

ELEKTRONIKA

dla wszystkich

nr 3/2024 (338) • marzec • www.elportal.pl

DIY PLUS
tylko dla prenumeratorów

Programowany hybrydowy zasilacz laboratoryjny z Wi-Fi

PROJEKTY dla elektroników

- ▶ Monitor prędkości na bazie GPS. Koniec z mandatami za przekroczenie prędkości
- ▶ Wielofunkcyjny monitor akumulatorów z dotykowym ekranem LCD, część 2
- ▶ Płytką treningowa SMD

DIY dla wszystkich

- ▶ Dzwonek bezprzewodowy
- ▶ Zabezpieczenie przeciwzwarciowe w pojazdach elektrycznych
- ▶ Generator tonów relaksujących na Arduino

TUTORIALE

- ▶ Audio OUT: Budowa radia tranzystorowego, część 2
- ▶ Chirurgia obwodowa: Transformatory i LTspice, część 1
- ▶ Elektronika
- ▶ Diagnostyka uszkodzeń aktywnych urządzeń elektroakustycznych przy pomocy programu komputerowego audioTester
- ▶ Piszemy prosty program dla ARM Cortex M4, czyli jak skomplikować prostą rzecz w przykładach
- ▶ Ekscytacje Maxa: Migające diody LED i śliniacy się inżynierowie



**Regulator prędkości obrotowej
silników DC 30 V/20 A**

ISSN 1425-1698 Indeks 33362X
0 3 >
9 771425 169245
16,90 zł (w tym 8% VAT)

EP.com.pl

Największy portal dla elektroników konstruktorów



**Król automatyki
jest w Tobie**

AutomatykaB2B.pl

FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

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

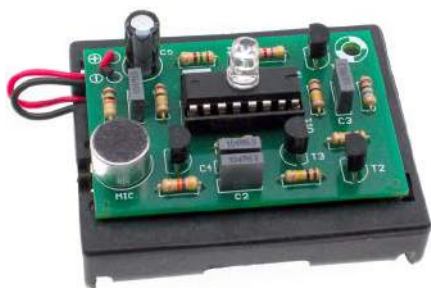
www.piekarz.pl





Najbardziej popularne kity AVT

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



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



AVT723 Uniwersalna gra zręcznościowa
<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>



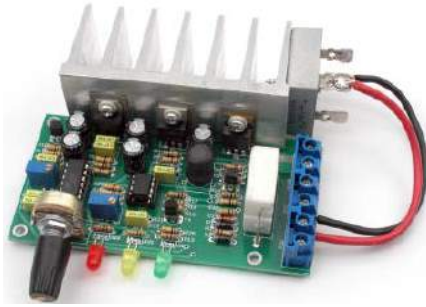
AVT990 Automatyczny włącznik świateł
<https://sklep.avt.pl/avt990.html>



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



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



AVT3120 Automatyczna ładowarka akumulatorów ołowiowych
<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>

PRENUMERATA

NA START
DO 6 WYDAŃ
GRATIS!

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

PO 5 LATACH
ZA PÓŁ CENY

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

PRENUMERATA EdW+

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

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

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

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

TYLKO prenumeratorzy* otrzymują pełny dostęp do:

ARCHIWUM

cyfrowego archiwum EdW na elportal.pl/archiwum



projektów w zbiorze DIY+ na elportal.pl/diy

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

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

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



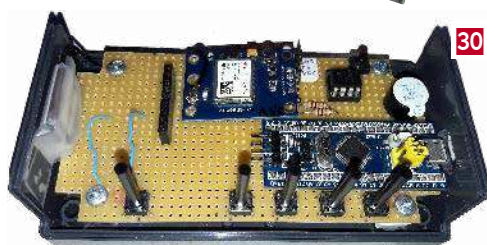
8

Projekty dla elektroników:

Programowany hybrydowy zasilacz laboratoryjny z modulem Wi-Fi, część 1	8
Regulator prędkości obrotowej silników DC 30 V/20 A	20
Monitor prędkości na bazie GPS.	
Koniec z mandatami za przekroczenie prędkości.....	30
Monitorowanie do 3 akumulatorów od 6 do 100 V – prąd do 10 A (lub 100 A+ z bocznikiem). Wielofunkcyjny monitor akumulatorów z dotykowym ekranem LCD, część 2.....	34
Płytką treningową SMD	42



20



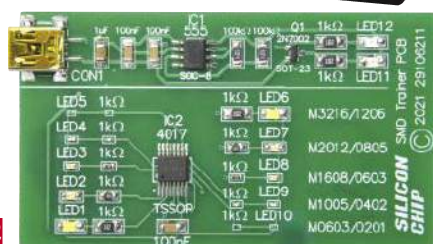
30

Tutoriale:

Ekscytacje Maxa:	
• Migające diody LED i śliniacy się inżynierowie, część 6	47
• Sprytnie porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania ..	51
Audio OUT: Budowa radia tranzystorowego, część 2.....	52
Chirurgia obwodowa: Transformatory i LTspice, część 1.....	58
Diagnostyka uszkodzeń aktywnych urządzeń elektroakustycznych przy pomocy programu komputerowego audioTester	62
Edukacja w EdW dla szkół i uczelni:	
Wykład 16: Wzmacniacze jedno- i dwutranzystorowe	66
Piszemy prosty program dla ARM Cortex M4, czyli jak skomplikować prostą rzecz w przykładach	72
Elektrony	81



34



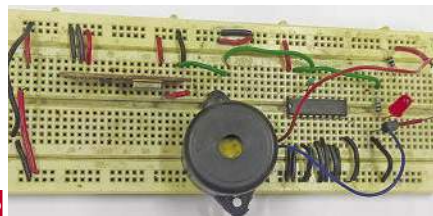
42

DIY dla wszystkich:

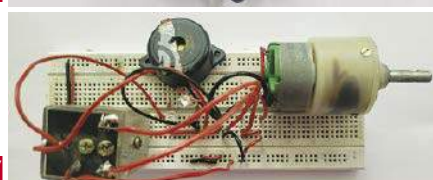
Dzwonek bezprzewodowy	85
Zabezpieczenie przeciwzwarceniowe w pojazdach elektrycznych	87
Generator tonów relaksujących na Arduino	89

DIY PLUS

Sterownik silnika krokowego z joystickiem.....	91
Programowalny kondycjoner sygnału z czujnika rezystancyjnego mostkowego	91



85



87

Rubryki stałe:

Prenumerata	3
Od wydawcy	5
Pocztka.....	6

A za miesiąc w kwietniowym EdW



* Jednookładowy cyfrowy odbiornik radiowy FM/AM/SW

Wykonać radio na układzie scalonym Si 4730 to łatwizna. Wystarczy dołączyć antenę, sterownik cyfrowy, wyświetlacz LCD i radio jest gotowe. Dla poprawyysterowania słuchawek do modułu dodano bufor w postaci pary wzmacniaczy operacyjnych.

* Programowany hybrydowy zasilacz laboratoryjny z Wi-Fi, część 2

Ten zasilacz dostarcza napięcie 0...27 V, przy napięciu do 16 V może być obciążony prądem 5 A i tylko nieco mniejszym przy wyższym napięciu. Może być zdalnie sterowany przez Wi-Fi. Można skonfigurować wiele jednostek do automatycznego sterowania i połączyć je szeregowo lub równolegle. W drugiej części artykułu pokazano, jak zmontować dwie płytki drukowane i podłączyć wszystko w niewielkiej plastikowej obudowie.

WSZYSTKO O USB – 3 artykuły

* Historia standardu USB

Rzeczywistość USB rozpoczął się w 1994 roku, gdy grupa firm mocno zaangażowanych w branżę PC porozumiała się i postanowiła ułatwić podłączanie urządzeń zewnętrznych do komputerów PC.

* Ładowarki USB

Opisano tanie moduły ładowarek, zasilaczy i adapterów sterowanych z gniazda USB.

* Jak pracuje złącze USB-C

Jeśli chcesz zasilać swoje projekty ze złącza USB-C, to masz kilka opcji opisanych w tym artykule.

* Lutowanie SMD – sztuczki i triki

Kolejny artykuł z serii o lutowaniu SMD.

* Plus kolejna porcja intrygujących projektów DIY.

* Plus wiele artykułów w Twoich ulubionych cyklach Tutoriali.

**W kioskach
od 29 marca**

Czym jest wzmacnienie i co to jest wzmacniacz?

Temat wykładu w tym numerze EdW prowokuje do zadania tego – pozornie tylko trywialnego – pytania. Odpowiedź zależy od adresata pytania. Oto trzy warianty odpowiedzi.

Dla Cycerona wzmacnianie było retorycznym zabiegiem wzbogacania wypowiedzi o metafory, obrazowość i hiperbole. Przykład. Wypowiedź prosta: „Zdałem sobie sprawę, że jestem głodny.” Wypowiedź wzmacniona: „Zdałem sobie sprawę, że mój żołądek woła o pożywienie z taką samą desperacją, jak ziemia pragnie deszczu po długim okresie suszy.”

Czat GPT w odpowiedzi na pytanie czym jest wzmacniacz dla elektronika snuje kwiecistą wypowiedź:

„W elektryzującym świecie elektroniki, wzmacniacze stanowią tytanów, kontrolując moc dźwięku z precyzją doświadczonego dyrygenta. Niezależnie od tego, czy jesteś początkującym hobbystą, czy doświadczonym weteranem, zrozumienie wzmacniaczy jest jak posiadanie klucza do ogromnego królestwa czystości dźwięku, mocy i niuansu. W swojej istocie misja wzmacniacza jest prosta: wziąć sygnał audio i wzmacnić go na mocniejszy, na tyle silny, by napędzać głośniki i dostarczać dźwięk, który może sięgać od szeptu bryzy do ryku koncertu. Jednak prostota tej misji ukrywa złożoność i pomysłowość wbudowaną w projekt i funkcję wzmacniacza.”

Dla adepta elektroniki liczy się jednak co innego. Cytuję fragment z podręcznika mojego autorstwa.

„Wzmacniacz – jak wskazuje nazwa – jest przyrządem służącym do wzmacniania, czyli do powiększania czegoś. Przykładowo lupa lub lornetka teatralna powiększają obraz, dźwignia mechaniczna zwiększa siłę, transformator zwiększa napięcie (lub prąd).

Czy te przyrządy są wzmacniaczami? Nie. Nie są to wzmacniacze, gdyż **wzmacniacz jest przyrządem umożliwiającym sterowanie większej mocy mniejszą**. Aby nastąpił efekt wzmacnienia, są konieczne dwie rzeczy: źródło energii i przyrząd do sterowania przepływu tej energii – wzmacniacz. Jeżeli na przykład strumień wody z węża ogrodniczego skierujemy na łopatkę turbiny, która obracając się będzie wykonywała jakąś pracę, to przekręcając kran wodny (jest to czynność nie wymagająca dużych energii) można powodować znaczne zmiany energii obracającej się turbiny. W tym przykładzie kran jest przyrządem służącym do sterowania dużej mocy za pomocą malej, czyli jest wzmacniaczem. Wzmacniaczem w ogólnym sensie jest również wyłącznik sieci oświetleniowej, wyłącznik radiowy itp. Transformator natomiast, nawet 1000-krotnie podwyższający napięcie lub prąd, nie jest wzmacniaczem, gdyż moc wydzielana na jego wyjściu może być co najwyżej prawie równa mocy dostarczanej na wejście. Natomiast tranzystor jest wzmacniaczem stosowanym zarówno do liniowego (wprost proporcjonalnego) zwiększania mocy sygnału, jak również do nieliniowego, przy czym często dyskretnego (skokowego, kłuczającego) sterowania mocy.”

Poza elementem (tranzystorem lub lampą elektronową) sterującym przepływem energii ze źródła zasilającego, drugim składnikiem niezbędnym dla prawidłowej pracy wzmacniacza jest **sprężenie zwrotne**, które określa wzmacnienie, stabilizuje pracę wzmacniacza i zmniejsza zniekształcenia z poziomu około 1% nawet do 0,001%. Rozumiejąc fundamentalną rolę tranzystora i sprzężenia zwrotnego możemy przystąpić do zgłębiania wykładu o wzmacniaczach tranzystorowych.

Wiesław Marciniak

W redakcyjnej szufladzie (głównie w skrzynce pocztowej) sporo miejsc zajmują listy czytelników skłonnych do uprawiania swojej filozofii obwodowej i fizycznej. Gdy chodzi o ujawnianie paradoksów, to podejmujemy dyskusję, ale często są to problemy „wydumane”, podobne do „zdroworozsądkowych” poglądów płaskoziemców. Na takie poglądy reagujemy niechętnie, chyba że jest ku temu specjalna okazja. Publikowany w tym wydaniu EdW wykład o wzmacniaczach tranzystorowych stwarza okazję do zareagowania na kilka listów prezentujących filozoficzne podejście do zasady działania tranzystora, w szczególności do jego funkcji wzmacniania sygnału. Najbardziej wyraziście są poglądy czytelnika, który wątpi w działanie wzmacniające tranzystora i przesłał nam cytaty z Internetu, prezentujący opinię niejakiego Cyrila Mieszkowa, wyrażającego podobne, a nawet tożsame poglądy. Oto ten cytat:

Czy tranzystor jest rzeczywiście elementem wzmacniającym? Czy jest to urządzenie aktywne czy pasywne? Czy istnieją elementy wzmacniające? Czy w ogóle możliwe jest wzmacnianie energii?

Powszechnie wiadomo, że tranzystor jest elementem aktywnym wykorzystywanym do budowy wzmacniaczy. Ale to nie jest prawda. Tranzystor nie jest elementem aktywnym, lecz pasywnym; jedyną rzeczą, jaką może zrobić tranzystor, jest rozpraszanie energii. Nie jest to więc element wzmacniający, lecz tłumiący. Jest to po prostu rezystor (nieliniowy, sterowany elektrycznie, ale wciąż rezystor), który zmniejsza natężenie prądu.

Prawdziwe wzmocnienie jest niemożliwe; więc nie ma prawdziwych wzmacniaczy. Tak zwane „wzmocnienie” to tylko iluzja, sprytna sztuczka, a „wzmacniacz” to tylko „magiczne pudełko”, w którym widzimy większą moc wyjściową, ale nie jest to wzmocniona mała moc wejściowa. To jest inna moc.

W elektronice analogowej realizujemy takie „wzmocnienie” w najbardziej paradoksalny, absurdalny i głupi sposób – aby uzyskać moc wyjściową większą niż wejściowa, uzyskujemy duże źródło zasilania, a następnie wyrzucamy jego część (od zera do całej mocy). Dla porównania, w energetyce nie mogą sobie na to pozwolić...

Mam rację?

Red. Nie masz racji. Czybyś oczekiwał, że tranzystor „sam z siebie” zwiększy moc, czyli z mniejszej mocy na wejściu „wyprodukuje” większą moc na wyjściu? Nie ma perpetuum mobile. Energia nie powstaje z niczego, tylko może zachodzić jej transformacja z jednej postaci w inną. Funkcja wzmocnienia jest realizowana przez tranzystor poprzez regulację przepływu energii ze źródła zasilającego w sposób modulowany sygnałem wejściowym. Tego tematu dotyczy też Wstępniak.

Więcej o SMD

Kilkanaście lat temu wróżyono koniec hobbistycznych projektów, a jedną z ważkich przyczyn miało być wyparcie elementów przewlekanych przez komponenty SMD, których montaż w warunkach amatorskich wydawał się niemożliwy. Na szczęście, rzeczywistość nie potwierdziła tych czarnych przewidywań. Okazało się, że w pewnym stopniu lutowanie ręczne SMD jest możliwe. Dziękuję, że podejmujecie tę tematykę. Zabrałem się do wykonania pęsety do testów SMD i zainteresował mnie zapowiadany na marzec trenażer SMD. Czy ta tematyka będzie kontynuowana?

N.Cz.

Red. Tak. W numerze kwietniowym jest planowany naszpikowany poradami praktycznymi artykuł „SMD Soldering Tips & Tricks” z magazynu Silicon Chip.

Patronat AVT

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

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

Subscribe to Elektor's newsletter and get the chance to

WIN

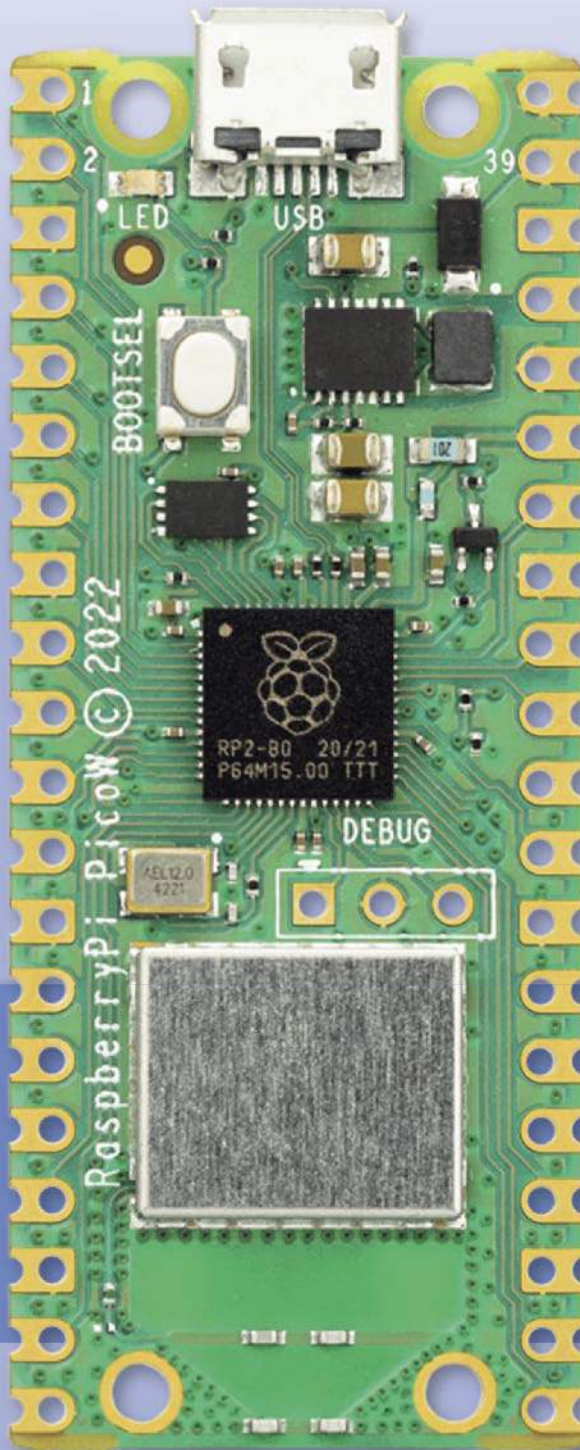
a Raspberry Pi Pico W board



www.elektor.com/eda



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



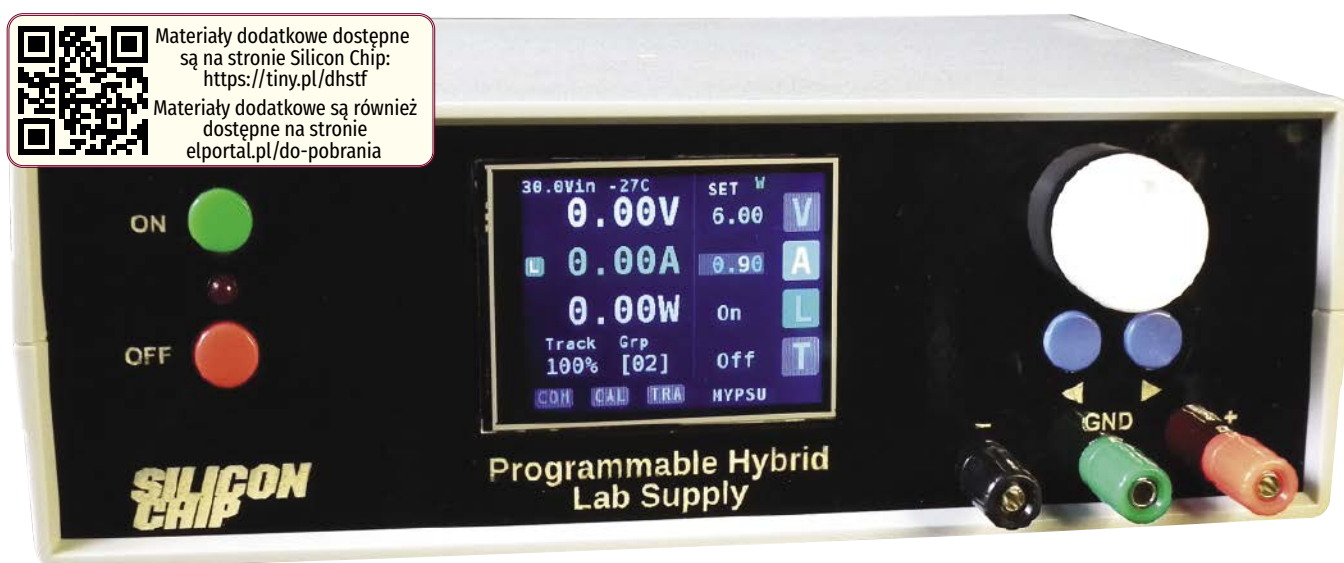
Be one of the 10 fortunate winners!



elektor
design > share > earn



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



Programowany hybrydowy zasilacz laboratoryjny z modułem Wi-Fi, część 1

Ten zasilacz laboratoryjny ma wbudowane bezprzewodowe sterowanie przez Wi-Fi a także obsługę za pomocą kolorowego ekranu dotykowego i enkodera obrotowego, z możliwością bezprzewodowej synchronizacji kilku zasilaczy. Jak na swoją wydajność, jest kompaktowy i niedrogi, dostarczając napięcie 0...27 V, z dopuszczalnym prądem 0...5 A przy napięciu do 18 V i nieco niższymi prądami powyżej. Posiada ograniczenie prądu i monitorowanie napięcia/prądu, łagodny rozruch, a jego końcowy stopień regulacji jest liniowy, co zapewnia dobre sterowanie i płynną regulację wyjścia DC.

Prezentowana konstrukcja, dzięki zastosowaniu przetwornicy impulsowej AC-DC i wstępnej regulacji napięcia, również w trybie przełączania, pozwala uniknąć nieporęcznych transformatorów mocy i znacznego wytwarzania ciepła. Końcowy stopień regulacji jest liniowy, co zapewnia lepsze sterowanie linii wyjścia i obciążenia, a także niższe tętnienia i szumy.

Dzięki niewielkiemu wytwarzaniu ciepła, zasilacz mieści się w kompaktowej obudowie z tworzywa sztucznego, a cała jednostka ma masę zaledwie 1,5 kg – mniej niż sam transformator mocy w konwencjonalnej konstrukcji.

Zasilacz jest programowalny dzięki czemu znakomicie sprawdzi się jako część zestawu przyrządów laboratoryjnych. Można go, na przykład, wykorzystać do zautomatyzowanego testowania. Interfejs Wi-Fi umożliwia zdalne monitorowanie za pośrednictwem łącza internetowego i zdalne sterowanie za pomocą

standardowego protokołu SCPI (Standard Commands for Programmable Instruments).

Napięcie i prąd są ustawiane w krokach co 10 mV i 10 mA, a napięcie jest sprawdzane z dokładnością do miliwoltów. W celu zachowania parametrów między sesjami, ustawienia urządzenia są przechowywane we wbudowanej pamięci flash. Ograniczenie prądu, zabezpieczenie przeciwzwarciowe i termiczne są sterowane programowo.

Ograniczenia bezpiecznego obszaru roboczego dla urządzeń wyjściowych są egzekwowane przez oprogramowanie, zapewniając dodatkową warstwę ochrony przed przeciążeniem zasilacza, oprócz zabezpieczeń wbudowanych dla trzech regulatorów.

Rysunek 1 przedstawia schemat działania prezentowanego zasilacza laboratoryjnego. Składa się on z trzech modułów: modułu sterującego na górze, modułu regulatora na dole oraz przetwornicy impulsowej AC-DC (gotowy

moduł), który zapewnia zasilanie prądem stałym dla wszystkich obwodów. Moduł sterujący jest zasilany z szyn zasilających o niższym napięciu, które pochodzą z modułu regulatora.

Więcej funkcji

Tradycyjnie zasilacze laboratoryjne, gdy wyjście jest podłączone za pomocą przełącznika lub przekaźnika, „uruchamiają się awaryjnie” (poza zakresem nominalnego obciążenia), w przeciwieństwie do zachowania większości zasilaczy wbudowanych w sprzęt, w których napięcie narasta przez dziesiątki milisekund. Ten zasilacz laboratoryjny ma funkcję łagodnego rozruchu, która podnosi napięcie od zera do ustawionej wartości w tempie 100 V na sekundę, gdy wyjście jest włączone.

Zdalne sterowanie obejmuje regulację napięcia wyjściowego i maksymalnego prądu za pośrednictwem Wi-Fi (TCP) i izolowanego modułu łącza USB. Zasilacz może łatwo

wykonywać sekwencje skryptowe, takie jak skokowe zmiany napięcia i dowolne przyrosty napięć.

Na przykład, można napisać skrypty SCPI w EEZ Studio (do pobrania za darmo ze strony <https://github.com/eez-open/studio>), aby ustawić napięcie wyjściowe naprzemiennie na dwie różne wartości, aby przetestować regulację obciążenia urządzenia lub reakcję na skokową zmianę napięcia wejściowego.

Po uruchomieniu urządzenia nie jest zalecana bezpośrednia komunikacja szeregową USB. Lokalne uzimienie USB jest bezpośrednio podłączone do ujemnego zacisku zasilacza, który zwykle jest na potencjale pływającym. Dlatego podłączenie ujemnego zacisku wyjściowego do źródła napięcia może spowodować uszkodzenie komputera. Znacznie bezpieczniej jest korzystać ze sterowania Wi-Fi lub użyć izolatora USB.

Zamiast podłączać przyrząd do istniejącej sieci Wi-Fi LAN, można również skonfigurować go tak, aby uruchamiał własną sieć chronioną hasłem z identyfikatorem SSID ESPINST.

Po włączeniu zasilania zasilacz najpierw próbuje połączyć się z istniejącą siecią Wi-Fi, jeśli dane uwierzytelniające zostały wcześniej podane za pośrednictwem menu ekranowego. Jeśli to się nie powiedzie, próbuje połączyć się z istniejącą siecią Wi-Fi ESPINST. Jeśli to się nie powiedzie, urządzenie samo skonfiguruje sieć Wi-Fi ESPINST.

Jeśli używana jest istniejąca sieć Wi-Fi, dostęp do zasilacza można uzyskać za pomocą jego adresu IP lub lokalnej nazwy urządzenia (domyślnie MYPSU.local) przy użyciu protokołu mDNS.

Urządzenie udostępnia stronę internetową, która wyświetla ustawienia i zmierzane wartości, wraz z „dużym czerwonym przyciskiem” do zdalnego wyłączenia wyjścia. Na stronie internetowej nie są dostępne żadne inne elementy sterujące, ponieważ nie jest ona zabezpieczona.

Kilka programowanych zasilaczy można skonfigurować jako grupę, komunikującą się przez Wi-Fi, co umożliwia zapewnienie normalnych funkcji śledzenia działania zasilaczy, tj. powiązanych ustawień napięcia i zsynchronizowanego ograniczania prądu bez konieczności korzystania z komputera hosta.

Ponieważ każdy zasilacz ma w pełni „pływające” wyjście, można je również łączyć szeregowo w celu zapewnienia wyższego napięcia wyjściowego lub równolegle w celu uzyskania wyższego prądu.

Ponieważ ograniczona ilość miejsca uniemożliwia tu pełne omówienie wszystkich funkcji urządzenia i sposobu ich wykorzystania, dlatego pełne opisy znajdują się w podręczniku

Właściwości i parametry:

- Hybrydowy zasilacz laboratoryjny z sieciowym zasilaczem przetwarzającym AC-DC, wstępnym regulatorem przetwarzającym DC-DC i końcowym regulatorem liniowym
- Zdalne monitorowanie i sterowanie przez Wi-Fi
- Kompaktowy, lekki i o niskim rozpraszaniu ciepła
- Możliwość skoordynowania kilku jednostek w celu elastycznego współdziałania i śledzenia parametrów
- Wysokowydajna konstrukcja z niskim poziomem tętnień i szumów na wyjściu
- Dostarcza do 24 V @ 0...3,5 A, 0...18 V @ 0...5 A.
- Rozdzielczość ustawień: 10 mV i 10 mA
- Zgrubna i precyzyjna regulacja napięcia i prądu wyjściowego
- Precyzja regulacji lepsza niż 1 mV i 1 mA
- Ograniczenie prądu, zabezpieczenie nadnapięciowe i nadprądowe
- Doskonała regulacja linii i obciążenia oraz dobra odpowiedź przejściowa bez przeregulowania
- Miękki start po włączeniu wyjścia, zapobiegający „awaryjnym” startom
- Obsługa HTTP, Telnet (TCP) i izolowanego sterowania szeregowego USB za pomocą uniwersalnych poleceń SCPI
- Uniwersalne wejście AC (100...240 V AC, 50...60 Hz)

dla tego projektu dostarczonym jako część plików do pobrania na stronie [siliconchip.com.au/link/ab72](https://www.siliconchip.com.au/link/ab72).

Przegląd funkcji

Napięcie wyjściowe i maksymalny prąd można ustawić za pomocą ekranu dotykowego, używając kombinacji przycisków dotykowych po prawej stronie ekranu (V i A), które wybierają ustawienie do zmiany, dwóch przełączników chwilowych wybierających, która cyfra jest zmieniana, oraz enkodera obrotowego do zmiany rzeczywistej wartości. Zapewnia to płynne przejście od regulacji zgrubnej do precyzyjnej.

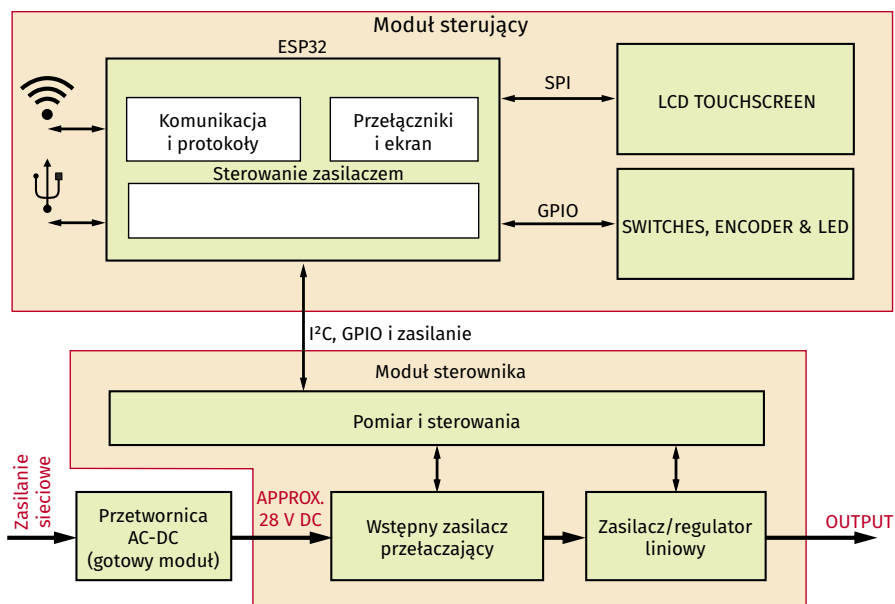
Ograniczenie prądu można włączyć za pomocą przycisku ekranowego (L), podobnie jak

funkcje śledzenia (T), gdy dostępne jest więcej niż jedno źródło zasilania.

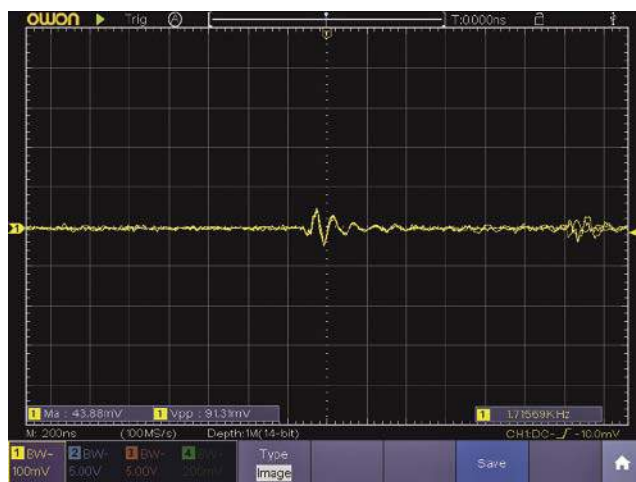
Rzeczywiste napięcie wyjściowe, prąd i moc są wyświetlane po lewej stronie ekranu głównego. Wzdłuż górnej krawędzi ekranu wyświetlane jest również napięcie wejściowe, temperatura radiatora, stan wentylatora i Wi-Fi, a także wskaźnik [E] (dla pamięci EEPROM), który pokazuje, kiedy trwa zapisywanie parametrów w pamięci FLASH.

Zapis do pamięci FLASH następuje z 30...40 sekundowym opóźnieniem po zmianie ostatniego ustawienia, ponieważ gwarantowana żywotność pamięci wynosi mniej niż 100 000 cykli kasowania/zapisu.

Dostęp do podmenu ustawień parametrów komunikacji (COM), funkcji kalibracji



Rysunek 1. Zasilacz laboratoryjny zbudowany jest z trzech modułów: zasilacza sieciowego AC-DC, hybrydowego regulatora impulsowego/liniowego i modułu pomiarowego oraz płytki sterującej Wi-Fi wykorzystującej gotowy moduł mikroprocesora ESP-32, z kolorowym ekranem dotykowym



Oscylogram 1. Na wyjściu nie występują wykrywalne tętnienia sieci. Obecna jest niewielka ilość (35 mV RMS) szumów przełączania, głównie o częstotliwości przełączania regulatora wstępnego



Oscylogram 2. Przy pełnym obciążeniu 5 A, tętnienia o częstotliwości przełączania 260 kHz są bardziej wyraźne, ale nadal mniejsze niż 100 mV międzyszczytowe (pomarańczowy przebieg). Żółty przebieg pokazuje napięcie przed wyjściową cewką toroidalną, aby uwidocznić skuteczność w redukcji szumów nawet przez kilku zwojów dławika

(CAL) i śledzenia (TRA) uzyskuje się za pomocą przycisków umieszczonych w dolnej części ekranu. Po uruchomieniu zasilacza dostęp do podmenu rzadko będzie konieczny.

Dwa dedykowane przyciski po lewej stronie panelu włączają i wyłączają przełącznik wyjściowy. Te przyciski panelu sterowania są podłączone na stałe do płyty zasilacza, aby zapewnić, że wyjście może zostać odłączone nawet w mało prawdopodobnym przypadku, gdy jednostka centralna zostanie odłączona.

Wyjście zasilania jest pływające, więc trzeci zacisk GND jest przewidziany dla sytuacji, w których wymagane jest uziemienie zasilania.

Wydajność

Konwersję AC-DC zapewnia gotowy, dostępny na rynku moduł przetwornicy o napięciu znamionowym 24 V i prądzie 4,5 A (nominalnie 108 W). Dopóki jednak nie przekroczymy ogólnego limitu mocy, możemy pobierać nieco więcej prądu przy niższych napięciach lub nieco wyższe napięcie.

Przy pełnym skróceniu potencjometru na konwerterze, zasilanie AC-DC prototypu zapewnia napięcie nieco poniżej 30 V. Przy lekkich i umiarkowanych obciążeniach, regulator wstępny i regulator końcowy mają spadek napięcia poniżej 2 V, co daje teoretyczne maksymalne napięcie wyjściowe do 27 V z zasilania DC 30 V.

Gdy obciążenie wzrasta, a para tranzystorów Q1/Q2 zaczyna przewodzić (sytuacja opisana bardziej szczegółowo poniżej), spadek napięcia wzrasta, co ogranicza maksymalne napięcie wyjściowe do nieco poniżej 24 V przy pełnej mocy. Charakterystyka ta wypada korzystnie w porównaniu ze spadkiem napięcia obserwowanym przy

dużym obciążeniu w konstrukcjach opartych na transformatorach.

Kilka czynników ogranicza maksymalny prąd wyjściowy zasilacza: całkowity zakres mocy przetwornicy AC-DC, wartość jej prądu znamionowego równa 4,5 A przy pełnej mocy oraz wartość znamionowego prądu wejściowego równa 5 A dla stopnia regulatora wstępnego (czyli konwertera DC-DC).

Czerwona linia na **rysunku 2**, krzywa bezpiecznego obszaru roboczego (SOA) zasilacza, pokazuje jego ograniczenia. Regulator wstępny może obsłużyć do 5 A, definiując górną linię. Narożnik odcięcia odpowiada mocy wyjściowej 90 W, ponieważ 18 W z 108 W mocy przetwornicy AC-DC jest przy pełnej mocy przekształcane w ciepło przez stopień liniowy. Prawa linia to maksymalne napięcie wyjściowe 27 V.

Stopień mocy może dostarczać przy wyższych napięciach nieco mniej niż absolutna maksymalna moc, a czerwona linia wskazuje jego zmniejszoną wydajność. Ograniczenia prądu SOA są wymuszane przez oprogramowanie: nawet jeśli ustawisz 5 A jako punkt ograniczenia prądu przy 20 V, ograniczenie rozpocznie się przy około 4,5 A, aby zapewnić, że maksymalna moc konwertera nie zostanie przekroczona.

Tętnienia i szumy

Tętnienia napięcia wyjściowego są niewielkie, a najbardziej znaczące wahania występują przy częstotliwości przełączania regulatora wstępnego wynoszącej 260 kHz (patrz **oscylogram 1**). Wyjściowy (pomarańczowy) przebieg na **oscylogramie 2** wskazuje, że 37 mV RMS (150 mV peak-to-peak) niepożądanych zakłóceń wyjściowych obejmuje tętnienia 100 mV

peak-to-peak, nałożone na stany przejściowe przełączania 50 mV.

Żółty przebieg, pokazujący napięcie na wyjściu regulatora liniowego, jest prawie identyczny z tym na jego wejściu, potwierdzając, że jego zdolność do tłumienia tętnień przy wysokich częstotliwościach nie jest rewelacyjna. Poprawa tłumienia szumów RF jest spowodowana dławikiem między płytką drukowaną a zaciskiem wyjściowym. Zwiększenie jego indukcyjności jeszcze bardziej zmniejszyłoby niepożądany sygnał, choć prawdopodobnie spowodowałoby niestabilność na wyjściu przy niektórych obciążeniach pojemnościowych.

Wpływ obciążenia

Oscylogram 3 pokazuje, że skokowa zmiana obciążenia od 0 do 2 A nie powoduje prawie żadnej mierzalnej zmiany napięcia wyjściowego. Krótkie skoki są spowodowane bardzo krótkimi czasami narastania i opadania prądu, ponieważ obciążenie było zasilane z tranzystorów wyjściowych sterowanych falą prostokątną.

Po wyłączeniuysterowania występuje niewielki dodatni skok przy spadku prądu do zera (około 100 mV), spowodowany reakcją oprogramowania na stan nieustalony. Zjawisko to ustabilizowało się w ciągu 10 ms. Po ponownym włączeniu obciążenia napięcie chwilowo spada o podobną wartość i stabilizuje się w czasie krótszym niż 5 ms.

Na **oscylogramie 4** tranzystory mocy obciążenia są sterowane falą trójkątną wytwarzającą impuls prądowy przełączania tranzystorów wyjściowych (przebieg zielony) o czasie narastania 1 ms. Na żółtym przebiegu napięcia wyjściowego nie ma zauważalnego skoku przełączania ani wahań napięcia.

Oscylogram 3. Zachowanie wyjścia (żółty przebieg), gdy obciążenie 2 A szybko się włącza i wyłącza (pomarańczowy przebieg) przy użyciu tranzystorów wyjściowych sterowanych falą prostokątną

Oscylogram 4. Stany nieustalone na Oscylogramie 3 są eliminowane, gdy zmiana obciążenia ma dłuższy czas narastania, ponieważ tranzystory wyjściowe są sterowane falą trójkątną. Warto porównać przebieg zielony z z przebiegiem pomarańczowym na oscylogramie 3

Regulacja napięcia

Chociaż zasilacz jest zdolny do dokładniejszej regulacji napięcia, do algorytmu sterowania napięciem została wprowadzona histereza 1 mV, aby zapobiec „samooscylacjom”. W prawie wszystkich praktycznych sytuacjach stabilność napięcia wyjściowego zasilacza jest znacznie bardziej krytyczna niż rzeczywista wartość napięcia. W końcu jest to zasilacz laboratoryjny, a nie zasilacz napięcia wzorcowego.

W przypadku konstrukcji sterowanej pokrętle, głównym kryterium jest możliwość zmiany napięcia tak szybko, jak można obracać pokrętle. Jednak w przypadku sterowania cyfrowego, w szczególności zdalnego sterowania cyfrowego, praktyczne staje się użycie zasilacza do zapewnienia skokowych zmian napięcia, a nawet generowania przyrostów napięcia lub fal prostokątnych. W takich warunkach czas ustalania się napięcia jest bardziej istotny.

Zdolność zasilacza do radzenia sobie ze zmianami napięcia wyjściowego pod obciążeniem jest dość znaczna. Oscylogram 5 pokazuje reakcję samego regulatora na krótki skok napięcia od 1 V do 10 V przy obciążeniu 20 Ω . Pokazuje przeregulowanie po czasie narastania wynoszącym 35 ms, przy czym napięcie ustabilizowało się w czasie krótszym niż 100 ms.

Ponieważ niepożądane jest przeregulowanie napięcia, szybkość narastania napięcia jest celowo ograniczona programowo do 100 V/s (Oscylogram 6), bez znaczącego pozostałego przeregulowania. Przy spadającym napięciu (Oscylogram 7), sygnał wyjściowy ustabilizował się w ciągu 25 ms przy minimalnym niedostosowaniu. Szybkość spadku napięcia nie jest ograniczona programowo i zależy głównie od stałej czasowej rezystancji obciążenia i kondensatora wyjściowego 10 μF .

Ograniczenie szybkości wzrostu napięcia, w połączeniu ze wzrostem napięcia wyjściowego od 0 V, gdy wyjście jest włączone, jest zasadniczym elementem funkcji miękkiego startu.

Konstrukcja sprzętowa

Układ blokowy zasilacza laboratoryjnego pokazano na uproszczonym schemacie, **rysunek 3**. Zasilanie AC jest konwertowane na 28...30 V DC przy użyciu dostępnego na rynku, gotowego modułu przetwornicy impulsowej o mocy 100 W. Moduł AC-DC został wybrany w celu zmniejszenia rozmiaru, kosztów oraz ciężaru budowanego zasilacza.

Następnie regulator obniżający DC-DC zbudowany na LM2679 zmniejsza napięcie DC do poziomu o 3,6 V powyżej wymaganego napięcia wyjściowego. Wreszcie, liniowy regulator mocy, wykorzystujący LM317, obniża napięcie do prawidłowej wartości wyjściowej.

