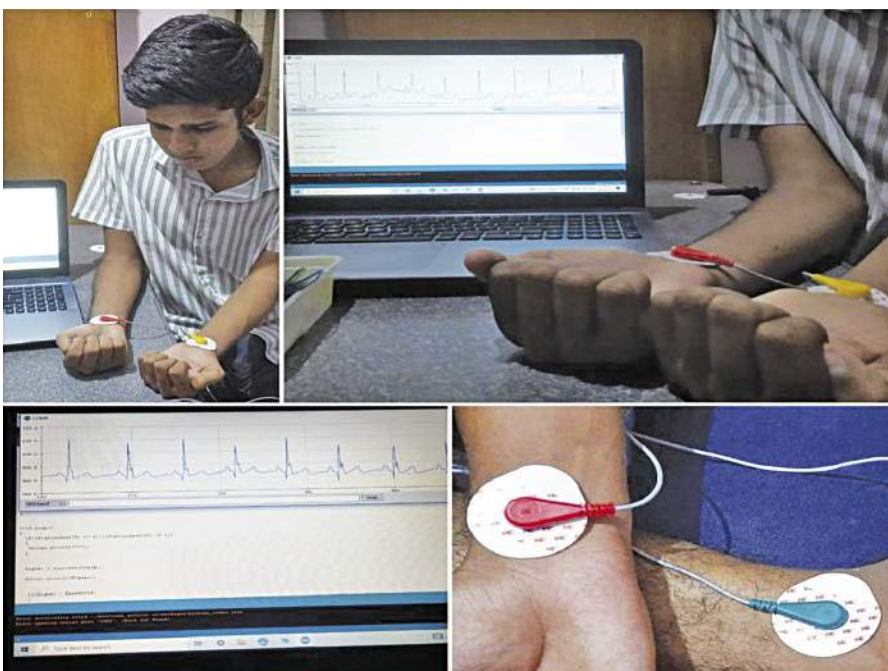
buzzer. Niezbędny software ECG_code.ino należy wgrać przy pomocy oprogramowania Arduino IDE.



Rysunek 6.



Rysunek 7.



Rysunek 8.

Konstrukcja i przetestowanie urządzenia

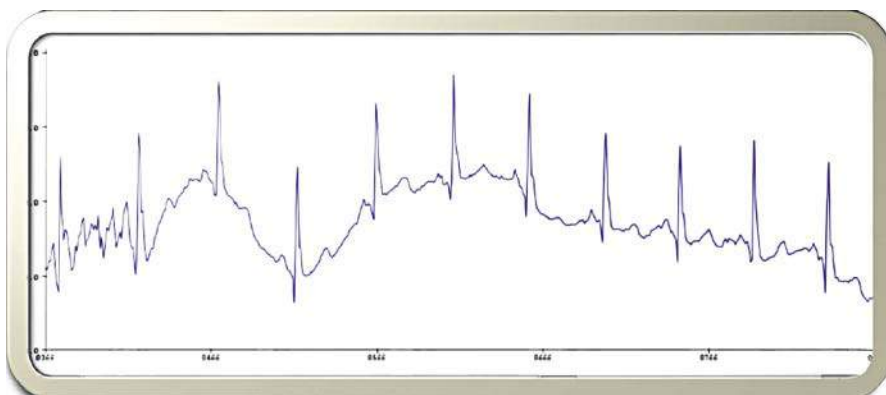
W celu pozyskania przebiegu EKG, należy poprawnie podłączyć elektrody. Aczkolwiek w tym zakresie istnieje kilka konfiguracji, zalecamy kierować się **rysunkiem 6**, na którym zaznaczono punkty na ciele oraz kolory elektrod. Poprawność podłączenia i jakość elektrod jest kluczowa dla pozyskania czystego sygnału elektrokardiogramu. Dostępne są niedrogie elektrody z końcówkami wypełnionymi chlorkiem srebra i żelom zapewniającym niską rezystywność złącza. Zewnętrzna średnica zawiera substancję kleistą pomocną w umocowaniu elektrody w wyznaczonym punkcie na ciele pacjenta.

Orientacyjny kształt przebiegu EKG pokazano na **rysunku 7**. To przebieg okresowy, w którym można wyróżnić kilka charakterystycznych faz. Standardowo oznacza się je

literami P, Q, R, S, T i U. Fragment przebiegu oznaczony P niesie informację o pracy węzła zatokowo-przedsionkowego serca. Przedział wyszczególniony jako QRS w rytmie pracy serca pochodzi z węzła przedsionkowo-komorowego. Kształt przebiegu w fragmencie P może świadczyć o migotaniu przedsionków. Zniekształcenia fragmentu QRS mogą świadczyć o depolaryzacji komór serca, zaś fragment T o repolaryzacji komór serca. Jak widać na **rysunku 7** w przedziale QRS dominuje amplituda impulsu opisanego jako R z fragmentami Q i S o przeciwnej polaryzacji. Na podstawie segmentu QRS można wnioskować o skurczach komór serca. Przedziały wyszczególnione jako R i S niosą także informację o transmisji impulsów do dalszych tkanek.

Testowanie urządzenia

W celu przetestowania Elektrokardiogramu, można podłączyć elektrody w nadgarstkach zgodnie ze zdjęciem pokazanym na **rysunku 8**. Moduł EKG zasilany jest z płytki Arduino.



Rysunek 9.



Rysunek 10.

Połączenia między modułem AD8232 i płytką Arduino

- GND – potencjał masy
- 3,3 V – zasilanie dla MOD1
- Output (na MOD1) – A0 (na Arduino)
 - analogowy sygnał przesłany z AD8232 do Arduino
- LO- (MOD1) do D8 (Arduino) oraz LO+ do D9
 - detekcja „leads-off” braku poprawnego podłączenia sond
- SDN – Shutdown (opcjonalnie sygnał z Arduino do MOD1)

Transmisja między Arduino i komputerem jest szeregową i należy ustawić odpowiedni port i prędkość 9600 bodów. Sercem modułu EKG jest układ AD8232 firmy Analog Devices, który jest wyposażony w detekcję poprawności podłączenia elektrod. To linie podłączone do wyprowadzeń 9 i 10 płytki Arduino.

Od Red. EdW: W konfiguracji 3-elektrodowej układ jest w stanie zidentyfikować którą linię sygnałową (sonda) jest „luźna”; dlatego wykorzystano dwie linie cyfrowe wyjścia, przekazujące informację do Arduino.