Prąd wyjściowy jest konwertowany na napięcie za pomocą przetwornika prądu na napięcie

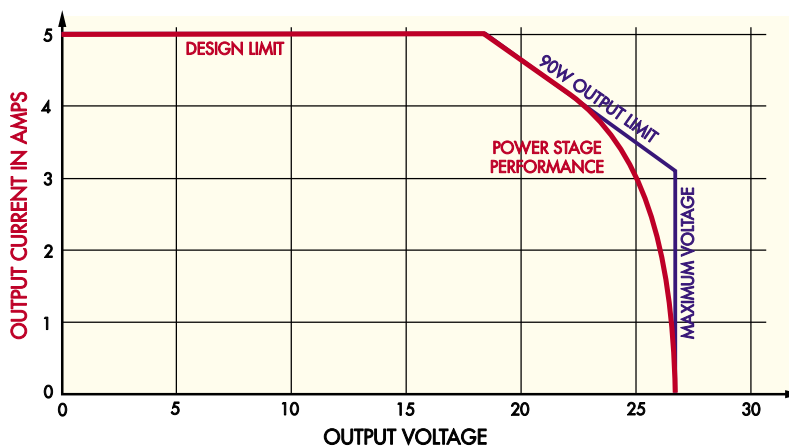
i układu INA282 mierzącego spadek napięcia na rezystorze bocznikowym 0,01 Ω . Napięcie i prąd wyjściowy są następnie mierzone za pomocą 16-bitowego przetwornika analogowo-cyfrowego (ADC).

Napięcie wyjściowe LM317 jest regulowane przez cyfrowy potencjometr IC3 i ustawiane za pomocą wyjścia przetwornika cyfrowo-analogowego (DAC).

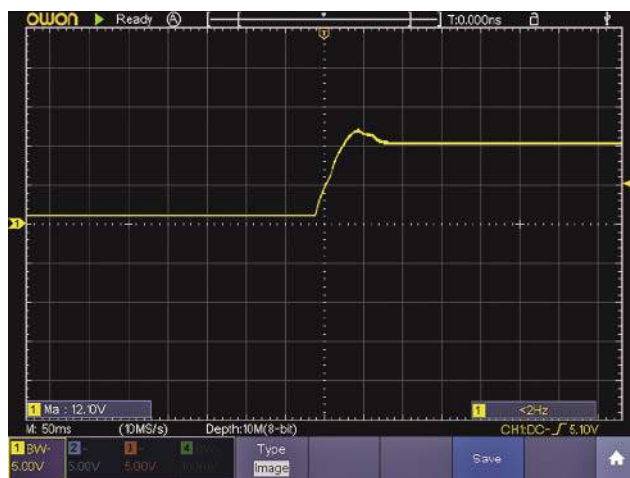
Wszystkie cyfrowe funkcje sterowania wykorzystują magistralę szeregową I²C, a dwa moduły mogą współdzielić jeden sterownik, poprzez zmianę za pomocą zworki jednego bitu adresu I²C każdego regulatora.

Chociaż trzystopniowe podejście do regulacji napięcia może wydawać się skomplikowane, zapewnia ono najlepszą równowagę między wydajnością a prostotą, spośród kilku testowanych konfiguracji.

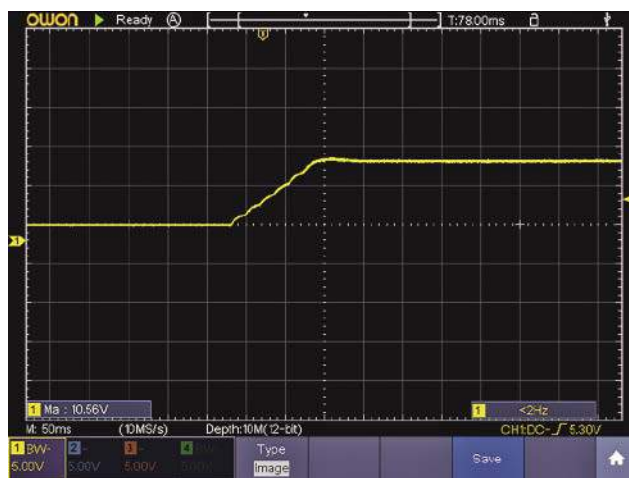
Jednym z kluczowych wyzwań projektowych w każdej konstrukcji impulsowej jest minimalizacja szumów przełączania na wyjściu.



Rysunek 2. Bezpieczny obszar działania (SOA) zasilacza laboratoryjnego. Może on dostarczyć 5 A przy napięciu do 18 V. Powyżej tej wartości, limit 108 W konwertera AC-DC i 18 W rozproszona w stopniu liniowym powodują, że maksymalny prąd obniża się, aż do osiągnięcia maksymalnego napięcia, które można wytworzyć, biorąc pod uwagę spadek napięcia regulatora liniowego



Oscylogram 5. Przy znacznym skoku napięcia o 9 V z obciążeniem 20 Ω , nieprzetworzone wyjście wykazuje niepożądane przeregulowanie, pomimo krótkiego czasu ustalania wynoszącego około 75 ms



Oscylogram 6. Ograniczenie czasu narastania napięcia do 100 V/s prawie eliminuje przeregulowanie

Na każdym etapie regulacji zwrócono szczególną uwagę, aby zminimalizować generowanie i przenoszenie tych zakłóceń.

Zasilacz jest zbudowany na dwóch płytkach drukowanych: jednej, która zawiera wszystkie elementy regulacyjne, oraz drugiej płytki, sterującej, która zawiera moduł mikrosterownika z Wi-Fi, ekran dotykowy, przyciski i pokrętkę. Są one połączone ze sobą kablem taśmowym.

Ponieważ kilka ważnych układów na płycie zasilacza jest dostępnych tylko jako części SMD, zdecydowaliśmy się na projekt w pełni SMD. Rozmiary części utrzymaliśmy na poziomie 2,0x1,2 mm (0805) lub większym, aby ułatwić montaż ręczny.

Dla tych Czytelników, którzy jeszcze nie odważyli się na montaż projektu SMD, mamy propozycję rozważenia budowy naszego pieca rozplwowego DIY SMD (EdW, maj i czerwiec 2023; w oryginale siliconchip.com.au/Series/343), aby zrealizować ten projekt. Jeśli jednak chcesz, możesz również złożyć

go ręcznie; w takim przypadku oprócz lutownicy z regulacją temperatury potrzebujesz strzykawki z pastą topnikową, miedzianej plecionki lutowniczej i pęsety z cienką końcówką oraz opcjonalnie lupy z podświetleniem.

Poziom rozpraszania ciepła

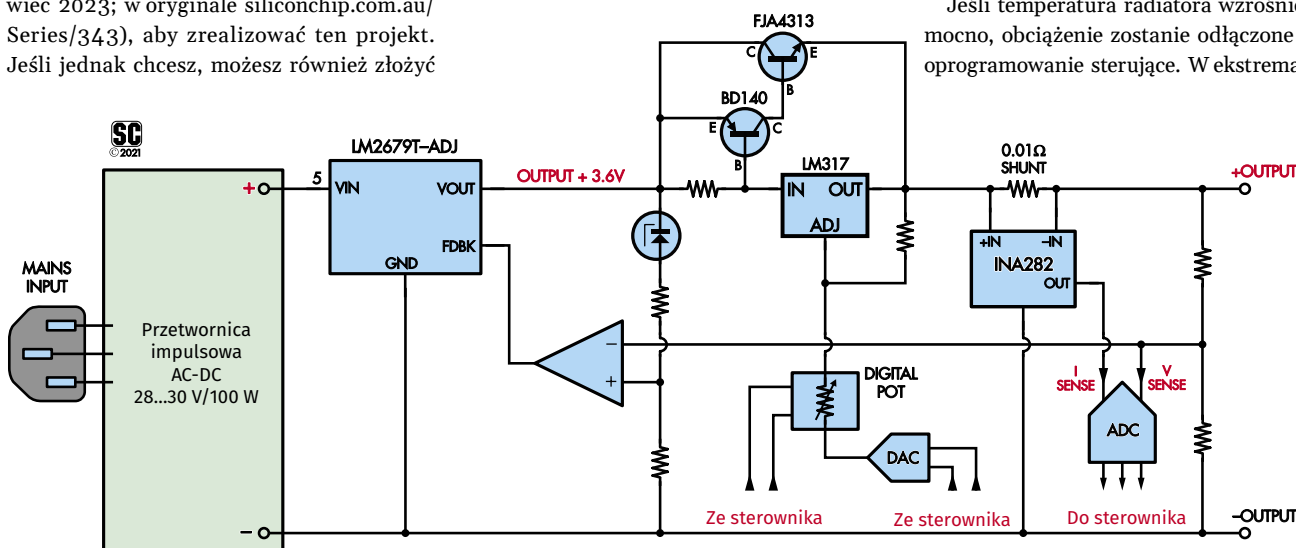
Większość ciepła odpadowego w zasilaczu jest generowana przez jeden tranzystor, Q2. Regulator wstępny utrzymuje na nim stałe napięcie, wyższe na kolektorze o 3,6 V niż na emiterze, więc jego moc cieplna jest wprost proporcjonalna do prądu obciążenia. Przy pełnym prądzie, Q2 wygeneruje 18 W ciepła odpadowego. Natomiast regulator LM317 działa przy niskim prądzie i niskim napięciu różnicowym.

Sprawność regulatora wstępnego w całym zakresie napięcia i prądu wynosi co najmniej 84%. Przy pełnym obciążeniu i sprawności

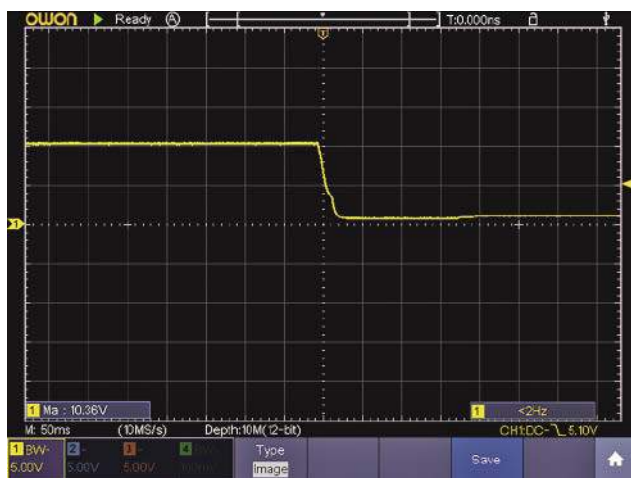
znamionowej może on wygenerować 18 W strat ciepłych, dzielonych między układ scalony regulatora, diodę Schottky'ego i cewkę indukcyjną. W praktyce ciepło generowane w tej sekcji naszego prototypu było znacznie mniejsze niż w przypadku szeregowego tranzystora Q2.

Przy potencjalnym maksymalnym poziomie 36 W strat ciepła do rozproszenia, ta hybrydowa konstrukcja stanowi znaczną poprawę w stosunku do 108 W strat, które byłyby generowane przez w pełni liniową konstrukcję dostarczającą pełny prąd przy napięciu wyjściowym w pobliżu zera woltów. Umiarkowana moc cieplna strat pozwala na zamontowanie radiatora o niewielkich rozmiarach w kompaktowym plastikowym pudełku, z małym wentylatorem, który wymusza ruch powietrza, gdy temperatura radiatora wzrasta.

Jeśli temperatura radiatora wzrośnie zbyt mocno, obciążenie zostanie odłączone przez oprogramowanie sterujące. W ekstremalnych



Rysunek 3. Uproszczony schemat blokowy obwodu demonstrujący działanie zasilacza laboratoryjnego. Po konwerterze AC-DC następuje regulator wstępny wykorzystujący regulator przetwórczy LM2679 5A, a następnie stopień liniowy składający się z LM317 z parą tranzystorów zwiększających prąd wyjściowy. Mikrokontroler monitoruje napięcie i prąd wyjściowy oraz steruje zaciskiem ADJ na LM317 za pomocą zmiennej rezystancji (za pośrednictwem potencjometru cyfrowego) i niewielkiego napięcia, dostarczanego przez przetwornik cyfrowo-analogowy w celu precyzyjnej regulacji



Oscylogram 7. Prawie nie występuje przeregulowanie przy napięciu opadającym na wyjściu (10 V, 20 Ω), więc ograniczenie czasu spadku napięcia nie jest wymagane

okolicznościach układy LM317 i LM2679 uruchomią swoje wewnętrzne obwody wyłączania termicznego, zapewniając ostatnią warstwę ochrony.

Istnieją dwie opcje radiatora dla tego projektu: można użyć komercyjnego radiatora (Cincon M-B012) lub można go złożyć z dwóch kawałków blachy aluminiowej o grubości 1,6 mm. Ponieważ rozpraszanie mocy nie jest aż tak wysokie, oba rozwiązania będą wystarczające. Plany radiatora DIY zostaną pokazane później.

Płytki sterująca

Regulator wyposażony jest w potężny układ ESP32 Wi-Fi typu „system-on-a-chip” (SoC), dużego brata modułu ESP8266 opisanego w naszym D1 Mini Backpack-u (EdW – kwiecień 2023, oryginalny tekst w SC w październiku 2020; siliconchip.com.au/Article/14599). Ma on na pokładzie dwa procesory, dzięki czemu jeden z nich może być dedykowany funkcjom komunikacyjnym.

Chociaż może się to wydawać nieistotne, ponieważ 32-bitowy procesor 180 MHz ma znacznie większą wydajność niż jest to potrzebne w przypadku najbardziej ambitnych projektów, funkcje Wi-Fi w projekcie jednoprosesorowym mają priorytet przed kodem użytkownika, czasami powodując niedopuszczalne opóźnienia przetwarzania w aplikacjach czasu rzeczywistego, takich jak ta.

ESP32 ma 520 kB pamięci RAM, w porównaniu do 80 kB w ESP8266. Jest to szczególnie ważne w przypadku przeprogramowywania „over-the-air” (OTA), ponieważ zarówno oryginalny, jak i nowy program muszą zmieścić się w pamięci jednocześnie.

Sterownik komunikuje się przez Wi-Fi, łącząc się z lokalną siecią LAN lub konfiguruje własną. Obsługiwana jest również

komunikacja Bluetooth, zarówno tradycyjna, jak i niskoenerygetyczna (BLE), a także komunikacja szeregową przez USB.

Moduł ESP32 podłącza się do gniazda na płytce sterującej. Wybrany przez nas moduł DevKit C ma znaczne możliwości rozbudowy (32 styki w porównaniu do 16 w D1). Jest to projekt referencyjny Espressif, który został wdrożony przez wielu producentów płyt, zapewniając szeroką dostępność i konkurencyjne ceny.

Ekran dotykowy LCD o przekątnej 2,8 cala lub 3,5 cala jest zamontowany z przodu płytki sterującej, wraz z dwoma przełącznikami chwilowymi i enkoderem obrotowym. W tym projekcie są one używane (wraz z menu dotykowym na ekranie) do ustawiania konfiguracji urządzenia i wartości sterujących. Po lewej stronie znajdują się dwa kolejne przełączniki i jedna dioda LED, używane jako przyciski włączania/wyłączania i wskaźnik wyjścia.

Możliwości rozbudowy sterownika są dostępne na 20-szpilkowym wtyku typu IDC Z-WS20 i obejmują porty I²C, SPI, szeregowo, GPIO, ADC, DAC i zasilania (3,3 V i 5 V). Sterownik może być zasilany za pomocą kabla USB, zewnętrznego zasilacza 5...12 V lub poprzez złącze IDC. Płytki zasilacza będzie zasilac sterownik w gotowym projekcie przez kabel taśmowy, podczas gdy zasilanie USB jest używane do prac przy uruchomieniu.

Pełny zakres funkcji płytki sterującej jest zawarty w instrukcji PDF, którą można uzyskać za pośrednictwem następującego łącza: siliconchip.com.au/link/ab72

Obwód regulatora

Pełny schemat ideowy obwodu płytki regulatora pokazano na **rysunku 4**. Przychodzący prąd stały z zasilacza impulsowego AC-DC jest podawany w lewym górnym rogu, a zaciski wyjściowe znajdują się w prawym górnym rogu.

Jest to konstrukcja uniwersalna...

Chociaż płytka sterująca opisana w tych artykułach została zaprojektowana głównie do sterowania tym konkretnym zasilaczem, jest to zasadniczo styl kompatybilny z Arduino „BackPack-iem” z dwoma potężnymi 32-bitowymi mikroprocesorami, dużą ilością pamięci FLASH i pamięci RAM oraz obsługą Wi-Fi i Bluetooth. Może więc być wykorzystywany do szerokiej gamy różnych projektów i zadań, i został zaprojektowany z taką myślą i zamierzeniem.

Sekcje po obu stronach, w których zamontowane są przyciski i enkoder obrotowy, można odłączyć, jeśli nie są one wymagane w danym projekcie. Można je również podłączyć z powrotem do głównej części płytki sterującej, jeśli ich funkcje są pożądane, ale trzeba zmienić ich lokalizację. Alternatywnie, w tych lokalizacjach można zamontować złącza, aby zapewnić więcej styków We/Wy niż jest dostępnych na 20-stykowym złączu IDC CON2.

Układ zasilania jest również elastyczny. Regulator może być zasilany napięciem około 7...15 V DC podanym na gniazdo zasilacza wtyczkowego, przez gniazdo USB w module ESP-32 lub przez styki CON2.

Nie możemy też zapominać o opcjonalnym gnieździe kart microSD. Podsumowując, jest to bardzo wydajny i elastyczny moduł sterujący, który zasługuje na wykorzystanie w innych aplikacjach!

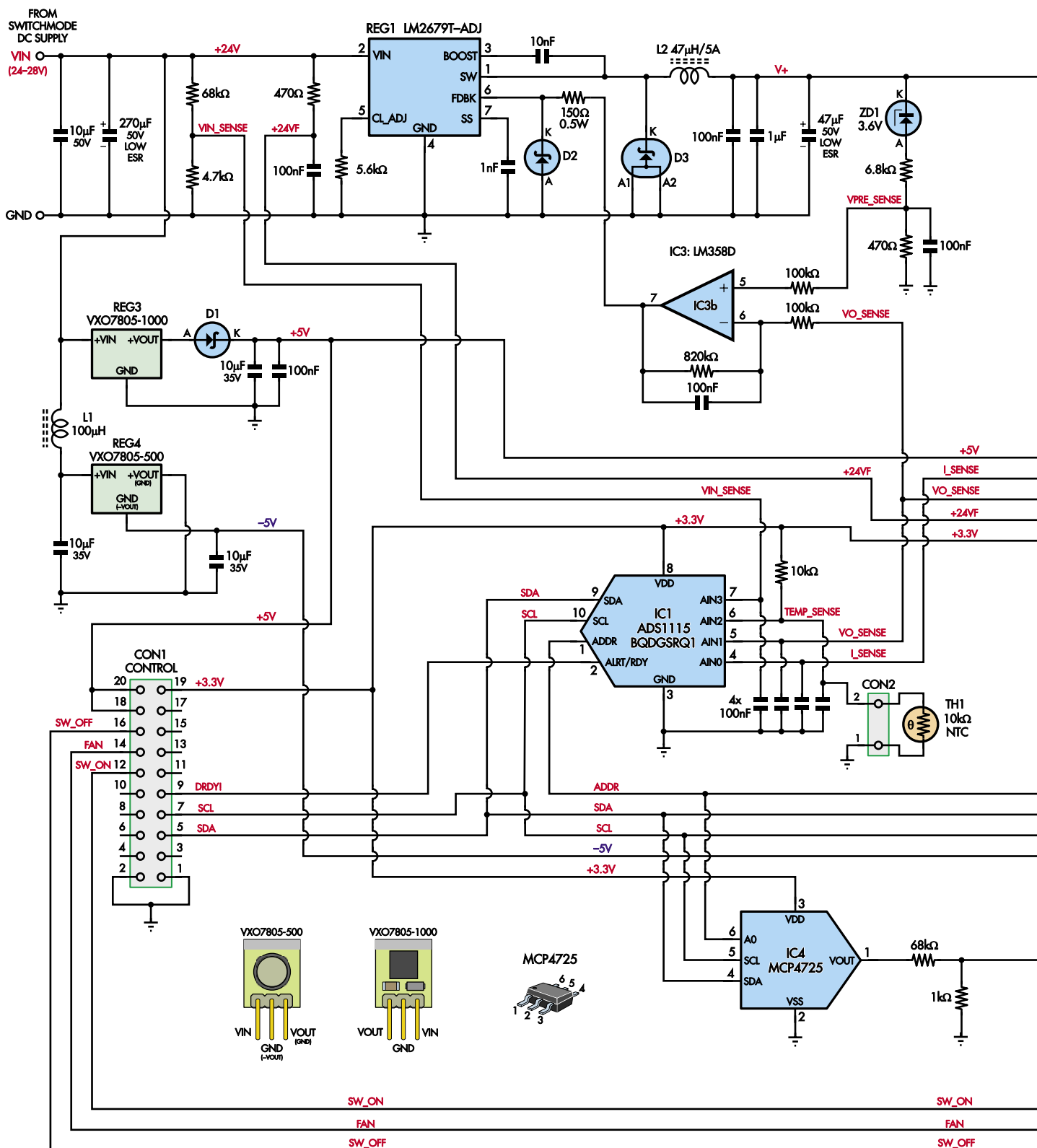
Napięcie wejściowe zasilania stopień wstępnego regulatora LM2679 (REG1), który jest nadzorowany przez wzmacniacz operacyjny IC3b. Końcowe napięcie wyjściowe i napięcie regulatora wstępnego są dzielone przez współczynnik 15 (68 k Ω /4,7 k Ω dla VO_SENSE i 6,8 k Ω /470 Ω dla VPRE_SENSE) przed odjęciem przez IC3b, działający jako wzmacniacz różnicowy. Różnica jest podawana do zacisku sprzężenia zwrotnego (FDBK) układu LM2979.

Napięcie regulatora wstępnego musi być o 3,6 V wyższe niż napięcie wyjściowe, aby umożliwić maksymalny spadek napięcia na końcowym stopniu regulacji liniowej. Dlatego na górze dzielnika VPRE_SENSE jest umieszczona dioda Zenera ZD1.

Wzmacniacz operacyjny ma umiarkowane wzmocnienie DC, aby zapewnić dokładne śledzenie, mimo że wejście FDBK REG1 ma punkt pracy 1,2 V. Wzmacniacz operacyjny jest silnie tłumiony zmiennoprądowo przez kondensator 100 nF równoległy z rezystorem sprzężenia zwrotnego, więc jego wzmocnienie AC jest bliskie jedności, zapewniając stabilność konfiguracji.

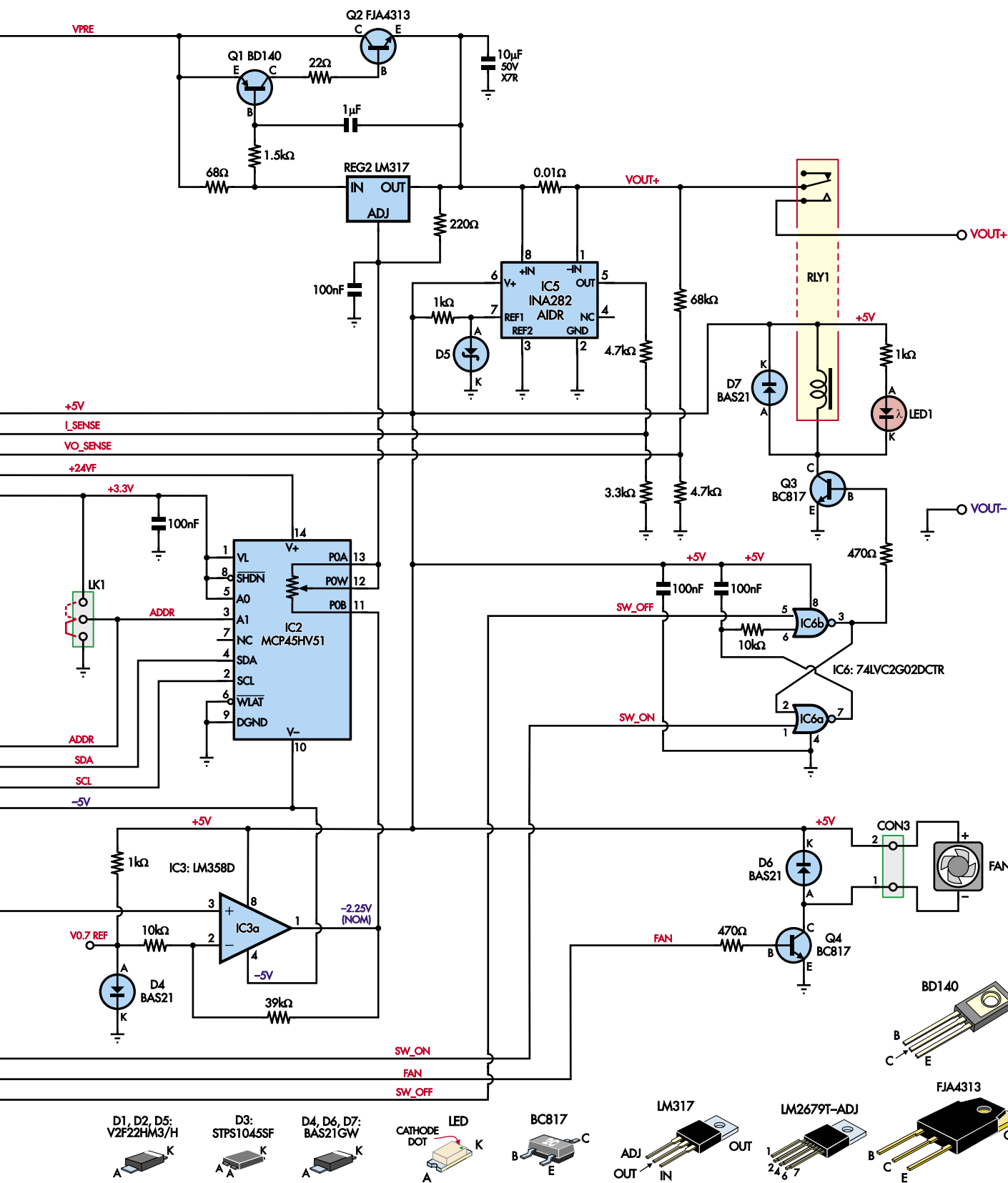
Dioda Schottky’ego (D2) na wejściu FDBK zapewnia, że napięcie nie będzie zbyt mocno ujemne podczas rozruchu, co potencjalnie mogłoby uszkodzić regulator. Funkcje miękkiego startu i ograniczania prądu LM2679 są włączone, a rezystor 5,6 k Ω od styku CL_ADJ do GND został wybrany w celu ograniczenia maksymalnego prądu przełączania tranzystora wyjściowego do 6,3 A.

Wybór napięcia 3,6 V dla diody Zenera ZD1 był kluczową decyzją projektową. Zwiększenie spadku napięcia na stopniu liniowym zwiększa straty na ciepło odpadowe. Jeśli jednak różnica napięć na LM317 stanie się zbyt mała, przestanie on pracować i może wpaść w oscylacje w połączeniu z tranzystorami liniowej regulacji prądu wyjściowego.



Laboratoryjny zasilacz hybrydowy –sekcja sterowania

Rysunek 4. Schemat ideowy płytki regulatora zawiera wstępny regulator impulsowy, zbudowany na REG1, końcowy stopień regulatora liniowego (REG2, Q1 i Q2) oraz obwody sterujące i monitorujące. Potencjometr cyfrowy IC2, przetwornik cyfrowo-analogowy IC4 i wzmacniacz operacyjny IC3a są używane do sterowania napięciem wyjściowym. Regulator wstępny ustawia napięcie wyższe o 3,6 V odżądanego napięcia wyjściowego dzięki działaniu wzmacnia-



cza różnicowego zbudowanego wokół IC3b, który steruje stykiem sprzężenia zwrotnego REG1. Monitor IC5 na boczniku rezystancyjnym podaje napięcie proporcjonalne do prądu wyjściowego do przetwornika ADC IC1, który monitoruje również napięcie wejściowe i wyjściowe oraz temperaturę radiatora za pomocą termistora NTC TH1 o rezystancji 10 kΩ

Ustawienie napięcia regulatora wstępnego na 3,6 V powyżej napięcia wyjściowego zapewnia kilkaset miliwoltów nadwyżki dla LM317 przy pełnym obciążeniu, gwarantując stabilność przy jednoczesnym ograniczeniu strat ciepła.

Ponieważ częstotliwość przełączania wynosi 260 kHz, kondensatory wyjściowe o małej wartości pojemności dla stopnia regulatora wstępnego odpowiednio ograniczają tętnienia; jednak kondensator elektrolityczny 47 μ F musi być typu low-ESR. Szum RF jest redukowany przez równoległe dodanie kondensatora ceramicznego 10 μ F typu MLC, który musi zawierać dielektryk X7R lub X5R, aby zapewnić dobrą odpowiedź przy wysokich częstotliwościach.

Płaszczyzna uziemienia dla wstępnego regulatora przełączającego jest oddzielona szczyliną od reszty obwodu, stykając się tylko we wspólnym punkcie uziemienia. L2 to dławik toroidalny, aby zminimalizować promieniowanie RF, ponieważ pole magnetyczne wstępnego regulatora przełączającego jest w większości generowane w urządzeniu.

Wprawni Czytelnicy zauważą, że liniowy stopień wyjściowy wykazuje uderzające podobieństwo do liniowego zasilacza laboratoryjnego 45 V/8 A Tima Blythmana (październik-grudzień 2019; siliconchip.com.au/Series/339). Główna różnica polega na tym, że napięcie wyjściowe jest sterowane

komputerowo za pomocą cyfrowego potencjometru 5 k Ω (IC2) i przetwornika cyfrowo-analogowego (IC4), wykorzystując wartości zmierzone przez 4-kanalowy, 16-bitowy przetwornik ADC (IC1).

Pozwala to na znaczną elastyczność oprogramowania w zakresie ograniczania prądu, ochrony obwodu, zdalnego sterowania, a nawet pozwala kilku oddzielnym jednostkom działać za pośrednictwem połączeń Wi-Fi jako pojedyncza jednostka.

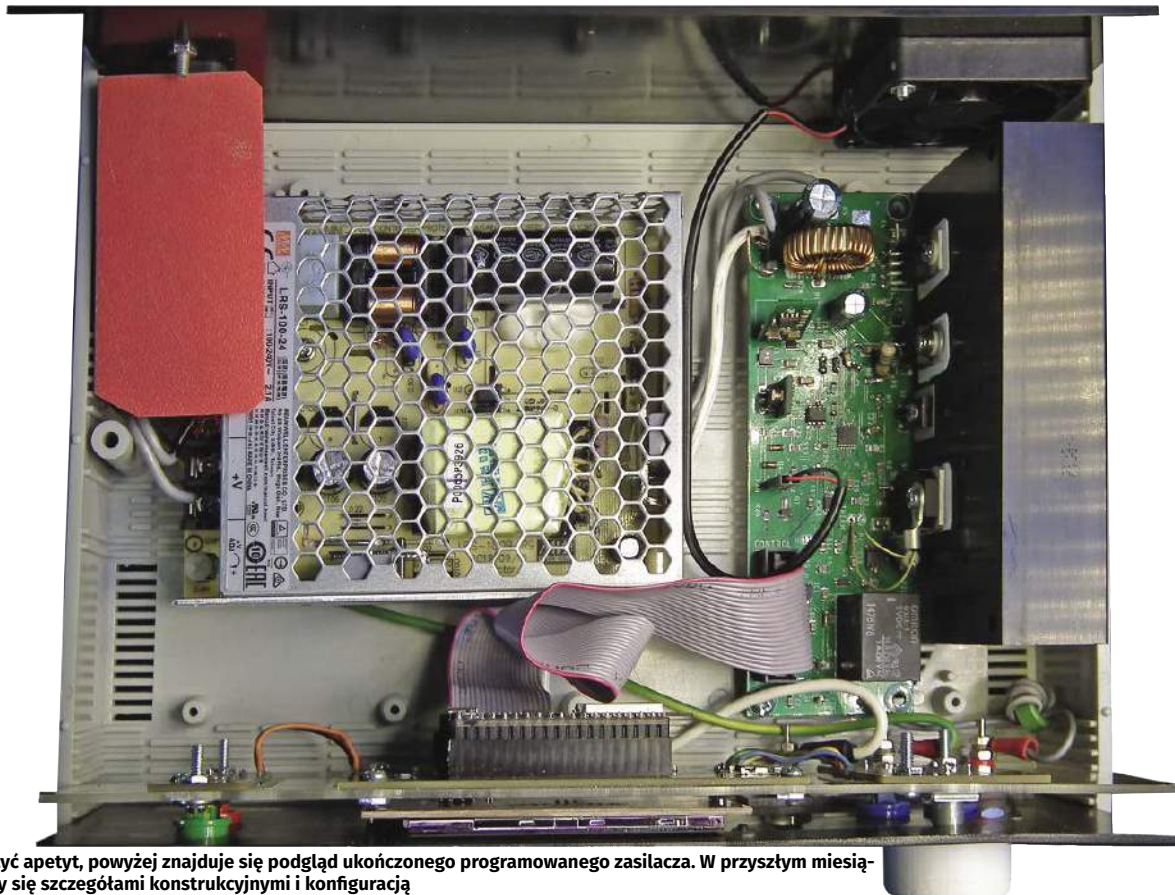
Zgrubne napięcie wyjściowe LM317 jest ustawiane przez stosunek rezystora 220 Ω między jego stykiem wyjściowym i stykiem ADJ oraz rezystancji potencjometru cyfrowego IC2. Napięcie wyjściowe ustabilizuje się, gdy napięcie między wyjściem LM317 (OUT) a stykiem regulacji (ADJ) wyniesie 1,25 V. Maksymalna rezystancja potencjometru cyfrowego wynosi 5 k Ω , zapewniając maksymalne napięcie wyjściowe 30 V.

Rozdzielczość potencjometru cyfrowego wynosi osiem bitów, co zapewnia kroki regulacji wynoszące około 120 mV. Nie jest to wystarczająco precyzyjna regulacja dla naszych celów, więc 12-bitowy przetwornik cyfrowo-analogowy i wzmacniacz operacyjny IC3a zapewniają podwójną funkcję precyzyjnej regulacji i zapewniają ujemne przesunięcie dla dolnego zakresu potencjometru cyfrowego, dzięki czemu napięcie na wyjściu LM317 może spaść do 0 V.

Napięcie na wejściu odwracającym IC3a ma wartość 0,7 V, ustawioną przez diodę D4. Przy wzmacnieniu wzmacniacza operacyjnego ustawionym na -3,9, przekłada się to na około -2,8 V na jego wyjściu. Przetwornik cyfrowo-analogowy dostarcza napięcie wyjściowe w zakresie 0...3,3 V, które jest dzielone przez rezystory 68 k Ω i 1 k Ω , co daje około 47,8 mV w pełnej skali i 186,5 mV po wzmacnieniu.

Przy przetworniku cyfrowo-analogowym ustawionym w punkcie środkowym, wzmacniacz operacyjny IC2a dostarcza około -2,35 V, co jest napięciem wymaganym do obniżenia napięcia na wyjściu LM317 do zera. Ujemne napięcie większe niż -1,25 V jest potrzebne, ponieważ cyfrowy potencjometr ma skończoną minimalną rezystancję (rezystancja ślizgacza) wynoszącą około 200 Ω . Każdy z 4096 kroków przetwornika DAC odpowiada zmianie napięcia wyjściowego o 45,5 μ V – to więcej niż wystarczająca rozdzielczość.

Po ustawieniu nowego napięcia wyjściowego, oprogramowanie oblicza najbardziej prawdopodobne ustawienie dla potencjometru i przetwornika DAC na jeden z dwóch sposobów. Jeśli zmiana jest niewielka, musi zostać zmieniona tylko wartość przetwornika DAC, aby uwzględnić różnicę. Początkowy skok jest nieco zachowawczy (delikatny), aby uniknąć przeregulowania, a ostateczne ustawienie jest



Aby zaostrić apetyt, powyżej znajduje się podgląd ukończonego programowanego zasilacza. W przyszłym miesiącu zajmiemy się szczegółami konstrukcyjnymi i konfiguracją

osiągane w ciągu 4...5 cykli poprzez powtórzenie procesu.

Jeśli zmiana jest duża, obliczane i ustawiane jest prawidłowe ustawienie dla potencjometru cyfrowego, przetwornik cyfrowo-analogowy jest ustawiany na wartość środkową i wywoływany jest algorytm precyzyjnego sterowania. Ponieważ każda iteracja sterowania zajmuje tylko 4 ms, czas ustalania jest rzędu 20 ms. Kondensator 100 nF od styku ADJ układu REG2 do masy poprawia regulację poprzez stabilizację napięcia na tym kontakcie, bez zwiększania czasu reakcji.

Zakres regulacji przetwornika cyfrowo-analogowego jest celowo ustawiony na około cztery przyrosty potencjometru cyfrowego, aby uniknąć wywołania mechanizmu regulacji zgrubnej przy małych zmianach napięcia i wynikających z tego zakłóceń napięcia wyjściowego.

Ograniczenie prądu jest realizowane w podobny sposób, wykorzystując do sterowania potencjometrem cyfrowym i ustawieniami

przetwornika cyfrowo-analogowego stosunek żądanego i rzeczywistego prądu wyjściowego.

Chociaż na panelu sterowania można wyłączyć ograniczanie prądu, oprogramowanie nadal monitoruje prąd wyjściowy, aby zapewnić ochronę przed przeciążeniem i zwarcie oraz utrzymuje zasilacz w bezpiecznym obszarze roboczym (SOA).

Prąd wyjściowy REG2 jest wzmacniany przez tranzystory Q1/Q2 działające jako para w układzie Sziklai'ego. Gdy prąd przepływający przez LM317 przekracza 100 mA, napięcie na rezystorze 68 Ω wzrasta powyżej 0,7 V, powodując przewodzenie Q1 i włączenie Q2, który przepuszcza większość prądu wyjściowego.

Kombinacja Q1 i Q2 ma potencjalne wzmocnienie prądowe ponad 10 000, więc należy zwrócić szczególną uwagę, aby zapewnić stabilność. Kondensator 1 µF zapewnia sprzężenie zwrotne AC do bazy Q1, a rezystor bazowy Q1 o wartości 1,5 kΩ jest tak dobrany, aby maksymalny prąd przepływający przez Q2 wynosił nieco ponad

5 A. Rezystor bazowy 22 Ω dla Q2 zapewnia, że prąd płynący przez Q1 jest ograniczony do kilkuset miliamperów.

Kondensator wyjściowy 10 µF jest typem wybranym ze względu na skuteczność przy wysokich częstotliwościach, redukując szumy RF. Dławik toroidalny L3 (nie pokazany na rysunku 4) dodatkowo obniża szumy HF.

Napięcia wejściowe i wyjściowe, prąd wyjściowy i temperatura radiatora są monitorowane przez 16-bitowy przetwornik analogowo-cyfrowy (ADC) ADS1115. Każdy sygnał wejściowy jest przetwarzany w zakresie, który przetwornik może obsłużyć, czyli 0...2,048 V.

Proste dzielniki napięcia są odpowiednie do dostosowania wartości napięcia i temperatury do zakresu przetwornika ADC. Jednak odczyty prądu przy braku obciążenia okazały się niewiarygodne, więc wyjście czujnika prądu INA282 jest kompensowane przez diodę Schottky'ego D5 w celu polaryzacji wejścia REF1 na styku 7 (dioda Schottky'ego ma około połowy

Wykaz elementów:

Lista części – programowany hybrydowy zasilacz laboratoryjny

- 1 obudowa przyrządu z ABS, 260×190×80 mm [Altronics H0482, Jaycar HB5910, Pro'skit 203-115B].
- 1 etykieta na panel przedni
- 1 przetwornika AC-DC MeanWell LRS-100-24 [Mouser, RS] lub analogiczna
- 1 moduł regulatora (patrz poniżej)
- 1 moduł panelu sterowania (patrz poniżej)
- 1 gniazdo zasilania IEC [Jaycar PP4005]
- 1 czerwony przewód zasilający
- 1 czarny przewód zasilający
- 1 zielony zacisk
- 1 wentylator niskoprądowy 40...60 mm 5 V DC [np. Altronics F1110]
- 16 śrub M3×15 z łbem walcowym i nakrętkami sześciokątnymi (do wentylatora, radiatora i panelu przedniego)
- 2 śruby M3×15 z łbem stożkowym i nakrętki sześciokątne (do złącza IEC)
- 3 śruby M3×25 z łbem stożkowym (do zasilacza MeanWell i radiatora)
- 3 wkręty samogwintujące 4G×8 (do PCB i konwertera AC-DC)
- 1 tuleja dystansowa M3×6 (do montażu zasilacza MeanWell)
- 1 przewód zasilający IEC z 3-stykową wtyczką
- 1 20-żyłowy kabel taśmowy 10cm+ z gniazdem Z-FC20
- 1 przewód zasilający o długości 1 m
- 1 odcinek przewodu zasilającego 5 A o długości 1 m
- 1 rurka termokurczliwa o długości 50 mm i średnicy 6 mm (do połączeń sieciowych)
- 3 zaciskane końcówki oczkowe o średnicy 3 mm ID (opcjonalnie)
- 3 podkładki izolacyjne TO-220 (mika lub guma silikonowa)
- 1 podkładka izolacyjna TO-3P (mika lub guma silikonowa)
- 1 mała tubka pasty termoprzewodzącej (wymagana tylko w przypadku stosowania podkładek izolacyjnych z miki)
- 1 toroid ferrytowy o średnicy 15 mm (lub większej) [Jaycar L01242]
- 1 wtyk 403-2 + 1 gniazdo 402-2 z kotkami (dla wentylatora)
- 1 osłona gniazda sieciowego

Lista części – moduł regulatora

- 1 dwustronna płyta drukowana o kodzie 18104212 i wymiarach 136×44,5 mm
- 1 wtyk IDC Z-WS20 prosty (CON1)
- 1 wtyk 403-2 + 1 gniazdo 402-2 z kotkami (dla wentylatora) (CON3)
- 1 cewka indukcyjna SMD 10 µH 1 A, 4×4 mm (L1) [np. Taiyo Yuden NRS4012T100MDG]
- 1 cewka toroidalna 47 µH 5 A (L2) [Altronics L6617]
- 1 przełącznik SPDT G5LE z cewką 5 V DC 10 A [np. Omron G5LE-1-DC5]
- 1 mały radiator [CINCON M-B012 lub wycięty i wygięty z blachy aluminiowej 1,6 mm]
- 1 termistor NTC 10 kΩ, montaż oczkowy, z wyprowadzeniami [Altronics R4112]

Półprzewodniki:

- 1 ADS1115DGSR ADC, MSOP-10 (IC1)
- 1 8-bitowy potencjometr cyfrowy MCP45HV51-502 5 kΩ I²C, TSSOP-14 (IC2)
- 1 podwójny wzmacniacz operacyjny z pojedynczym zasilaniem LM358D, SOIC-8 (IC3)
- 1 12-bitowy przetwornik cyfrowo-analogowy MCP4725AOT-E/CH, SOT-23-6 (IC4)
- 1 dwukierunkowy czujnik prądu INA282AIDR, SOIC-8 (IC5)
- 1 podwójna 2-wejściowa bramka NOR SN74LVC2G02DCTR, SSOP-8 (IC6; raster 0,65 mm)
- 1 regulator impulsowy LM2679T-ADJ, TO-220-7 (REG1)
- 1 regulator liniowy LM317, TO-220-3 (REG2)
- 1 moduł regulatora przełączającego CUI VX07805-1000 5 V 1 A, TO-220-3 (REG3)
- 1 moduł regulatora przełączającego CUI VX07805-500 5V 500 mA, TO-220-3 (REG4)
- 1 tranzystor PNP BD140 80 V 1,5 A, TO-126 (Q1)
- 1 tranzystor mocy NPN FJA4313 250 V 17 A, TO-3P (Q2)
- 2 tranzystory NPN BC817 lub odpowiednik 45 V, 500 mA, SOT-23 (Q3, Q4)

- 1 dioda LED SMD 0805 (LED1)
- 3 diody Schottky'ego V2F22HM3_H 1 A 20 V, DO219-AB-2 (D1, D2, D5)
- 1 dioda Schottky'ego STPS1045SF 15 A 60 V, TO-227A (D3)
- 3 diody niskosygnałowe BAS21 lub równoważne, SOD-123 (D4, D6, D7)
- 1 dioda Zenera BZV55 3,6 V, SOD-323/mini-MELF (ZD1)

Kondensatory: (wszystkie SMD 1210, o ile nie podano inaczej)

- 1 kondensator elektrolityczny low-ESR 270 µF 50 V (raster wyprowadzeń 3,5 mm, maksymalna średnica 8 mm)
- 1 kondensator elektrolityczny low-ESR 47 µF 50 V (raster wyprowadzeń 3,5 mm, maksymalna średnica 8 mm)
- 2 kondensatory ceramiczne 10 µF 50 V X7R SMD 1210
- 3 kondensatory ceramiczne 10 µF 35 V X7R SMD 1206
- 2 kondensatory ceramiczne 1 µF 50 V X7R SMD 0805
- 13 kondensatorów ceramicznych 100 nF 50 V X7R SMD 0805
- 1 kondensator ceramiczny 10 nF 50 V X7R SMD 0805
- 1 kondensator ceramiczny 1 nF 50 V NP0/C0G SMD 0805

Rezystory: (wszystkie 1% SMD 0805, o ile nie określono inaczej)

- 1 szt. 820 kΩ 2 szt. 100 kΩ 3 szt. 68 kΩ 1 szt. 39 kΩ 3 szt. 10 kΩ
- 1 szt. 6,8 kΩ 1 szt. 5,6 kΩ 3 szt. 4,7 kΩ 1 szt. 3,3 kΩ 1 szt. 1,5 kΩ
- 4 szt. 1 kΩ 4 szt. 470 Ω 1 szt. 220 Ω 1 szt. 150 Ω
- 1 szt. 68 Ω 1/2 W 1% montaż przewlekany osiowy
- 1 szt. 22 Ω 1/2 W 1% SMD 1206
- 1 szt. 10 mΩ 1 W 1% drutowy montaż przewlekany osiowy

Lista części – moduł panelu sterowania

- 1 dwustronna płyta PCB o kodzie 18104211 i wymiarach 167,5×56 mm
- 1 moduł Wi-Fi MCU WROOM-32 kompatybilny z Espressif ESP32-DEVKITC [Altronics Z6385A, Jaycar XC3800, NodeMCU-32S].
- 1 ekran dotykowy SPI LCD 2,8 cala ze sterownikiem ILI9341 [np. Silicon Chip Cat SC3410]
- 1 gniazdo zasilacza wtyczkowego DC 5,5/2,1 mm do montażu na płytce drukowanej (CON1; opcjonalnie) [Altronics P0620, Jaycar PS0519]
- 1 wtyk IDC Z-WS20 prosty (CON2) [WURTH 61202021621 lub podobne]
- 1 40-stykowa listwa żeńska (przycięta na dwie listwy po 19 styków) do modułu ESP-32
- 1 gniazdo karty SMD microSD (opcjonalnie) [Hirose DM3D-SF]
- 1 enkoder obrotowy (RE1) [Alps EC12E, np. Jaycar Cat SR1230]
- 1 pokrętko do enkodera obrotowego [np. Altronics H6514 (23 mm) lub Adafruit 2055 (35 mm)]
- 4 przełączniki dotykowe SPST 12 mm do montażu na płytce drukowanej z kwadratowymi przyciskami (S1-S4) [Altronics S1135, Jaycar SP0608].
- 2 czarne, białe lub szare nakładki na przełączniki [Altronics S1138]
- 1 czerwona nasadka przełącznika
- 1 zielona nasadka przełącznika

Półprzewodniki:

- 1 regulator liniowy 7805T 5 V/1 A (REG1; opcjonalnie)
- 1 czerwona lub zielona dioda LED 5 mm (LED1)

Kondensatory:

- 1 kondensator ceramiczny 47 µF 10 V X5R/X7R SMD 1210
- 1 kondensator ceramiczny 10 µF 25 V X5R/X7R SMD 1210
- 9 kondensatorów ceramicznych 100 nF 50 V X7R SMD 0805

Rezystory: (wszystkie 1% 1/10 W SMD 0805)

- 3 szt. 10 kΩ 2 szt. 1,8 kΩ 1 szt. 1 kΩ

spadku napięcia diody krzemowej), a następnie dzielone przez parę rezystorów 4,7 k Ω /3,3 k Ω .

Przy boczniku prądowym 0,01 Ω (10 m Ω) i wzmacnieniu 50 V/V, odpowiada to napięciu 2,5 V na wyjściu INA282 przy prądzie wyjściowym 5 A i 0,35...1,38 V na wejściu przetwornika ADC. Równoważne jest to rozdzielczości 150 μ A.

Temperatura Q2 jest mierzona przez termistorowy dzielnik napięcia, a linearyzacją zajmuje się oprogramowanie. Q4 włącza wentylator, gdy Q2 osiągnie temperaturę 35°C. Wentylator jest mały i cichy, więc wystarczające jest proste sterowanie włączaniem/wyłączaniem.

Wyjście jest przełączane przełącznikiem, sterowanym przez zatrask zbudowany z bramek logicznych (IC6a i IC6b) i tranzystora NPN Q3. Q3 steruje również wskaźnikiem LED1. Układ IC6 zapewnia, że wyjście jest zawsze wyłączone podczas rozruchu, niezależnie od stanu mikroprocesora.

Podwójna bramka NOR 74C02 jest skonfigurowana jako zatrask SR, z kondensatorem 100 nF zapewniającym krótki dodatni impuls po podłączeniu zasilania, resetującym ten zatrask.

Układ IC6 jest bezpośrednio sterowany przez przełączniki włączania/wyłączania na płycie sterującej, a także przez mikroprocesor, zapewniając, że naciśnięcie przycisku wyłączenia zawsze spowoduje natychmiastowe wyłączenie wyjścia, nawet jeśli mikroprocesor jest zajęty innymi zadaniami.

Zasilanie pomocnicze ± 5 V zapewnia zasilanie dla układu logicznego i wzmacniacza operacyjnego, a także płyty sterownika. Obie te szyny są zasilane przez 3-stykowe moduły przetwornic DC-DC, które mają taki sam układ wyprowadzeń jak standardowe regulatory liniowe.

Opublikowaliśmy podobne projekty w wydaniu SC z sierpnia 2020 r. (siliconchip.com.au/Article/14533), ale ich maksymalne napięcie wejściowe wynoszące 30 V jest tutaj po prostu niewystarczające. Dlatego wybraliśmy moduły dostępne na rynku, które mają wyższe wartości znamionowe.

Element 500 mA wybrany dla regulatora -5 V (VR4) ma maksymalne napięcie wejściowe 31 V dla ujemnych konfiguracji wyjściowych. Nie można go zastąpić wersją 1 A używaną dla regulatora +5 V (VR3), który może obsługiwać tylko 27 V w tym trybie.

Płyta regulatora łączy się z płytą sterowania/wyświetlacza poprzez CON1, 20-stykowe złącze i pasujący kabel taśmowy z gniazdami Z-FC20 na obu końcach. Zasilanie 3,3 V dla układów IC1, IC2 i IC4 pochodzi z regulatora na płycie sterującej i doprowadzone jest

za pośrednictwem złącza CON1. Zasilanie płytki sterującej jest dostarczane z szyny 5 V na tej płycie, przez styki 18 i 20 CON1.

Obwód sterowania

Schemat ideowy obwodu płytki sterującej pokazano na **rysunku 5**, z kablem taśmowym od CON1 na płytce regulatora zakończonym pasującym złączem IDC CON2.

Dwa główne komponenty na tej płycie to mikroprocesor ESP-32 i moduł Wi-Fi oraz 2,8-calowy lub 3,5-calowy ekran dotykowy. Są one połączone w zwykły sposób magistralą SPI i kilkoma cyfrowymi liniami sterującymi, umożliwiając mikroprocesorowi aktualizację zawartości ekranu i wykrywanie zdarzeń dotykowych.

Dostępne jest również opcjonalne gniazdo karty SD współdzielące tę samą magistralę SPI, choć w tym projekcie jest ono zbędne. Jest ono dostarczane, ponieważ płytka sterująca może być używana do innych celów, w których posiadanie wbudowanej pamięci może być przydatne.

Połączenia między ESP-32 i CON2 obejmują współdzieloną magistralę SPI, dwie magistrale I²C, magistralę szeregową oraz kilka cyfrowych styków We/Wy. Należy zauważyć, że wiele z nich nie jest podłączonych na drugim końcu i są one przewidziane do przyszłej rozbudowy.

Wykorzystywane funkcje to pierwsza magistrala I²C (SDA/SCL), do sterowania przetwornikiem ADC (IC1), cyfrowym potencjometrem (IC2) i przetwornikiem DAC (IC4) oraz czterema cyfrowymi liniami I/O. Są to styk 9, który jest sygnałem przerwania DRDY z przetwornika ADC, który wskazuje, że konwersja została zakończona, linie wykrywania włączania i wyłączenia na końcówkach 12 i 16 oraz linia sterowania wentylatorem na styku 14.

Moduł może być zasilany przez USB, na przykład, do celów związanych z rozwojem i uruchamianiem. Do projektów, w których wymagane jest zewnętrzne zasilanie, dołączono gniazdo zasilacza wtyczkowego i regulator 5 V. Zasilanie 5 V może być również dostarczane za pośrednictwem 20-stykowego złącza rozszerzeń (CON2), co jest rozwiązaniem zastosowanym w tym projekcie. Moduł ESP zużywa 225 mA przy dostarczaniu pełnej mocy wyjściowej Wi-Fi.

Moduł może dostarczyć do 50 mA przy napięciu 3,3 V dla dodatkowej logiki z wbudowanego regulatora ESP-32 i jak wspomniano wcześniej, jest to wykorzystywane przez płytke regulatora.

Przełączniki S1 i S2 mają rezystory polaryzujące do masy, kondensatory eliminujące drgania zestyków i są skonfigurowane jako aktywne w stanie wysokim. Podczas gdy wspomniane

funkcje mogą być zapewnione przez konfigurację portu i oprogramowanie, dodanie ich metodą sprzętową powoduje minimalny wzrost komplikacji i kosztów.

SW_ON i SW_OFF w tym projekcie przełączają wyjście zasilania. Oprócz prowadzenia do styków GPIO, są one również podłączone na stałe do złącza rozszerzeń.

Układ jest nieco nietypowy, ponieważ SW_ON jest wejściem, gdy wyjście zasilacza jest wyłączone, ale staje się wyjściem (o wysokim poziomie) po kliknięciu. Jest on ponownie konfigurowany, aby stał się wejściem, po naciśnięciu SW_OFF. Zapewnia to, że dioda LED1 pozostaje zapalona po zwolnieniu SW_ON.

SW_L i SW_R współpracują z enkoderem obrotowym, umożliwiając łatwe ustawianie wartości numerycznych. Enkoder obrotowy zmienia wartość o jedną „jednostkę” w górę (obróć w prawo) lub w dół (obróć w lewo) na kliknięcie. SW_L i SW_R wybierają wielkość tej jednostki, która jest również podświetlona na ekranie.

SW_L przesuwa cyfrę wybraną przez enkoder obrotowy do następnej cyfry po lewej stronie. Powoduje to dziesięciokrotne zwiększenie wartości dodawanej lub odejmowanej dla każdego kliknięcia enkodera. SW_R ma odwrotny efekt. Ten układ jest powszechny w instrumentach cyfrowych, ponieważ umożliwia szybką i dokładną regulację wartości i jest łatwy do opanowania.

Enkoder obrotowy i jego przełącznik są aktywne w stanie niskim. Mikroprocesor zapewnia odpowiednią polaryzację enkodera. Przełącznik przyciskowy enkodera w tym projekcie nie jest używany. W razie potrzeby można go podłączyć do IO26 w module ESP-32 za pośrednictwem zworki JP3.

Obecne oprogramowanie nie korzysta z przerwania ekranu dotykowego; można je jednak podłączyć do IO2 za pośrednictwem zworki JP1. Należy zachować ostrożność podczas używania IO2 do innych celów, ponieważ jego stan po włączeniu zasilania (wraz z IO0) określa sposób uruchamiania ESP-32.

W przyszłym miesiącu

W naszym kolejnym wydaniu EdW Czytelnicy będą mieli podane pełne szczegóły konstrukcyjne programowanego hybrydowego zasilacza laboratoryjnego oraz więcej informacji na temat jego konfiguracji i użytkowania.

Richard Palmer

Adaptacja do wydania
polskiego – Andrzej Nowicki

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

Regulator prędkości obrotowej silników DC 30 V/20 A

Ten niewielki, ale wydajny regulator prędkości silnika DC ma prąd znamionowy 20 A i jest wyposażony w wiele funkcji. Nadaje się do szerokiego zakresu zastosowań oraz jest prosty w budowie i obudwie. Funkcje obejmują ochronę przed niskim poziomem naładowania akumulatora, łagodny rozruch i regulowaną częstotliwość impulsów. Może obsługiwać silniki prądu stałego o napięciu od bliskiego 0 V do 30 V.

Istnieje wiele zastosowań silników prądu stałego, w których regulacja prędkości jest pożądana lub konieczna. Ponieważ silniki DC mogą być zasilane bezpośrednio z akumulatorów, są one używane w wózkach golfowych, elektrycznych skuterach, rowerach, hulajnogach i deskorolkach, zdalnie sterowanych samochodach i łodziach – lista jest długa.

W większości tych zastosowań potrzebny jest sposób regulacji prędkości obrotowej silnika. Jazda na maksymalnej prędkości przez cały czas nie zawsze jest dobrym pomysłem!

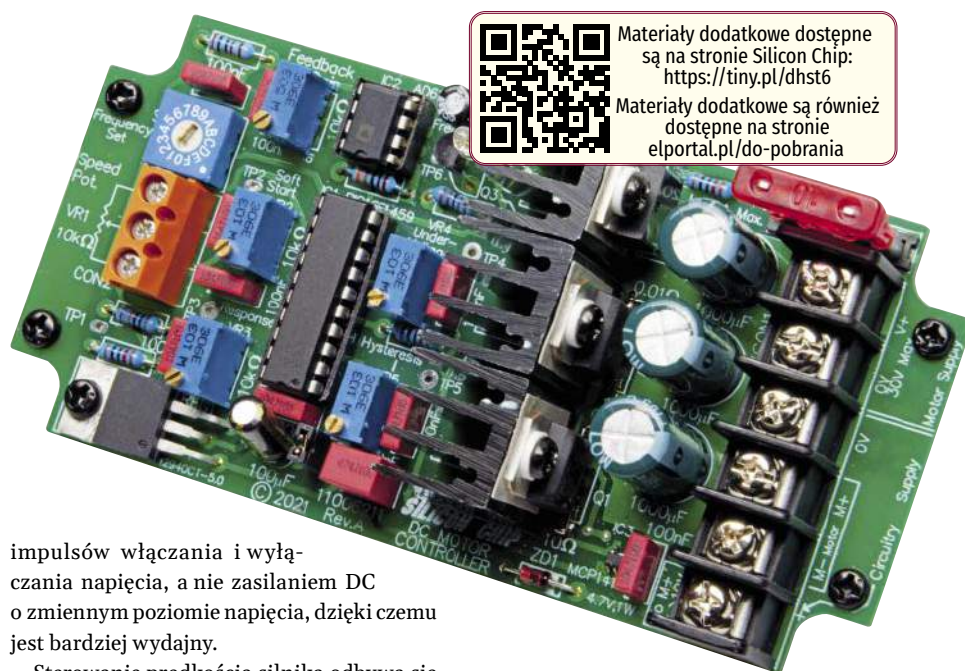
Regulator prędkości, taki jak ten zaprezentowany, jest idealnym rozwiązaniem. Może on obsługiwać silniki prądu stałego o napięciu znamionowym do 24 V (maksymalnie 30 V) i prądzie o ciągłym natężeniu do 20 A.

Regulator jest prezentowany jako niesłonięty moduł elektroniczny zbudowany na płytce drukowanej, który w razie potrzeby można umieścić w standardowej plastikowej obudowie UB3. Zawiera on wytrzymałe zaciski dla zasilania i połączeń silnika, a także dodatkowe zaciski dla potencjometru regulacji prędkości, który jest montowany poza płytką drukowaną.

W celu chłodzenia komponenty zasilające silnik są zamontowane na sporych radiatorach. Regulowane funkcje, takie jak czas trwania miękkiego startu i wzmocnienie sprzężenia zwrotnego, są ustawiane za pomocą wbudowanych wieloobrotowych potencjometrów nastawnych, z punktami testowymi do pomiaru napięcia. Wbudowana dioda LED wskazuje ustawienie prędkości, a także błędy, takie jak niski poziom naładowania akumulatora lub odłączenie silnika.

Konstrukcja regulatora prędkości

Chociaż w przeszłości publikowaliśmy opisy wielu regulatorów prędkości silnika DC, ta wersja ma więcej funkcji i lepszą wydajność. Prędkość silnika jest sterowana za pomocą modulacji szerokości impulsu (PWM). Oznacza to, że silnik jest napędzany serią



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/dhst6>
Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania

impulsów włączania i wyłączenia napięcia, a nie zasilaniem DC o zmiennym poziomie napięcia, dzięki czemu jest bardziej wydajny.

Sterowanie prędkością silnika odbywa się poprzez zmianę szerokości impulsu. Stosunek szerokości impulsu do odstępu między impulsami to cykl pracy. Niskie wypełnienie cyklu pracy będzie dostarczać napięcie do silnika tylko przez niewielką część czasu, a silnik będzie pracował powoli.

Wraz ze wzrostem czasu trwania impulsu, cykl pracy o większym wypełnieniu sprawia, że silnik pracuje szybciej, aż osiągnie 100% wypełnienia cyklu i będzie zasilany w sposób ciągły.

Przebiegi oscyloskopowe: **oscylogram 1 i oscylogram 2** pokazują, jak działa prezentowany układ zasilania w bazujący na metodzie PWM.

Na oscylogramie 1, górny (żółty) przebieg jest sygnałem sterującym dla MOSFET-ów Q1 i Q2. Gdy jest on wysoki, silnik jest zasilany. W tym przypadku wypełnienie cyklu pracy jest bardzo niskie i wynosi około 9,5%, więc silnik pracuje powoli. Dolny cyjanowy przebieg jest związany z prądem pobieranym przez silnik. Wartość tego prądu służy do utrzymywania stałej prędkości silnika przy zmiennym obciążeniu.

Oscylogram 2 ma te same dwa przebiegi, ale tym razem wypełnienie cyklu pracy jest znacznie większe, a silnik obraca się szybciej. Silnik jest mniej obciążony niż na oscylogramie 1, więc odczyt prądu jest niższy pomimo wyższego wypełnienia cyklu pracy.

Co nowego w układzie

Jednym z problemów związanych ze sterowaniem silnikami prądu stałego za pomocą PWM jest to, że silnik może generować dodatkowy hałas z powodu wibracji uzwojeń silnika i innych części mechanicznych przy częstotliwości PWM. Można to do pewnego stopnia złagodzić, dobierając częstotliwość PWM tak, aby generowała minimalny hałas.

Hałas ten ma tendencję do zmniejszania się wraz ze wzrostem częstotliwości PWM i jest w większości eliminowany przy częstotliwościach PWM powyżej 20 kHz (przy górnej granicy ludzkiego słuchu).

Zwiększanie częstotliwości może jednak powodować problemy. Utrzymanie prędkości silnika przy zmiennym obciążeniu staje się

trudniejsze przy użyciu tradycyjnego systemu sprzężenia zwrotnego. Bardzo wysokie częstotliwości PWM mogą również powodować utratę momentu obrotowego silnika.

Problemy te i ich rozwiązania opisano bardziej szczegółowo w oddzielnej sekcji zatytułowanej „Pułapki sterowania PWM silnikiem przy wyższych częstotliwościach”.

Sterownik ten umożliwia regulację częstotliwości PWM poza zakresem słyszalności, jednocześnie rozwiązując problemy związane z ograniczonym momentem obrotowym silnika przy niskich prędkościach i sterowaniem przy wyższych częstotliwościach.

Inne wbudowane funkcje obejmują łagodny rozruch, odcięcie zasilania przy niskim napięciu akumulatora, wskaźnik stanu LED i opcjonalne wykrywanie odłączenia silnika. Funkcje te można łatwo skonfigurować i ustawić za pomocą potencjometrów.

Łagodny rozruch

Polega on na powolnym zwiększaniu prędkości obrotowej silnika, aż do ustawienia prędkości na maksimum. Łagodny rozruch zmniejsza skok natężenia prądu i szybki wzrost momentu obrotowego silnika w porównaniu z nagłym włączeniem zasilania. Wypełnienie cyklu pracy PWM jest zwiększane przez dłuższy czas, dzięki czemu silnik uruchamia się płynnie.

Maksymalny okres łagodnego rozruchu wynosi dwie sekundy dla pełnego zakresu sterowania od 0% do 100%. Okres ten można regulować w zakresie od zera do dwóch sekund w 255 krokach.

Łagodny rozruch można zainicjować na kilka sposobów. Ma on zastosowanie, gdy regulator jest początkowo zasilany, lub gdy

Funkcje:

- Regulator PWM silnika prądu stałego
- Może zasilac silniki o napięciu znamionowym do 24 V i 20 A DC
- Napięcie zasilania silnika i regulatora może być oddzielone
- Możliwość wyboru 16 częstotliwości PWM
- Regulacja sprzężenia zwrotnego obciążenia silnika i regulacja wzmocnienia
- Regulowana szybkość łagodnego rozruchu
- Regulacja krzywej prędkości silnika
- Odcięcie zasilania przy zbyt niskim napięciu ze wskaźnikiem LED i regulowaną histerezą
- Wskaźnik LED cyklu pracy
- Opcjonalne wykrywanie odłączenia silnika

Specyfikacje:

- Zakres regulacji prędkości: 0% do 100% wypełnienia cyklu pracy
- Napięcie zasilania silnika: od niemal zerowego do maksymalnie 30 V
- Zasilanie regulatora: 10,5 V do maksymalnie 30 V (5,5...26 V z odłączoną, zwartą diodą ZD1)
- Wskazanie prędkości: jasność LED1 zmienia się wraz z wypełnieniem cyklu pracy PWM
- Częstotliwość PWM: 16 kroków od 30,6 Hz do 32,4 kHz (patrz tabela 1)
- Zakres łagodnego startu: 0...2 sekund w 255 krokach dla wypełnienia cyklu pracy od 0% do 100%
- Regulacja krzywej prędkości: minimalna prędkość może być ustawiona na 0...33% wypełnienia cyklu pracy
- Próg odłączenia przy niskim napięciu zasilania (under voltage – UV): 0...30 V w krokach co 29,6 mV
- Histeresa UV: 0...5 V w krokach co 29,6 mV
- Wskazanie UV: dioda LED1 miga przez 65 ms z częstotliwością 1 Hz
- Wykrywanie odłączenia silnika: silnik jest wyłączany, jeśli monitorowany prąd spadnie do zera podczas zasilania silnika; sygnalizowane miganiem diody LED 2 Hz/50% wypełnienia cyklu świecenia
- Wykrywanie odłączenia potencjometru prędkości: sygnalizowane słabo świecącą diodą LED

regulacja prędkości jest uruchamiana z pozycji całkowitego wyłączenia, a także po powrocie do normalnej pracy po wyłączeniu niskiego napięcia.

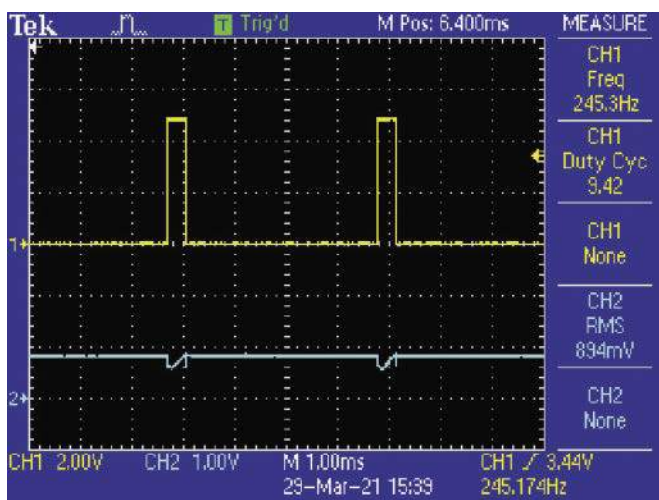
Wykrywanie niskiego napięcia

Funkcja wykrywania niskiego napięcia zapobiega nadmiernemu rozładowaniu akumulatora zasilającego silnik. Większość

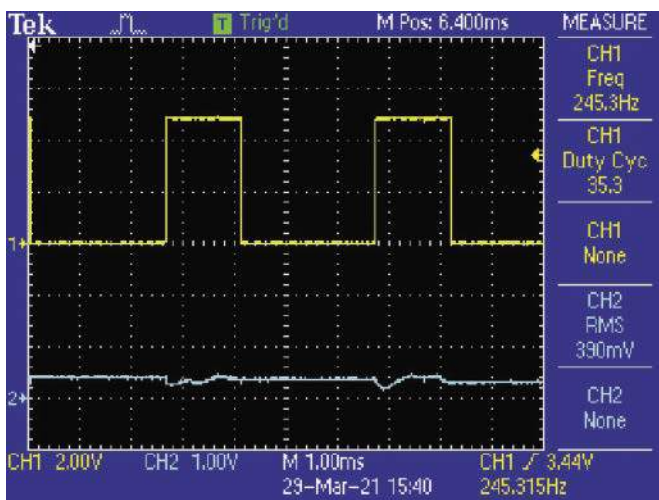
akumulatorów, w tym ołowiowo-kwasowe i litowe, ulega uszkodzeniu po rozładowaniu poniżej określonego napięcia.

Ta funkcja wyłącza napęd silnika przy wstępnie ustawionym napięciu progowym. Jest to sygnalizowane miganiem wskaźnika LED.

Napięcie musi znajdować się poniżej progu przez ponad dziesięć sekund, zanim napęd silnika zostanie wyłączony. Zapobiega



Oscylogram 1. Sygnał sterowania z modulacją szerokości impulsu (PWM) przy niskim wypełnieniu cyklu pracy, około 9,5%. Czas narastania prądu podczas każdego impulsu jest krótki, więc silnik pracuje powoli



Oscylogram 2. Kolejny sygnał PWM, tym razem z wypełnieniem cyklu pracy równym 35,5%. Jest to w przybliżeniu równoważne zasilaniu silnika przy 1/3 napięcia zasilania, więc będzie on obracał się szybciej, ale nie z pełną prędkością

to uciążliwym wyzwoleniom zadziałania układu detekcji niskiego napięcia, które w przeciwnym razie wyłączyłyby sterownik z powodu krótkotrwałego spadku napięcia podczas uruchamiania silnika.

Po wyłączeniu napięcie musi wzrosnąć o określoną wartość powyżej progu wykrywania niskiego napięcia, zanim zasilanie zostanie ponownie uruchomione. Ta histereza zapobiega ciągłemu włączaniu i wyłączaniu, gdy napięcie akumulatora wraca do normy po usunięciu obciążenia silnika, tylko po to, aby ponownie wyłączyć się po ponownym uruchomieniu silnika.

Odlączenie silnika

Opcjonalne wykrywanie odlączenia silnika zapobiega uruchomieniu silnika, jeśli zostanie on odlączony, a następnie ponownie podłączony, gdy ustawienie prędkości jest powyżej zera. Gdy nastąpi wykrycie odlączenia silnika, potencjometr prędkości musi zostać całkowicie obrócony w lewo, a silnik ponownie podłączony, zanim będzie on mógł ponownie pracować. Stan odlączenia sygnalizowany jest miganiem diody LED z częstotliwością 2 Hz.

Oddzielne napięcie zasilania

Kolejną funkcją jest możliwość oddzielenia napięcia zasilania sterownika od napięcia zasilania silnika. Oznacza to, że silnik może być zasilany znacznie niższym napięciem niż wymagane do działania regulatora prędkości silnika DC.

Tak więc, podczas gdy regulator prędkości silnika prądu stałego wymaga do działania zasilania o napięciu co najmniej 10,5 V (do 30 V), silnik może być zasilany oddzielnym napięciem od blisko 0 V do 30 V. Limit 30 V jest wystarczający, aby umożliwić korzystanie z dowolnego akumulatora 24 V; np. w pełni naładowany 12-ogniowy akumulator kwasowo-ołowiowy ma napięcie około 29 V.

Możesz użyć tego samego źródła zasilania zarówno dla regulatora, jak i silnika, pod warunkiem, że napięcie to mieści się w zakresie 10,5...30 V i jest odpowiednie dla silnika.

Szczegóły układu

Pełny schemat ideowy układu sterownika silnika DC pokazano na **rysunku 1**. Wykorzystuje on 8-bitowy mikrokontroler PIC16F1459, IC1, który zapewnia generowanie sygnału PWM i monitoruje napięcie akumulatora, prąd silnika oraz napięcie z potencjometru prędkości i z kilku potencjometrów nastawnych. IC1 monitoruje również przełącznik obrotowy S1, który wybiera częstotliwość PWM.

IC1 ma dwa wyjścia PWM i używamy obu. Jedno znajduje się na wyprowadzeniu 5 (PWM1), a drugie na wyprowadzeniu 8 (PWM2). Te wyjścia PWM mają różne funkcje,

ale zapewniają tę samą częstotliwość PWM i wypełnienie cyklu pracy przez większość czasu, gdy silnik jest napędzany.

Wyjście PWM1 jest używane do sterowania MOSFET-ami Q1 i Q2 poprzez sterownik bramek IC3. Układ IC3 to MCP1416, zaprojektowany w celu zapewnienia wysokoprądowego sterowania bramek MOSFET-ów z szybkimi czasami narastania i opadania. Zapewnia to szybkie ich włączanie i wyłączanie. Każda bramka MOSFET-a jest odizolowana od drugiej za pomocą rezystora 10 Ω . Rezystory zapobiegają również oscylacjom przełączania MOSFET-ów na progu wyzwolenia bramki.

Wykorzystano MOSFET-y wyzwolane poziomem logicznym, przez co w pełni przewodzą przy napięciu 5 V przyłożonym do bramki. Należy przypomnieć, że standardowe MOSFET-y zazwyczaj wymagają napięcia co najmniej 10 V do wyzwolenia pełnego ich przewodzenia. Dwa MOSFET-y są połączone równolegle i dzielą prąd obciążenia (silnika).

Rezystory o niskiej wartości są umieszczone między źródłem każdego MOSFET-a a masą, przy czym rezystor źródłowy Q1 jest używany do monitorowania prądu. Rezystor źródłowy Q2, choć nie jest używany do pomiaru prądu obciążenia, jest nadal niezbędny. Dzieje się tak, aby całkowita rezystancja włączenia MOSFET-a Q2 i jego rezystora źródłowego odpowiadała Q1 i jego rezystorowi źródłowemu.

Ponieważ rezystancja włączenia MOSFET-a wynosi zwykle 0,014 Ω , rezystor źródłowy 0,01 Ω dla Q2 pomaga utrzymać równomierny podział prądu obciążenia między dwoma MOSFET-ami. Bez niego Q2 przenosiłby około 2/3 prądu obciążenia, a Q1 tylko 1/3.

Dioda D1 znajduje się między dodatnim zasilaniem a drenami MOSFET-ów, aby minimalizować indukowany skok napięcia, gdy napęd silnika jest wyłączany. Dioda ta jest efektywnie podłączona do zacisków silnika. Jest to podwójna dioda Schottky'ego 10 A, która może przewodzić 20 A w sposób ciągły, gdy diody są połączone równolegle.

Równoległe połączenie diod zapewnia niemal równy podział prądu. Jest to możliwe, ponieważ obie diody znajdują się na tej samej strukturze krzemowej, a zatem mają taką samą charakterystykę i temperaturę pracy.

Zasilanie silnika jest podłączone do zacisków GND i zasilania silnika „+” na złączu śrubowym CON1. To dodatnie zasilanie jest doprowadzane do silnika za pośrednictwem bezpiecznika F1, bezpiecznika samochodowego o wartości znamionowej dobranej do silnika. Trzy kondensatory elektrolityczne 470 μ F 35 V o niskim ESR bocznikują zasilanie silnika za bezpiecznikiem. Mają one zapewnić wysoki krótkotrwały prąd szczytowy.

Regulacja sprzężenia zwrotnego

Wiele regulatorów prędkości silnika prądu stałego monitoruje wsteczną siłę elektromotoryczną (back-EMF) silnika, aby określić, kiedy zmiany obciążenia mogą zmniejszyć prędkość silnika. Ta wsteczna siła elektromotoryczna to napięcie generowane przez silnik, gdy zasilanie jest wyłączone, a silnik nadal się obraca. Napięcie indukowane zmniejsza się, gdy silnik zwalnia pod obciążeniem.

Kontrola prędkości jest utrzymywana poprzez zwiększenie wypełnienia cyklu pracy PWM w celu zwiększenia momentu obrotowego i prędkości silnika, gdy jego prędkość spada. Nie używamy jednak metody wykrywania napięcia wstecznego z powodów opisanych w sekcji „Pułapki sterowania silnikiem PWM przy wyższych częstotliwościach”. Zamiast tego monitorujemy pobór prądu.

Gdy MOSFET-y Q1 i Q2 przewodzą, napięcie na rezystorze źródłowym Q1 o rezystancji 0,01 Ω jest proporcjonalne do prądu pobieranego przez silnik. Gdy MOSFET jest wyłączony, na tym rezystorze nie ma napięcia. Używamy więc obwodu próbkowania i podtrzymania, aby uchwycić napięcie, gdy Q1 przewodzi.

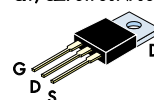
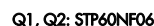
MOSFET Q3 i kondensator 100 μ F tworzą bufor próbkowania i podtrzymania. Bramka Q3 jest sterowana przez wyjście PWM2 układu IC1, które naśladuje wyjście PWM1. Tak więc, gdy Q1 i Q2 są włączone, tak samo jest z Q3, a kondensator 100 μ F ładuje się lub rozładowuje tak, że jego napięcie zbliża się do napięcia na rezystorze 0,01 Ω wykrywającym prąd.

Gdy MOSFET-y Q1 i Q2 wyłączają się, to samo dzieje się z Q3, izolując kondensator 100 μ F od rezystora 0,01 Ω , aby zapobiec jego rozładowaniu w czasie wyłączenia.

Powód, dla którego używamy oddzielnego wyjścia PWM2 do sterowania Q3 ma związek z przypadkiem, gdy silnik jest wyłączony. W tym przypadku wyjście PWM1 ma wypełnienie cyklu pracy 0% (tj. jest utrzymywane w stanie niskim), ale PWM2 jest zaprogramowane do generowania impulsu 60 μ s co 13,4 s. Powoduje to chwilowe włączenie Q3, rozładowując kondensator 100 μ F przez rezystor 0,01 Ω .

Ten czas włączenia jest wydłużany, jeśli kondensator musi zostać rozładowany z wyższego napięcia, zwłaszcza gdy silnik jest wyłączony przez zmniejszenie regulacji prędkości. Bez tego, kondensator 100 μ F powoli ładuje się poprzez prąd upływu ze wzmacniacza IC2, powodując dość gwałtowne uruchomienie silnika.

IC2 jest wzmacniaczem instrumentalnym i zapewnia wzmocnienie małego napięcia na boczniku do pomiaru prądu. Jego wzmocnienie można regulować w zakresie od 611,



23

Pałapki sterowania silnikiem PWM przy wyższych częstotliwościach

Podczas korzystania z sygnału PWM do zasilania silnika prądu stałego, średni prąd uzwojenia silnika zmienia się w zależności od wypełnienia cyklu pracy. Ponieważ moment obrotowy jest proporcjonalny do prądu uzwojenia, prędkość silnika można łatwo zmieniać.

Teoretycznie częstotliwość PWM nie ma wpływu na prędkość silnika; znaczenie ma tylko wypełnienie cyklu pracy, ponieważ określa ono średni prąd płynący przez uzwojenia silnika. Wyższe częstotliwości PWM spowodują mniejsze tętnienie prądu silnika, ale nie wpłyną znacząco na średnią.

Istnieją jednak przypadki, w których wyższe częstotliwości mogą wpływać na prąd przy niższych wypełnieniach cyklu pracy, do tego stopnia, że silnik w ogóle nie będzie się obracał przy niższych wypełnieniach cyklu pracy. Istnieje wiele niejasności co do przyczyn takiego stanu rzeczy i tego, co należy z tym zrobić.

Przeszukaliśmy Internet, próbując znaleźć dobre wyjaśnienie tego zjawiska, ale większość informacji, które znaleźliśmy, była myląca lub niepoprawna. Przeprowadziliśmy więc kilka eksperymentów, by przekonać się o tym na własnej skórze.

Najważniejsze jest to: jeśli używasz półmostka lub pełnego mostka do zasilania silnika prądu stałego, będzie on zachowywał się prawie tak, jak przewiduje teoria. Prąd silnika zmienia się niemal dokładnie liniowo wraz z wypełnieniem cyklu pracy PWM, niezależnie od częstotliwości.

Tego właśnie można oczekiwać, zastępując silnik modelem w postaci indukcyjności połączonej szeregowo z rezystancją. Jeśli indukcyjność wynosi L , a rezystancja szeregowo wynosi R , impedancja uzwojenia silnika wynosi $R + 2\pi f \times L$. Prąd dla fali sinusoidalnej przy dowolnej częstotliwości f wynosi $V / (R + 2\pi f \times L)$.

Sygnał PWM zawiera składową stałoprądową (średni poziom, $V \times$ wypełnienie cyklu pracy) plus składowe zmienneprądowe przy częstotliwości przełączania f oraz harmoniczne fali prostokątnej przy $3f$, $5f$, $7f$ itd. Dokładny skład mieszanki harmonicznych różni się w zależności od wypełnienia cyklu pracy.

Ponieważ prąd maleje wraz ze wzrostem częstotliwości, indukcyjność uzwojenia tłumi składowe AC sygnału PWM. Uzwojenia silnika wygładzają te tętnienia, ale

indukcyjność nie ma wpływu na poziom prądu stałego; jest on określany wyłącznie przez napięcie zasilania, wypełnienie cyklu pracy i rezystancję uzwojenia silnika.

Potwierdzają to nasze testy. Jednak podobnie jak wiele prostszych konstrukcji, nasz regulator prędkości obrotowej silnika nie wykorzystuje konstrukcji półmostkowej lub pełnego mostka, a zatem nie wytwarza fali prostokątnej na uzwojeniach silnika.

Dodatni zacisk silnika jest podłączony do $V+$, a ujemny koniec jest okresowo podłączany do $0V$, gdy włączają się MOSFET-y $Q1$ i $Q2$.

Przez pewien czas mamy napięcie $V+$ na silniku. Ale przez resztę czasu, gdy MOSFET-y $Q1$ i $Q2$ są wyłączone, indukcyjność uzwojenia i napięcie wsteczne EMF powodują wzrost napięcia na ujemnym zacisku silnika powyżej potencjału na zacisku dodatnim. Napięcie to jest tłumione przez diodę $D1$ do potencjału około $0,5V$ powyżej napięcia dodatniego.

Tak więc, gdy MOSFET-y są wyłączone, na silniku występuje napięcie ujemne, a nie $0V$, a przez diodę $D1$ przepływa znaczny prąd recyrkulacyjny. Powoduje to, że prąd uzwojenia silnika zanika znacznie szybciej niż w przypadku półmostka lub pełnego mostka, opisanych powyżej i pokazanych na schematach na dole.

Można to zobaczyć porównując oscylogramy 3 i 4. Pokazują one ten sam nieobciążony silnik prądu stałego zasilany przy tej samej częstotliwości PWM ($3,92kHz$) i przy tym samym wypełnieniu cyklu pracy (10%), ale z zasilaniem półmostkowym na oscylogramie 3 i zasilaniem „po jednej stronie” [single-ended] na oscylogramie 4. Żółty przebieg pokazuje przyłożone napięcie, podczas gdy zielony wykres pokazuje prąd płynący przez uzwojenia silnika (a raczej odczyt napięcia na boczniku).

Szybkość narastania prądu i prąd szczytowy są podobne w obu przypadkach. Ale kiedy MOSFET po wysokiej stronie zasilania [high-side] wyłącza się, a MOSFET po niskiej stronie zasilania [low-side] włącza się, na Oscylogramie 3 można zaobserwować wykładniczy spadek prądu w uzwojeniach silnika. Prąd płynie przez cały cykl, aż zacznie ponownie narastać w następnym cyklu.

Na oscylogramie 4, gdy prąd recyrkuluje przez diodę w czasie wyłączenia, spada

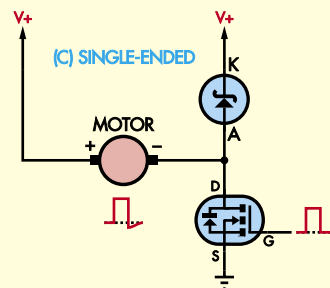
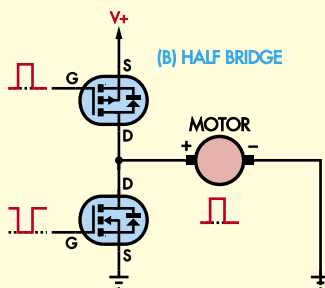
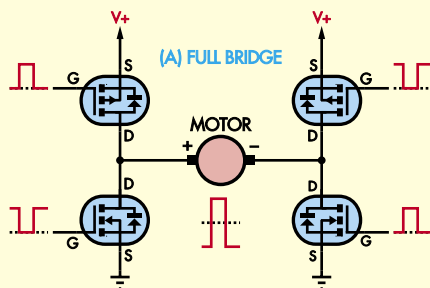
on wykładniczo (ale szybciej), a następnie liniowo, osiągając napięcie zero przed następnym cyklem. Dlatego średni prąd jest znacznie niższy, o około połowę (odczyt na boczniku $400mV$ vs $800mV$), mimo że cykl pracy jest taki sam. Oscylogram 5 pokazuje ten sam schemat zasilania półmostkowego użyty na oscylogramie 3, ponownie z 10% wypełnieniem cyklu pracy, ale przy znacznie wyższej częstotliwości PWM wynoszącej $31,4kHz$. Średni prąd jest tylko nieco niższy, odczyt na boczniku wynosi około $750mV$, w porównaniu do około $800mV$, ze względu na „czas martwy” MOSFET-a, który jest bardziej znaczący przy tej wyższej częstotliwości przełączania.

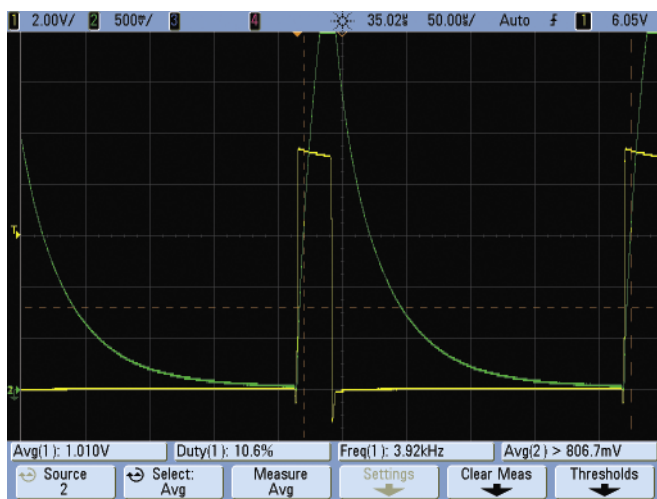
Oscylogram 6 pokazuje ten sam schemat zasilania „single-ended”, co na oscylogramie 4, ale tym razem przy częstotliwości $31,4kHz$. Różnica w prądzie wzrosła jeszcze bardziej – średni prąd płynący przez uzwojenia generuje teraz odczyt na boczniku tylko $286mV$. Tak więc schemat zasilania „single-ended” powoduje dodatkowe zmniejszenie prądu silnika przy wyższych częstotliwościach.