Dodatkowo wyjście pin 11 wykorzystano do podłączenia buzzera. Sygnalizacja dźwiękowa jest uruchamiana, gdy amplituda sygnału przekroczy poziom 590 mikrowoltów. Może to świadczyć i sygnalizować nadpobudliwość. Przykładowe przebiegi pozyskane monitorem EKG pokazuje rysunek 9. ■

dr Geetali Saha, Kushal Gadhvi

Ten artykuł był wcześniej opublikowany na łamach „EFY”, listopad 2022 (efymag.com)

Quiz: Półprzewodnikowe elementy mocy

Półprzewodnikowe elementy mocy są wytwarzane z następujących materiałów:

- ☐ Se, Ge, Si
- ☐ Ge, Si, SiC
- ☐ Si, SiC, GaN

Półprzewodniki szerokopasmowe SiC, GaN w porównaniu z Si mają szerokość pasma zabronionego większą:

- ☐ 2×
- ☐ 3×
- ☐ 5×

Natężenie pola elektrycznego wywołującego przebicie dla SiC, GaN jest większe niż dla Si:

- ☐ 2×
- ☐ 5×
- ☐ 10×

W diodach PIN obecność słabo domieszkowanej warstwy I jest niezbędna dla osiągnięcia?

- ☐ wysokich napięć pracy w kierunku zaporowym
- ☐ dużych prądów przewodzenia
- ☐ dobrej stabilności temperaturowej charakterystyk

Dioda Schottky'ego w porównaniu z diodą PIN ma szybkość działania:

- ☐ większą
- ☐ mniejszą
- ☐ porównywalną

Tyrystor triodowy ma strukturę naprzemiennych warstw typu p i typu n, przy czym liczba tych warstw wynosi

- ☐ 3
- ☐ 4
- ☐ 5

Tyrystor triodowy można wyłączyć ze stanu przewodzenia do stanu blokowania jednym z następujących sposobów:

- ☐ podając na bramkę napięcie ujemne względem katody
- ☐ zwierając bramkę z katodą
- ☐ zmniejszając prąd anodowy poniżej prądu trzymania IH

Tyrystor dwukierunkowy (diak, triak) składa się z naprzemiennych warstw typu p i typu n. których jest:

- ☐ 4
- ☐ 5
- ☐ 6

Tyrystor GTO wyłącza się prądem bramki, który osiąga wartość:

- ☐ ok. 25% prądu anoda – katoda
- ☐ ok. 1% prądu anoda – katoda
- ☐ ok. 0,1% prądu anoda – katoda

Tranzystor IGBT to struktura łącząca cechy budowy i zalety:

- ☐ tranzystorów bipolarnych i MOSFET
- ☐ tranzystorów bipolarnych i MESFET
- ☐ tranzystorów bipolarnych i JFET

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 24.03.2023.

Bezprzewodowy Power-Bank

Technologie bezprzewodowe są domeną naszych czasów, a ich szerokie stosowanie wynika nie tylko z mody. Zapewne również nie z chęci oszczędności drutu-miedzi. Ostatnim bastionem w tym „sektorze” jest bezprzewodowy przesył energii. Aczkolwiek możliwy i nie dziwi nikogo, kto przynajmniej widział transformator, trzeba stwierdzić – możliwy jedynie na bliską odległość. W każdym razie, bezprzewodowe ładowanie akumulatora telefonu komórkowego już nikogo nie dziwi.

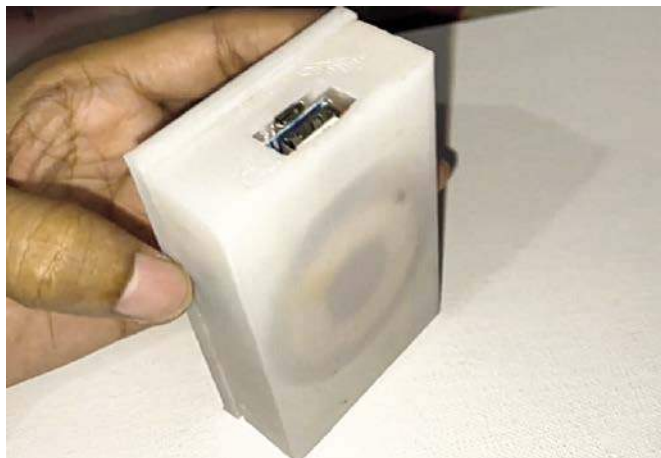
Bieżący DIY to Power Bank, który potrafi wysłać energię „wireless”. Ładowarki z tzw. technologią MagSafe są dostępne, a wiele urządzeń jest z MagSafe kompatybilnych. Okazuje się, że samodzielne wykonanie bezprzewodowego Power Banku nie jest trudne, i oprócz oszczędności finansowej dostarczy satysfakcji z jego wykonania.

Projekt i łączenie podzespołów

Korzystając z gotowych podzespołów, większość prac sprowadza się do montażu mechanicznego i niewielu połączeń elektrycznych.

Wielkość obudowy zdeterminowana jest głównie rozmiarem baterii. Zakładamy, że nasz powerbank zmieści dwa akumulatory rozmiaru 18650. Poza „aku”, w obudowie muszą się zmieścić dwa „moduliki” małych rozmiarów oraz płaska cewka, która będzie emitować energię. Jeden z tych modułów, to obecny w każdym powerbanu BMS (Battery Management System), który zawiera liniowo pracującą ładowarkę oraz przetworniczkę step-up podnoszącą napięcie akumulatora (ok. 3,2 V do 4,2 V) do względnie stabilnej wartości 5 V. Powerbank ma być bezprzewodowy,

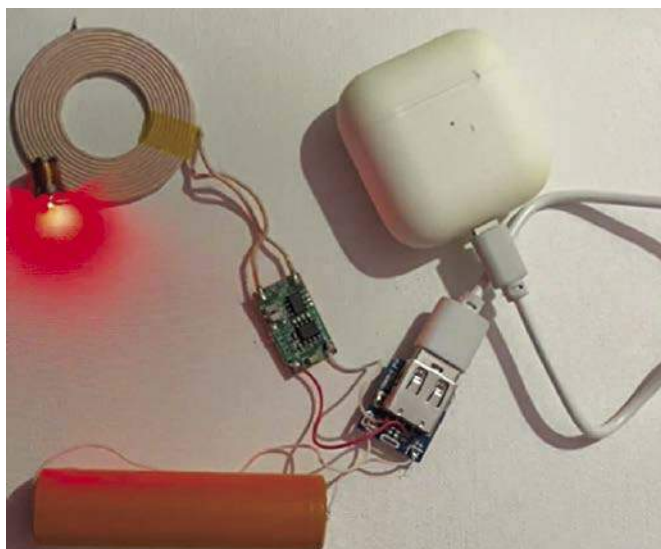
nie mniej jego ładowanie jest przewodowe. Dlatego w obudowie trzeba przewidzieć i stabilnie umocować złącze USB. W bieżącym projekcie wykorzystano typowy moduł BMS z układem scalonym TP4056. To ładowarka i step-up converter o zdolności jedno-amperowej. Dla zamontowania i podłączenia baterii można wykorzystać dostępne złącza dla rozmiaru 18650 lub wprost przylutować przewody do akumulatora. Połączenia należy wykonać wg schematu na rysunku 10. A więc: złącze mikro-USB na wejście +5V-IN i GND modułu BMS. Do punktów oznaczonych BAT+