W przypadku schematu zasilania „single-ended” średni prąd silnika dla niskich wypełnień cyklu pracy jest mniejszy niż oczekiwano, a efekt ten zwiększa się przy wyższych częstotliwościach. Dlatego dobrym pomysłem jest zwiększenie minimalnego wypełnienia cyklu pracy przy wyższych częstotliwościach PWM w celu kompensacji tego efektu, co jest powodem zastosowania potencjometru nastawnego $VR3$ w tym projekcie.

Wielkość tego efektu może się również różnić w zależności od silnika. Większe silniki o wyższej indukcyjności będą bardziej narażone z powodu zmniejszonego prądu (i momentu obrotowego) przy niskich wypełnieniach cyklu pracy z wyższymi częstotliwościami PWM.

W praktyce najłatwiejszym sposobem skompensowania tego efektu jest dostosowanie ustawienia minimalnego wypełnienia cyklu pracy (poprzez regulację $VR3$), aż do uzyskania zadowalającej regulacji prędkości na dolnym końcu zakresu potencjometru prędkości $VR1$. Jeśli nie można tego osiągnąć dla danego silnika, należy wypróbować niższą częstotliwość PWM.





Oscylogram 3. Napięcie na silniku (żółty przebieg) i prąd (zielony przebieg) z półmostkiem zasilającym przy 10% wypełnieniu cyklu pracy. Indukcyjność silnika ogranicza czas narastania i opadania prądu. Prąd nie spada do zera przed następnym impulsem, pomimo stosunkowo niskiego wypełnienia cyklu pracy; indukcyjność uzwojenia podtrzymuje go

10 k Ω /2 k Ω . Tak więc dla zasilania silnika w zakresie 0...30 V napięcie na wejściu AN10 mieści się w zakresie 0...5 V.

Napięcie to jest filtrowane za pomocą kondensatora 100 nF, aby zapobiec zmianie wyniku konwersji ADC przez zakłócenia.

Regulacja ustawień

Napięcie zasilania jest porównywane z napięciem ustawienia minimalnego progu napięciowego na wejściu AN7 (wyprowadzenie 7), ustawianym przez potencjometr nastawny VR4. Ten potencjometr jest podłączony do stabilizowanego napięcia zasilającego 5 V, umożliwiając regulację zakresu napięcia od 0 do 5 V. Punkt testowy TP4 pozwala zmierzyć ustawiony próg.

Aby ułatwić konfigurację, napięcie w punkcie TP4 wynosi jedną dziesiątą minimalnego progu napięcia. Jeśli więc próg minimalnego napięcia

ma wynosić 11,5 V, należy ustawić napięcie na punkcie TP4 na 1,15 V.

Napięcie na wejściu AN7 jest konwertowane na wartość cyfrową i mnożone przez 1,6666, więc skala odpowiada napięciu zasilania silnika podzielonemu przez sześć.

Napięcie zasilanie silnika musi spaść poniżej minimalnego progu zasilania przez 10 sekund, zanim napęd silnika zostanie wyłączony. Gdy to nastąpi, dioda LED1 miga przez chwilę co sekundę.

Zazwyczaj akumulator odzyskuje nieco energii po wyłączeniu napędu silnika; napięcie akumulatora wzrośnie, gdy nie będzie obciążenia. Aby zapobiec ponownemu włączeniu silnika z powodu tego efektu, dodajemy histerezę.

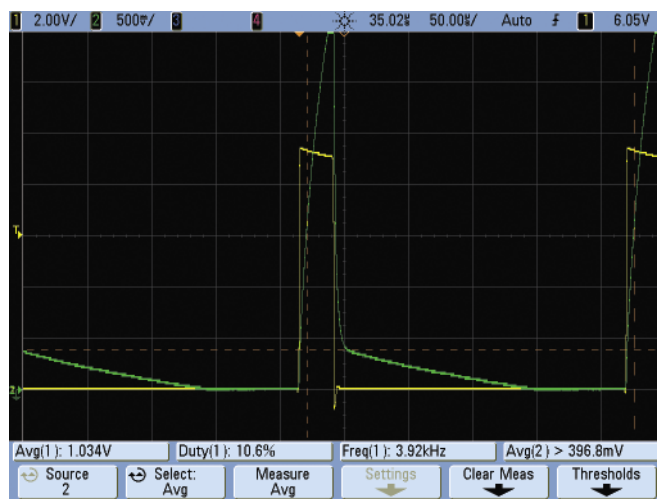
Napięcie zasilania silnika będzie musiało przekroczyć próg niskiego napięcia plus

napięcie histerezy, zanim silnik zostanie ponownie włączony. W praktyce akumulator musi zostać naładowany przed ponownym uruchomieniem silnika.

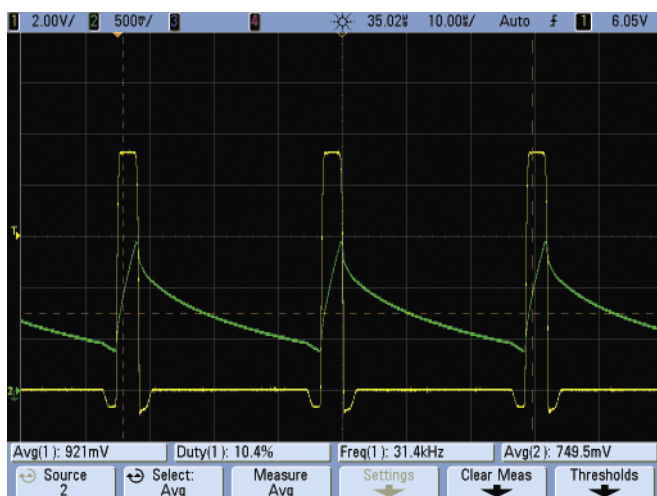
Histeresa jest ustawiana za pomocą potencjometru nastawnego VR5 i może być monitorowana w punkcie TP5. Odczyt TP5 to pełne napięcie histerezy (nie 1/10, jak w przypadku pomiaru minimalnego progu na TP4). Jeśli więc chcesz uzyskać histerezę 1 V, wyreguluj VR5, aż w punkcie TP5 odczytasz napięcie 1 V.

Regulacja okresu miękkiego startu odbywa się za pomocą potencjometru nastawnego VR2, z napięciem mierzonym w punkcie TP2. Napięcie to jest monitorowane na wejściu AN6 i ustawia maksymalną szybkość, z jaką wzrasta prędkość silnika.

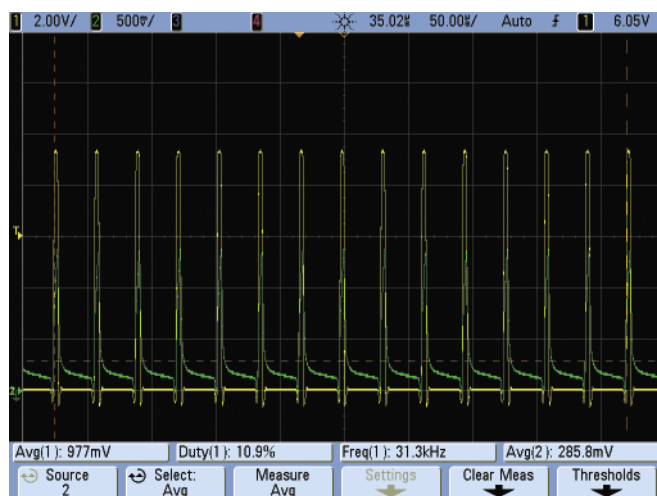
Maksymalny czas osiągnięcia od zera do 100% wypełnienia cyklu pracy wynosi dwie



Oscylogram 4. Podobnie jak Oscylogram 3, ale zmieniliśmy sterownik z półmostka na pojedynczy MOSFET z diodą recyrkulacyjną, jak w tym (i wielu innych) regulatorach prędkości. Ma to ogromny wpływ na to, jak prąd zmniejsza się na końcu każdego impulsu, więc prąd silnika jest znacznie mniejszy przy niskich wypełnieniach cyklu pracy



Oscylogram 5. Przetwarzając się z powrotem na sterowanie półmostkowe, ale zwiększając częstotliwość PWM do 31,4 kHz, można zauważyć, że średnia wartość prądu prawie nie ulega zmianie. Średni poziom prądu jest wyższy w czasie wyłączenia ze względu na krótszy czas wyłączenia



Oscylogram 6. Zasilanie „po jednej stronie” [single-ended] z wyższą częstotliwością wykazuje ten sam szybki spadek prądu, jak pokazano na oscylogramie 4, z tą różnicą, że tym razem średni prąd jest jeszcze niższy, ponieważ ma mniej czasu na narastanie podczas krótszych impulsów włączenia



Regulator prędkości silnika prądu stałego z potencjometrem VR1 podłączonym do testowania

sekundy, przy napięciu 5 V w punkcie TP2. Ustawienie 2,5 V zapewni jednosekundowy okres łagodnego rozruchu i tak dalej.

Potencjometr nastawny VR3 służy do regulacji krzywej prędkości, z odpowiadającym punktem pomiarowym TP3. Jest on monitorowany na wejściu AN4 układu IC1, wyprowadzenie 16. Pozwala to na wykorzystanie potencjometru prędkości w całym jego zakresie, gdy częstotliwość PWM jest ustawiona stosunkowo wysoko, a także może zrekomensować fakt, że silniki mogą wymagać cyklu pracy z wypełnieniem znacznie powyżej 0%, zanim zaczną się obracać.

Jak opisano w oddzielnym panelu zatytułowanym „Pułapki sterowania silnikiem PWM przy wyższych częstotliwościach”, w niektórych przypadkach sterowanie silnikiem z wysoką częstotliwością PWM może oznaczać, że silnik nie uruchomi się, dopóki wypełnienie cyklu pracy nie osiągnie 20% lub nawet więcej.

Regulacja krzywej prędkości ustawia początkowe wypełnienie cyklu pracy, gdy potencjometr prędkości jest obracany od położenia zerowego w prawo. Regulacja ta usuwa martwą strefę zakresu obrotu potencjometru prędkości. Zakres regulacji krzywej wynosi od prawie zera do 33% wypełnienia cyklu pracy.

Za każdym razem, gdy ustawienie krzywej prędkości jest niezerowe, oprogramowanie w IC1 rozszerza zakres regulacji prędkości tak, że maksymalne wypełnienie cyklu pracy jest nadal osiągalne, gdy VR1 jest całkowicie skrecony w prawo.

Praca przy niskich częstotliwościach może być również zoptymalizowana przy użyciu regulacji krzywej prędkości, z założoną zworką JP1 w celu obniżenia normalnie wysokiego poziomu wejścia cyfrowego RA5 (wyprowadzenie 2). Bez założonej zworki wejście RA5 jest polaryzowane do stanu wysokiego przez wewnętrzny prąd polaryzujący.

Regulacja krzywej prędkości przy założonej zworce JP1 pozwala na lepszą regulację sprzężenia zwrotnego przy bardzo niskich wypełnieniach cyklu pracy. Regulacja zmniejsza efekt

zaskoku prędkości obrotowej silnika, przy którym napięcie sprzężenia zwrotnego nagle wzrasta wraz ze wzrostem wypełnienia PWM przy starcie tuż od zera. Ta regulacja ustawia wartość przesunięcia sprzężenia zwrotnego, dzięki czemu sprzężenie zwrotne jest ignorowane poniżej określonego ustawienia prędkości.

Potencjometr nastawny VR3 służy również do włączania lub wyłączania wykrywania odłączenia silnika. Odbывается to poprzez podzielenie zakresu VR3 na połowy, 0...2,5 V i 2,5...5 V. Od 0 V do 2,5 V sprawdzanie odłączenia silnika jest wyłączone. Powyżej 2,5 V wykrywanie odłączenia silnika jest włączone, a regulacja krzywej prędkości jest odwrócona, przy czym całkowite obrócenie w prawo daje taki sam efekt jak całkowite obrócenie w lewo.

Gdy prąd sprzężenia zwrotnego silnika znajduje się poniżej ustawionej wartości przez ponad 200 ms, silnik jest deklarowany jako odłączony. W takim przypadku cykl pracy PWM jest ustawiony na zero, a dioda LED miga z częstotliwością 2 Hz.

Silnik uruchomi się ponownie dopiero po ponownym podłączeniu, a potencjometr prędkości zostanie najpierw całkowicie obrócony w lewo. Zapobiega to nieregularnemu działaniu z powodu luźnych przewodów itp.

Wykrywanie rozłączenia silnika jest opcjonalne, ponieważ jeśli silnik nie jest prawidłowo skonfigurowany do pracy z wysokimi częstotliwościami, fałszywe zdarzenia rozłączenia mogą powodować uciążliwe wyłączenia. Może się to zdarzyć, jeśli krzywa nie jest prawidłowo ustawiona, z wystarczająco wysokim wypełnieniem cyklu pracy na początku obrotu potencjometru prędkości.

Opcje częstotliwości PWM

Przełącznik S1 służy do wyboru częstotliwości PWM dla silnika. Jest to 16-pozycyjny obrotowy przełącznik BCD (kodowanie szesnastkowe). Istnieją cztery zaciski przełącznika oznaczone jako 8, 4, 2 i 1 oraz wspólne połączenie, które podłączyliśmy do masy.

Pozostałe zaciski przełącznika łączą się odpowiednio z wejściami cyfrowymi RA1, RB6, RB7 i RB5 układu IC1. Wszystkie te wejścia z wyjątkiem RA1 są skonfigurowane w IC1 w celu zapewnienia prądu polaryzacji. Wejście RA1 nie ma takiej opcji, więc zewnętrzny rezystor polaryzujący 10 kΩ łączy to wejście z napięciem 5 V.

Dzięki polaryzacji wejścia znajdują się w stanie wysokim (przy 5 V), gdy przełącznik nie łączy tego zacisku z masą. 16 możliwych kombinacji jest dekodowanych w IC1 i wybierana jest wymagana częstotliwość PWM (patrz tabela 1).

Zasilanie

Zasilanie regulatora jest podłączone poprzez zaciski CON1 pomiędzy GND i wejście dodatniego zasilania układu. Prąd zasilania przepływa przez diodę Zenera ZD1, a wejście regulatora REG1 jest zbocznikowane kondensatorem 470 nF.

REG1 to samochodowy stabilizator napięcia 5 V o niskim spadku napięcia. Jest on w stanie wytrzymać napięcie o odwrotnej polaryzacji, więc zapewnia układowi ochronę przed odwrotnie podłączonym zasilaniem. Maksymalne zalecane napięcie robocze na wejściu REG1 wynosi 26 V. Tak więc do użytku przy napięciu do 30 V, dioda ZD1 obniża napięcie na wejściu o około 4,7 V.

Napięcie zaniku pracy dla REG1 wynosi zazwyczaj 0,5 V powyżej znamionowego. Oznacza to, że potrzebuje on 5,5 V na wejściu, aby zapewnić regulację na wyjściu. Z dodatkowym napięciem dla regulatora wynosi to 5,5 V + 4,7 V = 10,2 V. Dla bezpieczeństwa zaokrąglamy tę wartość do 10,5 V.

Należy pamiętać, że połączenia dodatniego zasilania regulatora i silnika są oddzielne, więc w razie potrzeby silnik może być zasilany innym napięciem.

Oznacza to, że zasilanie silnika może znajdować się poza zakresem sterownika, a obwód

Tabela 1. Opcje częstotliwości PWM

Ustawienie przełącznika BCD (S1)	Częstotliwość PWM
0	30,6 Hz
1	61,3 Hz
2	122,5 Hz
3	245 Hz
4	367,6 Hz
5	490 Hz
6	980 Hz
7	1,96 kHz
8	2,97 kHz
9	3,92 kHz
A	5,88 kHz
B	7,84 kHz
C	11,8 kHz
D	15,7 kHz
E	23,5 kHz
F	32,4 kHz



nadal będzie działał, o ile odpowiednie napięcie zasilania sterownika zostanie połączone, gdy napięcie zasilania silnika mieści się w odpowiednim zakresie sterownika.

Budowa

Regulator prędkości silnika DC jest zbudowany przy użyciu dwustronnej płytki PCB z metalizacją otworów (technologia PTH) o kodzie 11006211 i wymiarach 122×58 mm. **Rysunek 2** przedstawia szczegółowy schemat montażowy.

Zacznij od zainstalowania dwóch rezystorów SMD 10 Ω i dwóch rezystorów SMD 0,01 Ω, wszystkie w pobliżu Q1 i Q2. Teraz zamontuj IC3, sterownik SMD bramek MOSFET-ów. Zachowaj ostrożność podczas lutowania; możesz potrzebować szkła powiększającego i oddzielnej lampy roboczej. Przylutuj najpierw wyprowadzenie 1 i sprawdź, czy pozostałe wyprowadzenia są prawidłowo wyrównane, przed przylutowaniem pozostałych.

Teraz można zainstalować diodę Zenera ZD1, zwracając uwagę na jej orientację. Następnie należy przylutować siedem rezystorów do montażu przewlekane. Przed montażem należy sprawdzić każdy z nich za pomocą multimetru cyfrowego.

Gdy te części znajdą się na miejscu, zainstaluj podstawkę dla IC1. Układ IC2 można zamontować za pomocą podstawki lub przylutować bezpośrednio do płytki. Upewnij się, że każdy z nich jest prawidłowo zorientowany.

W tym momencie warto zamontować MOSFET Q3, diodę LED i dwuszpilkową listwę kołkową pod zworkę JP1. Upewnij się, że dłuższe wyprowadzenie (anoda) diody LED1 znajduje się w otworze po lewej stronie, oznaczonym literą „A”. Zamiast tego można zamontować tam dwuszpilkowy odcinek listwy kołkowej lub bezpośrednio przylutować do płytki dwużyłowy kabel, aby dioda LED mogła być zamontowana na obudowie.

Następnie można przylutować kondensatory poliestrowe; natomiast te elektrolityczne najłatwiej jest zainstalować po uprzednim zamontowaniu wszystkich półprzewodników. Postępuj analogicznie z wieloobrotowymi potencjometrami nastawnymi. Ustaw je tak, aby śruby regulacyjne znajdowały się w pozycji, jak pokazano na rysunku. Teraz można zainstalować przełącznik BCD S1, z kropką wskazującą pierwsze wyprowadzenie w prawym dolnym rogu.

Następny na liście jest 3-drożny śrubowy blok zacisków (CON2). Upewnij się przed lutowaniem jego wyprowadzeń, że jest prawidłowo osadzony na płytce i że jego otwory są skierowane na zewnątrz. Następnie można podłączyć CON1, 6-stykowy blok zacisków śrubowych. Należy pamiętać, że Altronics twierdzi, że są one przewidziane na obciążenie 15 A; jednak dane

Dinkle dla tych terminali DT-35-B07W-XX dopuszczają prąd 20 A, więc są odpowiednie dla tego regulatora 20 A.

Następny jest uchwyt bezpiecznika. Można zamontować oprawkę monolityczną lub dwa oddzielne zaciski styków bezpiecznika. Jeśli używasz pojedynczych klipsów, dobrym pomysłem może być włożenie bezpiecznika przed lutowaniem, aby upewnić się, że są one prawidłowo ustawione.

W punktach testowych TP1-TP5 i TP GND można zainstalować kołki PC lub pozostawić je puste i kontaktować do punktów testowych PCB bezpośrednio za pomocą sond multimetru.

Instalacja półprzewodników

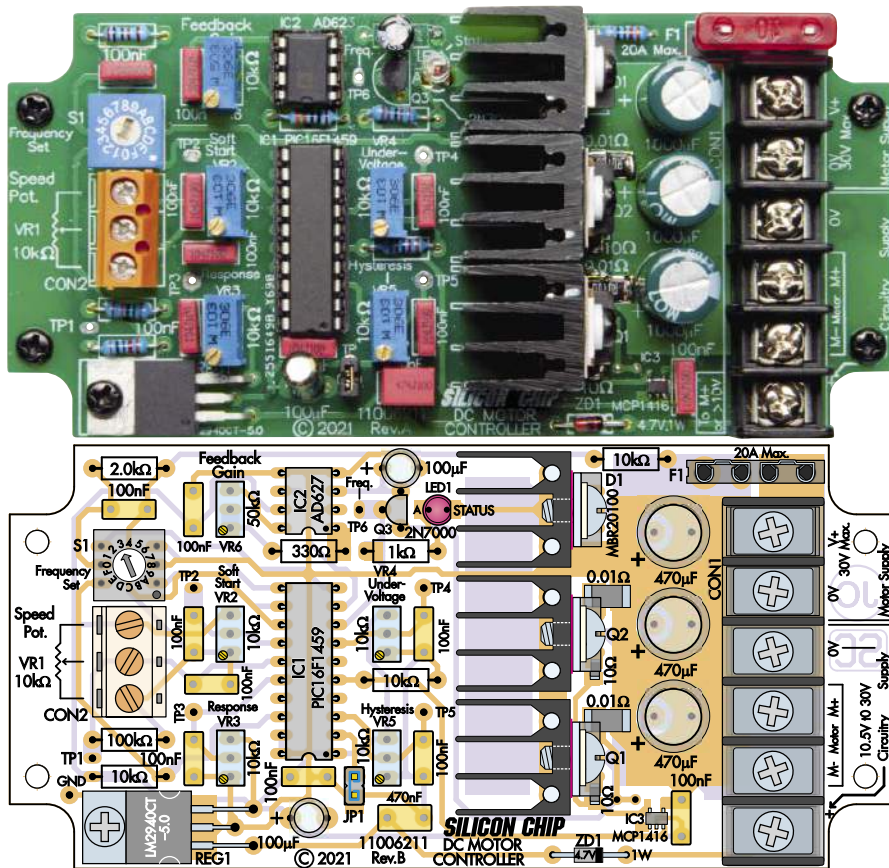
Stabilizator napięcia REG1 jest montowany poziomo na płytce poziomo. Instaluje się go, najpierw wyginając jego wyprowadzenia, aby przeszły przez otwory montażowe. Radiator REG1 jest następnie mocowany do płytki drukowanej za pomocą śruby M3×6 i nakrętki, po czym wyprowadzenia są lutowane.

MOSFETy Q1 i Q2 oraz dioda Shottky'ego D1 są montowane pionowo i mocowane do oddzielnych małych radiatorów. Trzy radiatory muszą być zainstalowane jako pierwsze, poprzez przylutowanie ich kołków pozycjonujących

do odpowiednich pól lutowniczych PCB. Przed przylutowaniem radiatorów należy upewnić się, że są one prawidłowo osadzone na płytce drukowanej.

MOSFETy Q1 i Q2 należy zamontować wyprowadzeniami w przygotowanych dla nich otworach na płytce PCB, pamiętając jednocześnie o zastosowaniu podkładek silikonowych i tulei izolacyjnych podczas montażu tych elementów (patrz rysunek 3), aby odizolować każdy z nich od radiatora. Przycmocyj je za pomocą śrub M3×10 do gwintowanych otworów w radiatorach. Dokręć mocno śruby, a następnie przylutuj ich wyprowadzenia. Dioda D1 jest montowana w podobny sposób. **Od Red. EdW:** alternatywnie można przycmocyjować półprzewodniki, z użyciem podkładek, tulei izolacyjnych i śrub M3×10, do luźnych radiatorów, a następnie cały taki blok wpasować do PCB, rozpoczynając lutowanie od wypustek radiatorów. Unika się w ten sposób niewygodnego manipulowania półprzewodnikami przy radiatorach osadzonych na PCB.

Teraz zainstaluj pozostałe kondensatory elektrolityczne, zwracając uwagę na ich prawidłową orientację. Na koniec użyj multimetru, aby potwierdzić, że metalowe radiatory



Rysunek 2. Schemat montażowy regulatora. Płytkę PCB regulatora prędkości jest stosunkowo kompaktowa i wykorzystuje tylko pięć części SMD: cztery rezystory i sterownik MOSFET-ów IC3. MOSFET-y Q1 i Q2 oraz dioda D1 są przycmocyjowane do radiatorów montowanych w celu chłodzenia na płycie drukowanej. Podczas montażu należy zwrócić uwagę na polaryzację trzech układów scalonych, diody ZD1, kondensatorów elektrolitycznych i przełącznika BCD S1

Wada sprzężenia zwrotnego regulatora prędkości korzystającego ze wstecznej siły elektromagnetycznej

Zazwyczaj silnik prądu stałego działa jak generator, gdy zasilanie jest wyłączone. W przypadku korzystania z napędu PWM, to generowane napięcie lub wsteczna siła elektromotoryczna (EMF) występuje powtarzalnie, gdy zasilające MOSFET-y są wyłączone. Indukowane napięcie nie powstaje jednak natychmiast po wyłączeniu; nie pojawia się ono, dopóki nie rozproszy się energia zgromadzona w indukcyjności uzwojeń silnika.

W wielu regulatorach prędkości napięcie wstecznej EMF jest wykorzystywane do stabilizacji prędkości przy zmiennym obciążeniu. Gdy silnik jest obciążony, prędkość i wsteczna EMF zmniejszają się, a zmiana ta jest wykorzystywana do zapewnienia sprzężenia zwrotnego, które zwiększa wypełnienie cyklu pracy PWM w celu utrzymania stałej prędkości pod obciążeniem.

Jednak przy wyższych częstotliwościach PWM napięcie wstecznej EMF pojawia się w cyklu PWM znacznie później; czasami powstaje dopiero po ponownym włączeniu MOSFET-ów, więc niemożliwe jest wykrycie wstecznej EMF.

Porównaj przebiegi na Oscylogramie 7 i Oscylogramie 8. Są takie same, z wyjątkiem tego, że częstotliwość PWM wynosi nieco poniżej 3 kHz na Oscylogramie 7 i prawie 12 kHz na Oscylogramie 8. Można zobaczyć „półkę” tylnej części przebiegu EMF pojawiającą się około 80 µs po wyłączeniu na Oscylogramie 7, ale jest ona ledwo widoczna na Oscylogramie 8 i nie byłaby w ogóle obecna przy wyższej częstotliwości przełączania.

Brak wstecznej siły elektromagnetycznej przy wysokich częstotliwościach PWM oznacza, że musimy użyć innego sposobu wykrywania obciążenia silnika. Najprostszą alternatywą jest pomiar prądu silnika. Robimy to tylko wtedy, gdy silnik jest napędzany, wzmacniając napięcie na rezystorze bocznym o niskiej wartości połączonym szeregowo z silnikiem.

Korzystając z regulacji sprzężenia zwrotnego wykorzystującej pomiar prądu, wypełnienie cyklu pracy PWM można zwiększyć, gdy silnik jest obciążony. Pozwala to przynajmniej w pewnym stopniu przezwyciężyć niedociągnięcia związane z niskim momentem obrotowym przy wysokich częstotliwościach i niższych stopniach wypełnienia cyklu pracy.

diody D1 oraz tranzystorów Q1 i Q2 są odizolowane od radiatorów chłodzących.

Testowanie

Przed włożeniem układu IC1 do podstawki należy sprawdzić działanie stabilizatora napięcia, przykładając napięcie 10,5...30 V między zaciskami 0 V i dodatniego zasilania regulatora na CON1.

Zmierzyć napięcie między metalowym radiatorem REG1 a jego prawym skrajnym wyprowadzeniem. Odczyt powinien być zbliżony do 5 V (4,75 do 5,25 V). Jeśli tak nie jest, sprawdź, czy napięcie wejściowe na lewym wyprowadzeniu REG1 wynosi co najmniej 5,5 V.

Jeśli odczyt jest prawidłowy, wyłącz zasilanie i zainstaluj układ IC1, upewniając się, że jest prawidłowo zorientowany, a żadne

z jego wyprowadzeń nie zagina się pod korpusem. **Od Red. EdW: Napięcie 5 V warto zmierzyć bezpośrednio na wyprowadzeniach 1 (Vdd) oraz 20 (Vss) podstawki pod mikrokontroler, przed jego zamontowaniem. Obecność poprawnego napięcia bezpośrednio na stabilizatorze nie gwarantuje, że dociera ono (zarówno Vdd jak i Vss) również do samego mikrokontrolera. Można będzie w ten sposób uniknąć przykrych niespodzianek na skutek ewentualnych niedokładności lutowania, na przykład w obwodzie masy samego stabilizatora. Jeśli użyłeś podstawki dla IC2, włóż do niej ten układ.**

Na tym etapie warto podłączyć potencjometr VR1 do CON2. Będziesz także musiał włożyć bezpiecznik, aby kontynuować testowanie. Bezpiecznik powinien być dopasowany

do silnika; jeśli jest to silnik o prądzie 1 A, zainstaluj bezpiecznik 1 A; w przypadku silnika 20 A użyj bezpiecznika 20 A itp.

Następnie należy obrócić pokrętko regulacji krzywej prędkości VR3 całkowicie w lewo. Można znaleźć tę pozycję, obracając go o co najmniej 20 obrotów w lewo lub do momentu usłyszenia słabego dźwięku kliknięcia. Gdy układ jest zasilany, odczyt napięcia między TP3 i GND powinien być bardzo zbliżony do 0 V.

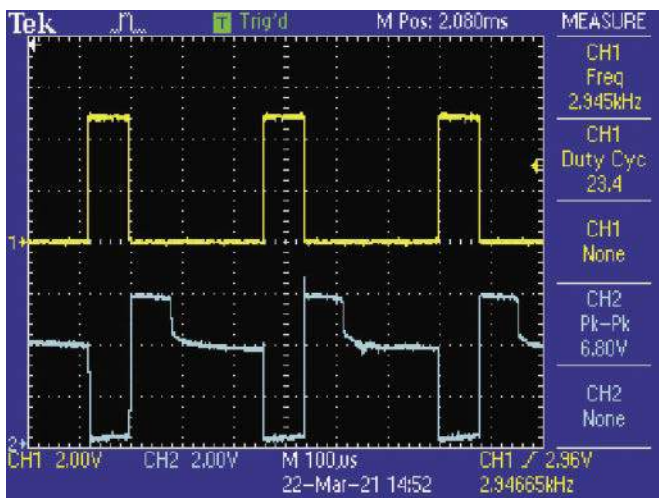
Testowanie wyłącznika niskiego napięcia

Po podłączeniu zasilania dioda LED będzie migać z częstotliwością 1 Hz, ponieważ do zacisków zasilania silnika nie jest podłączone napięcie.

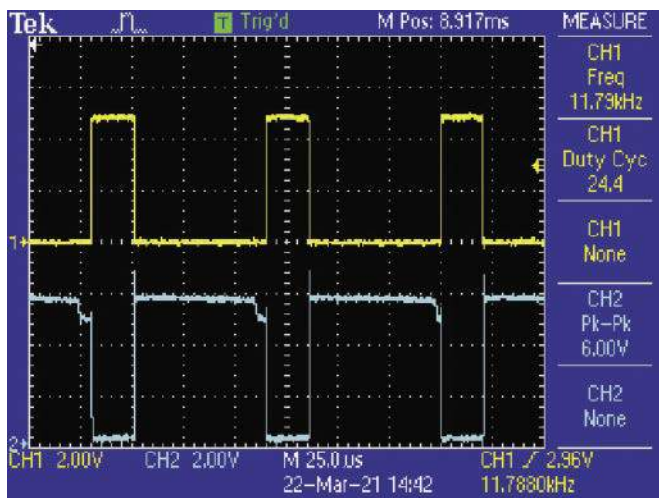
Potencjometr nastawny VR4 ustawia odciecie przy niskim napięciu. Za pomocą multimetru podłączonego między TP4 i TP GND, wyreguluj VR4 na jedną dziesiątążądanego niskiego napięcia odcięcia. Tak więc, aby uzyskać odciecie przy niskim napięciu równym 11,5 V (bezpieczny poziom dla większości akumulatorów kwasowo-ołowiowych 12 V), wyreguluj TP4, aż uzyskasz odczyt 1,15 V.

Regulacja histerezy jest podobna, przy użyciu potencjometru nastawnego VR5 i pomiaru na TP5. Histereza to napięcie zmierzone na TP5 (nie 1/10 napięcia jak poprzednio). Aby uzyskać histerezę 1 V, należy ustawić napięcie na TP5 na 1 V. Histerezę można ustawić na maksymalnie 5 V, ale 1 V jest rozsądnym punktem wyjścia. Przy zalecanym napięciu odcięcia 11,5 V oznacza to, że napięcie akumulatora musi wzrosnąć powyżej 12,5 V (około połowy naładowania) przed wznowieniem pracy.

Jeśli posiadasz laboratoryjny zasilacz regulowany, możesz przetestować punkt odcięcia



Oscylogram 7. Przy częstotliwości PWM nieco poniżej 3 kHz jest wystarczająco dużo czasu na wykrywanie wstecznej siły elektromagnetycznej EMF. Napięcie silnika wzrasta natychmiast po wyłączeniu MOSFET-ów, a następnie spada do niższego plateau, gdy pole magnetyczne zanika, a EMF zaczyna dominować



Oscylogram 8. Przy częstotliwości PWM wynoszącej prawie 12 kHz, napięcie back-EMF jest ledwo widoczne tuż przed rozpoczęciem następnego impulsu. Próbkowanie wstecznej siły elektromagnetycznej przy tej częstotliwości byłoby dla tego silnika niepraktyczne i niemożliwe przy wyższych częstotliwościach

Wykaz elementów:

1 dwustronna płytka drukowana o kodzie 11006211 i wymiarach 122x58 mm
1 obudowa UB3 Jiffy box (opcjonalnie) [Jaycar HB6013, HB6023, Altronics H0203].
16-drożny zacisk śrubowy barierowy 20 A* do montażu na PCB, raster 8,25 mm (CON1) [Altronics P2106]
13-drożny zacisk śrubowy, raster 5,08 mm (CON2)
1 potencjometr liniowy 10 kΩ (VR1)
1 pokrętko pasujące do VR1
1 dwuszpilkowy odcinek listwy kołkowej prostej, raster 2,54 mm plus zworka (JP1)
3 silikonowe podkładki i tuleje izolacyjne TO-220
1 podstawka DIL-20 IC dla IC1
1 podstawka DIL-8 IC dla IC2 (opcjonalnie)
3 radiatory TO-220 do montażu na PCB [Jaycar HH8516, Altronics H0650]
1 4-bitowy przełącznik BCD (S1) typu „jeden z 16” [Jaycar SR1220, Altronics S3001A]
1 oprawka bezpiecznika płytkowego 20 A (F1) [Altronics S6040]
1 bezpiecznik topikowy pasujący do silnika (do 20 A)
4 śruby z łbem walcowym M3x10
1 nakrętka M3
4 kołki dystansowe M3 o długości 6,3 mm i 8 śrub M3x6 (opcjonalnie; do montażu płyty)
6 kołków PC (opcjonalnie)
* Dinkle określa je jako 20 A; Altronics podaje 15 A

Półprzewodniki

1 mikrokontroler PIC16F1459-1/P, DIP-20, zaprogramowany kodem 1100621A.hex (IC1)
1 wzmacniacz instrumentalny AD627ANZ, DIP-8 (IC2) [element 14, RS]
1 sterownik MOSFET-ów MCP1416T-E/OT, SOT-23-5 (IC3) [RS Components 668-4216]
1 regulator napięcia LM2940CT-5.0, TO-220 (REG1) [Jaycar ZV1560, Altronics Z0592]
2 N-kanalowe MOSFET-y STP60NF06, TO-220 (Q1, Q2) [Jaycar ZT2450]
1 N-kanalowy MOSFET matosygnatowy 2N7000, TO-92 (Q3) [Jaycar ZT2400, Altronics Z1555]
1 dioda LED 3 mm o wysokiej jasności (LED1)
1 dioda Zenera 4,7 V 1 W (ZD1)
1 podwójna dioda Schottky'ego MBR20100 10 A, TO-220 (D1) [Jaycar ZR1039].

Kondensatory

3 kondensatory elektrolityczne 470 µF 35 V low-ESR
2 kondensatory elektrolityczne 100 µF 16 V
1 kondensator poliestrowy 470 nF 63 V MKT
9 kondensatorów poliestrowych 100 nF 63 V MKT

Rezystory (wszystkie 1/4 W, 1% metalizowane osiowe, chyba że podano inaczej)

1 szt. 100 kΩ 3 szt. 10 kΩ 1 szt. 2 kΩ 1 szt. 1 kΩ
1 szt. 330 Ω 2 szt. 10 Ω SMD 1206
2 szt. 0,01 Ω SMD 2512 3 W [RS Components Cat 188-0753, Vishay WFMA25120100FEA lub odpowiednik]
4 wieloobrotowe potencjometry nastawne 10 kΩ (styl 3296W) (VR2-VR5)
1 wieloobrotowy potencjometr nastawny 50 kΩ (styl 3296W) (VR6)

przy niskim poziomie naładowania baterii. Podłącz to zasilanie między dodatnim zasilaniem silnika a 0 V i obróć VR1 całkowicie w prawo. Dioda LED zaświeci się, gdy napięcie zasilania znajduje się w zakresie roboczym i zacznie migać, gdy wykryte zostanie zbyt niskie napięcie.

Ustaw zasilanie na wartość wyższą niż ustawienie odcięcia przy niskim napięciu plus ustawienie histerezy, aby odcięcie niskiego napięcia nie zostało początkowo aktywowane. Następnie należy zmniejszyć napięcie do wartości napięcia odcięcia. Należy pamiętać, że zadziałanie zabezpieczenia przy zbyt niskim napięciu zasilania nastąpi po czasie około 10 sekund, od momentu gdy napięcie zasilania spadnie poniżej ustawionego progu. Dioda LED1 powinna wtedy migać z częstotliwością 1 Hz.

Powoli zwiększaj napięcie zasilania do wartości nieco powyżej progu plus wartość ustawienia histerezy (12,5 V w naszym przykładzie), a dioda LED1 powinna się zaświecić w sposób ciągły. W razie potrzeby wyreguluj VR4 i VR5, aby uzyskać odcięcie i włączenie dokładnie przy wymaganych napięciach.

Ustawienie miękkiego startu

Wyreguluj VR2 dla wymaganego czasu trwania miękkiego startu. Zazwyczaj odpowiednio jest ustawienie 5 V na TP2,

co zapewni maksymalnie dwusekundowy okres łagodnego rozruchu. Można zmniejszyć tę wartość, aby uzyskać szybszy rozruch, lub wyłączyć łagodny rozruch przy napięciu 0 V mierzonym na TP2.

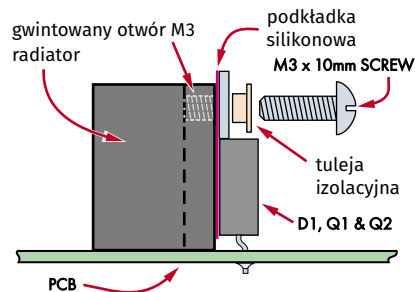
Regulacja krzywej prędkości

VR3 ustawia regulację krzywej prędkości. Jest ona wyłączona, gdy VR3 jest obrócony całkowicie w lewo, przy 0 V na TP3. Obrót VR3 w prawo zwiększy regulację krzywej prędkości. W przypadku ustawień powyżej 2,5 V, patrz opcjonalna sekcja wykrywania odłączenia silnika poniżej.

Jak wspomniano wcześniej, ustawienie krzywej prędkości zapewnia poprawę działania przy wysokiej częstotliwości PWM, gdy JP1 jest otwarty lub poprawę działania przy niskiej częstotliwości, gdy JP1 jest zwarty.

Gdy JP1 jest otwarty, VR3 zwiększa minimalny cykl pracy dla niskich ustawień VR1. Aby dokonać regulacji, obróć potencjometr prędkości VR1 lekko w prawo od skrajnej lewej pozycji, uzyskując odczyt nieco ponad 20 mV na TP1. Następnie ustaw VR3 w prawo, aż silnik zacznie pracować.

Wyreguluj regulator wzmocnienia (VR6), aby uzyskać najlepsze sterowanie silnikiem w celu utrzymania prędkości silnika pod obciążeniem. Ustawienie w prawo zwiększy wzmocnienie, a ustawienie w lewo



Rysunek 3. Ten widok z boku pokazuje szczegóły montażu elementów w obudowach TO-220 do radiatorów. Otwór w radiatorze jest fabrycznie gwintowany. Radiatory są połączone do masy za pośrednictwem płytki drukowanej i kołków montażowych, więc potrzebne są podkładki i tuleje izolacyjne

zmniejszy wzmocnienie. Ustawienie zbyt wysokiego wzmocnienia może spowodować niestabilność prędkości silnika.

Ustaw częstotliwość PWM na wartość, przy której silnik pracuje najlepiej. Będzie to kompromis między wydajnością sterowania silnikiem a ilością szumu PWM wytwarzanego przez silnik. Bardzo niskie częstotliwości mogą powodować nierówną pracę silnika. Bardzo wysokie częstotliwości poprawią płynność, ale mogą zmniejszyć moment obrotowy przy niższych prędkościach, chyba że sterowanie sprzężeniem zwrotnym zostanie dostosowane w celu zapewnienia lepszej wydajności pod obciążeniem.

Wyreguluj potencjometr VR3, aby uzyskać najlepszy zakres regulacji prędkości dla VR1. Gdy częstotliwość PWM jest niska, może się okazać, że prędkość silnika może gwałtownie wzrosnąć po zwiększeniu VR1 od zera, zwłaszcza przy wysokim wzmocnieniu sprzężenia zwrotnego. Regulacja VR3 przy zwartym JP1 może zmniejszyć ten efekt zatrząskiwania. Zaczniij od 0 V (na TP3) i reguluj VR3, aż silnik będzie działał dobrze przy niskich wypełnieniach cyklu pracy, bez efektu zaskoku.

Wykrywanie odłączenia silnika

Jeśli chcesz skorzystać z tej opcji, potencjometr nastawny regulacji krzywej prędkości (VR3) jest ustawiany w odwrotny sposób. Nie ma regulacji krzywej prędkości, gdy VR3 jest skrecony całkowicie w prawo (5 V na TP3), a regulacja krzywej prędkości zwiększa się, gdy VR3 jest skrecony od tego punktu w lewo. Regulacja jest użyteczna do osiągnięcia 2,5 V na TP3. ■

John Clarke

Adaptacja do wydania
polskiego – Andrzej Nowicki

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



Monitor prędkości na bazie GPS

Koniec z mandatami za przekroczenie prędkości

W Internecie można znaleźć wiele przykładów pokazujących zastosowanie modułu GPS w projektach z mikrokontrolerami. Ale co można zrobić ze współzrędnymi GPS, poza naniesieniem ich na mapę? Chciałem zrobić coś bardziej przydatnego i wymyśliłem ten monitor prędkości wykorzystujący GPS. Urządzenia tego można używać w dowolnym pojeździe, takim jak samochód lub łódź. Pozwala ono na jazdę bez patrzenia na prędkościomierz i bez niemiłych finansowych „niespodzianek” w poczcie.

Prezentowane urządzenie nie jest ogranicznikiem prędkości, który wymagałby dokonania modyfikacji w pojeździe, ale alarmem sygnalizującym osiągnięcie (i przekroczenie) zadanej prędkości. Aby nie absorbować wzroku, urządzenie ma brzęczyk, który wyda dwa długie sygnały dźwiękowe w przypadku przekroczenia wybranej prędkości, natomiast dwa krótkie sygnały dźwiękowe zasygnalizują, że pojazd zwolnił do akceptowalnej prędkości.