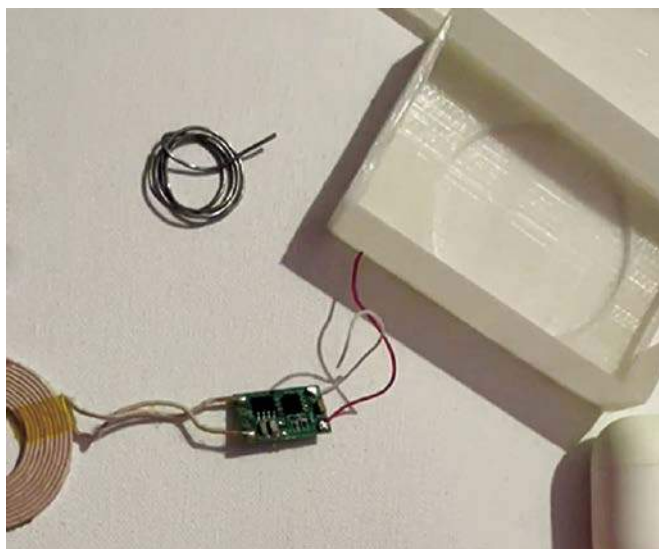
Rysunek 1.



Rysunek 3.



Rysunek 2.



Rysunek 4.

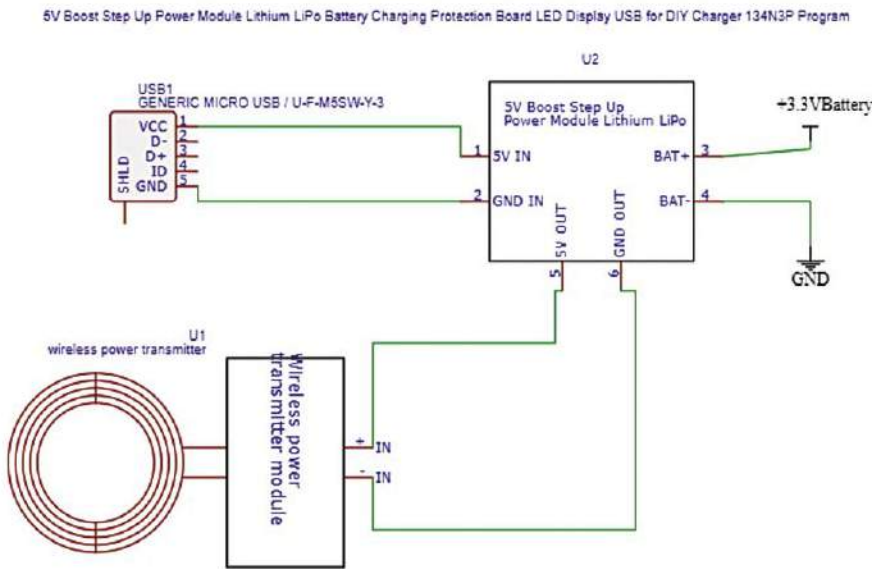
i BAT– podłączamy dwa akumulatorki równolegle. Wyjście przetwornicy step-up jest podłączone do złącza USB typ A i tradycyjnie

służy do ładowania przewodowego. Tę funkcję warto pozostawić, tym bardziej że „duże USB” może stanowić element „nośny” umocowania

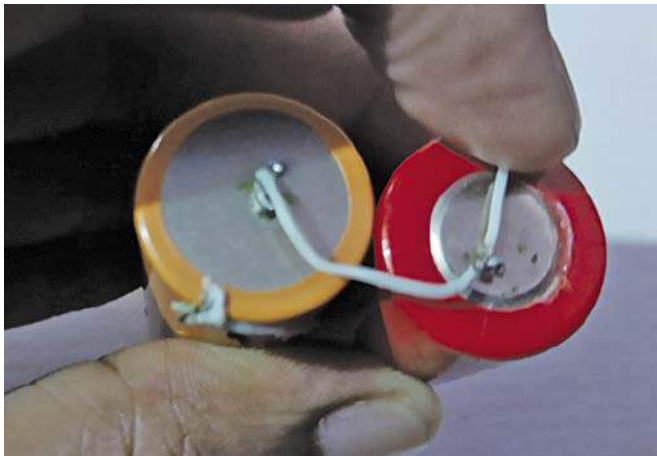
płytki. Do wyjścia +5V-OUT i GND-OUT podłączamy drugą płytkę będącą nadajnikiem bezprzewodowym. Wyjście nadajnika pracuje na cewkę indukcyjną, która jest elementem dość sporych rozmiarów. Jednak ponieważ jest płaska, nietrudno ją przykleić do obudowy. Wraz z cewką, od wewnętrznej strony obudowy należy przykleić magnesik. Niewielkich rozmiarów, aczkolwiek dość silny magnes jest też niezbędnym elementem standardu MagSafe.

Testowanie działania powerbanku

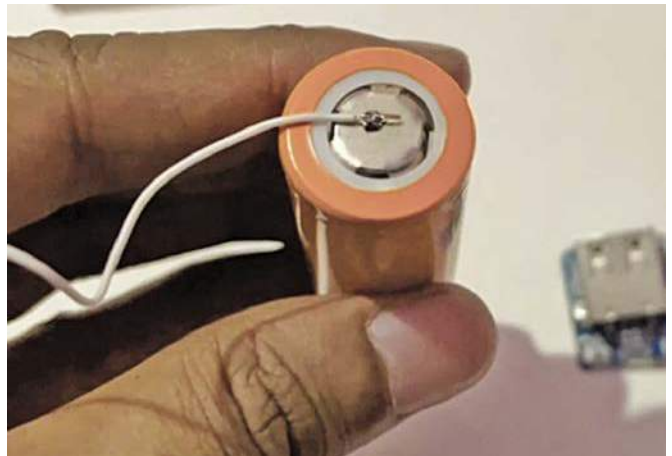
Ładowanie akumulatorów odbywa się w tradycyjny (przewodowy) sposób. I to należy sprawdzić w pierwszej kolejności podłączając do złącza mikro-USB dowolną pięcio-woltową ładowarkę. Sprawdzenie, czy powerbank transmituje energię bezprzewodowo, można wykonać „bezprzewodową diodą LED”, jak pokazują rysunki 17 i 19. Finalnie, można ładować telefon, który oczywiście też musi być wyposażony w funkcję ładowania bezprzewodowego.



Rysunek 5.



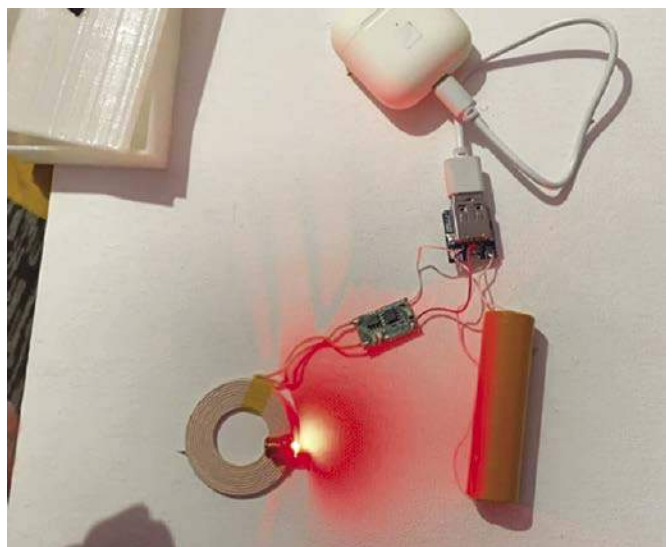
Rysunek 6.



Rysunek 8.

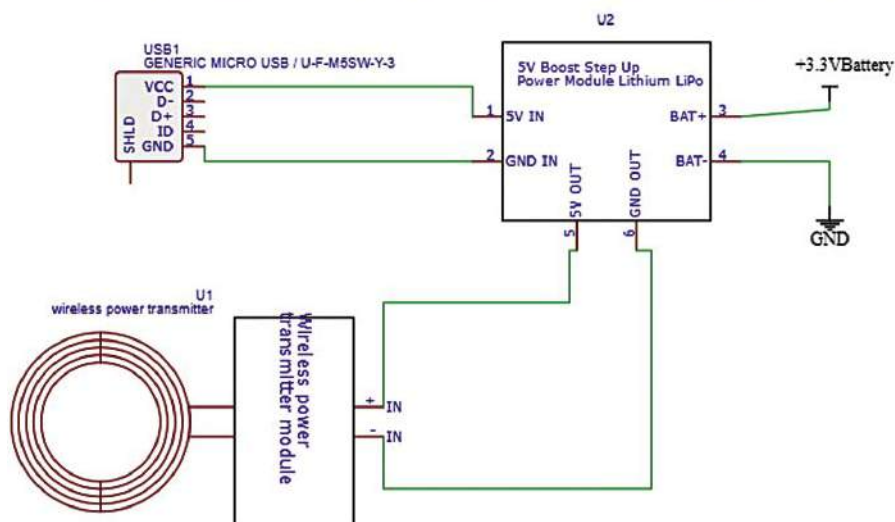


Rysunek 7.



Rysunek 9.

5V Boost Step Up Power Module Lithium LiPo Battery Charging Protection Board LED Display USB for DIY Charger 134N3P Program



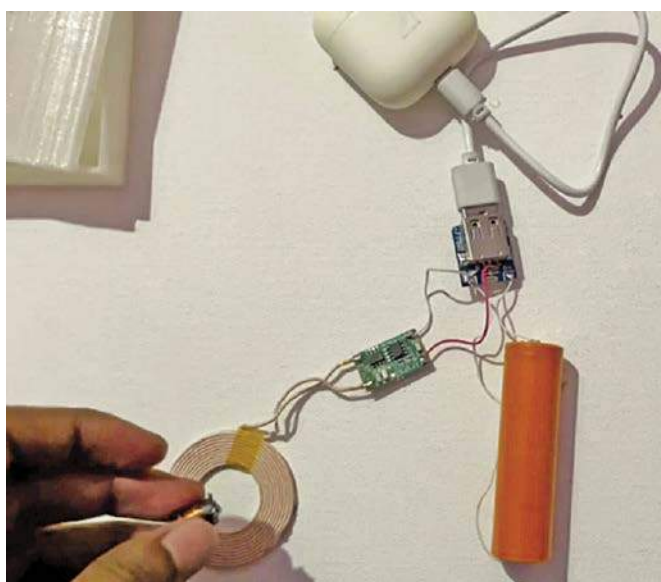
Rysunek 10.

Wykaz elementów, kupuj w sklep.avt.pl (W-wa, ul. Leszczynowa 11, tel. +48222578451, e-mail: handlowy@avt.pl):
 cewka indukcyjna („promieniująca” energię)
 moduł BMS (ładowarka i +5 V step-up converter)
 moduł transmitera bezprzewodowego TX
 akumulatory 18650 (2 sztuki 1000 mAh lub 2000 mAh)
 obudowa

Od strony elektrycznej cały układ jest prosty (jeśli nie wnikać w działanie BMS-u i płytki transmitera) i wymaga jedynie kilku połączeń. Dlatego więcej informacji (niż tekst) niosą zdjęcia zamieszczone na rysunkach od 1 do 20. Autor pokazał na nich konstrukcję mechaniczną swojego prototypu, a także połączenia i końcowy efekt pracy bezprzewodowego powerbanku. ■

Ashwini Kumar Sinha

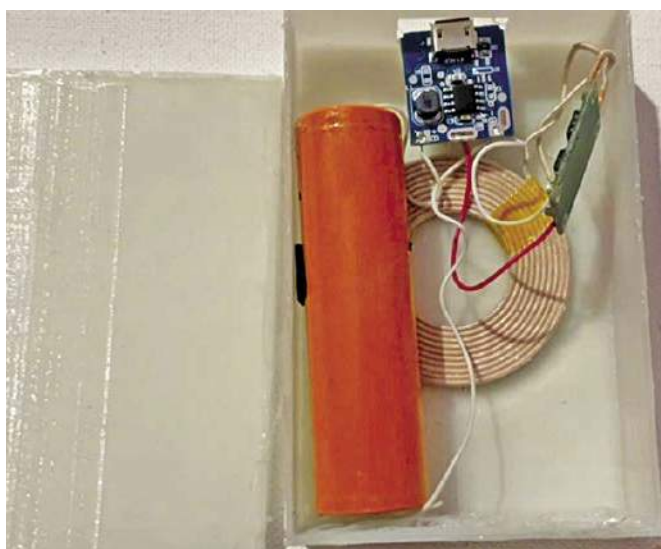
Ten artykuł był wcześniej opublikowany na łamach „EFY”, lipiec 2022 (efymag.com)



Rysunek 11.