Cztery ustawienia wstępne

Monitor prędkości ma cztery programowalne ustawienia ograniczeń prędkości, które powinny wystarczyć w przypadku większości, napotkanych w podróży, przypadków. Co więcej, ekran praktycznie nie pozwala na wyświetlenie więcej niż czterech limitów. Podczas jazdy kierowca powinien mieć możliwość szybkiego wyboru ustawienia wstępnego, w miarę możliwości bez patrzenia na urządzenie. Nie chciałem używać ekranu dotykowego – kierowca powinien fizycznie poczuć naciśnięcie przycisku. Ze względów ergonomicznych przyciski umieszczono nad czterema ograniczeniami prędkości wyświetlanymi na ekranie.

Dobry pretekst, aby spróbować wersji 32-bitowej

Moje pierwsze programy wykorzystujące 8-bitowy mikrokontroler ATmega328 skojarzony z modułem GPS szybko pokazały,

że jego 32 kB pamięci flash nie wystarczą do bardziej zaawansowanych aplikacji, szczególnie jeśli dane mają być wyświetlane graficznie na ekranie TFT. Dlatego też, pomimo wielu zalet mikrokontrolera ATmega328, zdecydowałem się przejść na wyższy poziom i wypróbować 32-bitowy mikrokontroler STM32. Najbardziej odpowiednią formą dla majsterkowiczów, takich jak ja, jest płytka BluePill z jej mikrokontrolerem STM32F103C8. Ta niewielka płytka znacznie ułatwia obsługę tego mikrokontrolera. Można ją łatwo zaprogramować poprzez port USB za pomocą Arduino IDE (pod warunkiem, że płytka zostanie wcześniej zaprogramowana za pomocą odpowiedniego bootloadera), jest kompaktowa i niedroga (cena to tylko kilka euro).

Rozmiary pamięci flash

Projekt ten jest pierwszą próbą zastosowania STM32F103C8. Należy zauważyć, że oficjalnie C8 jest wyposażony w 64 kB pamięci flash, podczas gdy wersja CB ma 128 kB, jednak czasami C8 ma aż 128 kB (podrabiane części?). Dzięki temu chętni będą mogli kontynuować naukę poprzez dodanie kolejnych urządzeń peryferyjnych, takich jak karta SD, moduł SIM, transmisja danych drogą radiową itp.

Moduł GPS

Moduł GPS zastosowany w tym projekcie to moduł NEO-6 firmy u-blox. Korzystanie z niego jest łatwe, ponieważ biblioteka TinyGPS+ [2] wykonuje całą pracę polegającą na dekodowaniu strumienia danych

w formacie NMEA 0183 wysyłanego przez moduł GPS. Główne dane jakie w ten sposób uzyskujemy to:

- data i godzina UTC;
- długość i szerokość geograficzna w stopniach dziesiętnych;
- prędkość;
- kierunek;
- wysokość;
- liczba widocznych satelitów;
- wartość HDOP.

Poziome rozmycie precyzji, czyli HDOP [2] zależy od pozycji satelitów widzianych przez odbiornik GPS. Im niższa jest ta wartość (blisko 1), tym lepsza jest precyzja współrzędnych. Wartość HDOP równa 10 wskazuje, że współrzędne nie są zbyt dokładne lub nawet nieprawidłowe. Monitor prędkości nie używa tego parametru.

Wyświetlacz TFT

Jako wyświetlacza użyłem zwykłego modułu TFT 2,2" z układem sterownika ILI 9341 i interfejsem SPI. Ma rozdzielczość 320×240 pikseli, co w zupełności wystarczy na dane, które chcemy wyświetlić (cztery wartości ograniczenia prędkości z krótkim komunikatem informującym o ewentualnym osiągnięciu (lub przekroczeniu) ograniczenia prędkości).

Monitor prędkości ma dwa główne tryby prezentacji wyników podczas normalnego użytkowania. Bieżący tryb wyświetlania jest zapisywany w pamięci EEPROM i przywoływany po ponownym włączeniu systemu. Wstępnie ustawione ograniczenia prędkości są również przechowywane w pamięci EEPROM. Używanych jest tylko sześć bajtów tej pamięci EEPROM: cztery bajty na ograniczenia prędkości, jeden bajt na ostatni wybrany ekran i jeden bajt na ostatnio wybrane ograniczenie prędkości.

Schemat układu

Schemat monitora prędkości pokazano na **rysunku 1**. Centrum sterowania stanowi tu moduł BluePill pokazany na jednym ze zdjęć poniżej. Zasilanie podłącza się poprzez złącze USB modułu BluePill. Nowsze pojazdy są często wyposażone w złącze USB; w przeciwnym razie wymagany jest adapter gniazda zapalniczkowego na USB (lub użycie powerbanku). Moduł BluePill zasilany jest z napięcia 5 V i ma własny stabilizator 3,3 V. Ten ostatni zasilą pamięć EEPROM napięciem 3,3 V. Ekran TFT i moduł GPS zasilane są napięciem 5 V. Całkowity pobór prądu przez układ wynosi około 200 mA.

Sterowanie wyświetlaczem odbywa się za pośrednictwem magistrali SPI, moduł GPS jest sterowany za pomocą łącza szeregowego, a pamięć EEPROM jest podłączona do magistrali I²C (z dwoma rezystorami podciągającymi). W moim prototypie użyłem pamięci EEPROM 24C256 o pojemności 32 kB, ale nawet 128-bajtowa pamięć 24C01 byłaby nawet więcej niż wystarczająco duża.

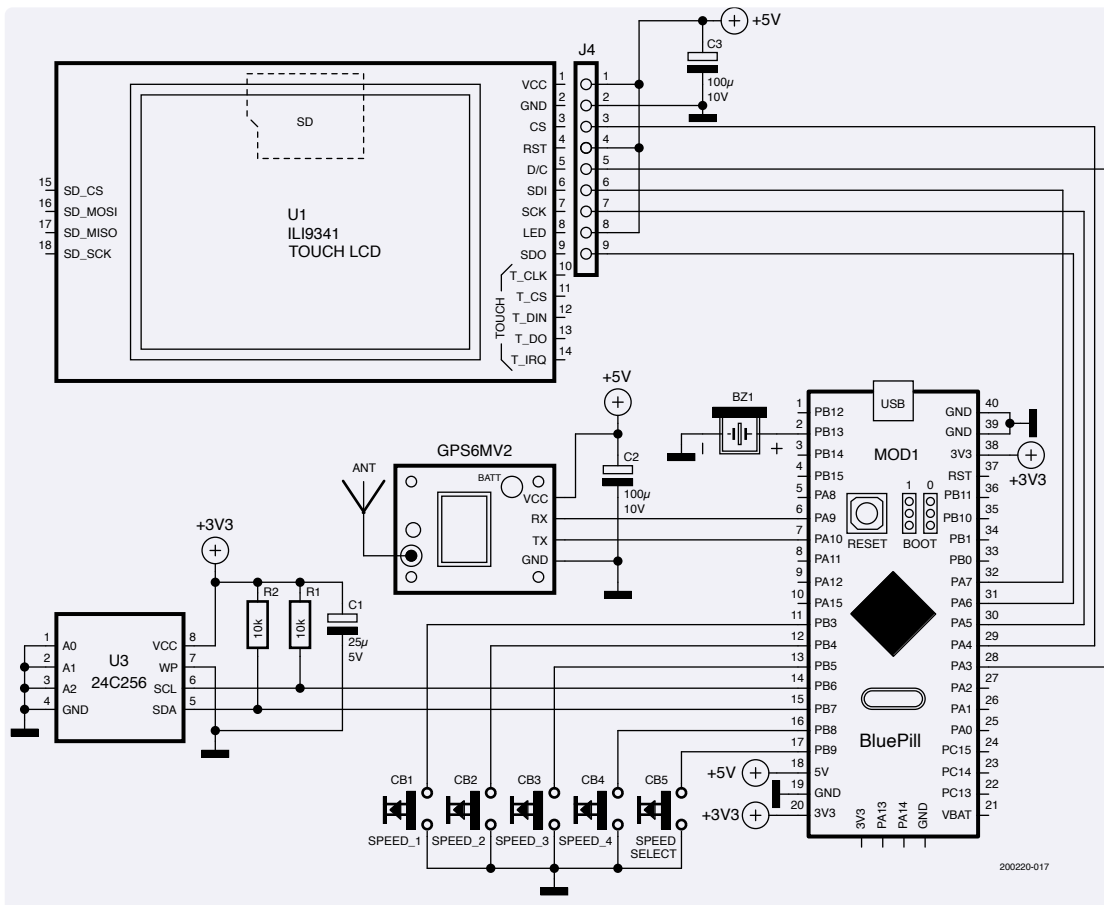
Wolnych pozostaje dwanaście portów GPIO, dzięki czemu możliwe jest dodanie większej liczby urządzeń peryferyjnych, możliwości i funkcji, o ile tylko program aplikacyjny zmieści się w pamięci flash. Skompilowane oprogramowanie tego projektu zajmuje 64 020 bajtów, co stanowi prawie całą pamięć flash mikrokontrolera.

Budowa monitora prędkości

Najważniejszym kryterium, które należy wziąć pod uwagę przy montażu tego projektu, jest jego ergonomia. Kierowca musi być w stanie znaleźć przyciski bez dwuznaczności. Przyciski mają popychacze o długości 25 mm – najdłuższe, jakie udało mi się znaleźć.

Prototyp ten wykonałem na płytce prototypowej (**rysunek 2**). Płytką drukowaną (PCB) umożliwiłaby bardziej zwartą realizację

Rysunek 1. Monitor prędkości łączy w sobie łatwo dostępne moduły z kilkoma elementami elektronicznymi



LINKI INTERNETOWE
 [1] TinyGPS++ library: <http://arduinoiana.org/libraries/tinygpsplus/>
 [2] Co to dokładnie jest HDOP?: [https://en.wikipedia.org/wiki/Dilution_of_precision_\(navigation\)](https://en.wikipedia.org/wiki/Dilution_of_precision_(navigation))
 [3] Pliki projektowe na stronie Elektor Labs: <https://elektormagazine.com/labs/save-money-with-this-speed-monitoring-by-gps>
 [4] STM32 Pakiet bibliotek płytek dla Arduino: https://github.com/stm32duino/BoardManagerFiles/raw/main/package_stmicroelectronics_index.json

w cieńszej obudowie. Złącze USB musi pozostać dostępne, ponieważ podczas normalnego użytkowania służy do zasilania układu. Z czasem konieczna może okazać się również aktualizacja oprogramowania.

Antena GPS powinna znajdować się jak najdalej od złącza wyświetlacza TFT (co najmniej kilka centymetrów).

Uwagi dotyczące modułu GPS

Moduł GPS wyposażony jest w małą baterię podtrzymującą dane. Jeśli GPS nie będzie używany przez kilka dni z rzędu, bateria rozładuje się, a dane zostaną utracone. Kiedy ponownie włączysz GPS, to przywrócenie danych (i naładowanie baterii) może zająć kilka minut. Dioda modułu GPS miga, gdy poprawnie dekoduje on dane i po chwili wyświetlana jest liczba satelitów.

Antena GPS musi „widzieć” satelity. Dlatego odbiór sygnałów musi być jak najmniej utrudniony. W obszarach zabudowanych np. pomiędzy dwoma budynkami, odbiór może być utrudniony. Niemniej jednak czujnik prędkości umieszczony u stóp pasażera w moim samochodzie działa prawidłowo.

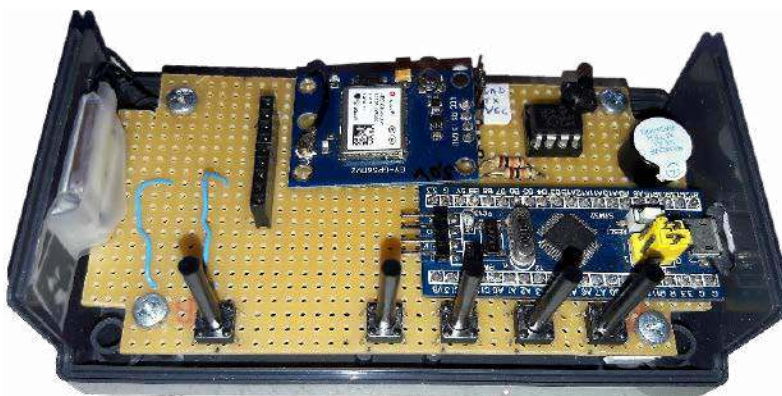
Przygotowanie ograniczeń prędkości

Czynność tę należy wykonywać wyłącznie na postoju pojazdu! Przy pierwszym włączeniu monitora prędkości wszystkie ograniczenia prędkości są ustawione na 255. Odpowiada to FF w formacie szesnastkowym, czyli zawartości pustej pamięci EEPROM.

Na ekranie z **rysunku 3**, który pokazuje dane z satelitów, naciśnij trzeci przycisk, aby aktywować ustawianie ograniczenia prędkości. Następnie za pomocą pierwszego i drugiego przycisku ustaw żądany limit. Naciśnij ponownie trzeci przycisk, aby zatwierdzić ograniczenie prędkości i przejdź do następnego, i tak dalej. Upewnij się, że limity są ułożone w logicznej kolejności, ponieważ oprogramowanie ich nie sortuje.

Korzystanie z monitora prędkości

Monitor prędkości jest łatwy w użyciu. Jednak upewnij się, że urządzenie, a w szczególności jego kabel zasilający, nie poruszają się podczas jazdy. Po umieszczeniu go w odpowiednim miejscu wybierz żądany tryb wyświetlania. **Rysunek 4** pokazuje ekran z okrągłymi znakami drogowymi ograniczeń prędkości. **Rysunek 5** pokazuje ekran, który



Rysunek 2. Prototyp zbudowany na płycie uniwersalnej. Wymiary obudowy to 12×6,5×4 cm

wyświetla aktualną prędkość w postaci cyfr 7-segmentowych. Wybierz ograniczenie prędkości, którego chcesz użyć. Jeśli jedziesz za szybko, zabrzmi brzęczyk.

Pasażer może wyświetlić przebieg zmian prędkości z ostatniego czasu – 4,5 minuty (270 sekund), wybierając ekran pokazany na **rysunku 6**. Na tym ekranie brzęczyk się nie włącza. Linie poziome (żółte) pokazują cztery ograniczenia prędkości. Pionowa linia (zielona) wskazująca „teraz” pokazuje aktualną prędkość. Przewijany wykres mógłby być ładniejszy, ale wymaga pełnego odświeżenia dla każdego nowego punktu danych i jest powolne. Poruszanie kursorem jest dużo łatwiejsze.

Ekran danych GPS

Na **rysunku 7** pokazano ekran danych GPS, który pozwala zlokalizować się na mapie za pomocą pokazywanych współrzędnych. Dokładność wynosi około 30 metrów, co jest całkiem niezłe jak na nasze małe i niedrogie urządzenie. W tym trybie wyświetlania brzęczyk się nie włącza.

Należy pamiętać, że data i godzina GPS są wartościami UTC, tj. na południku Greenwich. Jeśli chcesz z nich skorzystać, powinieneś przekonwertować te wartości na lokalną strefę czasową i czas letni/zimowy. Ponieważ Monitor prędkości nie jest zegarem, nie zapewnia do tego wygodnego menu.

Oprogramowanie

Program jest łatwy w modyfikacji, a kod źródłowy można pobrać online [3]. Jest to szkic Arduino i wymaga pakietu płytek STM32 dla Arduino (użyliśmy wersji 2.3.0) [4]. Lwia część kodu przeznaczona jest na interfejs graficzny użytkownika, reszta zajmuje się prostym zadaniem

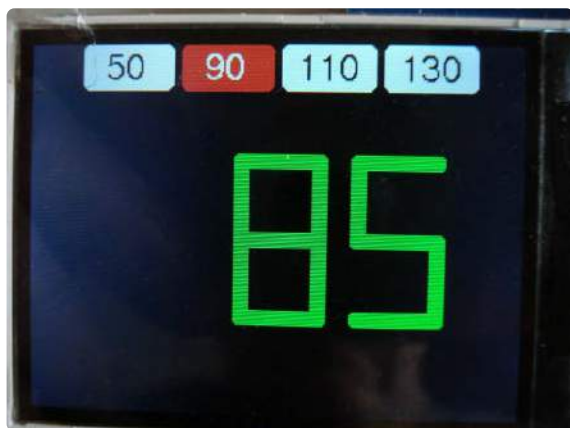


Rysunek 3. Dane GPS są dostępne na tym ekranie. Jest to również ekran zapewniający dostęp do definicji ograniczeń prędkości



Rysunek 4. Ekran, który pokazuje ograniczenia prędkości w postaci okrągłych znaków drogowych wraz z krótkim komunikatem

Tabela 1. Różne ekrany i nawigacja między nimi					
Ekran	Przycisk 1	Przycisk 2	Przycisk 3	Przycisk 4	Przycisk 5
1. okrągłe znaki ograniczeń	Ustaw Limit 1	Ustaw Limit 2	Ustaw Limit 3	Ustaw Limit 4	Idź do ekr. 2
2. wyświetlacz 7-segm. prędkości	Ustaw Limit 1	Ustaw Limit 2	Ustaw Limit 3	Ustaw Limit 4	Idź do ekr. 3
3. dane GPS	Ustaw Limit 1	Ustaw Limit 2	Idź do ekr. 4		Idź do ekr. 1
4. ustawianie limitów	Limit - (5 km/h)	Limit + (5 km/h)	Zatwierdzenie, przejście do następnego, powrót do ekr. 3	Idź do ekr. 5	
5. wykres prędkości				Idź do ekr. 2	



Rysunek 5. Na tym ekranie prędkość rzeczywista jest prezentowana na symulowanym wyświetlaczu 7-segmentowym



Rysunek 6. Ekran, który pokazuje przebieg zmian prędkości. Cztery żółte linie poziome pokazują cztery zaprogramowane limity, a zielona linia pionowa pokazuje aktualną prędkość



Rysunek 7. Ten ekran umożliwia ustawienie ograniczeń prędkości. Tutaj ustawiane jest pierwsze ograniczenie „Sp. 1” (na czerwono, jeśli się przyjrzyjiesz)

odbioru danych GPS, wyodrębnieniem prędkości i porównania jej z aktualnym limitem.

Jak zwykle w szkicu Arduino, funkcja `setup()` zajmuje się inicjalizacją wszystkich urządzeń peryferyjnych. Po zakończeniu wyświetla ekran powitalny przez pięć sekund. Funkcja `loop()` rozpoczyna się od odczytania stanu przycisków przed sprawdzeniem GPS pod kątem świeżych danych (w funkcji `dataDecode()`). Następnie, w zależności od tego, który ekran jest aktywny, wyświetlane są odpowiednie dane. Jeśli aktualna prędkość jest wyższa niż wybrane ograniczenie prędkości, włącza się alarm.

Należy pamiętać, że logi systemowe są wysyłane przez port szeregowy (115200,N,8,1).

Dalsze możliwości

Oto kilka rzeczy, które możesz chcieć dodać lub zmienić:

Użyj innych jednostek prędkości. Biblioteka GPS umożliwia korzystanie z innych jednostek prędkości, takich jak węzły lub mile na godzinę (MPH). Będzie to wymagało także zmiany ekranów (np. zamiany km/h na węzły) i wielkości kroku przyrostów limitu (w funkcji `SpeedSettings()`). Może to być przydatne w przypadku łodzi.

Jeżeli Twój moduł BluePill ma 128 kB pamięci flash, można dodać dodatkowe funkcje, takie jak:

Nagrywanie na kartę SD. Na tylnej ścianie wyświetlacza TFT znajduje się gniazdo kart SD. Można być użyte do rejestrowania, na przykład, danych GPS w celu uzyskania dostępu do nich z poziomu komputera PC;

Ekran dotykowy dla nowych funkcji. W tym przypadku, pamiętaj o bezpieczeństwie użytkownika!

Ustawianie czasu lokalnego. ■

Olivier Croiset (Francja)

Pytania lub komentarze?

Masz jakieś pytania techniczne lub uwagi do tego artykułu? Skontaktuj się z redakcją Elektora pod adresem editor@elektor.com lub z redakcją EdW edw@elportal.pl

REKLAMA

Certyfikat Underwriters Laboratories
UL 94V-0 E480140 TYPE 1

Zakład produkcyjny:
05-660 Warka
ul. M. Ropielewskiej 17
tel. 22 781 63 95
22 761 95 40
fax. 22 781 63 95 w.23
www.elmax.waw.pl
elmax@elmax.waw.pl

OBWODY DRUKOWANE

Produkcja, Projektowanie, Montaż

Płytki jednostronne	Serie dowolne	Dokumentacja technologiczna	Montaż elektroniczny
Płytki dwustronne	Prototypy	Dokumentacja konstrukcyjna	Ilości modelowe produkcyjne
Płytki na podłożu aluminium	Najmniejszy wymiar płytek 1w 630 mm		
Aktywny kalkulator prototypów na stronie internetowej	Pokrycie Sn lub SnPb linie na życzenie	Płyty czolowe FR4	Krótkie terminy
	Maki, opisy montażowe w różnych kolorach	Trawione szablony SMD	Wykonania super ekspresowe

Monitorowanie do 3 akumulatorów od 6 do 100 V
– prąd do 10 A (lub 100 A+ z bocznikiem)

Wielofunkcyjny monitor akumulatorów z dotykowym ekranem LCD, część 2

W zeszłym miesiącu, w 1 części naszego nowego Battery Multi-Logger'a (czyli wszechstronnego monitora akumulatorów, w skrócie MB) opisaliśmy, w jaki sposób łączy on funkcje modułu Micromite LCD BackPack'a z elementami do pomiaru napięcia i prądu oraz technikami oszczędzania energii, a wszystko to na jednej płytce drukowanej. Teraz omówimy budowę, testowanie, konfigurację i procedury kalibracji, abyś mógł go zbudować i używać.

Zanim przejdziemy do instrukcji montażu, przejrzymy szybko możliwości rejestratora.

- Może on obsługiwać baterie akumulatorów o napięciu od 6 do 100 V i monitorować do trzech dwukierunkowych prądów o natężeniu do 10 A za pomocą wbudowanych boczników lub znacznie więcej (do 100 A lub ponad) za pomocą zewnętrznych boczników.
- Zużycie energii podczas aktywnego rejestrowania wynosi przy wyłączonym ekranie mniej niż 1 mA.
- Może wyświetlać na 2,8-calowym podświetlanym ekranie dotykowym LCD bieżące i historyczne dane, a dane można również pobrać do komputera przez port USB w celu dalszej analizy.
- Śledzi aktualny stan naładowania akumulatora zarówno w amperogodzinach (Ah), jak i watogodzinach (Wh), a rozdzielczość pomiaru prądu wynosi około 0,1% pełnej skali, co odpowiada krokom 10 mA przy użyciu wewnętrznych boczników.

Wszystkie te funkcje zostały umieszczone na niewielkiej płytce drukowanej. Ponieważ wszystkie funkcje interfejsu użytkownika są dostępne za pośrednictwem ekranu dotykowego, moduł MB można łatwo zintegrować z innymi urządzeniami poprzez dodanie prostokątnego wycięcia w obudowie.

Budowa

Monitor baterii – MB – jest zbudowany na dwustronnej płytce drukowanej o kodzie 11106201 i wymiarach 86×50 mm. **Rysunek 5** pokazuje rozmieszczenie komponentów po obu stronach płytki.



Jak zwykle przy montażu płytki z wieloma elementami SMD, warto mieć pod ręką: pastę topnikową, miedzianą plecionkę lutowniczą, łupę, pęsetę i lutownicę o regulowanej temperaturze. Najmniejsze części mają odstępy między stykami poniżej 1 mm, więc mostki lutownicze są prawie nieuniknione, stąd potrzeba pasty topnikowej i plecionki lutowniczej.

Ponieważ topnik ma tendencję do generowania dymu, użyj odciągu oparów lub pracuj na zewnątrz, gdzie dym może łatwiej się rozprzyszczyć.

Jedną z najbardziej skomplikowanych części jest gniazdo USB, CON5, więc zacznij

od jego zamontowania. Nałóż topnik na pola lutownicze, a następnie umieść gniazdo USB na miejscu; powinno zablokować się w otworach w płytce drukowanej. Dodaj więcej topnika na wierzchołki styków.

Dodaj lut na czysty grot lutownicy, a następnie dociśnij grot do małych styków i pól montażowych. Metalowa osłona gniazda trochę przy tym przeszkadza.

Gdy upewnisz się, że przylutowałeś wszystkie wyprowadzenia, sprawdź, czy nie ma mostków i usuń je, jeśli to konieczne, a następnie przylutuj większe wypustki na osłonie.

Układy scalone

Następnie należy przylutować układy scalone (IC1-IC6 i REF1, z tyłu PCB). Sugerujemy, aby najpierw zamontować układ IC5, ponieważ ma on najdrobniejszy rozstaw styków.

Dla każdego z układów scalonych należy sprawdzić orientację wyprowadzenia 1 w stosunku do sitodruku na PCB, dopasowując kropkę przed przylutowaniem jakichkolwiek styków.

Układ IC6 jest asymetryczny, więc chociaż ta część jest mała, łatwo jest ją prawidłowo zorientować. Należy pamiętać, że niektóre układy scalone mogą nie mieć kropki wskazującej styk 1. Zamiast tego będą miały fazę wzdłuż jednej krawędzi lub linię na jednym końcu; w każdym przypadku to wyróżnienie znajduje się najbliżej styku 1. W przypadku REF1 wskaźnik styku 1 może być nawet małym krzyżykiem wytrawionym laserowo.

Podczas lutowania układów scalonych nałóż topnik na pola stykowe, a następnie umieść układ scalony na miejscu i przylutuj jedno wyprowadzenie. Sprawdź pozycjonowanie, upewniając się, że część jest płaska i wyrównana na podłożu. Jeśli nie, ponownie roztop lut i wyreguluj część za pomocą pęsety.

Po prawidłowym umieszczeniu części przylutuj pozostałe styki. Nie przejmuj się mostkami lutowniczymi, ponieważ łatwiej jest usunąć wiele mostków później, wszystkie w tym samym czasie. W razie potrzeby, podczas lutowania, nałóż dodatkową porcję topnika.

Aby usunąć mostki, nałóż świeży topnik i dociśnij miedziany oplot lutowniczy za pomocą lutownicy do nadmiaru lutu. Gdy spoiwo się roztopi, pozwól miedzianej plecionce wchłonąć lut, a następnie delikatnie odciągnij ją od elementu.

Napięcie powierzchniowe między komponentem a polem lutowniczym powinno utrzymać wystarczającą ilość lutu, aby zachować dobre połączenie, nawet jeśli oplot lutowniczy usunie większość lutu.

Teraz jest dobry moment na dokładne sprawdzenie Twojej pracy za pomocą lupy, ponieważ wprowadzanie zmian będzie trudniejsze w miarę dodawania kolejnych części.

Dobrym pomysłem jest najpierw wyczyszczenie nadmiaru topnika; alkohol izopropylowy jest dobrym rozpuszczalnikiem, ale specjalistyczne produkty do usuwania topnika często wykonują lepszą robotę.

Tranzystor i regulatory

Kolejnymi najtrudniejszymi częściami są tranzystory i regulatory napięcia w obudowach SOT-23. Istnieje sześć takich części w trzech typach obudowy: Q1 i Q3 (P-kanalowe



Multi-Logger może być zamontowany w obudowie UB5 Jiffy Box, jak w wielu projektach wykorzystujących Micromite i jak pokazano tutaj. Możesz jednak użyć ramki do zamontowania Multi-Loggера na przednim panelu obudowy sprzętu; możesz wtedy użyć obudowy Jiffy Box do ochrony tylnej części urządzenia

MOSFET-y), Q2 i Q4 (N-kanalowe MOSFET-y) oraz REG1 i REG2 (regulatory LDO).

Na szczęście pasują one tylko w jeden sposób, więc użyj podobnej techniki jak w przypadku układów scalonych. Przylutuj jedno wyprowadzenie i sprawdź jego położenie przed przylutowaniem pozostałych styków. Pozostałe układy SMD mają znacznie większe podkładki lutownicze, więc znacznie łatwiej poradzisz sobie z nimi.

Rezystory i kondensatory

Wiele z pozostałych części to rezystory i kondensatory o rozmiarze 1206 (3,2×1,6 mm). Rezystory powinny być oznaczone ich wartościami, podczas gdy kondensatory zazwyczaj nie są, więc zachowaj szczególną ostrożność z kondensatorami i nie mieszaj ich. Zalecamy pracę z jedną wartością na raz.

Tam, gdzie to możliwe, oznaczyliśmy na PCB wartości rezystorów i kondensatorów poniżej nadruku samej części; wyjątkiem są części wokół IC4.

Pamiętaj, że jeśli używasz zewnętrznych boczników do wykrywania prądu, pomijasz trzy rezystory boczników 15 mΩ. Większe rezystory bocznikowe należy na razie odłożyć na bok, nawet jeśli zamierzasz je zamontować. W przypadku pozostałych części porównaj wartość wydrukowaną na sitodruku z wartością na części, która będzie najczęściej podana w postaci kodu numerycznego, który można dopasować do elementów z naszej listy części.

Dla każdej części nałóż topnik na pole lutownicze, przylutuj jedno wyprowadzenie, sprawdź i wyreguluj w razie potrzeby, a następnie przylutuj drugie wyprowadzenie. W razie potrzeby odśwież pierwszy styk.

Większość kondensatorów to typy 100 μF, 10 μF lub 100 nF, więc zalecamy umieszczenie ich w pierwszej kolejności. Kondensatory 100 μF i 10 μF będą najprawdopodobniej większe, więc ich rozróżnienie nie będzie zbyt trudne. Wszystkie cztery typy 100 μF są zamontowane z tyłu płytki drukowanej.

Użyj tej samej metody, co w przypadku rezystorów. Postępuj zgodnie z pozostałymi kondensatorami, zwracając uwagę na ich wartość przed wyjęciem z opakowania i pracując pojedynczo.

Istnieją dwie małe cewki (L2 i L3), które również mają wymiary 1206; są one lutowane w podobny sposób.

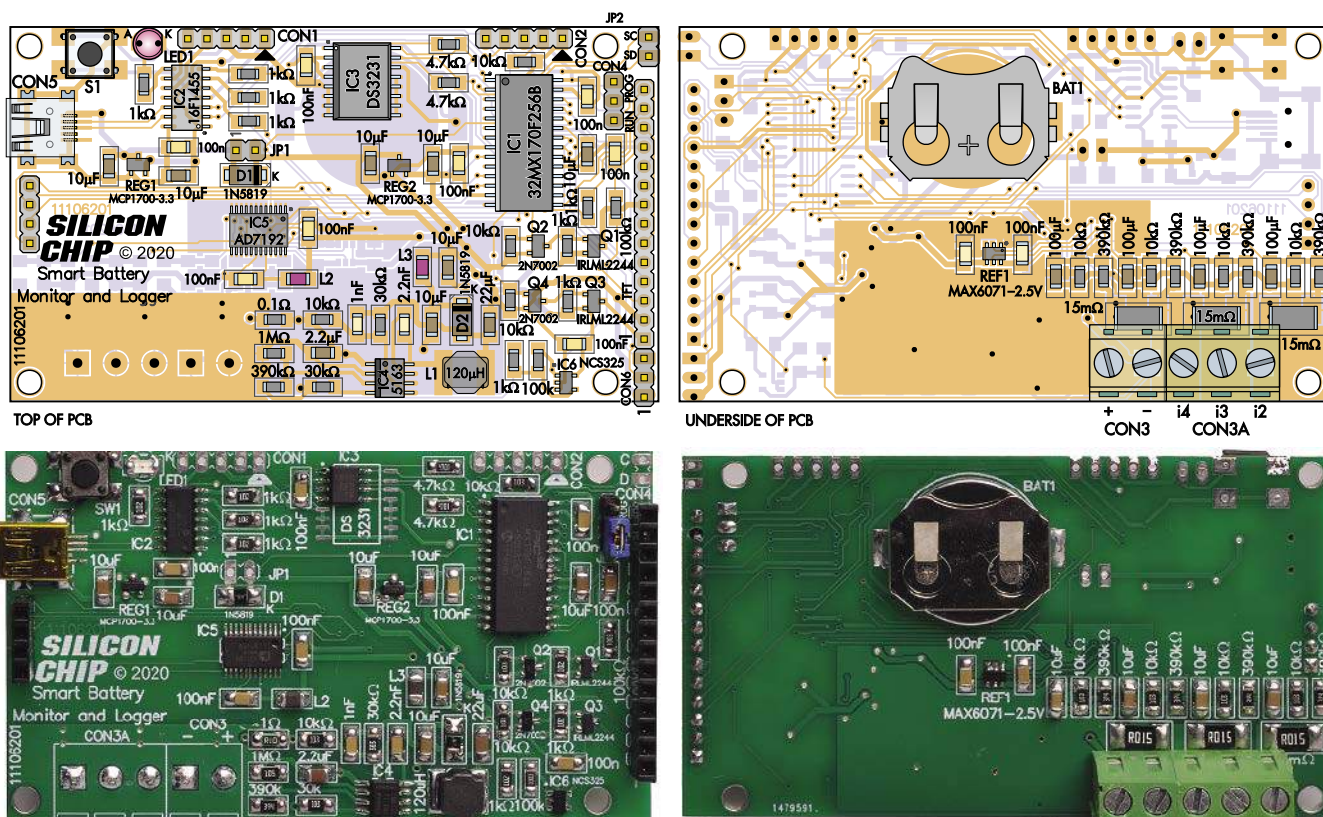
Większy dławik 120 μH (L1) może wymagać gorętszej lutownicy do lutowania. Użyj tej samej techniki pracy na jednym styku na raz. Czasami można uzyskać lepsze przenoszenie ciepła, naciskając na zestyk długą krawędzią grotu lutownicy. Następnie przylutuj drugie wyprowadzenie.

Kolejno przylutuj uchwyt ogniwa pastylkowego. Ponownie, może być konieczne podkreślenie temperatury lutownicy, aby dostarczyć więcej ciepła. Dodaj topnik do styków i umieść uchwyt tak, aby ogniwo można było włożyć od krawędzi płytki drukowanej.

Przylutuj jeden styk, a gdy będziesz z niego zadowolony, przylutuj drugi. Odśwież pierwszy styk, aby zmniejszyć naprężenia na polach lutowniczych PCB. Sprawdź nasze zdjęcia, aby zobaczyć, jak to powinno wyglądać.

A reszta elementów?

Istnieją dwie diody montowane powierzchniowo; obie są zamontowane z katodami skierowanymi w stronę REG2 (ponieważ to właśnie one są zasilane!).



Rysunek 5. Schemat montażowy rejestratora. Zdjęcia PCB pokazane powyżej pochodzą z wczesnego prototypu, więc różnią się nieco od płytek w naszej ostatecznej wersji, w tym aktualnymi wartościami komponentów. Komponenty znajdują się po obu stronach płytki, chociaż tylna część płytki jest znacznie mniej zabudowana. Należy zwrócić szczególną uwagę na orientację wszystkich układów scalonych, dwóch diod i diody LED. Większość pozostałych komponentów nie jest kierunkowa

W przypadku diod LED1 i S1 można użyć elementów do montażu powierzchniowego lub przewlekane. Zamontuj je w następnej kolejności. Katoda diody LED1 jest skierowana w prawo, w stronę CON1. Większość diod LED do montażu powierzchniowego ma katodę oznaczoną zieloną kropką, ale należy to dokładnie sprawdzić, ponieważ niektóre nie mają takiego oznaczenia.

Na tym etapie praktycznie wszystkie diody SMD zostały zamontowane, więc jest to dobra okazja do usunięcia nadmiaru topnika pozostałego na płytce drukowanej.

JP1 nie jest zwykle potrzebny, więc można go pominąć (użyliśmy go w naszych testach), ale JP2 jest wymagany. Załóż zwórkę na kołki, aby ułatwić manipulację i przylutuj jedną szpilkę. Sprawdź, czy listwa jest prostopadła do PCB i przylega doń, a następnie przylutuj pozostałe szpilki.

Jeśli masz wstępnie zaprogramowane mikroprocesory (IC1 i IC2), załóż zwórkę JP2 na dwóch dolnych szpilkach (jak widać na naszym zdjęciu). Jest to pozycja „RUN”. Jeśli chcesz zaprogramować IC1, dopasuj zwórkę do dwóch górnych szpilek (w pobliżu otworu montażowego PCB).

Do programowania wystarczy zamontować CON1, ponieważ IC2 może zaprogramować

IC1. Ale jeśli masz programator, może okazać się szybsze i łatwiejsze bezpośrednie programowanie obu mikroprocesorów.

Aby ułatwić debugowanie, użyliśmy kątowych listew kołkowych dla CON1 i CON2 (brak na zamieszczonych fotografiach), ale listwy proste również będą działać i zmieszczą się pod wyświetlaczem LCD.

Połączenia dla 2,8-calowego wyświetlacza LCD składają się z żeńskich gniazd jednorzędowych: 4-stykowego i 14-stykowego. W aktualnej wersji oprogramowania potrzebne jest tylko gniazdo 14-stykowe, chociaż wlutowanie obu gniazd sprawi, że montaż będzie bardziej wytrzymały.

Użyj 2,8-calowego wyświetlacza LCD jako przyrządu do montażu gniazd. Może być konieczne przylutowanie krótszych styków listew kołkowych do wyświetlacza LCD, jeśli nie są one wstępnie zainstalowane; większość ekranów nie jest dostarczana z zamontowaną listwą 4-kołkową. W takim przypadku należy podłączyć listwy kołkowe do gniazd i włożyć je do odpowiednich pól płytek drukowanych. Gniazda znajdują się na naszej płytce drukowanej, a listwy kołkowe po stronie LCD.

Przylutuj gniazda na miejscu, utrzymując PCB równolegle do ekranu. Następnie

delikatnie oddziel wyświetlacz LCD od płytki drukowanej, poruszając nim w razie potrzeby.

Ostatnim krokiem w montażu PCB jest dopasowanie gniazd zaciskowych CON3 i CON3A, połączeń baterii i obciążenia. Zamontuj je od spodu PCB, aby umożliwić dostęp do nich nawet po zmontowaniu stosu z ekranem. Upewnij się, że zamontowałeś trzy większe boczniki 15 mΩ, jeśli nie będziesz używać zewnętrznych boczników.

Programowanie

Jeśli masz wstępnie zaprogramowane układy scalone, nie musisz się martwić o ten krok i powinieneś przejść do sekcji konfiguracji.

Zarówno IC1, jak i IC2 wymagają do działania oprogramowania układowego. Jedynym sposobem na zaprogramowanie IC2 w obwodzie jest użycie złącza ICSP CON1 i programatora, takiego jak PICkit 3 lub PICkit 4.

Można użyć MPLAB X IDE (zintegrowane środowisko programowania), które jest dostępne do pobrania jako część pakietu MPLAB X ze strony www.microchip.com/en-us/tools-resources/develop/mplab-x-ide.

Wybierz PIC16F1455 jako chip docelowy i programator z rozwijanej listy Tool. Podłącz programator do CON1 zgodnie z jego instrukcjami i wyszukaj plik .hex Microbridge

(2410417A.hex). Następnie naciśnij przycisk Programm, aby go załadować.

Po otwarciu IPE można go również użyć do przesłania oprogramowania układowego dla IC1. Podłącz programator do CON2, wybierz PIC32MX170F256B jako chip docelowy i wyszukaj plik 1110620A.hex. Prześlij ten plik za pomocą przycisku Programm.

Po zakończeniu programowania nie zapomnij przesunąć zworki JP2 do pozycji RUN (dolnej).

Microbridge i MMBasic

Jeśli masz ochotę majstrować przy kodzie BASIC, możesz zaprogramować IC1 również za pomocą plików MMBasic, chociaż jest to nieco bardziej skomplikowane.

Przedstawimy kroki, zakładając, że masz już pewne doświadczenie ze środowiskiem Micromite, znasz się dość dobrze z MMBasic'em i czujesz się komfortowo przesyłając pliki do Micromite. Jeśli nie chcesz tego robić, przejdź do następnej sekcji.

Będziesz potrzebował oprogramowania Microbridge na IC2, więc zacznij od zworki JP2 w pozycji PROGRAMM, ponieważ wymaga (przynajmniej) pliku .hex dla środowiska BASIC, który należy najpierw przesłać do IC1.

Można to zrobić za pomocą PICKit i IPE (jak opisano powyżej), ale zamiast oprogramowania układowego rejestratora baterii należy wybrać najnowszy plik oprogramowania układowego Micromite MMBasic.

Alternatywnie, firmware MMBasic można wgrać przez Microbridge, naciskając S1 (aby wejść w tryb programowania). Następnie użyj programu takiego jak pic32prog lub P32P GUI, aby załadować plik .hex Micromite MMBasic. Użyliśmy wersji 5.5.2.

JP2 można teraz przesunąć do pozycji RUN. Ze środowiska BASIC (port szeregowy działający z prędkością 38 400 bodów) należy uruchomić polecenia, aby skonfigurować 2,8-calowy wyświetlacz LCD i panel dotykowy, jak zwykle w przypadku V2 Micromite Backpack.

OPTION LCDPANEL ILI9341,

LANDSCAPE, 2, 23, 6

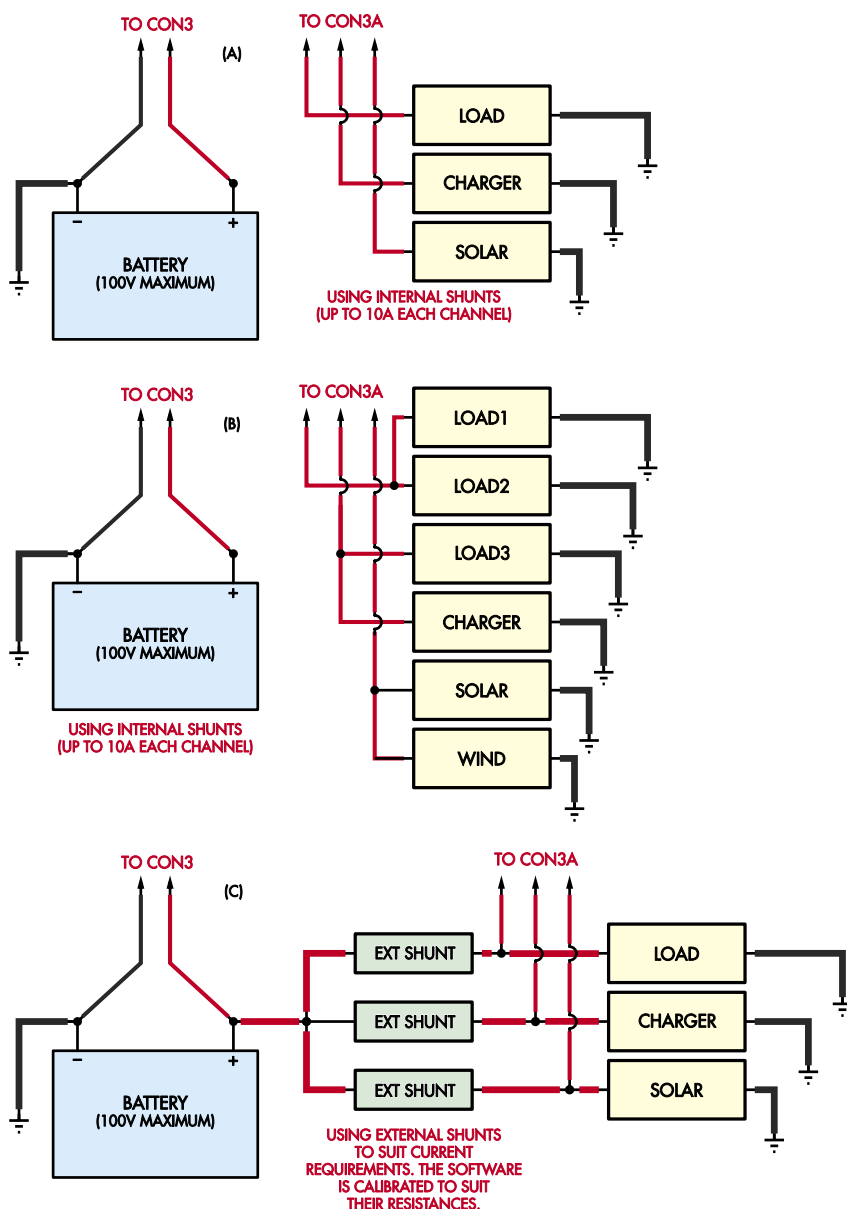
OPTION TOUCH 7, 15

GUI CALIBRATE

Pliki BASIC są ułożone jako plik biblioteki uzupełniający główny kod źródłowy. Dzięki temu Micromite może kompresować niektóre z używanych danych. Załaduj plik library.bas, a następnie uruchom polecenie:

LIBRARY SAVE

Spowoduje to zapisanie i skompresowanie pliku biblioteki. Następnie załaduj główny plik Battery Logger.bas i uruchom go. Instrukcje te znajdują się w pliku library.bas.