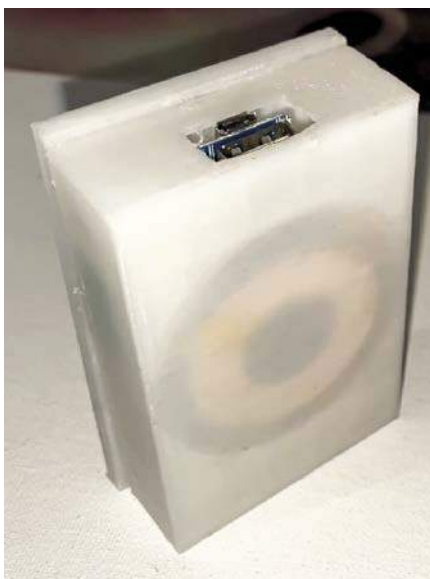
Rysunek 13.



Rysunek 12.



Rysunek 14.



Rysunek 15.



Rysunek 18.



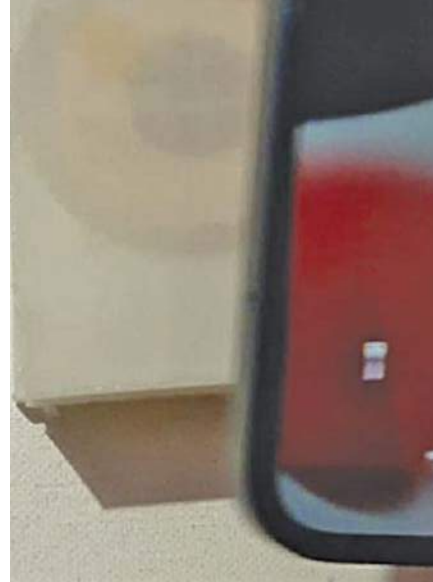
Rysunek 20.



Rysunek 16.



Rysunek 17.



Rysunek 17.

REKLAMA



Certyfikat Underwriters Laboratories

UL 94V-0 E480148 TYPE 1

Zakład produkcyjny:

05-660 Warka
ul. M. Ropolewskiej 17
tel. 22 781 63 95
22 761 95 80
fax. 22 781 63 95 w. 23
www.elmax.waw.pl
elmax@elmax.waw.pl



OBWODY Drukowane

Produkcja, Projektowanie, Montaż

Płytki jednostronne	Serie dowolne	Dokumentacja technologiczna	Montaż elektroniki
Płytki dwustronne	Prototypy	Dokumentacja konstrukcyjna	ilości modelowe produkcyjne
Płytki na podłożu aluminium	Maksymalny wymiar płytek 1w 630 mm	Płyty czołowe FR4	Krótkie terminy
Aktywny kalkulator prototypów na stronie internetowej	Pokrycie Sn lub SnPb inne na życzenie	Trawione szablony SMD	Wykonania super expresowe
	Maski, opisy montażowe w różnych kolorach		

Autonomiczny robot-samochodzik omijający przeszkody

Temat robotów autonomicznych jest bardzo na czasie. Zakres prac dla, których użyteczne są tego rodzaju urządzenia jest szeroki, a ich użyteczność trudno przecenić. Funkcja rozpoznawania i omijania przeszkód jest istotną cechą, w którą robot autonomiczny powinien być wyposażony. Dzięki niej robot-samochodzik powinien dotrzeć do celu unikając nieprzewidzianych kolizji. W tym celu musi być wyposażony w funkcję rozpoznawania przeszkody, a także krawędzi lub obramowania drogi, w ramach której ma się poruszać. Układy wykonawcze muszą mieć możliwość wyboru drogi alternatywnej. Oba obwody muszą współpracować tworząc logiczną pętlę sprzężenia zwrotnego działającego w czasie rzeczywistym.

Prezentowany tu projekt robotu-samochodzika ma wbudowaną inteligencję spełniającą wyżej nakreślone wymagania. Algorytm działania wbudowany jest w software, którego pracą steruje obwód rozpoznawania przeszkody. Prototyp wykonany przez autora widzimy na rysunku 1. Aczkolwiek to zabawka, algorytm można przenieść dla robota wykonującego poważniejsze zadania. Robota potrafiącego poruszać się w nieznanym terenie bez ciągłego nadzoru i sterowania przez operatora.

Spis materiałów:

- Arduino UNO R3 – 1 szt.
- Płytki drivera silników – 1 szt.
- Koła – 4 szt.
- Silniki DC z wbudowanymi przekładnikami – 4 szt.
- Servo-motor – 1 szt.
- Czujnik ultradźwiękowy – 1 szt.
- Stopka mocująca czujnik – 1 szt.
- Bateria LiPO – 1 szt.
- „Pokrywa” akrylowa – 1 szt.
- Przewody oraz jumpery mędkie i żeńskie

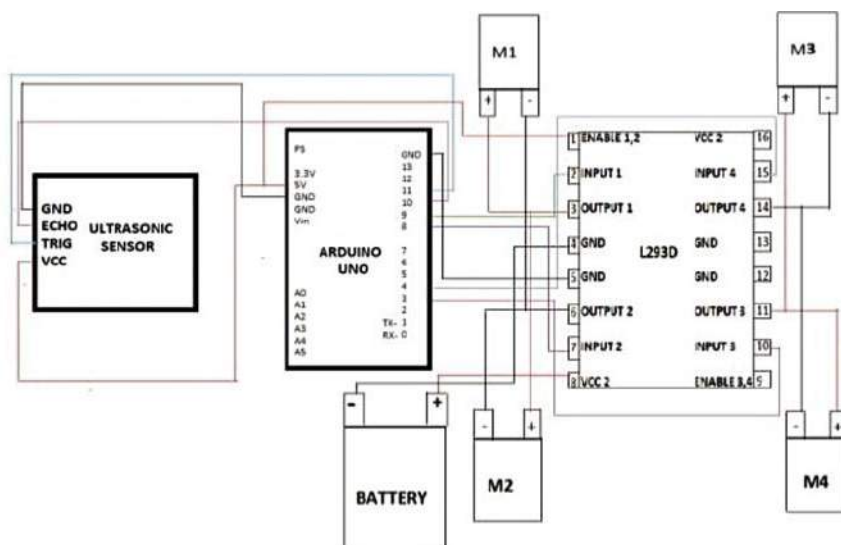


Rysunek 1.