Rysunek 6. Ten schemat przedstawiliśmy w zeszłym miesiącu, aby pokazać możliwości rejestratora. Powtarzamy go teraz, ponieważ możesz chcieć użyć go jako przewodnika podczas podłączania. Podczas korzystania z wewnętrznych boczników, akumulator łączy się przez CON3, a dodatnie końce obciążen lub ładowarek idą do zacisków CON3a. Wszystkie ujemne bieguny obciążenia i ładowarki są podłączone bezpośrednio do akumulatora. W przypadku korzystania z zewnętrznych boczników należy postępować zgodnie ze schematem (C) i upewnić się, że przewody od akumulatora do boczników są krótkie i grube, aby zapewnić maksymalną precyzję

Konfiguracja i obsługa

Jeśli jeszcze tego nie zrobiłeś, zamontuj ogniwo CR2032 w uchwycie BAT1, zamontuj panel LCD i podłącz rejestrator do komputera lub zasilacza USB przez CON5. Jeśli zaprogramowałeś IC1 za pomocą pliku .hex specyficznego dla tego projektu, oprogramowanie rejestratora powinno uruchomić się od razu. Jeśli samodzielnie załadowałeś pliki BASIC, może być konieczne ręczne uruchomienie programu po raz pierwszy.

Podczas uruchamiania powinien pojawić się **ekran 1**. Komunikat o błędzie może pojawić się przez pierwsze kilka sekund, podczas gdy

program czeka na prawidłowy odczyt baterii; jeśli nie zniknie po około dziesięciu sekundach, może to oznaczać problem z układem IC5. Napięcie wyświetlane po „V=” powinno wynosić zero, ponieważ bateria nie jest jeszcze podłączona.

Możesz jednak zobaczyć kilka odczytów wartości prądu, ponieważ nie zakończyliśmy jeszcze kalibracji. I1 odpowiada poborowi prądu przez rejestrator, podczas gdy I2-I4 to prądy mierzone na zaciskach CON3A, jak pokazano na rysunku 5. Wartości te mogą nieco skakać, ale najważniejsze są średnie długoterminowe.



Ekran 1. Ekran główny zapewnia prezentację wszystkich krytycznych statystyk dotyczących akumulatora, a także trzy proste opcje menu umożliwiające dostęp do innych funkcji. Wartości pokazane na szaro to obliczenia pojemności akumulatora, które nie są jeszcze prawidłowe, ponieważ rejestrator nie wykrył pełnego cyklu ładowania i rozładowania; zaświecą się jaśniej, gdy to nastąpi

Po prawej stronie znajdują się pomiary pojemności i stanu naładowania. CHGv% to prosta liniowa zależność pomiędzy nominalnym napięciem pełnego i rozładowanego akumulatora, podczas gdy CHGm% wykorzystuje zmierzony prąd od ostatniego stanu pełnego naładowania i rozładowania.

Odczyt CHGm% nie będzie w pełni dokładny, dopóki akumulator nie przejdzie pełnego cyklu ładowania i rozładowania. Podobnie odczyty pojemności nie będą od razu miarodajne.

W prawym górnym rogu znajduje się licznik czasu; gdy osiągnie zero, wyświetlacz zgaśnie. Jest to normalny tryb, w którym rejestrator zapisuje dane, ale nie musi niczego wyświetlać, oszczędzając w ten sposób energię. Licznik czasu świecenia ekranu można zresetować, dotykając dowolnego miejsca na ekranie głównym.

Ten limit czasu ma miejsce tylko na ekranie głównym pokazanym na Ekranie 1, więc należy do niego powrócić za każdym razem, gdy kończy się dostęp do interfejsu graficznego rejestratora.

Aby ponownie aktywować ekran, naciśnij i przytrzymaj panel dotykowy, aż podświetlenie się zaświeci. Aby zapewnić maksymalną wydajność energetyczną, Micromite sprawdza panel tylko w jedno-sekundowych odstępach, więc wybudzenie może zająć około sekundy. Rejestrator czeka na zwolnienie dotyku przed wyświetleniem ekranu głównego, więc nie można przypadkowo nacisnąć jakiegoś przycisku podczas wybudzania.

Interfejs jest dość intuicyjny, ale i tak przejdziemy przez różne ekrany. **Ekran 2** jest dostępny po naciśnięciu przycisku Dane

i wyświetla wykres napięcia i prądów. Skalę prądu (po lewej stronie) można ustawić ręcznie, podczas gdy skala napięcia wykorzystuje nominalne wartości pełnego naładowania i rozładowania. Domyślnie są one ustawione na 14,4 V i 11,0 V, aby pasowały do akumulatora kwasowo-ołowiowego 12 V.

Przyciski znajdujące się na dole strony umożliwiają wyświetlanie różnych skal, z ramami czasowymi wyświetlanymi na dole ekranu. Na każdej skali przycisk Export powoduje rzut danych do portu szeregowego. Dane te są tworzone w taki sposób, aby można je było zapisać jako plik CSV (wartości rozdzielana przecinkami), a następnie otworzyć w większości programów do obsługi arkuszy kalkulacyjnych. Naciśnięcie przycisku Exit powoduje powrót do ekranu głównego.

Ekran 3 jest dostępny za pomocą przycisku Ustawienia. Każdą wyświetlaną wartość można zmienić, naciskając odpowiedni przycisk.

Ekran 4 pokazuje wprowadzaną liczbę, w tym przykładzie w celu aktualizacji bieżącego roku. Jeśli wprowadzona liczba jest nieprawidłowa, wyświetlany jest komunikat. Naciśnięcie przycisku OK spowoduje wyświetlenie monitu o potwierdzenie nowej wartości (patrz **ekran 5**).

Ustawienia godziny i daty są natychmiast zapisywane w zegarze czasu rzeczywistego i wyświetlane na tym ekranie oraz na ekranie głównym. Dwie wartości B/L określają jasność podświetlenia w procentach, w zakresie 1...100. Pierwsza wartość (B/L) jest używana przez większość czasu.

Druga wartość (B/L dim) jest używana przez ostatnie pięć sekund przed wyłączeniem



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

ekranu, aby wskazać, że ma to nastąpić. Minimalna wartość 1% jest narzucona dla obu ustawień, aby zapewnić, że wyświetlacz jest zawsze widoczny.

Wartości V(full) i V(empty) powinny być ustawione tak, aby pasowały do konkretnej baterii/akumulatora. Nie można ustawić wartości V(empty) wyższej niż wartość V(full).

Wartość Timeout określa, jak długo wyświetlacz pozostaje włączony przed wygaszeniem ekranu głównego. Wartość ta wynosi co najmniej pięć sekund, ponieważ jest to okres ściemniania, który występuje przed wygaszeniem.

Duża wartość może być użyta do zatrzymania wygaszania wyświetlacza; np. zapis 99 999 999 sekund oznacza, że czas świecenia wyświetlacza wynosi około trzech lat.

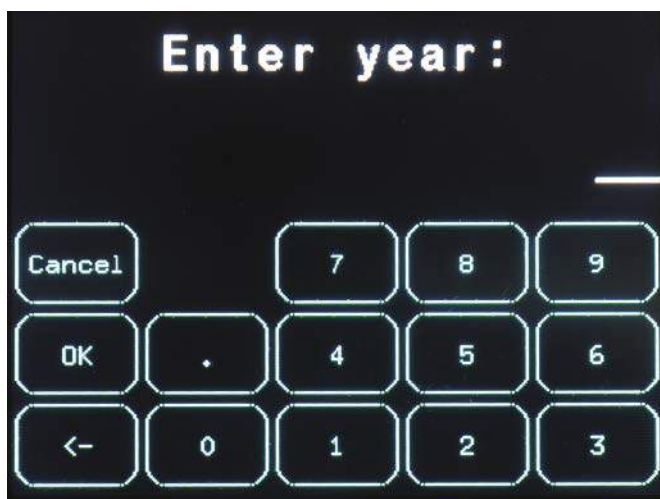
Wartość „I scale” ustawia granice wykresu tylko na stronie danych. Ustawienie wartości 20 spowoduje, że wykres będzie obejmował zakres od -20 A do 20 A.

Wartość „V(sdown)” ustawia krytyczny limit baterii. Poniżej tego poziomu rejestrator przechodzi w stan uśpienia na znacznie dłuższe okresy pomiędzy kolejnymi aktywnościami. Kod MMBasic ustawia tę wartość na 15 sekund. Ponieważ przetwornik ADC (IC5) przechodzi w stan uśpienia po każdej konwersji, pobór prądu spada nawet poniżej normalnego trybu „wyłączonego ekranu”.

To ustawienie ma na celu oszczędzanie baterii, która jest już mocno rozładowana. Nadal możesz korzystać z rejestratora, choć będziesz musiał dotykać ekranu nawet przez 15 sekund, aby go wybudzić, a dane będą znacznie rzadsze, ponieważ nie będą rejestrowane tak często. Mimo to, powinieneś być w stanie



Ekran 3. Ekran ustawień zawiera najczęściej używane opcje konfiguracji rejestratora, w tym napięcia baterii, godzinę i datę oraz sterowanie podświetleniem. Każdy wpis jest sprawdzany, aby upewnić się, że nie koliduje z innymi wartościami (takimi jak napięcie „Rozładowany” jest wyższe niż napięcie „W pełni naładowany”), a następnie natychmiast zapisywany w pamięci FLASH



Ekran 4. Ekran wprowadzania jest wyświetlany za każdym razem, gdy trzeba wprowadzić liczbę. Symbol w lewym dolnym rogu umożliwia usunięcie ostatnio wpisanego znaku. Ponieważ liczby ujemne nie są używane, nie ma symbolu minusa

szybko zidentyfikować problem z baterią i go rozwiązać.

Aby wyłączyć tę funkcję (np. do testowania bez podłączonej baterii), ustaw tę wartość na 0 V. W takim przypadku regulator impulsowy „buck” wyłączy się poniżej napięcia wejściowego wynoszącego około 5,5 V, powodując całkowite wyłączenie rejestratora, chyba że jest on zasilany z USB.

Kalibracja

Pozostały przycisk na stronie głównej powoduje przejście do strony kalibracji (ekran 6). Zawsze należy najpierw skalibrować współczynnik V, ponieważ zmierzony prąd zależy od dokładności zmierzonych napięć.

Wewnętrznie istnieje współczynnik V (stosunek między rzeczywistym napięciem a nieprzetworzonym 24-bitowym odczytem ADC) dla każdego z czterech dzielników, ale wyświetlany jest tylko jeden, ponieważ wszystkie powinny być podobne, uwzględniając tolerancję boczników. Wartość nominalna zapewni odczyt w pełnej skali przy 100 V.

Cztery współczynniki V umożliwiają kompensację zmian w dzielnikach, głównie z powodu tolerancji komponentów. Pozwalają one na wyzerowanie trzech dzielników prądu względem głównego dzielnika napięcia. Dlatego ten krok należy wykonać najpierw przed próbą kalibracji poszczególnych prądów; w przeciwnym razie wystąpi przesunięcie zera.

Konieczne będzie podłączenie akumulatora lub przynajmniej stabilnego źródła napięcia powyżej 6 V. Wyższe napięcie będzie oznaczać, że błąd kwantyzacji (spowodowany krokami między kolejnymi wartościami ADC) będzie

proporcjonalnie mniejszy, potencjalnie dając nieco lepszą kalibrację.

Nie podłączaj jednak niczego do CON3A, ponieważ nie chcemy, aby przepływ prądu wypaczał wyniki. Jeśli to możliwe, pozostaw również podłączone zasilanie USB, ponieważ zminimalizuje to obciążenie baterii, a wyświetlacz będzie zasilany z USB. W tym przypadku jedynym obciążeniem baterii będzie prąd spoczynkowy IC4 bez obciążenia, wynoszący około 10 μ A.

Podłącz woltomierz do zacisków akumulatora i pozwól urządzeniu ustabilizować się przez minutę. Odczyt musi być stabilny, aby uzyskać optymalne wyniki. Naciśnij przycisk „Volts” i potwierdź, że zaciski nie są obciążone.

Wprowadź napięcie akumulatora wyświetlane na woltomierzu. Na stronie wyświetlone zostaną różne współczynniki V oraz szacunkowe wartości ich zmian. Jeśli różnica wynosi więcej niż kilka procent (z powodu tolerancji komponentów), może to oznaczać problem z dzielnikami, taki jak nieprawidłowa wartość komponentu lub fałszywe obciążenie akumulatora.

Nowe wartości można potwierdzić, naciskając przycisk OK, lub użyć przycisku Anuluj, aby przeprowadzić dalszą analizę. Kalibracja zostanie zapisana w pamięci FLASH i natychmiast użyta. Wróć, aby sprawdzić, czy wyświetlane prądy (I2-I4) ustabilizowały się w pobliżu zera. Oznacza to, że kalibracja jest prawidłowa.

Pozostałe kalibracje nie są tak krytyczne, ponieważ nie spowodują przesunięcia w wynikach, ale po prostu dadzą nieprawidłowe skalowanie prądu. Wartości domyślne są obliczane na podstawie nominalnych wartości komponentów; będziesz musiał je zmienić, jeśli używasz zewnętrznych boczników.

Kalibracja prądu

Metoda kalibracji prądu jest prosta. Do każdego zacisku przykładane jest znane obciążenie, prąd jest mierzony i wprowadzany do rejestratora, który następnie oblicza współczynnik konwersji.

W przypadku I2-I4 są to obciążenia zewnętrzne na CON3A, podczas gdy I1 jest prądem własnym rejestratora. Tak więc dla I2-I4, obciążenie powinno być przyłożone pomiędzy CON3A a ujemnym biegunem akumulatora (masą).

W takim przypadku rzeczywisty prąd wyświetlany na ekranie głównym będzie ujemny (akumulator się rozładowuje). Mimo to można wprowadzić tylko wartość dodatnią, więc należy wprowadzić tylko wielkość prądu.

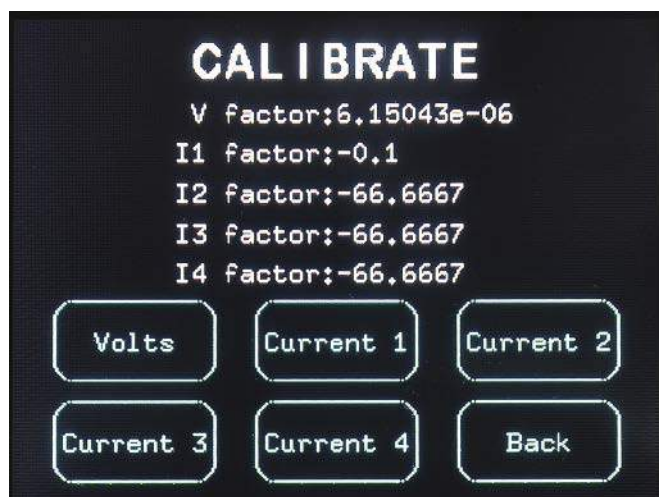
Wartości początkowe są ustawione w programie MMBasic, ale można je również zmienić tutaj, co należy zrobić, jeśli używasz boczników o wartościach innych niż 15 m Ω . Wartości kalibracji prądu są po prostu odwrotnością rezystancji bocznika w omach, więc domyślne boczniki 15 m Ω mają współczynnik kalibracji 66,67.

W przypadku I1 prawdopodobnie konieczne będzie odłączenie akumulatora, aby umożliwić podłączenie amperomierza do zasilania rejestratora. W tym celu należy odłączyć kabel USB i upewnić się, że żaden z zacisków CON3A nie jest obciążony.

Nominalna wartość współczynnika używanego dla I1 jest odwrotnością rezystancji rezystora bocznikowego (w omach) podzieloną przez wzmocnienie obwodu wzmacniacza operacyjnego. Należy wziąć pod uwagę, że zmierzone napięcie bocznika byłoby takie samo, jak gdyby rezystancja bocznika została



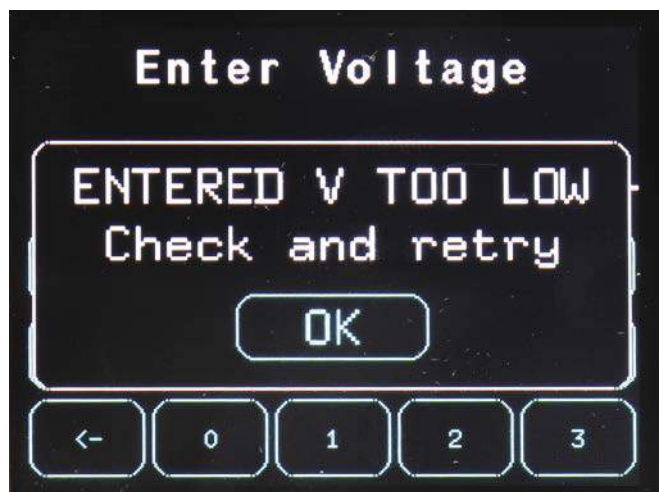
Ekran 5. Każda wprowadzona wartość jest sprawdzana przed przetworzeniem i zapisaniem, aby zapewnić możliwość bezpiecznego wprowadzania zmian



Ekran 6. Ekran kalibracji zapewnia w większości zautomatyzowany sposób regulacji rejestratora w celu uwzględnienia tolerancji komponentów. Operator musi po prostu wprowadzić odczyt miernika (woltu lub ampery), a rejestrator oblicza współczynniki kalibracji, aby uzyskać żadaną wartość



Ekran 7. Wszelkie warunki, które muszą być spełnione w celu dokładnej kalibracji, są wyświetlane przed rozpoczęciem kalibracji. Chociaż jest to dodatkowy krok, oznacza to, że istnieje niewielka szansa na niepowodzenie kalibracji



Ekran 8. Jak wspomniano, wszystkie wartości są sprawdzane pod kątem poprawności przed zapisaniem i użyciem przez MB. W tym przypadku wyświetlany jest krótki, ale pomocny komunikat, aby umożliwić użytkownikowi ustalenie, co poszło nie tak

pomnożona przez wzmocnienie. Domyślną wartością jest więc odwrotność $0,1 \Omega$ (tj. $1/0,1 = 10$, podzielone przez 100, lub $0,1$).

Montaż i zakończenie

Mając wszystko skalibrowane i skonfigurowane, można zamontować i podłączyć MB. Będąc podobnym co do rozmiaru i kształtu do panelu V2 Micromite, rejestrator baterii może być wyposażony w wycięty laserowo akrylowy panel przedni zaprojektowany dla pudełek UB3 Jiffy.

W tej formie można go zamontować w obudowie. Mimo to, spodziewamy się, że większość Czytelników użyje akrylowego panelu jako ramki do zamontowania rejestratora w obudowie sprzętu, z przewodami podłączonymi wewnętrznie i panelem dotykowym dostępnym z zewnątrz.



Po zamontowaniu wewnątrz obudowy urządzenia, ważne funkcje są dostępne dla konserwacji, w tym zakończenia kabli i bateria zapasowa RTC

Aby to zrobić, oddziel wyświetlacz LCD i płytkę drukowaną rejestratora, delikatnie poruszając. Zdecyduj, która strona ramki ma być widoczna; my preferujemy matową powierzchnię, ale jest ona odwracalna, więc jeśli chcesz, możesz umieścić błyszczącą stronę na zewnątrz.

Wkręć cztery śruby M3 przez przód maskownicy, umieść podkładki na gwintach, a następnie zamocuj wyświetlacz LCD. Podkładki dystansowe zapewniają prześwit dla przewodów, które wystają z tyłu złączy.

Zabezpiecz śruby M3 za pomocą gwintowanych podkładek dystansowych. Podłącz ponownie płytkę drukowaną rejestratora i przymocuj ją do zestawu za pomocą pozostałych śrub M3.

Ten kompletny zespół można teraz przymocować na przykład do przednich drzwi szafki na sprzęt, używając śrub M3 i nakrętki w każdym rogu, aby go zabezpieczyć. Po otwarciu szafki dostęp do połączeń akumulatora można uzyskać od tyłu.

Zabezpieczenie tylnej części rejestratora można łatwo wykonać za pomocą obudowy UB3 Jiffy. Dołączone śruby mogą być zbyt krótkie, jeśli trzeba będzie je wkręcić przez panel, ale kołki dystansowe pokryją się z otworami w ramce.

W takim przypadku wystarczy tylko kilka otworów z boku lub z tyłu obudowy, aby doprowadzić przewody.

Aby zakończyć okablowanie, można postępować zgodnie z trzema przykładami pokazanymi na **rysunku 6** (reprodukcja z zeszłego miesiąca). Pokazuje to opcje użycia z wewnętrznymi lub zewnętrznymi bocznikami, w tym jedną możliwość współdzielenia zaciśków na CON3A, jeśli masz więcej niż trzy całkowite obciążenia plus źródła ładowania.

Idealnie byłoby, gdyby na każdym przewodzie wychodzącym z CON3A (lub wysokoprądowym okablowaniu prowadzącym do boczników) znajdował się bezpiecznik.

Bezpiecznik powinien również znajdować się w przewodzie prowadzącym od dodatniego bieguna akumulatora do dodatniego bieguna CON3. W ten sposób usterka rejestratora lub któregośkolwiek z podłączonych obciążeń nie może spowodować zwarcia akumulatora.

Okablowanie będzie specyficzne dla indywidualnych rozwiązań, więc możemy zaoferować tylko ogólne porady.

Podsumowanie

Podobnie jak w przypadku wielu naszych projektów, zwłaszcza tych napisanych w MMBasic, spodziewamy się, że Czytelnicy będą chcieli dostosowywać, majstrować i być może ulepszać oprogramowanie.

Z niecierpliwością czekamy na informacje, jakie funkcje Czytelnicy chcieliby dodać,

ponieważ już planujemy w przyszłości uzupełnienie rejestratora o dodatkowy sprzęt.

Zauważył zapewne, że nie pozostawiliśmy zbyt wielu niewykorzystanych portów mikroprocesora, ale jednak wyprowadziliśmy dwa porty do złącza CON4 (nie zamontowanego na PCB) w prawym górnym rogu płytki drukowanej. Są one podłączone do magistrali I²C Micromite, ponieważ uznaliśmy, że będzie to dobry sposób na rozszerzenie urządzenia (są one już używane przez zegar czasu rzeczywistego, ale I²C jest magistralą współdzieloną).

Połączenia zasilania 3,3 V i uziemienia są również dostępne na pobliskim CON2, podczas gdy CON6 łączy się z drugim portem COM Micromite (COM1), na stykach 21 i 22.

Zapewnia to dedykowany kanał komunikacyjny, który można wykorzystać w przyszłości do dodania większej liczby funkcji. ■

Tim Blythman

*Adaptacja do wydania
polskiego – Andrzej Nowicki*

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

REKLAMA

Wydawnictwo AVT nawiąże współpracę redakcyjną z osobami dobrze operującymi terminologią elektroniki i słowem pisanym. Propozycja szczególnie interesująca dla nauczycieli elektroniki, autorów artykułów, skryptów i książek.

Aplikacje prosimy kierować na adres: redakcja@elportal.pl



Czy jesteś zainteresowany nauką lutowania małych podzespołów do montażu powierzchniowego, ale nie chcesz zniszczyć drogiej płytki lub układu scalonego, zdobywając te umiejętności? Być może nie masz innego wyjścia, jak tylko się tego nauczyć, ponieważ tak wiele części produkowanych w dzisiejszych czasach jest dostępnych tylko w obudowach SMD. Ten prosty projekt płytki treningowej czy testowej (Trainer) to świetny sposób na ćwiczenie lutowania różnych części do montażu powierzchniowego. Jeśli zrobisz to poprawnie, zostaniesz nagrodzony serią migających sekwencyjnie diod LED.



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/dhs7h>
Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania

Płytką treningowa SMD

Podzespoły do montażu powierzchniowego (SMD), ze względu na ich kompaktowość, dobrą niezawodność, niski koszt, automatyzację montażu i powszechną dostępność są preferowanym typem części używanych w większości urządzeń komercyjnych. Niektórzy producenci nadal produkują nowe części w wersji do montażu przewlekane, jednak wybór komponentów staje się z dnia na dzień coraz bardziej ograniczony, jeśli nie możesz stosować podzespołów SMD.

Wiemy, że początkowo wydaje się to zniechęcające (dla nas, redakcji Silicon Chipa, też tak było), ale będziesz zaskoczony, jak łatwo przy odrobinie praktyki zacząć ich używać. I właśnie do tego została zaprojektowana ta płytka. Jest to działający obwód zaprojektowany przy użyciu szerokiej gamy różnych części SMD, umożliwiający trening ich lutowania. W ten sposób można opanować tę technikę montażu a zarazem zapoznać się z popularnymi rozmiarami i obudowami SMD.

Płytkę zaprojektowano tak, abyś mógł zacząć od większych części, a gdy nabierzesz pewności siebie, przejść do tych mniejszych. Po drodze możesz sprawdzić cały proces, dzięki czemu szybko dowiesz się, czy popełniłeś błąd i będziesz miał okazję go poprawić.

Ten artykuł zawiera podstawowe instrukcje dotyczące budowania i testowania płytki treningowej, wraz z opisem jej działania. Tematyka jest kontynuowana za miesiąc, w artykule zawierającym znacznie więcej szczegółów dotyczących niezbędnych narzędzi i technik.

Zalecamy Czytelnikowi zapoznanie się z tym artykułem teraz i zagłębienie do niego później,

jeśli napotka tutaj coś, czego w pełni nie zrozumie. Jest to szczególnie ważne, jeśli nie masz doświadczenia w lutowaniu lub masz wątpliwości co do swoich umiejętności postępowania z podzespołami SMD.

Zakładając, że przeczytałeś ten artykuł (przynajmniej częściowo) i zaczynasz mieć pojęcie, jak zabrać się za montaż tej płytki, przejdźmy do opisu projektu.

Szczegóły obwodu

Schemat ideowy płytki treningowej SMD pokazano na **rysunku 1**. Zanim przejdziemy dalej, wyjaśnimy, jak to działa. Ważne jest, aby wiedzieć, co układ powinien robić, zwłaszcza abyś mógł mieć jakieś wyczucie na temat tego, co mogło pójść nie tak, jeśli układ nie zadziała od razu.

Obwód składa się z dwóch głównych części, z których druga zależy od pierwszej. Pierwsza część obwodu jest również łatwiejsza do zmontowania, więc możesz wypróbować swoje umiejętności na niej, zanim przejdiesz do części o zwiększonym poziomie trudności.

Wspólne dla obu części jest zasilanie. Uchwyt litowego ogniwa pastylkowego BAT1 typu CR20xx jest połączony równolegle z gniazdem USB CON1. Należy zamontować tylko jedno zasilanie! Zalecamy uchwyt na ogniwo pastylkowe, ponieważ jest mniej prawdopodobne, że w przypadku popełnienia błędu podczas budowy ogniwo to dostarczy zbyt wysoki prąd.

Ze względu na obecność ogniwa pastylkowego, należy uważać, aby płytka treningowa SMD Trainer była przechowywana poza zasięgiem dzieci. Posiada migające światelka, więc

będzie przyciągała uwagę, niemniej nie jest zabawką i z tego powodu nie powinna dostać się w ręce małego dziecka, również z uwagi na samą baterię, która mogłaby zostać połknięta i wyrządzić szkody zagrażające zdrowiu a nawet życiu.

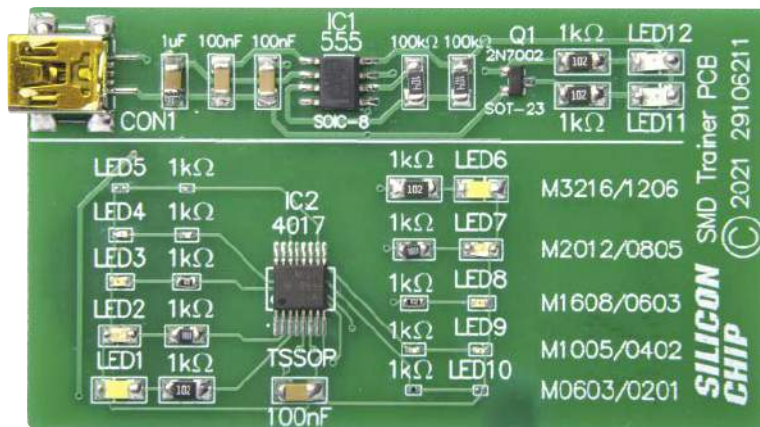
Pierwsza połowka

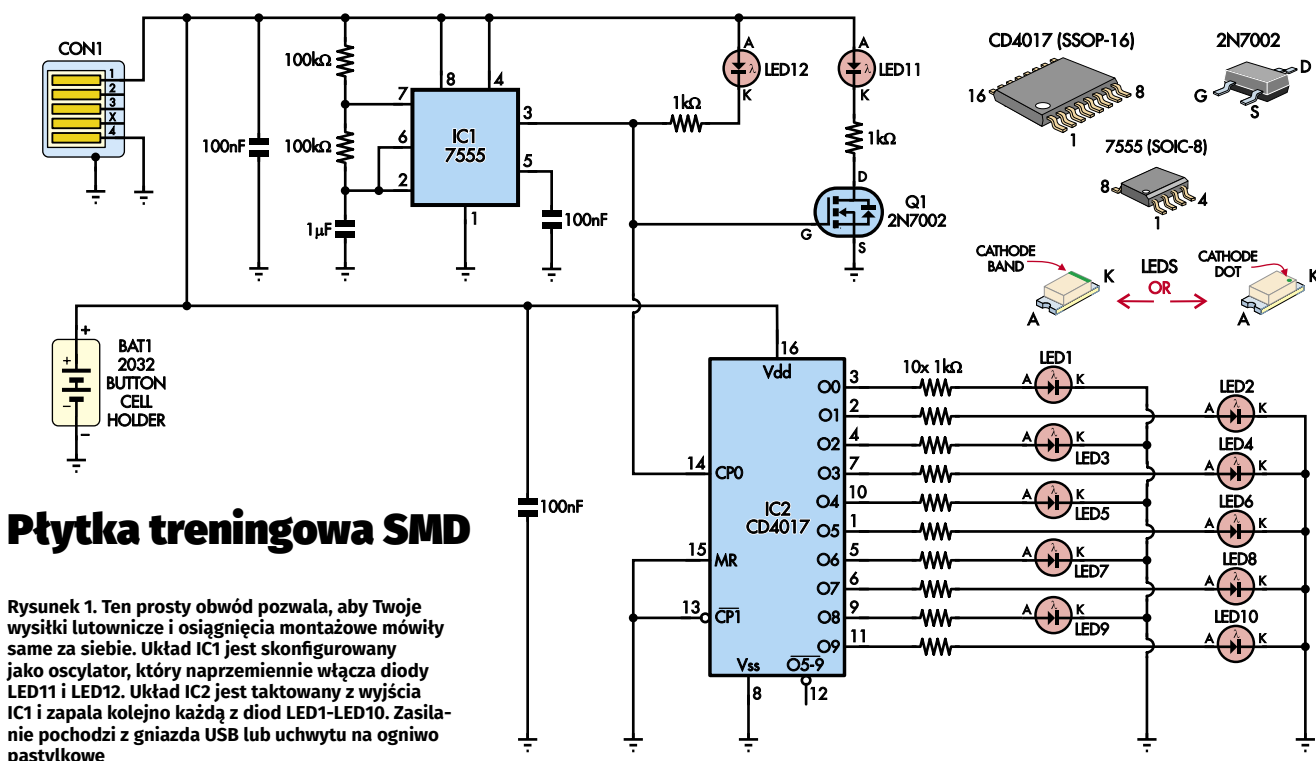
IC1 to scalony układ czasowy (timer 7555). Wybraliśmy ten wariant CMOS zamiast układu na tranzystorach bipolarnych 555, aby umożliwić pracę obwodu przy niskim napięciu ok. 3 V. Zasilanie dochodzi do styków 8 (dodatni) i 1 (ujemny) układu IC1. Styk 4 (RESET) jest utrzymywany w stanie wysokim, aby umożliwić działanie timera.

Układ IC1 ma zasilanie zbocznikowane kondensatorem 100 nF, a drugi kondensator 100 nF stabilizuje wewnętrzne napięcie na styku CV, końcówka 5. Układ IC1 jest zastosowany w dobrze znanej konfiguracji oscylatora astabilnego, z rezystorami 100 kΩ i kondensatorem 1 μF.

W tym układzie kondensator 1 μF łąduje się z zasilania poprzez dwa rezystory 100 kΩ; jego biegun dodatni jest podłączony do zwartych wyprowadzeń wejściowych 2 i 6. Gdy potencjał na wyprowadzeniu 2 wzrośnie powyżej 66% napięcia zasilania (do około 2 V), wewnętrzny przerzutnik przełączy się, a wyprowadzenie 7 jest podłączane do masy (przez tranzystor wewnątrz IC1). Jednocześnie wyprowadzenie 3 przechodzi w stan niski.

Powoduje to rozładowanie kondensatora 1 μF przez dolny rezystor 100 kΩ i poprzez wyprowadzenie 7, a napięcie





na kondensatorze spadnie do 33% napięcia zasilania (około 1 V). Przerzutnik resetuje się wtedy, wyprowadzenie 3 przechodzi w stan wysoki, wyprowadzenie 7 przestaje pobierać prąd z kondensatora, kondensator zaczyna się ponownie ładować i cykl się powtarza.

Przy podanych wartościach komponentów częstotliwość oscylatora wynosi około 4,8 Hz z 66% cyklem pracy na wyprowadzeniu 3 (tj. wyprowadzenie 3 jest w stanie wysokim przez około 2/3 czasu).

Gdy wyprowadzenie 3 jest w stanie niskim, prąd jest pobierany z zasilania przez diodę LED12 i jej szeregowy rezystor ograniczający prąd o rezystancji 1 kΩ, powodując jej świecenie. Gdy wyprowadzenie 3 jest w stanie wysokim, MOSFET Q1 jest włączany przez dodatkowo napięcie na jego bramce, a prąd przepływa przez diodę LED11 i jej rezystor szeregowy. W ten sposób te dwie diody LED migają naprzemiennie.

Ta pierwsza część obwodu jest zbudowana z większych części SMD, takich jak te, które zwykle uwzględniamy w naszych projektach, gdy części przewlekanych nie można zastosować. Może działać niezależnie od pozostałej części obwodu i może być zbudowana i przetestowana jako pierwsza część dwuczęściowego zadania.

Druga połowa

Pozioma linia na płytce PCB dzieli ją na dwie odrębne części; część druga znajduje się poniżej tej linii.

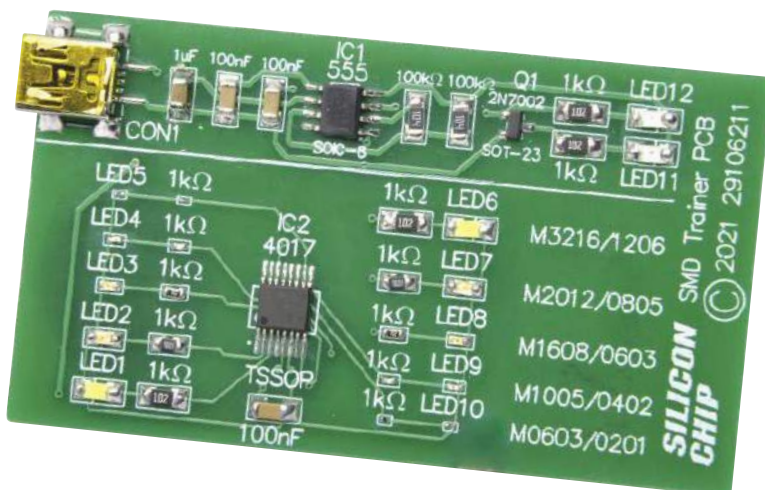
Sercem drugiej części obwodu jest IC2, licznik dekadowy typu 4017. Jest on zasilany z tego samego źródła co IC1, podłączonego do jego wyprowadzenia 16 (zasilanie dodatnie) i wyprowadzenia 8 (zasilanie ujemne). Jego zasilanie, dla zapewnienia stabilności, jest również zbocznikowane kondensatorem 100 nF.

Układ IC2 ma dziesięć wyjść dostępnych na wyprowadzeniach 3, 2, 4, 7, 10, 1, 5, 6, 9 i 11. Są one ustawiane w stanie wysokim, jeden po drugim, w odpowiedzi na sygnał zegarowy przyłożony do wyprowadzenia 14. Sygnał ten pochodzi

z wyprowadzenia 3 wspomnianego powyżej układu IC1. Wyprowadzenia 13 i 15 są spolaryzowane do stanu niskiego, aby umożliwić normalne działania zliczania. Styk 12 jest wyjściem przeniesienia, które może być połączone kaskadowo z innymi układami, ale w tym przypadku pozostaje odłączone.

Każde z wymienionych powyżej dziesięciu wyjść ma rezystor szeregowy 1 kΩ i diodę LED podłączone do niego. Tak więc sygnał zegarowy na styku 14 powoduje, że diody LED zapalają się kolejno, jedna po drugiej.

Komponenty wokół IC2 mają różne rozmiary, aby stanowić bardziej interesujące wyzwanie.

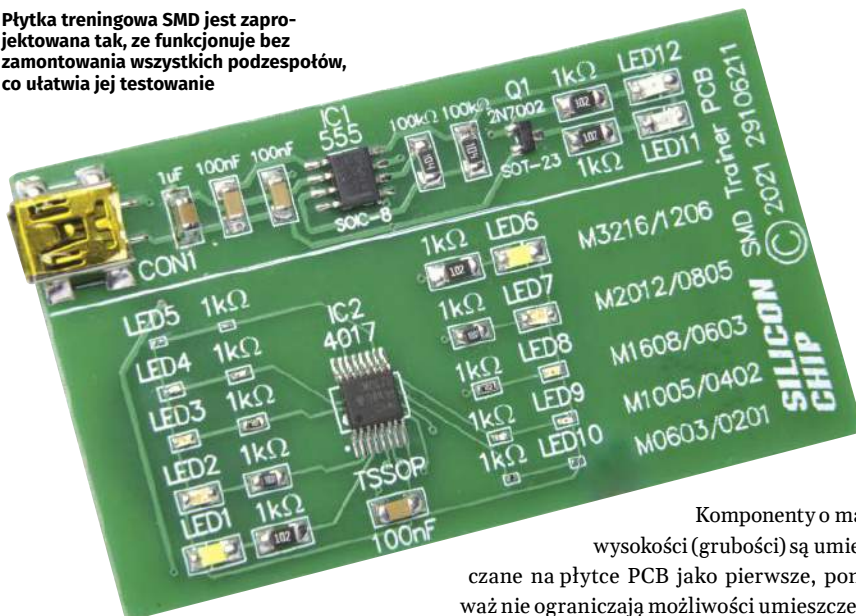


To jest płytka treningowa SMD Trainer, którą złożyliśmy (pokazana w powiększeniu, około 166% rzeczywistego rozmiaru). Jeśli masz problemy z dostrzeżeniem diod LED SMD 0201, może to być spowodowane tym, że nie są one zamontowane! Nie byliśmy w stanie przylutować ich ręcznie i nie będziemy udawać, że jest to łatwe

Tabela 1. Typowe rozmiary pasywnych elementów SMD

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

Płytkę treningową SMD jest zaprojektowana tak, że funkcjonuje bez zamontowania wszystkich podzespołów, co ułatwia jej testowanie



IC2 jest również w mniejszej obudowie SMD niż IC1. Więcej szczegółów można znaleźć w tabeli 1.

Umieszczenie i kolejność

Zalecana przez nas kolejność montażu dla większości projektów z częściami do montażu przewlekane wynika z kilku powodów. Praca według typu komponentu, na przykład zaczynając od rezystorów, następnie diod, kondensatorów, a potem kolejno układów scalonych, ułatwia śledzenie, na jakim etapie znajduje się montaż.

W większości przypadków kolejność ta jest podyktowana wysokością komponentów.

Komponenty o małej wysokości (grubości) są umieszczane na płytce PCB jako pierwsze, ponieważ nie ograniczają możliwości umieszczenia wyższych części. Oznacza to również, że płytkę drukowaną można odwrócić do góry nogami bez wypadania elementów do montażu przewlekane; są one utrzymywane na płytce drukowanej przez powierzchnię roboczą.

Pracę z częściami SMD determinują podobne przesłanki, ale istnieje znacznie mniejsza potrzeba odwracania PCB (montaż po jednej stronie), więc nie ma realnej szansy na wypadnięcie części. Ponadto większość części SMD ma niski profil.

Tak więc podstawową kwestią będzie zamontowanie trudniej dostępnych lub trudnych do lutowania części w pierwszej kolejności, tak aby ich montaż nie był utrudniony przez części zamontowane wcześniej.

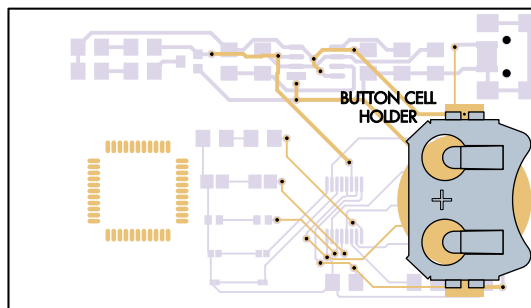
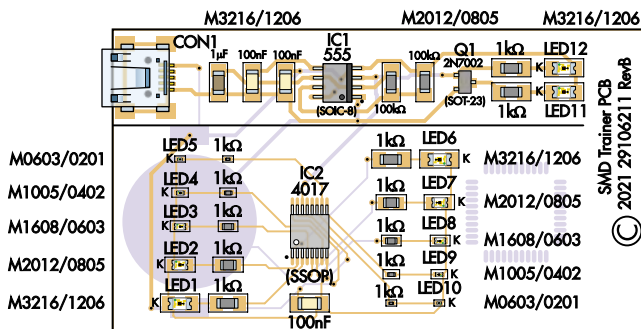
Mając to na uwadze, najlepszym sposobem konstruowania obwodów hybrydowych (zawierających zarówno części przewlekane, jak i SMD) jest zamontowanie części SMD w pierwszej kolejności. Niezależnie od tego, czy znajdują się one po tej samej stronie, czy nie, wyższe części do montażu przewlekane będą stanowić większą przeszkodę w budowie, jeśli zostaną wlutowane przed mniejszymi częściami SMD.