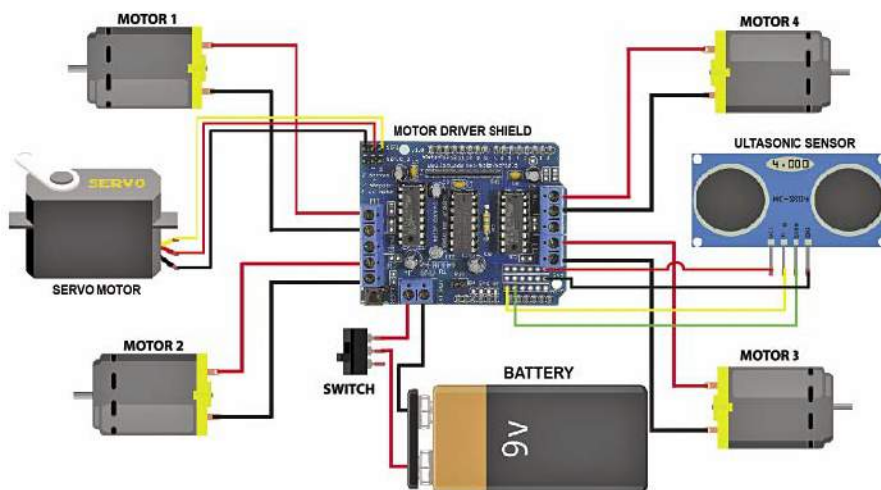
Budowa robota i działanie układu

Na rysunkach 2 i 3 pokazano schemat blokowy i połączenia między współpracującymi podzespołami. Układ wykonano na bazie mikrokontrolera rodziny ATmega 8 na płytce Arduino Uno R3. Kluczowym

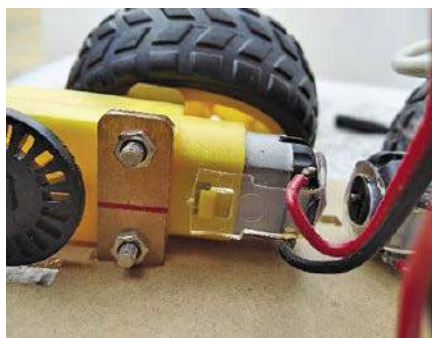
podzespołem jest czujnik ultradźwiękowy. Gdy rozpozna przeszkodę lub barierę na drodze robota, przesyła tę informację do mikroprocesora. To jest sygnał wejściowy, a wyjściowym jest sterowanie silnikami umieszczonymi w każdym kole. Silniki sprzężone są z mikrokontrolerem przez



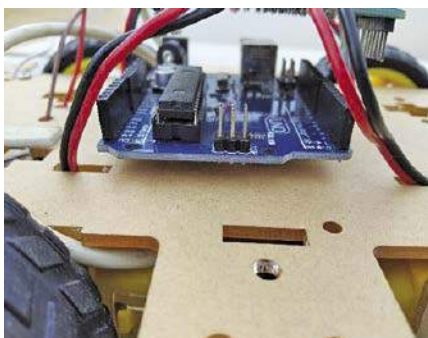
Rysunek 2.



Rysunek 3.



Rysunek 4.



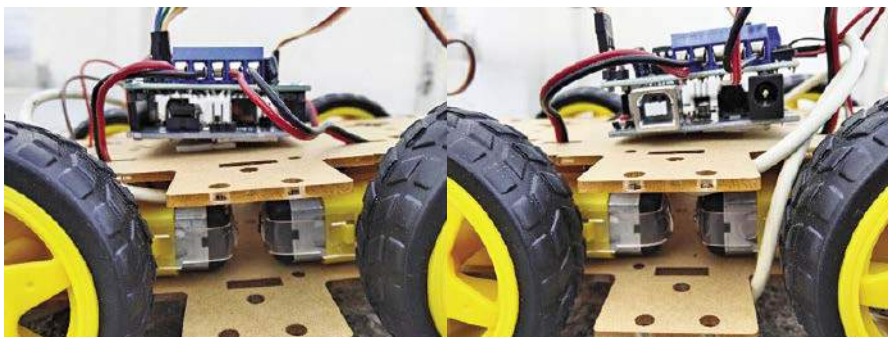
Rysunek 7.



Rysunek 8.



Rysunek 5.



Rysunek 9.



Rysunek 6.

wzmacniacze drivery. Mimo, iż na bieżąco nie jest znana droga omijająca przeszkody, sterowanie kierunkiem obrotów silników plus możliwość skrętów, pozwoli na znalezienie drogi omijającej kolizje. Proces

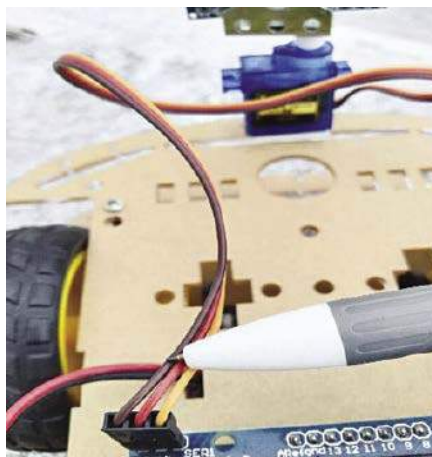
podejmowania decyzji nie jest narzucony z góry, aczkolwiek inteligencja oprogramowania jest równie ważna jak praca obwodów wykonawczych.

Podobnie, kluczowe jest działanie „oczu” robota. W tej roli pracuje czujnik ultradźwiękowy, składający się z dwóch sekcji, nadajnika i odbiornika. Nadajnik emituje paczki ultradźwięków o częstotliwości 40 kHz. Odbiornik odbiera echo odbite od ew. przeszkód. Układ potrafi obliczyć opóźnienie czasowe echa względem emitera. Zatem, ta prosta technika pozwala na bieżąco „wiedzieć” nie tylko obecność, ale i odległość przeszkody lub bariery na drodze poruszającego się robota. W ten sposób, sprzężenie zwrotne działa w czasie rzeczywistym. Decyzje są autonomiczne i polegają na zmianie kierunku lub wycofaniu się robota.

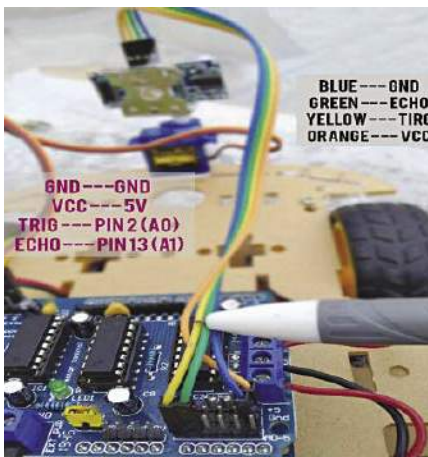
Oprogramowanie

Program napisano w języku zrozumiałym przez Arduino. Można go ściągnąć pod nazwą ARDUINO_OBSTACLE_AVOIDING_CAR. ino i przepisać z komputera do mikrokontrolera Arduino z użyciem oprogramowania Arduino IDE. Autor użył wersji IDE 1.8.11. W celu kompilacji i przepisania softwareu należy płytke Arduino UNO połączyć z komputerem kablem USB. Wymagane ustawienia sprawdzają się do wyboru typu mikrokontrolera obecnego na Arduino oraz portu użytego w procesie przepisania kodu źródłowego celem za-programowania mikroprocesora.

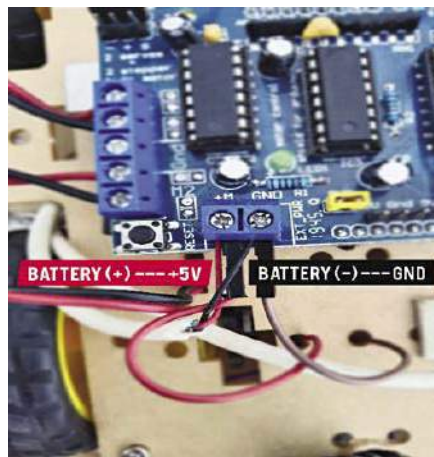
Od Red. EdW. Jak wynika ze schematu na rysunku 2, robot ma możliwość obrotu kołami w obu kierunkach. Mimo iż każde koło ma swój silnik wraz z przekładnią, praca kół lewy przód/



Rysunek 10.



Rysunek 11.



Rysunek 12.

DIY dla wszystkich

tył i prawy przód/tył jest współbieżna. Na schemacie nie pokazano także serwomotoru służącego do obrotu czujnika ultradźwiękowego. Jego lokalizacja i zakres pozycji serwomechanizmu mogą być kluczowe dla końcowego efektu wykrywania przeszkód na drodze robota.

Konstrukcja mechaniczna i testowanie pracy urządzenia

Najpierw załaduj oprogramowanie, następnie postępuj zgodnie z poniższymi krokami

Krok 1: Zmontowanie chassis robota (patrz rysunki 4, 5 i 6)

- przylutuj przewody plus i GND do silników
- umocuj mechanicznie silniki wraz z przekładnią od spodu chassis
- przymocuj koła na osie przekładni sprzężonych z silnikami

Krok 2: Umocowanie podzespołów

- przykręć płytkę Arduino od góry chassis
- załóż płytkę driverów silników przygotowaną jako „kanapka” mocowana na Arduino

- podłącz silniki do odpowiednich par wyprowadzeń driverów: lewy przód/tył i prawy przód/tył

Krok 3: Mocowanie serwo

- umocuj serwo w górnej części chassis
- podłącz serwo do 3-stykowego złącza SER1

Krok 4: Podłączenie czujnika ultradźwiękowego (patrz rysunek 11)

Umocuj czujnik przy pomocy przygotowanej „stopki”. Cztery przewody podłącz do czujnika przy pomocy przygotowanych jumperów żeńskich. Do Arduino podłącz przewody w następujący sposób:

GND do GND

Vcc do +5V

TRIG do pinu nr 2 (A0)

ECHO do pinu 13 (A1)

Krok 5: Podłączenie zasilania

Zastosowano baterię LiPO, którą należy podłączyć do drivera L293D, zwracając uwagę na polaryzację:

LiPo battery+ do +12V

LiPo battery- do GND

Kod źródłowy
tego projektu jest
dostępny do pobrania
ze strony
<https://bit.ly/3soblh2>

Inteligentny robot powinien być gotowy do pracy.

Prototyp wykonano w formie zabawki. Solidniej wykonana wersja może mieć zastosowanie w wielu aplikacjach:

- Od tak poważnych jak militarne czy prace poszukiwawcze i ratownicze z nawigacją w trudnym terenie.
- W gospodarstwie domowym jako np. samobieżne odkurzacze.
- Rozbudowa oprogramowania pozwoli na aplikację jako asystent parkowania lub rozbudowa funkcjonalności zdalnie sterowanych dźwigów itp. ■

Dr S. Rajan, Mythili M.

Ten artykuł był wcześniej opublikowany na łamach „EFY”, listopad 2022 (efymag.com)

REKLAMA

Sięgnij po archiwalne wydania ELEKTRONIKI dla WSZYSTKICH

Prenumeratorzy mają bezpłatny dostęp do e-wydań archiwalnych EdW starszych niż 24 miesiące.

Przesyłka
GRATIS

Zamów wygodnie na
www.UlubionyKiosk.pl

eprasa.pl 250f0d7be5

Przedstawiamy początkowe fragmenty dwóch projektów ze zbioru kilkudziesięciu projektów dostępnych wyłącznie dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl. W rubryce **DIY PLUS** zamieszczamy aktualnie najciekawsze projekty publikowane w Internecie w formacie open source. Prenumeratorów EdW zapraszamy do zapoznania się na www.elportal.pl z niezwykle inspirującymi zasobami rubryki **DIY PLUS**.

PALPi konsola do gier retro

W tym projekcie pokazano, jak można skonfigurować konsolę do gier i zrobić kompletnego ręcznego Gameboya, który może emulować każdą grę retro, jaką można sobie wyobrazić. Nazywa się PALPi, ponieważ wykorzystuje kompozytowy wyświetlacz PAL. Kiedy byłem mały, uwielbiałem grać w gry takie jak Pokemon, contra, Super Mario, Final Fantasy i inne, głównie na Gameboy Advance i konsole do gier, które podłączaliśmy do naszych telewizorów CRT, aby uruchamiać stare, wspaniałe gry. Obecnie możemy po prostu pobrać ROM z grą retro, otworzyć go w emulatorze i grać w nią na laptopie lub urządzeniach przenośnych. Ale jako twórca, chciałem zrobić coś innego, więc przygotowałem tę podręczną konsolę do gier retro, która jest zasilana przez Raspberry Pi Zero a system operacyjny, którego używam to Recalbox OS. Jest to przyzwyczajony emulator OS, który jest również dostarczany z kilkoma preinstalowanymi grami.