Oznacza to również, że proces umieszczania układów scalonych jako ostatnich nie jest zbyt odpowiedni. W dzisiejszych czasach układy scalone są zwykle bardziej wytrzymałe i mniej podatne na uszkodzenia spowodowane ładunkami elektrostatycznymi, co zwykle było motywacją do montowania ich tak późno, jak to tylko możliwe.

W projektach SMD (lub przynajmniej tych przeznaczonych do lutowania ręcznego), układy scalone mają zazwyczaj drobniejsze wyprowadzenia i są trudniejsze w montażu. Dlatego sensowne jest wykonanie ich montażu w pierwszej kolejności, a następnie kontynuacja pracy z otaczającymi je elementami pasywnymi, które często są większe.

Montaż zestawu płytki treningowej SMD

Zapoznaj się teraz ze schematami montażowymi PCB, na rysunkach 2 i 3, które pokazują, rozmieszczenie komponentów na odpowiednich lokalizacjach. Treningowa płytka drukowana SMD jest dwustronna, ma wymiary 70,5×40 mm i jest oznaczona kodem 29106211.



Rysunki 2 i 3. Zaczni od zamontowania komponentów w górnej połowie płytki drukowanej, która tworzy migacz z naprzemiennie świecącymi diodami LED11 i LED12. Te komponenty to większe SMD, które zazwyczaj nie są zbyt trudne do przylutowania. Po ich uruchomieniu można przejść do trudniejszego fragmentu poniżej, również utworzonego z diodami LED. Po zamontowaniu układu IC2 i jego kondensatora odsprężającego, zamontuj diody LED1, LED6 i ich rezystory szeregowo, a następnie przejdź do mniejszych części, testując je po każdym kroku, aby upewnić się, że lutowanie jest poprawne

Zalecamy rozpoczęcie lutowania od gniazda USB, jeśli zdecydujesz się je zamontować. Wyprowadzenia nie są zbyt małe, ale potrafią być trudno dostępne dla grota lutownicy. Na szczęście ta część ma kolki blokujące na spodzie PCB, które wchodzą w otwory w płytce drukowanej. Tak więc prawidłowe ustawienie elementu jest łatwe.

Umieść topnik na wszystkich polach stykowych gniazda USB i dociśnij część. W tym zastosowaniu tylko dwa zewnętrzne zestyki z pięciu są potrzebne do zasilania; dlatego są jedynymi, które mają przedłużone ścieżki. Można również dodać więcej topnika na wierzch wyprowadzeń.

Oczyść grot lutownicy, nałóż niewielką ilość lutowia i dociśnij grot jednocześnie do padu płytki PCB oraz leżącego nań wyprowadzenia złącza USB. Jeśli lut nie spływa na kontakt, zbliż grot, aż dotknie końcówki, jeśli to konieczne. Powtórz tę czynność dla drugiego zewnętrznego zestyku.

W przypadku tego złącza upewnij się, że podczas wykonywania tych połączeń zasilania nie dotykasz grotem lutownicy obudowy gniazda USB. Ciasny kąt między obudową a PCB sprawia, że jest to trudne. Jeśli utworzysz zwarcie pomiędzy sąsiednimi padami, podgrzejsz wszystkie styki, aby wylutować część i przywrócisz pierwotny wygląd zarówno gniazda, jak i płytki drukowanej, za pomocą miedzianej plecionki lutowniczej.

W przypadku większych styków, które zabezpieczają gniazdo USB mechanicznie, wystarczy przyłożyć grot, dodać trochę lutu, aż utworzy się schludny spływ lutowia, a następnie usunąć lutownicę. Nieco większa ilość lutu w tym miejscu zapewni bezpieczne połączenie mechaniczne.

Korzystając z podobnej procedury, umieść na płytce IC1 i Q1, upewniając się, że są prawidłowo zorientowane. Następnie przylutuj rezystory i kondensatory. Zwróć uwagę, że każdy z nich może mieć jedną z dwóch różnych wartości;

Wykaz elementów:

1 dwustronna płytka drukowana o kodzie 29106211 i wymiarach 71×40 mm
1 gniazdo mini-USB (CON1) LUB
1 uchwyt uniwersalny SMD na litową baterię pastylkową typu CR2020, CR2025 lub CR2032 (BAT1) [BAT-HLD-001; Digi-Key, Mouser itp.]
1 litowa bateria pastylkowa CR20xx

Półprzewodniki:

1 układ scalony timera CMOS 7555, SOIC-8 (IC1)
1 układ scalony licznika dekadowego 4017B, SSOP-16 (IC2)
1 N-kanalowy MOSFET 2N7002, SOT-23 (Q1)
4 diody LED SMD1206, dowolny kolor (LED1, LED6, LED11, LED12)
2 diody LED SMD 0805, dowolny kolor (LED2, LED7)
2 diody LED SMD 0603, dowolny kolor (LED3, LED8)
2 diody LED SMD 0402, dowolny kolor (LED4, LED9)
2 diody LED SMD 0201, dowolny kolor (LED5, LED10)

Kondensatory: (wszystkie SMD X7R 10 V+ ceramiczne)

1 szt. 1 µF SMD 1206
3 szt. 100 nF SMD 1206

Rezystory: (wszystkie SMD 1% lub 5%)

2 szt. 100 kΩ SMD 1206 4 szt. 1 kΩ SMD 1206 2 szt. 1 kΩ SMD 0805
2 szt. 1 kΩ SMD 0603 2 szt. 1 kΩ SMD 0402 2 szt. 1 kΩ SMD 0201

Zestaw Altronics będzie dostępny
Altronics ogłosił, że przygotowuje zestaw dla tego projektu, kod K2001

możesz również zapoznać się z naszymi zdjęciami. Te podzespoły nie są kierunkowe, możesz je lutować w jednej dowolnej z dwóch orientacji.

Natomiast diody LED są spolaryzowane i muszą być zamontowane anodami do rezystorów a katodami do wspólnej ścieżki masy.

Jeśli chcesz zamontować uchwyt ogniwa zamiast gniazda USB, zrób to teraz. Zwykle łatwiej jest zamontować wszystkie części po jednej stronie płytki na raz, ale umieszczenie gniazda ogniwa pastylkowego pozwoli przetestować pierwszą część obwodu, który właśnie zmontowałeś.

Odwróć płytkę drukowaną i nałóż trochę topnika na dwa mniejsze zewnętrzne pola lutownicze. Pozostaw duży wewnętrzny zestyk czysty, ponieważ samo pole lutownicze PCB staje się zaciskiem ujemnym i nie wymaga lutowania.

Upewnij się również, że otwór uchwytu jest skierowany w stronę krawędzi płytki drukowanej, aby można było łatwo włożyć ogniwo. Umieść uchwyt z grubsza na miejscu i dodaj trochę topnika na górę wyprowadzeń.

Należy pamiętać, że w przeciwieństwie do gniazda USB, nie ma nic, co mogłoby zablokować tę część na miejscu.

Prawdopodobnie będziesz musiał nieco zwiększyć temperaturę lutownicy (jeśli jest regulowana) i nałożyć trochę lutu na grot; nieco więcej niż w przypadku mniejszych części. Użyj pęsety, aby utrzymać uchwyt ogniwa na miejscu i przyłóż grot równocześnie do wyprowadzenia uchwytu jak i padu lutowniczego płytki PCB, na którym leży to wyprowadzenie.

Potrzebujesz trochę czasu na rozgrzanie; pamiętaj, że jest to jeden kawałek metalu, więc jest mało prawdopodobne, aby został uszkodzony przez zbyt wysoką temperaturę. Powinieneś zobaczyć dymiący topnik i wypływający lut. Odsuń lutownicę i pozwól części (i lutowi) ostygnąć przez kilka sekund przed zwolnieniem pęsety.

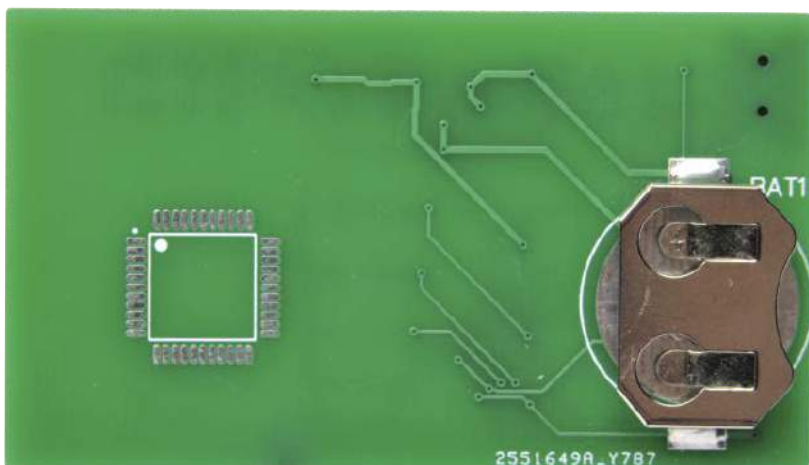
Pierwsze połączenie nie musi być idealne; najważniejsze jest, aby część była dokładnie umieszczona i mocno trzymana.

Do drugiego zestyku można podejść jak do większych pól lutowniczych w przypadku gniazda USB. Przyłóż grot lutownicy i podawaj lut, aż utworzy się poprawny rozpliw lutowia, dobrze spajający pad lutowniczy oraz wyprowadzenie złącza, a następnie odejmij grot. Daj uchwytowi kilka sekund na zastygnięcie lutu przed powrotem do pierwszego lutu, aby go poprawić. Możesz to zrobić, przykładając grot i lutując w ten sam sposób.

Wstępne testowanie

Pierwsza część obwodu powinna być już sprawna. Rozpocznij sprawdzenie obwodu przez podłączenie ogniwa pastylkowego lub zasilania z gniazda USB. Jeśli używasz ogniwa pastylkowego, upewnij się, że jego polaryzacja jest prawidłowa. Diody LED11 i LED12 powinny migać naprzemiennie.

Jeśli jedna dioda LED świeci stale, IC1 nie oscyluje i powinieneś sprawdzić poprawność jego przylutowania, jak również poprawność przylutowania podzespołów wokół IC1. Gdy jedna dioda LED miga, druga być może jest



Na spodzie płytki drukowanej znajduje się zestaw styków TQFP. Służy on do ćwiczenia lutowania i nie ma żadnego połączenia elektrycznego z obwodem

przylutowana nieprawidłowo; przyczyną może być też wadliwy montaż jednego z rezystorów 1 kΩ lub tranzystora Q1.

Możesz również zobaczyć coś, co wydaje się być dwiema diodami LED włączonymi w tym samym czasie. W takim przypadku prawdopodobnie migają one szybciej, niż jest to widoczne dla oka. Jednym z możliwych powodów jest to, że kondensator czasu oscylacji o pojemności 1 μF został zamieniony z jednym z kondensatorów 100 nF.

W tym momencie najlepiej jest sprawdzić, czy ta część obwodu działa poprawnie. W przeciwnym razie, jeśli druga część nie zadziała, trudniej będzie znaleźć przyczynę.

Pozostała część obwodu

Warto zauważyć, że komponenty w dolnej połowie płytki PCB są dość regularnie i luźno rozłożone. Jest to luksus mało prawdopodobny w docelowych projektach SMD.

Dzięki dużej ilości miejsca na płycie treningowej SMD, z pewnością możliwe jest lutowanie tych komponentów w dowolnej kolejności. Zalecamy jednak rozpoczęcie pracy od IC2 i sąsiedniego kondensatora, a następnie diod LED w kolejności od największej do najmniejszej. Pozwoli to na uruchomienie obwodu w dowolnym momencie po zamontowaniu którejkolwiek z większych diod LED i sprawdzenie, czy zamontowany fragment działa.

Zacznij od układu IC2. Nałóż topniki i umieść chip. Byliśmy dość rozrutni tutaj z długością padów lutowniczych z dwóch powodów. Po pierwsze, widzieliśmy warianty obudowy SOP tego układu dostępne z różnymi szerokościami korpusu. Tak więc ta konfiguracja padów zapewnia elastyczność umożliwiającą akceptację szeregu kompatybilnych części. Po drugie, ułatwia to lutowanie.

Wyczyść grot lutownicy i dodaj do niego niewielką ilość świeżego lutowia. Przytrzymaj IC2 pęsetą i przyłóż grot tylko do płytki

PCB. Powinieneś zobaczyć, jak lut płynie po ścieżce i tworzy połączenie wystarczająco mocne, aby utrzymać część na miejscu.

Sprawdź, czy wyprowadzenia IC2 są wyrównane z padami i przylutuj pozostałe wyprowadzenia w ten sam sposób. Te małe styki nie wymagają dużej ilości lutowia, więc może się okazać, że wystarczy tylko od czasu do czasu dodać trochę lutu na grot lutownicy.

Sprawdź, czy nie ma mostków i usuń je w razie potrzeby. Postępuj podobnie z sąsiednim kondensatorem 100 nF. Diody LED1 i LED6 to części o rozmiarze SMD 1206, więc powinieneś czuć się komfortowo montując je i odpowiadające im rezystory 1 kΩ. Zwróć uwagę, że wszystkie katody LED-ów znajdują się po stronie przeciwnej względem układu IC2. Do IC2 diody podłączone są anodami, za pośrednictwem szeregowo włączonych rezystorów.

Ponowne testowanie

Nasz projekt da się uruchomić i przetestować w dowolnym momencie w postaci częściowo ukończonej płytki. Diody LED1 i LED12 powinny nadal świecić naprzemiennie; jeśli tak się nie dzieje, może to oznaczać zwarcie, które powoduje odcięcie zasilania od układu IC1 i jego komponentów.

Diody LED1 do LED10, gdyś zamontowane, powinny migotać na przemian. Jeśli nic się nie dzieje, sprawdź, czy układ IC2 jest prawidłowo zamontowany, w prawidłowej orientacji i bez zwarć pomiędzy wyprowadzeniami. Gdy pojedyncze diody LED nie migają, prawdopodobnie oznacza to, że akurat ta dioda LED lub jej rezystor nie są poprawnie przylutowane.

Zakończenie

Nie spiesz się i sprawdź kolejno diody LED i rezystory o różnych rozmiarach. Nie rozczaruj się, jeśli nie możesz ręcznie przylutować części w rozmiarze SMD 0402 lub SMD 0603. W żadnym z naszych projektów nie



Ten element o rozmiarze SMD 0201, pokazany na czubku palca, mierzy zaledwie 0,6×0,3 mm, co sprawia, że łatwo go zgubić

używaliśmy elementów mniejszych niż SMD 0603, bo nawet dla nas elementy SMD 0402 i mniejsze stanowią wyzwanie.

Ostatni raz tak małe komponenty jak SMD 0603 zastosowaliśmy w radiu z ekranem dotykowym DAB+ (styczeń-marzec 2019 r.; www.siliconchip.com.au/Series/330). Nawet wtedy oferowaliśmy płytki PCB z tymi najmniejszymi częściami wstępnie zamontowanymi.

Wszystko, co jest tak małe, nie jest przeznaczone do lutowania ręcznego. Mniejsze diody LED często mają odsłonięte zestyki tylko od spodu, co bardzo utrudnia przenoszenie ciepła tam, gdzie jest ono potrzebne.

Istnieją pewne sztuczki, których można użyć, takie jak nałożenie niewielkiej ilości lutowia na pola lutownicze i próba przewodzenia ciepła przez ścieżkę PCB podgrzewaną lutownicą z lutowiem. Można też spróbować swoich sił w poprawianiu lutów przy użyciu gorącego powietrza lub podczerwieni.

Projekt DIY Solder Reflow Oven czyli rozplwowego pieca do lutowania elementów SMD redakcja Silicon Chipa opublikowała w wydaniach z kwietnia i maja 2020 roku (siliconchip.com.au/Series/343) – w EdW opis znalazł się w numerach majowym i czerwcowym z 2023 r. Możliwy jest również skuteczny montaż płytki za pomocą «narzędzi», takich jak patelnie elektryczne czy żelazka do ubrań!

Czyszczenie

Gdy będziesz zadowolony z postępów, wyczyść resztki topnika rozpuszczalnikiem i pozwól płycie całkowicie wyschnąć. Chociaż płytka nie robi nic niesamowicie użytecznego, nadal jest przydatnym narzędziem referencyjnym i przypomni ci o sztuczках i technikach, których nauczyłeś się podczas jej budowy.

Kompletny zestaw

Do wyczerpania zapasów internetowy sklep Silicon Chipa będzie sprzedawać kompletny zestaw części (siliconchip.com.au/Shop/20/5260) lub otrzymasz go od Altronics. ■

Tim Blythman

Adaptacja do wydania polskiego – Andrzej Nowicki

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

Więcej informacji

Oczywiście w przeszłości zamieszczaliśmy w Silicon Chip-ie artykuły na temat technologii montażu powierzchniowego, urządzeń i konstrukcji. Niektóre są wymienione poniżej:

- Make Your Own SMD Tools, Circuit Notebook July 2007 (siliconchip.com.au/Article/2289),
- How To Hand-Solder Very Small SMD ICs October 2009 (siliconchip.com.au/Article/1590),
- Soldering SMDs: it's becoming unavoidable December 2010 (siliconchip.com.au/Article/376),
- Simple DIY gizmos for SMD desoldering Circuit Notebook July 2014 (siliconchip.com.au/Article/7944),
- Publisher's Letter: SMDs present challenges and opportunities September 2015 (siliconchip.com.au/Article/8955),
- Third hand for soldering tiny surface mount devices Circuit Notebook April 2016 (siliconchip.com.au/Article/9901),
- Publisher's Letter: It's getting hard to avoid tiny SMDs January 2019 (siliconchip.com.au/Article/11361),
- Sterownik pieca rozplwowego DIY do nowoczesnego lutowania kwiecień i maj 2020 (siliconchip.com.au/Series/343), po polsku w EdW maj i czerwiec 2023 r.



Migające diody LED i śliniący się inżynierowie (6)

Jak wspominałem w poprzednim odcinku (EdW 02/2024), moim obecnym projektem hobbystycznym jest zbudowanie tablicy 12×12=144 piłeczek pingpongowych, z których każda zawiera trójkolorową diodę LED w postaci WS2818 („NeoPixel”). Jak być może pamiętasz, przedstawiłem również moduł Seeeduino XIAO, za pomocą którego planuję sterować moją tablicą. XIAO jest wielkości standardowego znaczka pocztowego, ale posiada 32-bitowy mikrokontroler Arm Cortex-M0+ pracujący z częstotliwością 48 MHz z 256 KB pamięci Flash (programu) i 32 KB pamięci SRAM. Pysznie! W rzeczywistości jestem tak zachwycony tym małym cudem, że postanowiłem stworzyć wideo (<https://bit.ly/373vPQC>).

Niewielką komplikacją jest to, że wejścia/wyjścia (I/O) XIAO w stanie wysokim oczekują/wystawiają napięcia 3,3 V, podczas gdy ja potrzebuję sterować moimi pikselami za pomocą stanów wysokich o napięciu 5 V. Na szczęście znalazłem dość fajny hack, który pozwala mi użyć diody 1N4001 i „ofiarnego” piksela, aby zaimplementować tani i wygodny konwerter napięcia 3,3 V na 5 V (<https://bit.ly/3cxMhcV>).

W oczekiwaniu na ukończenie budowy mojej tablicy z piłeczek pingpongowych, odwiedziłem stronę wiki XIAO (<https://bit.ly/3obCyGH>), aby dowiedzieć się, jak podłączyć wspomniane wcześniej maleństwo do Arduino IDE. Następnie stworzyłem naprawdę prosty szkic testowy (program), aby „połechtać” moje piksele, wybrałem XIAO jako płytkę docelową i nacisnąłem ikonę „Kompiluj”.

Niestety próba kompilacji zakończyła się wieloma błędami i ponurymi ostrzeżeniami, których nigdy wcześniej nie widziałem. „Ojej”, powiedziałem do siebie (lub coś w tym stylu). Zmieniłem płytkę docelową na Arduino Uno i kompilacja przebiegła zgodnie z oczekiwaniami. Powiodła się również, gdy przełączyłem się z powrotem na XIAO i zakomentowałem wszystkie instrukcje związane z NeoPixel. Hmmm.

W tym momencie miałem już smutną minę, więc wysłałem swój szkic do ludzi z Seeed Studio, wyjaśniając mój problem. Następnego ranka w skrzynce odbiorczej znalazłem wiadomość e-mail od inżyniera aplikacji terenowych (FAE) Ansona He, który powiedział, że mój program testowy skompilował się bez problemu. Anson zalecił mi sprawdzenie, czy mam najnowszą wersję Arduino IDE (miałem) i najnowszą wersję biblioteki NeoPixel firmy Adafruit (nie miałem).

Po usunięciu starej biblioteki NeoPixel, postępowałem zgodnie z instrukcjami na stronie Adafruit NeoPixel Überguide (<https://bit.ly/2XzuqOB>),

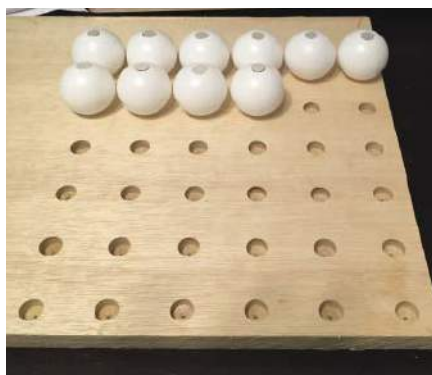
aby zainstalować najnowszą i najlepszą wersję ich biblioteki. Po tym wszystko skompilowało się wspaniale. Teraz pozostało tylko skonstruować tablicę, co brzmi łatwo, jeśli mówisz to szybko i wściekle gestykulujesz.

Gorący klej jest moim przyjacielem

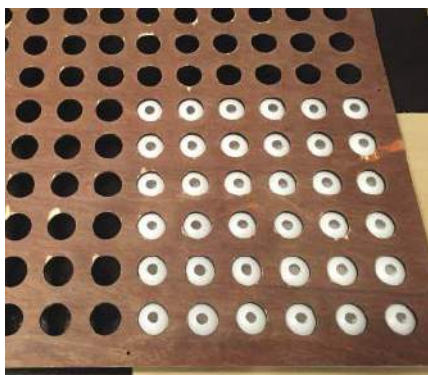
Budowa tablicy postępowała w szybkim tempie. Niestety, nie zastanawiałem się nad tym, jak precyzyjnie ustawić piłeczki pingpongowe. Jak być może pamiętacie z poprzedniego odcinka, zamiast wiercić otwory w piłeczkach, aby pomieścić NeoPiksele, wyciąłem je ręcznie za pomocą małych zakrzywionych nożyczek do paznokci. Rzecz w tym, że to co ułatwiło wszystko to, co było do zrobienia przed zamontowaniem piłeczek na tablicy, pozostawiło mnie teraz z problemem równego ich rozmieszczenia i zamocowania.

Przed wszystkim nie miałem pojęcia, jak trudne może być okiełznanie 144 piłeczek pingpongowych i sprawienie, by stały się posłuszne mojej konstruktorskiej myśli przewodniej. Tymczasem piłeczki zdawały się pozostawać zdeterminowane by ze swych miejsc, za wszelką cenę, uciekać. Po wielu nieudanych eksperymentach stworzyłem przyrząd, wierząc matrycę 6×6 mniejszych otworów w kawałku płyty przeznaczonej do utylizacji i następnie używając go do pozycjonowania i unieruchamiania piłek (**rysunek 1**).

Następnie położyłem jeden róg głównej płyty na górze (**rysunek 2**), użyłem kawałka drewnianego kołka, aby zmienić położenie otworów w piłeczkach pingpongowych, tak aby były skierowane prosto w górę, i użyłem pistoletu do klejenia na gorąco, aby przymocować wszystko na miejscu. Następnie powtórzyłem ten proces dla pozostałych trzech ćwiartek.



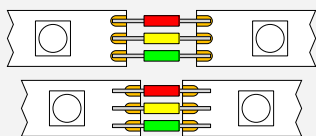
Rysunek 1. Przyrząd do przytrzymywania piłeczek pingpongowych



Rysunek 2. Gotowy do podłączenia pierwszych 36 kulek



Rysunek 3. Cięcie segmentów NeoPixel



Rysunek 4. Tworzenie przewodów łączących w celu dostosowania do różnych rozmiarów szczelin

Na marginesie, traktuję tę tablicę 12×12 jako prototyp do przetarcia szlaków w procesie montażu. W przyszłości mam nadzieję zbudować wyświetlacz wielkości ściany. Postanowiłem, że następnym razem stworzę mniejsze panele 8×8, a następnie połączę je ze sobą, w którym to przypadku przyrząd do wyrównywania będzie tego samego rozmiaru co pojedynczy panel.

Nudne fragmenty

Każdy projekt obejmuje jeden lub więcej nudnych fragmentów i właśnie w tym miejscu się teraz znalazłem. Zaczęło się od wycięcia 145 segmentów z mojego paska NeoPixel – 144 dla matrycy i „ofiarnego” piksela dla konwertera poziomu napięcia (**rysunek 3**). Jest to mniej przyjemne, niż można by się spodziewać, ponieważ pasek jest dostarczany w ochronnej plastikowej osłonie, dzięki czemu piksele są wodoodporne, aby można je było umieścić na zewnątrz, a ta osłona musi być starannie usunięta (nie żebym narzekał, ale wszystko to wymaga czasu).

Następnym krokiem było użycie większej ilości gorącego kleju do przymocowania tych segmentów do piłeczek pingpongowych. Pamiętając, że piksele będą połączone w łańcuch, zrobiłem to w ten sposób, że ustawiłem naprzemiennie rzędy w różnych kierunkach, najpierw od prawej do lewej, potem od lewej do prawej, potem od prawej do lewej i tak dalej. Dzięki temu przewody sygnałowe są krótkie podczas łączenia końca jednego rzędu z początkiem następnego.

Potrzebowałem trzech krótkich przewodów do połączenia sąsiednich pikseli: 5 V, 0 V i sygnał danych. Problem polegał na tym, że występowały różnice spowodowane takimi rzeczami jak przycinanie segmentów na nieco inne długości i niedokładne wyrównanie otworów w piłeczkach pingpongowych. W rezultacie niektóre przerwy były mniejsze, a inne większe. Zamiast docinać przewody indywidualnie, zdecydowałem się zrobić odcinki z drutu wystarczająco długie, aby wystarczały dla najszerszych szczelin, podczas gdy odcinki izolacyjne były wystarczająco krótkie, aby pasowały do najwęższych szczelin (**rysunek 4**).

Tak więc, pomijając końce rzędów, musimy wykonać 11 grup po 3 połączenia w 12 rzędach, co równa się wykonaniu 396 małych przewodów (na co straciłem kilka godzin). Pozostawiłem kawałek izolacji, ponieważ nie chciałem, aby podczas montażu tych przewodów moje szczytce z długimi noskami działały jak chłodzące radiator, pogorszając jakość lutowania. Z perspektywy czasu odkryłem, że to naprawdę nie był problem i mogłem zaoszczędzić sobie wiele wysiłku, po prostu używając krótkich odcinków nieizolowanego ocynowanego drutu miedzianego wyciętego z rolki. Tak czy inaczej, stworzyłem filmy pokazujące proces montażu przewodów (<https://bit.ly/2MCQ7aI>, <https://bit.ly/3dHzgPt>).

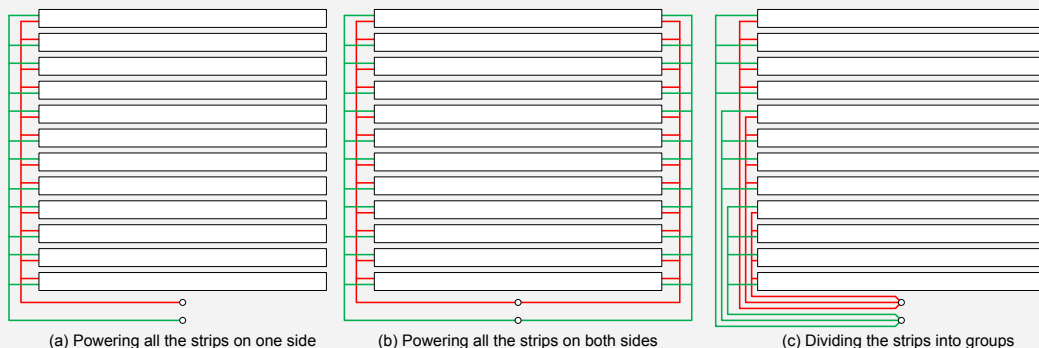
Poczuj moc!

W tym momencie sprawy znów zaczęły się robić interesujące, ponieważ byłem teraz na etapie podłączania zasilania.

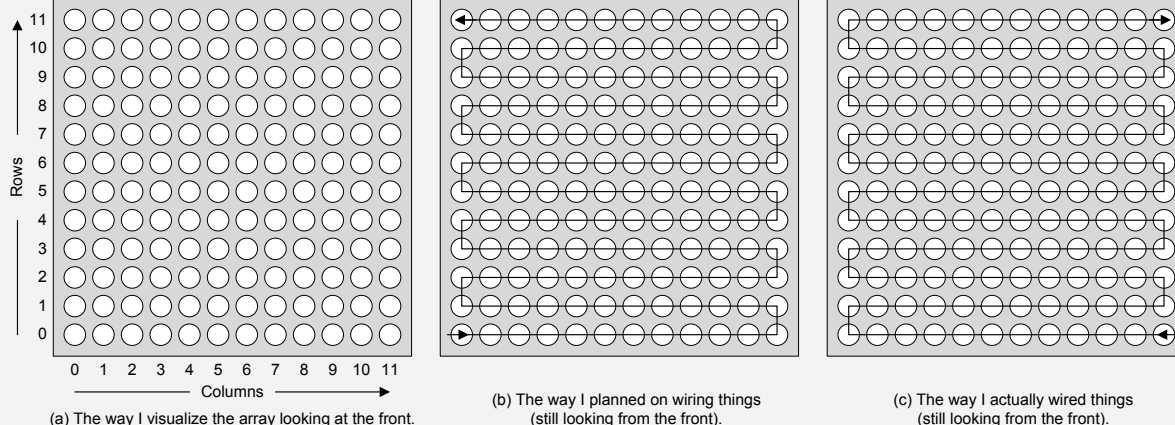
Każdy NeoPixel zawiera czerwone, zielone i niebieskie subpiksele. Dane katalogowe mówią, że przy pełnym włączeniu każdy subpiksel zużywa 20 mA, więc każdy NeoPixel będzie zużywał 60 mA. Szczepie mówią, kiedy mierzyłem to w rzeczywistości, nigdy nie widziałem więcej niż 45 mA dla w pełni włączonego NeoPixela, więc jest to wartość,

którą zwykle sam stosuję co wcale nie znaczy, że rekomenduję, tudzież narzucam taką praktykę Czytelnikowi.

Kolejną kwestią jest to, że podczas użytkowania rzadko zdarza się, aby wszystkie elementy tablicy były w pełni włączone. Byłbym zaskoczony, gdybyśmy mieli średnio przez większość czasu



Rysunek 5. Różne scenariusze okablowania zasilania



Rysunek 6. Kiedy plany spotykają się ze światem rzeczywistym

złączonych więcej niż 20% pikseli. Mówiąc to, mam jednocześnie świadomość, że zawsze powinniśmy projektować pod kątem najgorszego scenariusza (co jest sprzeczne z moim użyciem 45 mA vs. 60 mA).

Mamy 145 pikseli, jeśli uwzględnimy „ofiarny” piksel, nawet jeśli nie planujemy go do niczego używać, najgorszy scenariusz pełnego wykorzystania zgodnie z notą katalogową wyniósłby $145 \times 60 = 8,7$ A. Dla porównania, najgorszy scenariusz zakładający 45 mA na piksel wyniósłby $145 \times 45 = 6,6$ A.

Pierwotnie planowałem zasilanie wszystkich pasków po jednej stronie, jak pokazano na **rysunku 5a** (zauważ, że powodem, dla którego przewody zasilania i uziemienia zmieniają się od góry do dołu od paska do paska, jest to, że naprzemienne paski mają swoje ścieżki sygnałowe biegnące od prawej i od lewej do prawej, co również zamienia orientację sygnałów 5 V i 0 V).

Niestety, jedyny miedziany przewód, który miałem pod ręką w mojej skrzyni z drobiazgami, i którego używam do zasilania, ma prąd znamionowy 3,5 A. Oczywiście, jak w przypadku wszystkiego w elektronice, istnieją niezliczone sposoby rozwiązania tego problemu. Jednym ze sposobów byłoby trzymanie się scenariusza przedstawionego na **rysunku 5a**, ale użycie par przewodów. Dałoby nam to prąd znamionowy 7 A, co spełniłoby nasze wymagania 6,6 A.

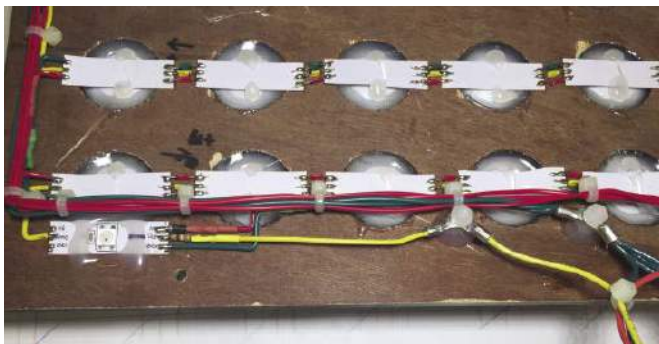
Alternatywą byłoby pozostanie przy pojedynczych przewodach, ale podłączenie jednego zestawu po lewej stronie pasków, a drugiego po prawej, jak pokazano na **rysunku 5b**. Ten scenariusz, który również zapewni wydajność 7 A, ma tę zaletę, że jeśli jedno z połączeń zasilania lub uziemienia na środku paska ulegnie awarii (np. z powodu złego połączenia lutowanego), wszystkie piksele na tym pasku będą nadal zasilane.

Podejście, które ostatecznie wybrałem i które udokumentowałem w filmie (<https://bit.ly/2Y6L9bh>), polegało na podzieleniu pasków na trzy grupy po cztery i zasilaniu każdej grupy niezależnie, jak pokazano na **rysunku 5c**, zapewniając mi w ten sposób wydajność 10,5 A, co z niewielką spełnia nawet najgorsze wymagania 8,7 A związane z pełnym włączeniem wszystkich pikseli przy zużyciu 60 mA każdy. Uff!

Twoja druga lewa!

Sposób, w jaki myślałem o mojej tablicy 12×12 , to macierz 12 wierszy i 12 kolumn. Jeśli patrzę na to od przodu, element (0,0) będzie znajdował się w lewym dolnym rogu, jak pokazano na **rysunku 6a**. Na tej podstawie planowałem okablowanie w taki sposób, aby – nadal patrząc od przodu – pierwszy piksel w łańcuchu znajdował się w lewym dolnym rogu, jak pokazano na **rysunku 6b**.

Niestety, kiedy usiadłem przy stole warsztatowym i zacząłem podłączać wszystkie przewody, zapomniałem, że patrzę teraz na tył matrycy i zlokalizowałem pierwszy piksel w łańcuchu w lewym dolnym rogu ale ale z tej konkretnej perspektywy (patrząc od tyłu) – **rysunek 7**. Zauważ, że samotny piksel znajdujący się najbliższej nas jest pikselem „ofiarnym”, który działa jako konwerter poziomu



Rysunek 7. To twoja druga lewa strona!

Y = Rows 0 to 11	11	133	134	135	136	137	138	139	140	141	142	143	144
	10	132	131	130	129	128	127	126	125	124	123	122	121
	9	109	110	111	112	113	114	115	116	117	118	119	120
	8	108	107	106	105	104	103	102	101	100	99	98	97
	7	85	86	87	88	89	90	91	92	93	94	95	96
	6	84	83	82	81	80	79	78	77	76	75	74	73
	5	61	62	63	64	65	66	67	68	69	70	71	72
	4	60	59	58	57	56	55	54	53	52	51	50	49
	3	37	38	39	40	41	42	43	44	45	46	47	48
	2	36	35	34	33	32	31	30	29	28	27	26	25
	1	13	14	15	16	17	18	19	20	21	22	23	24
	0	12	11	10	9	8	7	6	5	4	3	2	1
		0	1	2	3	4	5	6	7	8	9	10	11
		X = Columns 0 to 11											

Rysunek 8. Kolejność (numeracja) NeoPikseli w tablicy

napięcia. Można się tego domyślić widząc czarną diodę 1N4001 podłączoną między zasilaniem 5 V a wejściem zasilania tego piksela; a także rezystor 390 Ω buforujący sygnał danych z mikrokontrolera XIAO. Finał wszystkich moich szaleństw dotyczących okablowania matrycy pokazano na **rysunku 6c**.

Oczywiście strona lokalizacji pierwszego piksela nie ma większego znaczenia, bo i tak wszystko skorygujemy w oprogramowaniu, ale i tak irytuje mnie, że w tak głupi sposób pomyliłem kierunki.

Testowanie, testowanie...

Po zakończeniu okablowania byłem gotowy do przeprowadzenia wstępnych testów. Zawsze opłaca się, aby pierwszy test był tak prosty, jak to tylko możliwe, więc moim odpowiednikiem klasycznego programu „Hello World” było po prostu zapalenie każdego piksela po kolei. Normalnie, nasze 144 piksele byłyby ponumerowane od 0 do 143. Ponieważ jednak w rzeczywistości mamy 145 pikseli, z „ofiarnym” pikselem zajmującym lokalizację 0, piksele w naszej tablicy są ponumerowane od 1 do 144 (**rysunek 8**).

Sercem pierwszego programu testowego jest funkcja pokazana poniżej. Jak można się spodziewać, wynikiem jest podświetlenie pikseli w serpentynowym wzorze, zaczynając od piksela w prawym dolnym rogu i przesuwając się od prawej do lewej i od lewej do prawej, gdy przechodzimy przez łańcuch (**rysunek 9a**).

```
void LightOneAfterAnother (uint32_t thisColor)
{
    for (int iNeo = 1; iNeo < NUM_NEOS; iNeo++)
    {
        Neos.setPixelColor(iNeo, thisColor);
        Neos.show();
        delay(TestCycleTime);
    }
}
```

Główny program wywołuje tę funkcję w kółko, najpierw ustawiając kolor na czerwony, potem na zielony, a następnie na niebieski. Jeśli chcesz, możesz pobrać pełny program do przeglądania i rozważania (plik CB-Aug20-01.txt – dostępny na stronie PE z sierpnia 2020 r.). Stworzyłem również film pokazujący działanie programu (<https://bit.ly/3dIpr3G>).

OK, tutaj sprawy znów zaczynają się robić interesujące. Pamiętaj, że sposób, w jaki chcę wizualizować tablicę – i sposób, w jaki chcę ją traktować w moich programach – to 12 wierszy ponumerowanych od 0 na dole do 11 na górze i 12 kolumn ponumerowanych od 0 po lewej stronie do 11 po prawej stronie. W przyszłości chcemy być w stanie wydawać polecenia (programowo) takie jak „podświetl

piksel w kolumnie 4 w wierszu 2 kolorem czerwonym”.

W naszym drugim teście chcemy rozpocząć od wiersza 0 i podświetlić piksele w kolejności od kolumny 0 do 11, a następnie powtórzyć dla wiersza 1 i przejść do wiersza 11. Wynikowy skan rastrowy powinien wyglądać jak na ilustracji na **rysunku 9b**.

W zależności od doświadczenia, można myśleć o kombinacjach kolumn i wierszy jako o współrzędnych X-Y. Aby to osiągnąć, zmodyfikowałem moją główną funkcję, która teraz wygląda następująco:

```
void LightOneAfterAnother
(uint32_t thisColor)
```

```
{
    int iNeo;

    for (int yInd = 0; yInd < NUM_ROWS; yInd++)
    {
        for (int xInd = 0; xInd < NUM_COLS; xInd++)
        {
            iNeo = GetNeoNum(xInd, yInd);
            Neos.setPixelColor(iNeo, thisColor);
            Neos.show();
            delay(TestCycleTime);
        }
    }
}
```

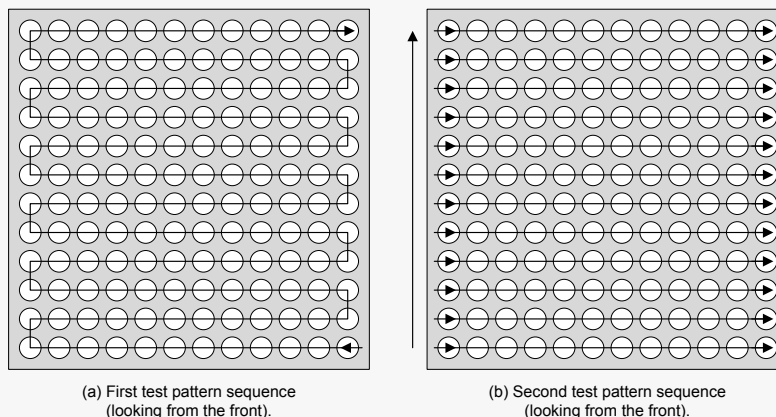
Jak widzimy, mamy zewnętrzną pętlę, która działa w górę wierszy (wartości Y) i wewnętrzną pętlę, która działa w poprzek kolumn (wartości X). Interesująca część to ta, w której wywołujemy funkcję GetNeoNum(), przekazując jej wartości (X,Y) i – miejmy nadzieję – otrzymując numer odpowiedniego NeoPixela w łańcuchu.

Jakiego algorytmu moglibyśmy użyć do zaimplementowania funkcji GetNeoNum()? Myślę, że dobrym ćwiczeniem byłoby zastanowienie się nad tym przed dalszą lekturą. Może naszkicuj coś ołówkiem na papierze.

Nie wiem jak Tobie, ale mnie takie rzeczy nie przychodzą naturalnie. Jestem pewien, że profesjonalni programiści mogliby to zrobić bez zastanowienia, ale ja jestem bardziej wizualnym „rozwiązywaczem” problemów, więc zacząłem od naszkicowania tablicy liczb przedstawionej na **rysunku 8**.

Po chwili zastanowienia zdecydowałem, że jeśli jestem w rzędzie Y, moim punktem wyjścia jest stwierdzenie, że mam (Y/12) pikseli. Następnym krokiem jest określenie, czy jestem w nieparzystym czy parzystym rzędzie. Jeśli jestem w rzędzie parzystym, muszę dodać (12-X) pikseli; dla porównania, jeśli jestem w rzędzie nieparzystym, muszę dodać (X+1) pikseli.

Sposób, w jaki określam, czy jestem w nieparzystym, czy parzystym wierszu, polega na użyciu operatora % (modulo), który zwraca resztę całkowitą z dzielenia liczb całkowitych. Jeśli więc podzielę wiersz Y przez % 2, a wynik będzie równy 0, będziemy w wierszu parzystym; jeśli wynik będzie równy 1, będziemy w wierszu nieparzystym. Kod tej funkcji wygląda następująco:



Rysunek 9. Wyniki pierwszego i drugiego testu

```
int GetNeoNum (int xInd, int yInd)
{
    int iNeo;

    iNeo = yInd * NUM_COLS;

    if ( (yInd % 2) == 0)
    { // Rząd parzysty
        iNeo = iNeo + (12 - xInd);
    }
    else
    { // Rząd nieparzysty
        iNeo = iNeo + (xInd + 1);
    }

    return iNeo;
}
```

Wypróbujmy to. Załóżmy, że chcemy podświetlić piksel znajdujący się w kolumnie 4 w wierszu 2, czyli (X,Y)=(4,2). Najpierw mnożymy wiersz przez liczbę pikseli, więc (Y*12)=(2*12)=24. Następnie dzielimy wiersz przez % 2, aby ustalić, czy jest parzysty, w którym to przypadku musimy dodać (12-X)=(12-4)=8. Tak więc liczba pikseli w łańcuchu odpowiadająca współrzędnym (X,Y) (4,2) wynosi 24+8=32. Wypróbuj to sam, używając kilku przykładowych wartości (X, Y) i sprawdzając wyniki za pomocą rysunku 8.