Dokończenie artykułu na stronie:
<https://bit.ly/3XITS08>



Wielokanałowy analizator FFT widma

Artykuł opisuje analizator FFT widma. Ma 8, 16, 24, 32 lub nawet 64 pojemniki częstotliwości (kanały) i można nawet podwoić tę liczbę, jeśli zmodyfikuje się oprogramowanie. Płyta PCB może sterować matrycą pikseli (WS2812 lub matryca led) lub można podłączyć jeden lub więcej wyświetlaczy HUB75E, ale będzie trzeba dokonać wyboru, którego użyć i odpowiednio dostosować parametry w ustawieniach. Można podłączyć sygnał audio za pomocą wejścia audio lub można użyć wejścia mikrofonu do podłączenia małego mikrofonu pojemnościowego. Używanie mikrofonu może jednak ograniczyć pasmo przenoszenia ze względu na jego ograniczenia. Czutość wejściowa jest zautomatyzowana i dostosowuje się do poziomów sygnału wejściowego. Można dostosować jasność i czas utrzymywania szczytu. Gdy układ nie otrzyma żadnego sygnału wejściowego, po chwili przejdzie w tryb pożarowy, w którym niektóre diody/wyświetlacz zapalą się jak ogień. W prototypie użyto dwóch paneli HUB75E, które zamontowano na drewnianym stojaku.

Dokończenie artykułu na stronie:
<https://bit.ly/412haR8>



Niektóre projekty aktualnie dostępne tylko dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl:

1. Uniwersalny wskaźnik włączanego biegu manualnej skrzyni biegów w motocyklu
2. Automatyk LED-owe oświetlenie klatki schodowej z wykorzystaniem Arduino
3. RPi – stacja pogodowa IoT
4. Niskobudżetowy monitor jakości powietrza IoT oparty o RaspberryPi 4
5. Automatyczny system ogrodniczy z NodeMCU i Blynk, ArduFarmBot 2
6. Wzmacniacz piezoelektryczny do gitary i skrzypiec
7. TinyML – Rozpoznawanie ruchu przy pomocy Raspberry Pi Pico
8. Wysokowydajny i niezawodny sterownik bipolarnego silnika krokowego
9. Sterownik silnika prądu stałego z wykorzystaniem przełącznika i mosfetu – interfejs Arduino
10. Przedwzmacniacz do mikrofonu MEMS
11. Izolowany obwód wykrywania napięcia 250 V AC z pojedynczym wyjściem (wejście 250 V prądu przemiennego, wyjście 5 V)
12. Generator sygnałów AD9833
13. Obserwacja charakterystyk tranzystora
14. Wyświetlacz EKG z użyciem Arduino
15. Łatwy do zbudowania robot kroczący
16. Sonarowy theremin MIDI
17. Zamek elektroniczny na kod
18. Prosty tester tranzystorów
19. Super prosty czuły wykrywacz metali
20. Stymulator czaszkowy Arduino (Bio-BrainTuner)
21. Zegar binarny z użyciem Microbit

Miesięcznik „Elektronika dla Wszystkich” (12 numerów w roku) jest wydawany we współpracy z kilkoma redakcjami zagranicznymi



Wydawnictwo:
AVT-Korporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Adres redakcji:
03-197 Warszawa, ul. Leszczyńska 11
e-mail: edw@elportal.pl, www.elportal.pl

Redaktor merytoryczny
Paweł Sujko

Dział Reklamy:
Katarzyna Guła
katarzyna.gula@elportal.pl, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobarski
jakub.sobanski@elportal.pl

Copyright AVT-Korporacja Sp. z o.o., Warszawa, ul. Leszczyńska 11. Projekty publikowane w „Elektronice dla Wszystkich” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki dla Wszystkich”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice dla Wszystkich” jest dozwolone wyłącznie po uzyskaniu pisemnej zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice dla Wszystkich”.

DTP, okładka, redakcja strony internetowej www.elportal.pl:
MAD Sp. z o.o.

Prenumerata:
W Wydawnictwie AVT, e-mail: prenumerata@avt.pl
tel. 22 257 84 22, (godz. 10:00–14:00)

W RUCH S.A., e-mail: prenumerata@ruch.com.pl
tel. 801 800 803, 22 717 59 59, www.prenumerata.ruch.com.pl

Nowość

**O CZYM LEKARZE
CI NIE POWIEDZĄ**
Twoje zdrowie w Twoich rękach
Licencja europejska WHAT DOCTORS
DON'T TELL YOU

**Wydanie
specjalne 1/2023**
17,50 zł (w tym 8% VAT)

Holistyczny weterynarz



Jak sobie radzić z... astmą • biegunką • chorobami serca • nerek i wątroby • cukrzycą • czyrakami i cystami między palcami • grzybicą • guzkami piersi • katarem i kichaniem • nadciśnością tarczycy • metaboliczną chorobą kości u jaszczurek • nietrzymaniem moczu • otyłością • pchłami • problemami z gruczołami okołoodbytowymi i kolanami • przeziębieniem u papug • robaczycą • ropomaciczem • roztoczami u jeży • suchym i podrażnionym okiem • świerzbie • wymiotami • zaburzeniami snu • zaciemą • zapaleniem stawów • zatwardzeniem

Jak zapobiegać przegrzaniu i udarowi cieplnemu • sikaniu kota w domu • atakom kleszczy i leczyć skutki ukąszeń

Ponadto: czworonożny wegetarianin • pułapki w pielęgnacji świnek morskich • ziołowe lekarstwa DIY na powszechne schorzenia u psów i kotów • pies w podróży • wzmacnianie odporności u psów chorych na raka

Numer specjalny 1/2023
cena 17,50 zł (w tym 8% VAT)
ISSN 2300-9985 | Ciepłota 411000
0 1 >
597836

Zadbaj o zdrowie swojego pupila!

Na łamach tego wydania specjalnego weterynarze i terapeuci zwierzęcy radzą, co robić w przypadku wielu nieprzyjemnych dolegliwości oraz jak postępować, by uniknąć poważnych problemów. Ponadto nasi eksperci podpowiadają, na co zwracać uwagę, by nie umknęły nam wczesne sygnały kłopotów zdrowotnych i co robić, aby pomóc naszym pieszczochom, jeśli – nie daj Boże – zachorują.

Przejrzyj i zamów na [UlubionyKiosk.pl](https://ulubionykiosk.pl)

eprasa.pl 250f0d7be5

DARMOWA
PRZESYŁKA