Jeśli chcesz, możesz pobrać pełny program do przeglądania i rozważania (plik CB-Aug20-02.txt – dostępny na stronie PE z sierpnia 2020 r.). Stworzyłem też krótki film pokazujący to wszystko w akcji (<https://bit.ly/3cFcfLM>).

Teraz jesteśmy naprawdę gotowi, by dać czadu. Co powinniśmy zrobić najpierw? Mam kilka pomysłów, które omówię i zademonstruję w następnym odcinku. Do tego czasu, jak zawsze, czekam na wasze komentarze, pytania i sugestie.



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

REKLAMA

EPcom.pl

Odwiedź stronę z mnóstwem doskonałych projektów



Chrupiące kawałki

Mam przyjaciela na emeryturze, który nazywa siebie „Crusty” i który uczy się programować w C. Kilka tygodni temu wysłał mi e-mail z problemem. Stworzył program dla swojego Arduino Uno z pętlą `for()`, która wyglądała mniej więcej tak:

```
for (i = 0; i <= 10, i++)  
{  
    // Zrób kilka rzeczy  
}
```

Wszystko działało zgodnie z oczekiwaniami, a pętla wykonała się 11 razy. To nie był problem, do którego nawiązałem. Problem pojawił się, gdy Crusty zmodyfikował swój kod, aby wyglądał jak poniżej:

```
for (i = 10; i >= 0, i--)  
{  
    // Zrób kilka rzeczy  
}
```

Problem Crusty polegał na tym, że ta pętla nigdy się nie skończyła. Zamiast tego wykonywała się w kółko. Każdy profesjonalny programista natychmiast odgadnie przyczynę. Z moim kilkudziesięcioletnim, boleśnie zdobytym doświadczeniem, było to oczywiste również dla mnie, ale biedny, stary Crusty po prostu nie mógł tego rozgryźć. A ty?

Nie jesteś w moim typie!

Kiedy ludzie po raz pierwszy zapoznają się z językiem programowania C, jednym z pierwszych typów danych, z którymi się spotykają, jest `int`, który oznacza liczbę całkowitą; na przykład:

```
int MyInt;
```

Zmienne zadeklarowane z tym typem, takie jak `MyInt` w tym przykładzie, mogą przechowywać dodatnie i ujemne liczby całkowite (wiem, że z definicji nie ma żadnych ujemnych liczb całkowitych, ale wiesz o co mi chodzi). Na przykład, `-7`, `0` i `42` są poprawnymi wartościami typu `int`. Kiedy widzimy liczbę taką jak `42`, zgodnie z konwencją zakładamy, że reprezentuje ona wartość dodatnią bez konieczności jawnego pisania `+42`. Podobnie, gdy widzimy zmienną zadeklarowaną przy użyciu typu danych `int`, zakładamy, że mówimy o „podpisanej” liczbie całkowitej, która może reprezentować zarówno wartości dodatnie, jak i ujemne. Możemy to również wyraźnie zaznaczyć, korzystając z poniższego zapisu:

```
signed int MyInt;
```

Oczywiście prowadzi nas to do faktu, że możemy również zadeklarować zmienną całkowitą jako nieoznaczoną, co oznacza, że nie ma ona znaku i może reprezentować tylko wartości dodatnie; na przykład:

```
unsigned int MyUInt;
```

W przeciwieństwie do liczb zapisanych ołówkiem na papierze, które mogą być tak duże, jak tylko chcemy, ograniczone jedynie trwałością naszego ołówka i wytrzymałością naszej ręki, rozmiar liczb przechowywanych w komputerze jest ograniczony ilością pamięci związanej z tym typem danych, których używamy do ich reprezentacji.

Co powiedzieć?

W tym miejscu sprawy zaczynają się nieco komplikować, ponieważ – wierz lub nie – standard C nie definiuje jednoznacznie rozmiaru `int`. Mówi jedynie, że `int` powinien mieć co najmniej dwa bajty (16 bitów). W przypadku Arduino Uno rozmiar `int` rzeczywiście wynosi dwa bajty; w innych komputerach może wynosić cztery bajty lub więcej.

2-bajtowe (16-bitowe) pole może być używane do reprezentowania $2^{16}=65\,536$ różnych wzorców zer i jedynek. W przypadku `int` ze znakiem, wzorce te mogą być używane do reprezentowania wartości ujemnych i dodatnich w zakresie od `-32 768` do `+32 767` (należy pamiętać, że musimy również reprezentować `0`, więc `32 768+32 767+1` [do reprezentowania `0`] równa się `65 536`). Dla porównania, w przypadku `unsigned int`, wzorce te mogą być używane do reprezentowania

tylko wartości dodatnich w zakresie od `0` do `65 535` (ponownie, musimy reprezentować `0`, więc `65 535+1` [do reprezentowania `0`] równa się `65 536`).

Wszystko zostało ujawnione!

Wracając do problemu Crusty’ego (i pamiętając, że wysłał mi tylko kod swojej pętli `for()`), było dla mnie oczywiste, że zadeklarował zmienną `i`, której używał do sterowania pętlą, jako typ danych bez znaku. Założyłem, że jest to `unsigned int` i rzeczywiście tak było.

Ponieważ Crusty myślał, że jego sterowanie pętlą będzie zawsze tylko dodatnie (tj. `>=0`), wpadł w pułapkę myślenia „większe jest lepsze”, decydując się na użycie `int` bez znaku, ponieważ może on przechowywać większe wartości dodatnie, mimo że tak naprawdę nigdy nie planował używać niczego większego niż `10`.

Przeprowadźmy teraz mały eksperyment myślowy. Wiemy, że zmienna typu `unsigned int` na Arduino Uno może być używana do reprezentowania dodatnich wartości w zakresie od `0` do `65 535`. Załóżmy, że ładujemy taką zmienną wartością `0`, a następnie będziemy ją inkrementować (dodawać jeden), aż osiągniemy maksymalną wartość `65 535`. Jak myślisz, co się stanie, jeśli spróbujemy inkrementować ją jeszcze raz? W rzeczywistości, ponieważ wynik przekroczy pojemność tego typu danych, zostanie on przepełniony i powróci do wartości `0`.

Odwrotnie dzieje się w drugą stronę. Oznacza to, że jeśli nasz `unsigned int` zawiera `0` i próbujemy odjąć `1` od tej wartości, wynikiem nie może być `-1`, ponieważ – z definicji – nasz `unsigned int` może zawierać tylko wartości dodatnie. Zamiast tego, `0 - 1` da wynik `65 535`. Chociaż może to nie wydawać się szczególnie intuicyjne, w rzeczywistości ma to doskonały sens, gdy zrozumie się, w jaki sposób te wartości są przechowywane, reprezentowane i manipulowane wewnątrz komputera, ale to dyskusja na inny dzień.

Wyjaśniłem to wszystko Crusty’emu. Poprosiłem go również o dodanie instrukcji `Serial.begin(9600)`; na początku jego funkcji `setup()` i zmodyfikowanie jego pętli `for()`, jak pokazano poniżej:

```
for (i = 10; i >= 0, i--)  
{  
    Serial.print(„i = „);  
    Serial.println(i);  
    // Zrób kilka rzeczy  
}
```

Kiedy Crusty załadował ten nowy szkic i uruchomił okno Serial Monitor, wyświetlona sekwencja zliczania była zgodna z przewidywaniami: „...3, 2, 1, 0, 65 535, 65 534...” Oczywiście wyjaśnia to, dlaczego pętla Crusty’ego nigdy się nie kończy, ponieważ i jest zawsze `>=0`.

Puszka z robakami

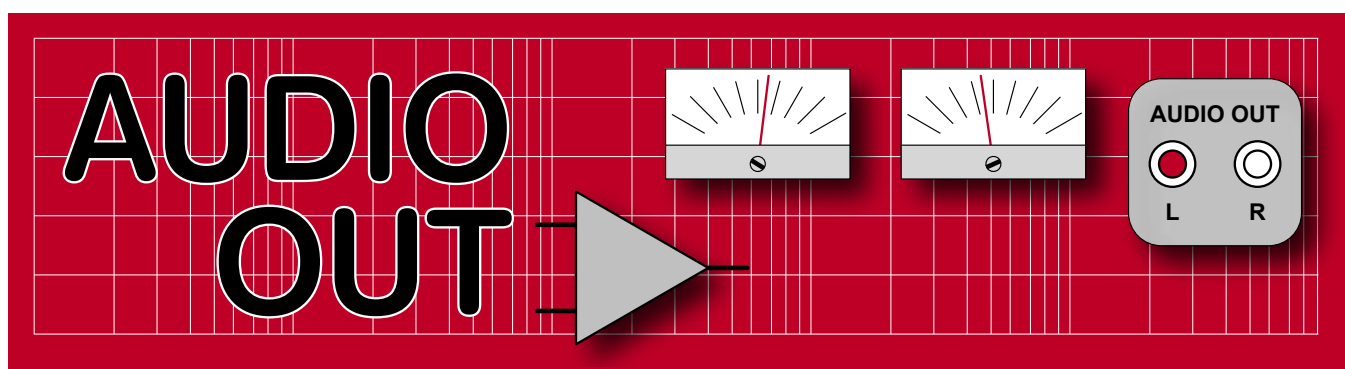
W rzeczywistości pytanie Crusty’ego otworzyło puszkę z robakami – interesującymi robakami, ale jednak robakami – ponieważ mamy również podpisane i niepodpisane wersje innych typów danych, takich jak `short` i `long`.

Mamy również typ danych `char`, którego podpisany lub niepodpisany status nie jest w rzeczywistości określony przez standard C, co oznacza, że może się różnić w zależności od komputera lub – bardziej poprawnie – kompilatora, ponieważ komputer po prostu robi to, co mu powiedziano (w przypadku Arduino Uno i jego kompilatora, `char` ma rozmiar jednego bajta i zachowuje się jak 8-bitowa liczba całkowita ze znakiem).

No i jest jeszcze typ danych `byte`, który tak naprawdę nie istnieje w standardowym C, ale który twórcy Arduino zdecydowali się wrzucić do miksu dla chichotu i uśmiechu. Jestem pewien, że będziesz zachwycony, jeśli ci powiem, że zagłębimy się w to wszystko w następnej części tego cyklu. ■

Clive „Max” Maxfield

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



Budowa radia tranzystorowego, część 2

W zeszłym miesiącu przedstawiłem „Biedronkę” – radio tranzystorowe i doszedłem do zbudowania prostego odbiornika kryształkowego. W tym miesiącu zbudujemy całość wersji tranzystorowej, a dodatkowo zasugerujemy ulepszenia, zachowujące ducha rozwiązania germanowego.

Aktywne radio

Kolejne znaczące ulepszenie to dodanie wzmacniacza tranzystorowego do odbiornika kryształkowego z **rysunku 14**, który uczyni odbiornik głośniejszym. Wystarczająco żebyśmy usłyszeli cokolwiek. Rozczarowałem się, gdy stwierdziłem, że oryginalne układy zasugerowane w książce nie działały właściwie, ze względu na złą polaryzację. Głównym problemem była dioda, było na niej napięcie stałe. Naprawiono to poprzez dodanie kondensatora sprzęgającego i ponowne ustawienie wartości polaryzacji, aby wyjście znajdowało się w połowie napięcia zasilania, jak pokazano to na **rysunku 15**. Układ można uczynić niewrażliwym na h_{FE} tranzystora poprzez dodanie rezystora emiterowego 1 kΩ. Ten należy zbocznikować kondensatorem o pojemności około 22 μF, jak pokazano na **rysunku 16**. Zastąpienie rezystora 2,7 kΩ i słuchawki piezoelektrycznej

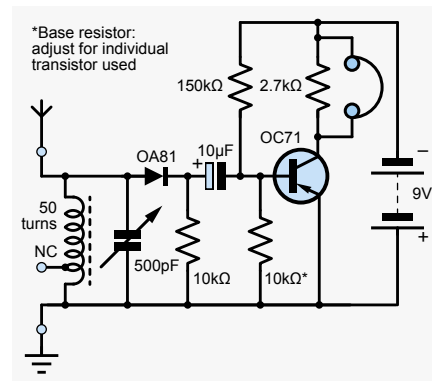
słuchawką telefoniczną polepszyło brzmienie (patrz obwód na **rysunku 17**). Słuchawki piezo mają raczej słabe, metaliczne brzmienie.

Radio refleksowe

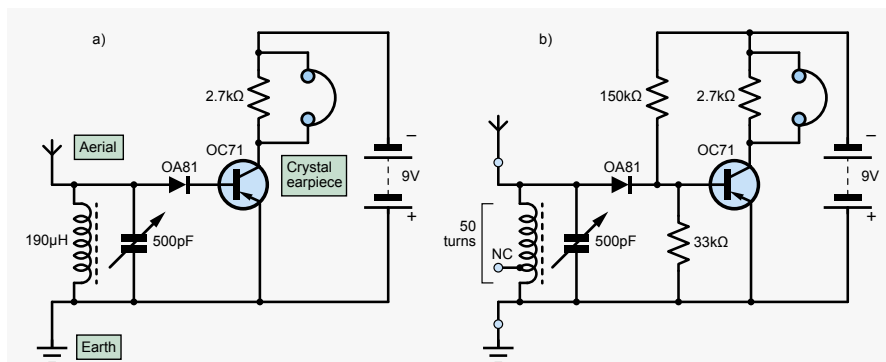
Biorąc pod uwagę prostą konstrukcję, nad którą pracujemy, istniał tylko jeden prosty sposób uzyskania użytecznego sygnału, a mianowicie zbudowanie obwodu radiowego z regeneracją lub refleksją, jak pokazano na **rysunku 18**. Chociaż ten projekt ma zarówno wzmacnienie, jak i selektywność, potrzebuje tylko jednego tranzystora, który jest efektywnie wykorzystywany dwukrotnie; łącząc w sobie zarówno wzmacniacz częstotliwości radiowych, jak i stopień wzmacniacza audio. Obciążeniem audio jest rezystor R2 w obwodzie kolektora, podczas gdy obciążenie RF to dławik L2.

W dawnych czasach tranzystor był znacznie droższy niż cewka indukcyjna. Dziś

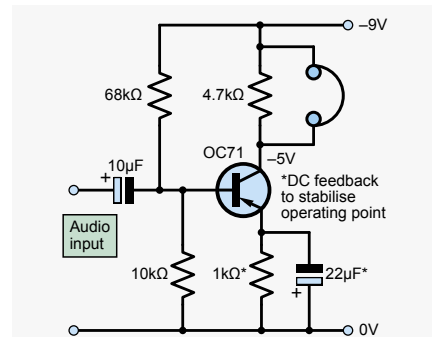
sytuacja jest odwrotna i zastosowano by co najmniej dwa tranzystory, ale przy podwojonym poborze prądu. Stosuje się również



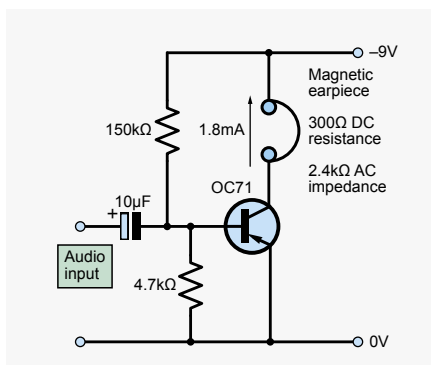
Rysunek 15. Dodanie kondensatora sprzęgającego i regulacja polaryzacji w układzie z **rysunku 14**



Rysunek 14. Dodanie jednostopniowego wzmacniacza do odbiornika kryształkowego: a) tranzystor bez polaryzacji b) tranzystor spolaryzowany ale napięcie na bazie blokuje przewodzenie diody detektora



Rysunek 16. Odbiornik kryształkowy z ulepszonym jednostopniowym wzmacniaczem. Dodanie rezystora emiterowego i bocznikującego go kondensatora do obwodu z **rysunku 15** daje lepszą tolerancję na rozrzut parametrów tranzystora

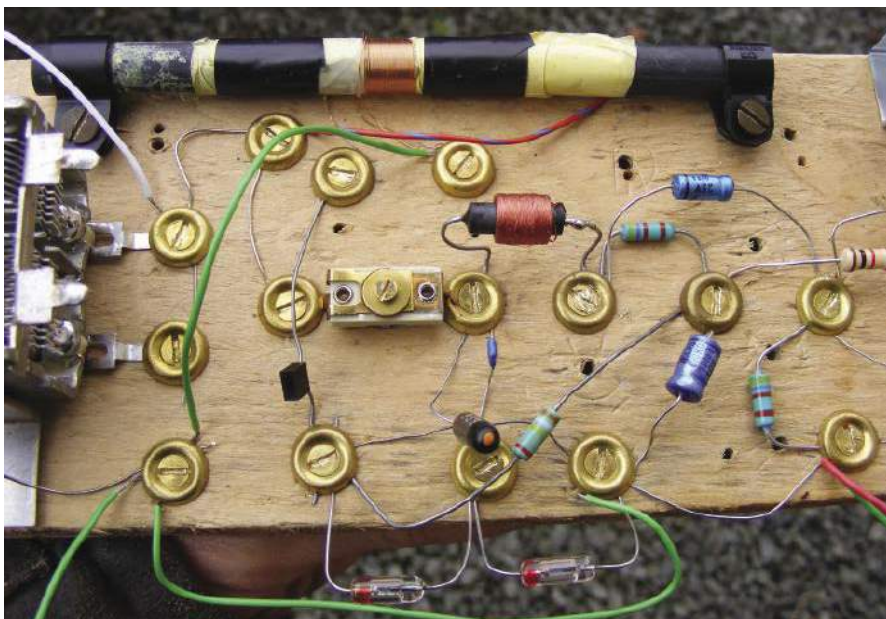


Rysunek 17. Dodanie wkładki telefonicznej; teraz dźwięk jest lepszy niż w samym odbiorniku kryształkowym

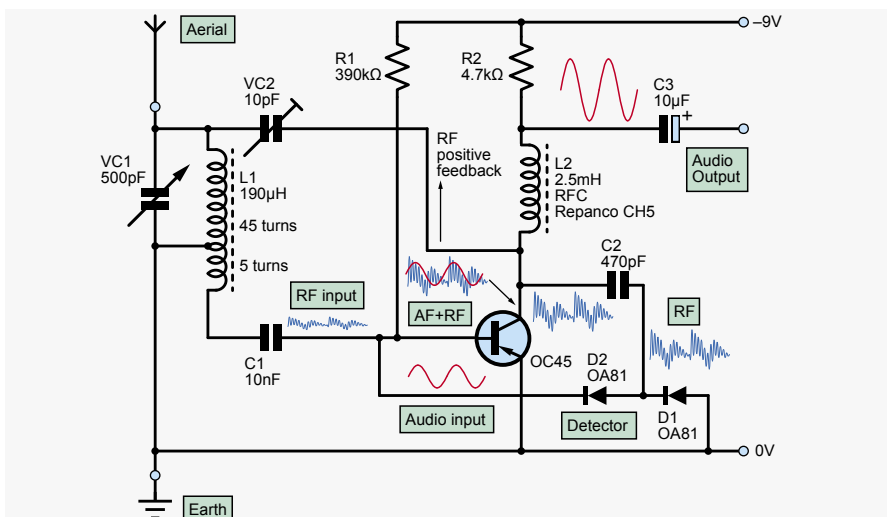
dotądnie sprzężenie zwrotne, „regenerację” lub „reakcję”, jak to się nazywa w żargonie „radiowców”. Służy do najlepszego dostrojenia układu. Podobnie jak w przypadku wszystkich dodatknych sprzężeń zwrotnych, istnieje realne ryzyko, że zbyt duże, może spowodować oscylacje – radio zacznie działać jak nadajnik! Układ dwóch połączonych ze sobą diod służy do wykrywania dźwięku i jest w zasadzie obwodem podwajacza napięcia. Jego wyjście podawane jest ponownie na wejście tranzystorów. Diody germanowe sprawdzają się tutaj najlepiej, ponieważ wersje krzemowe mogą powodować niestabilność. Oryginalną konstrukcję refleksowej Biedronki pokazano na rysunku 19.

Uwagi dotyczące komponentów

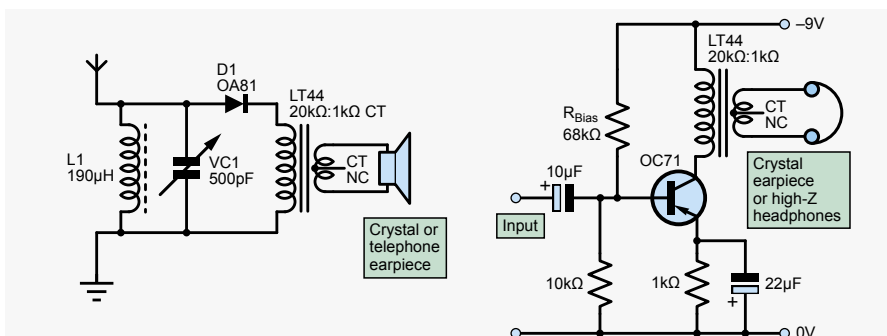
W przypadku takiego retrospektywnego projektu komponenty muszą wyglądać jak należy – nie ma tu żadnych anonimowych, czarnych elementów do montażu powierzchniowego. Dobrym początkiem będzie zakup klasycznych diod w obudowie szklanej.



Rysunek 19. Zmontowany odbiornik refleksowy



Rysunek 18. Odbiornik refleksowy – tranzystor jest używany dwukrotnie. zwróć uwagę na ścieżki sygnałów: częstotliwości audio (czerwony) i w.c. (niebieski)



Rysunek 20. Dodanie transformatora audio do radia kryształkowego zwiększa poziom wyjściowy

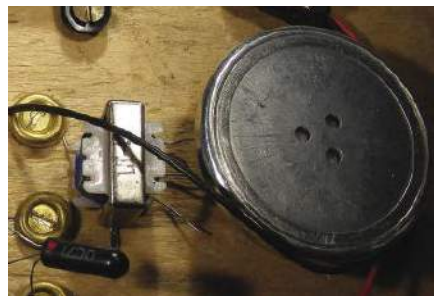
Śruby

Stosowane są wkręty z litego miedzi nr 4 lub nr 6, 3/8 cala (długość 10 mm × średnica 2 mm) z łbem stożkowym. Są one zwykle dostępne tylko z główką szczelinową. Misceczki

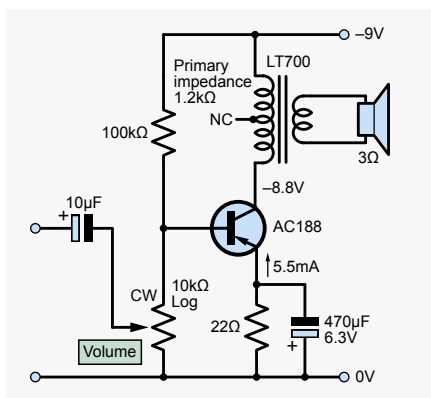
(podkładki stożkowe) mają numer 6 (średnica 11 mm i otwór 3,5 mm). Obawiałem się, że lakier zepsuje kontakt, ale tak się nie stało. Użyłem podkładek nr 5/6 ze śrubami nr 4 i nie było żadnych problemów. Typy pokryte miedzią działają, ale z czasem mogą się utleniać. Podkładki z litego miedzi są dostępne, ale drogie.

Transformatory

Jednym ze sposobów poprawy czułości odbiornika jest podłączenie słuchawki piezo przez transformator słuchawkowy, dopasowujący impedancję z 20 kΩ do 1 kΩ, taki jak



Rysunek 21. Słuchawka telefoniczna – zastosowanie transformatora LT44 również poprawia poziom wyjściowy



Rysunek 22. Ulepszony i z dodanym sprzężeniem zwrotnym stabilizującym punkt pracy, wzmacniacz wyjściowy, (dla uniknięcia płynięcia termicznego punktu pracy)

LT44 pokazany na **rysunku 20**, w celu zwiększenia impedancji obciążenia. Indukcyjność tego elementu tworzy obwód rezonansowy z pojemnością słuchawki, co jeszcze bardziej zwiększa napięcie wyjściowe. Transformator ten dobrze współpracuje również ze starymi wkładkami telefonicznymi o zrównoważonej armaturze, przedstawionymi na **rysunku 21**.

Mają one impedancję dla prądu przemiennego 2,4 kΩ i rezystancję dla prądu stałego 300 Ω. Zastosowanie transformatora w kolektorze wzmacniacza tranzystorowego ze słuchawką telefoniczną ostatecznie zwiększyło napięcie wyjściowe do użytecznego poziomu.

Transformator wyjściowy

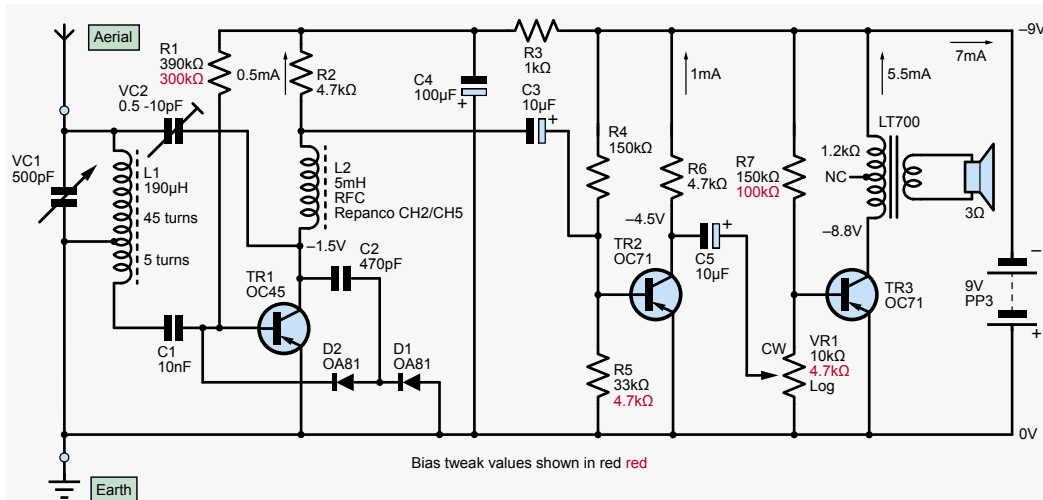
Na wyjściu zastosowano ulubiony transformator hobbystów Eagle LT700. Ma on stosunek impedancji od 1,2 kΩ do 3,2 Ω. Moc wyjściowa radia „Biedronka” jest bardzo niska, rzędu około 50 mW, nieco ponad 1 Vpp przy 3 Ω. Jeśli potrzebna jest podwójna moc wyjściowa, można zastosować LT726 lub większy LT730, które mają uzwojenia pierwotne 500 Ω. Ich uzwojenia wtórne mogą sterować głośniki 3 Ω lub 8 Ω. Jeśli możesz go zdobyć, najlepszy jest Repanco lub RS T/T4, ponieważ został zaprojektowany

tak, że przez rdzeń może przepływać prąd stały. Można też użyć Xicon 42TU400-RC z firmy Mouser. Prąd stały dla tych transformatorów będzie musiał zostać zwiększony do 14 mA poprzez zmniejszenie rezystora polaryzacji RBias (**rysunek 20**) do 51 kΩ, co jednak będzie sporym obciążeniem dla baterii PP3. Zdecydowałem, że oryginalny LT700 będzie ogólnie najlepszy do tego zadania i po wypróbowaniu innych ponownie użyłem właśnie tego. Wypróbowałem także 80 mm głośnik Philips o wysokiej impedancji 150 Ω podłączony bezpośrednio, co dało bardzo dobry rezultat.

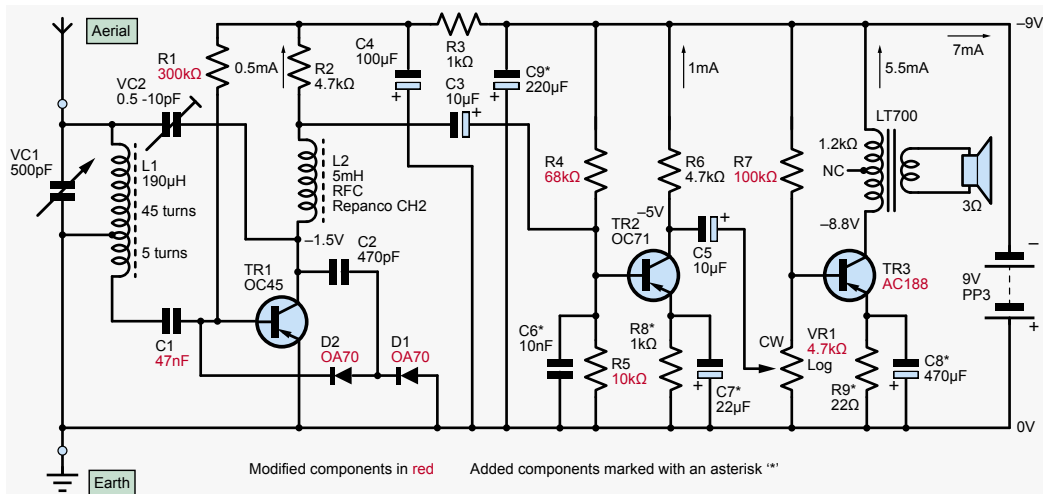
Tranzystory germanowe

Uważaj na typy, które są zbyt „dobre”. Używanie prostego układu polaryzacji (jak na **rysunku 23**) spowodowało, że większość tranzystorów była stale złączona. OC70 o niskim wzmocnieniu był

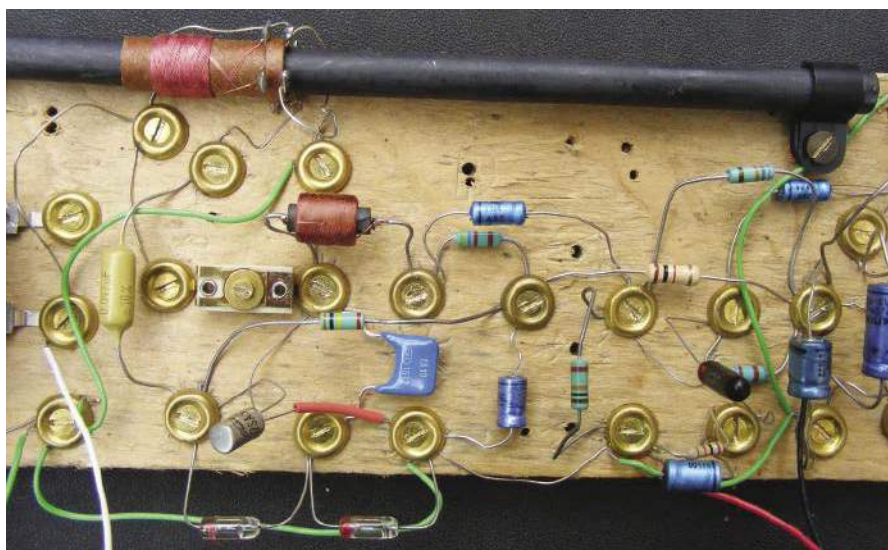
tu najodpowiedniejszy. Zmniejszyłem rezystor polaryzacji R5 z 33 kΩ do 4,7 kΩ, aby uzyskać połowę napięcia zasilania (4,5 V) na rezystorze obciążenia kolektora R6. Bardzo dobrze sprawdziło się tu zwiększenie odporności układu na zmiany wzmocnienia prądowego (h_{fe}), umieszczając zbocznikowany kondensator rezystor emiterowy: 1 kΩ i 22 µF w pierwszym stopniu audio oraz 22 Ω i 470 µF na wyjściu. Na stronie 44 książki „Learn About Simple Electronics.” zauważyłem, że we wzmacniaczu audio zastosowano stabilizację polaryzacji. We fragmencie dotyczącym radia odkryłem, że zmniejszenie rezystora 390 kΩ do 300 kΩ dało dokładniejszą polaryzację przy napięciu 3,5 V na rezystorze obciążenia kolektora R2. Po kilku godzinach zorientowałem się, że tranzystor wyjściowy OC71 wpadł w stan niestabilności termicznej i zaczął trwale przewodzić. Co ciekawe, proces ten został zapoczątkowany



Rysunek 23. Pełny oryginalny schemat odbiornika „Biedronka”. Zrobione tak prosto jak to możliwe, z prostą polaryzacją baz, przez co układ jest wrażliwy na parametry tranzystorów. Jest mało prawdopodobne, że układ ruszy bez dobierania rezystorów oznaczonych na czerwono



Rysunek 24. Zmodyfikowany układ odbiornika z rysunku 23. Pojemność C1 jest większa dla lepszego przenoszenia basów, C6 jest dodane dla filtrowania RF



Rysunek 25. Wersja długofalowa stopnia w.cz. przy użyciu tranzystora 2SA53

emisją RF w wyniku dotknięcia pobliskiej lampy LED. Jak przez mgłę pamiętam, jak mój OC71 przegrzewał się w 1972 roku, co oznaczało śmierć tranzystora germanowego. Jest całkiem prawdopodobne, że nastąpi niestabilność temperaturowa i tranzystor przegrzeje się, ponieważ uzwojenie wtórne transformatora ma rezystancję dla prądu stałego tylko 60 Ω . Wymieniłem więc OC71 na większej mocy AC188 (rysunek 22). Użyłem także potencjometru log 4,7 k Ω do regulacji głośności, ponieważ nie miałem żadnego typu 10 k Ω . Następnie musiałem zmniejszyć rezystor polaryzacji R7 do 100 k Ω . Spowodowało to ustawienie prądu stałego na około 5,5 mA. Moc wyjściowa wynosiła 42 mW przy 3 Ω . Te ulepszenia polaryzacji dodają sześć dodatkowych komponentów do całego radia. Teraz rozumiem, dlaczego George Dobbs pominął je w swoim pierwotnym obwodzie (rysunek 23) – układ był za bardzo złożony dla dzieci. Zmodyfikowany układ z „Biedronki” pokazano na rysunku 24.

Tranzystor na częstotliwości radiowe

Dla wzmocnienia częstotliwości radiowych wybrano OC45. Miałem tylko stary, który sprawiał wrażenie szumiącego. Wypróbowałem nowszy i mniej szumiący Toshiba 2SA53 (rysunek 25), co efekt szumienia zmniejszyło, ale szum wciąż pozostawał. Zwarcie obwodu strojonego ujawniło ciągły syk. Początkowo nie zwracałem sobie głowy dodaniem stabilizacji emiterowej do tego stopnia, ponieważ tranzystory częstotliwości pośredniej (IF) mają mniejszą rozpiętość h_{fe} , zazwyczaj 50, dla częstotliwość (Ft) 4...6 MHz. Właściwe typy RF, takie jak OC44, mogą znów być zbyt dobre ($h_{fe}=80$, $Ft=8...12$ MHz).

Myszę, że George Dobbs był zainspirowany kilkoma bardzo podobnymi układami autorstwa A Sapciyana „Pocket Reflex Receiver” w numerze pisma „Radio Constructor” z sierpnia 1968 r. (str. 27) i „Reflex-3 cale” w wydaniu z grudnia 1968 r. (str. 302). Wykorzystwały one tranzystor AF114 lub AF124, ale martwiłbym się o szum w paśmie akustycznym od tranzystorów germanowych stopowych.

Niekończące się modyfikacje

Wprowadzenie odpowiedniego dzielnika polaryzacji (R1, R12) na rysunku 26 zmniejszyło selektywność RF, ponieważ jego impedancja bocznikowała wejściowy obwód strojony. Rezystor obciążenia (R11) i kondensator separujący (C11) były potrzebne, aby uniknąć zakłócenia punktu pracy przez jego sygnał wyjściowy. Izolowanie diod od prądu bazowego spowodowało usunięcie ich przewodzenia i dalsze zmniejszenie czułości, ale także redukcję szumów (to nie był tranzystor mimo wszystko). Pięć dodatkowych komponentów i nieco gorsza selektywność RF, ale lepsza jakość dźwięku. Ostatni z modyfikowany układ pokazano na rysunku 26. W tym rozwiązaniu, z użyciem 8-calowego pręta ferrytowego LW (długofalowy) okazało się, że kondensator trymera nie zwiększał czułości, a nawet wydawało się, że go zmniejsza.

Ułożenie elementów było także bardziej krytyczne. Dławik L2 musiał być zamontowany pod kątem prostym do pręta ferrytowego, aby zmniejszyć sprzężenie (moc wyjściowa była mniejsza), gdy były w jednej linii. Zdałem sobie sprawę, że to coś wyjątkowego – synergia z oryginalnym prostym obwodem. Warto dobrać rezystor polaryzacji R1 dla konkretnego tranzystora.

Na koniec kondensator odsprężający 220 μ F (C9) należy dodać w poprzek szyny zasilającej by radio się nie wzbudzało w miarę rozładowywania się baterii. To dobra praktyka w dowolnym obwodzie. Ten nowy układ z innymi modyfikacjami w stopniach audio jest ograniczony konstrukcją śrubo-podkładkową (patrz rysunek 34). Być może potrzebna jest teraz porządna płytka PCB?

Kondensator strojeniowy

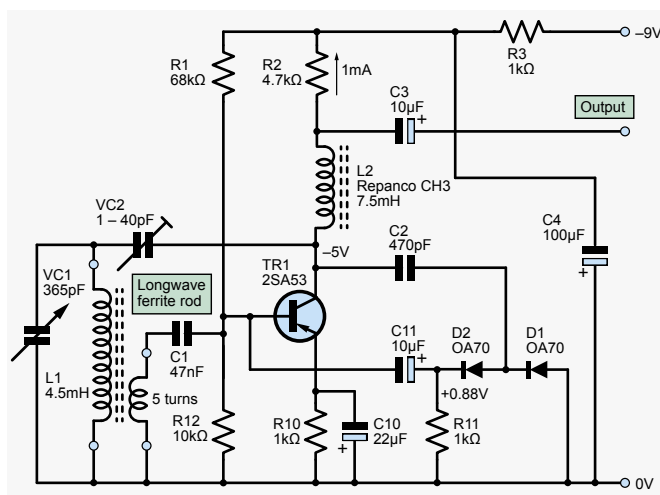
Kondensator strojeniowy ze stałym dielektrykiem Jacksona Dilecona nie jest już dostępny. W sprzedaży dostępny jest jeszcze jednoskrzydłowy typ 100 pF (rysunek 27). Jednakże nadal można jeszcze nabyć (za około 5 funtów) nastawny kondensator powrotny zmienny firmy Plessey, o łącznej pojemności 500 pF (2 sekcje 364 pF i 186 pF) pokazany na rysunku 28. Zaletą są tu napęd wolnobieżny i łożyska kulkowe.

Kondensatory powietrzne mają też niższe straty niż typy z dielektrykiem stałym.

Ponieważ radio może odebrać tylko kilka stacji, można zastosować równolegle kondensator dostrojczy 50 pF ze stałym kondensatorem 150 pF dla oszczędzania pieniędzy, jeśli chcesz słuchać tylko jednej stacji.

Kondensator dostrojczy

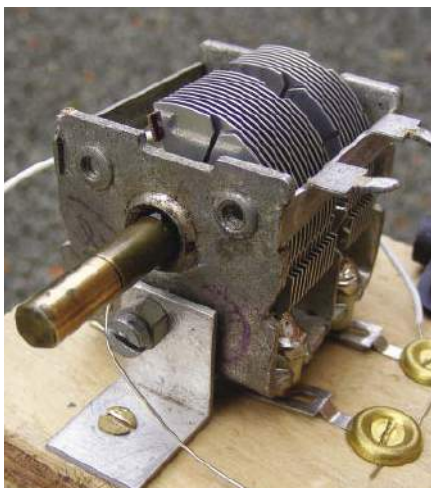
Użyłem trymera z miki 40 pF, chociaż określono 10 pF, egzemplarz, którego użyłem, mógł zejść wystarczająco nisko do około 1 pF. Końcówki muszą być wygięte pod kątem



Rysunek 26. Zmodyfikowany obwód RF z obniżonymi szumami



Rysunek 27. Oryginalny kondensator strojeniowy Jackson Dilecon



Rysunek 28. Polecany kondensator strojeniowy firmy Plessey. Nadal dostępny i mechaniczne dzieło sztuki. Zwróć uwagę na kątownik. Jest potrzebny z przodu ze śrubą 4BA celem dokonania montażu



Rysunek 29. Mouser sprzedaje odpowiedni dławik. Ten jest produkcji Bournsa. Zwróć uwagę na wiele warstw uzwojenia dla redukcji pojemności pasożytniczych



Rysunek 30. Repanco CH2 RFC. Ten wykorzystuje nawinięcie koszykowe i licę, aby zminimalizować straty



Rysunek 31. Gotowe radio – wersja oryginalna



Rysunek 32. Długofalowa cewka antenowa; zauważ jak dwa oddzielne uzwojenia są połączone pod spodem, żeby uzyskać odczep

Wykaz elementów:

Wszystkie komponenty oznaczone „*” przeznaczone są do modyfikacji w stosunku do pierwotnego projektu. Mogę dostarczyć wszystkie części lub zestawy: kontakt tel. 01597 829102 lub wyślij e-mail na adres jrothman1962@gmail.com. Bardzo dobrym dostawcą jest także firma Birketts, którzy wspierają „Biedronkę” i inne projekty George’a Dobbsa od lat. Najlepiej dzwonić w czwartki kiedy Judy Birkett często tu przychodzi. Nie działają online, co w przypadku tego projektu wydaje się całkowicie odpowiednim! J Birkett, dostawcy komponentów radiowych, 25 The Strait, Lincoln LN2 1JF, tel. 01522 520 767.

Rezystory

Używaj dużych, staromodnych typów o potówkowej lub jednowatowej mocy. Używałem głównie ElectroSilu typu TR5 z tlenku cyny z długimi przewodami. Dla autentyczności wizualnej używaj rezystorów węglowych, tolerancja 10% jest wystarczająca.

R1 390 kΩ (*300 kΩ oryginał, 68 kΩ ostateczna modyfikacja)
R2 4,7 kΩ
R3 1 kΩ
R4 150 kΩ (*68 kΩ)
R5 33 kΩ (*4,7 kΩ oryginał, wersja ostateczna 10 kΩ)
R6 4,7 kΩ
R7 150 kΩ (*100 kΩ oryginał)
R8 *1 kΩ
R9 *22 Ω
R10 *1 tys
R11 *10 kΩ
R12 *10 kΩ
VR1 log 10 kΩ (*4,7 kΩ log)

Kondensatory

Najlepiej pasują typy osiowe; Użyłem niebieskich elektrolitów Mullard 017 dla wyglądu. Tolerancja lub napięcie pracy bez ograniczeń.

C1 ceramiczny 10 nF (*47 nF poliester C296)
C2 srebrzona mika lub ceramiczny 470 pF
C3,5 10 μF 25 V

C4 100 μF 10 V

C6 ceramiczny *10 n

C7, 10* 22 μF 25 V

C8* 470 μF 6,3 V

C9* 220 μF 10 V

C11* 10 μF 10 V tantalowy

VC1 365 + 186 pF podwójny Plessey z przekładnią spowalniającą
VC2 10 pF trymer mikowy lub z folii z tworzywa sztucznego

Tranzystory germanowe

TR1 OC45 lub *2SA53 w.cz.

TR2, TR3 OC71 (TR3*AC188 lub AC153) m.cz.

Diody germanowe

D1, D2 OA81 lub * OA70, OA85, CG92,
OA91, IN60 ostrzowe

Elementy indukcyjne

L1 Pręt ferrytowy o średnicy 10 mm, od 7,5 do 20 cm (im dłuższy, tym lepiej)

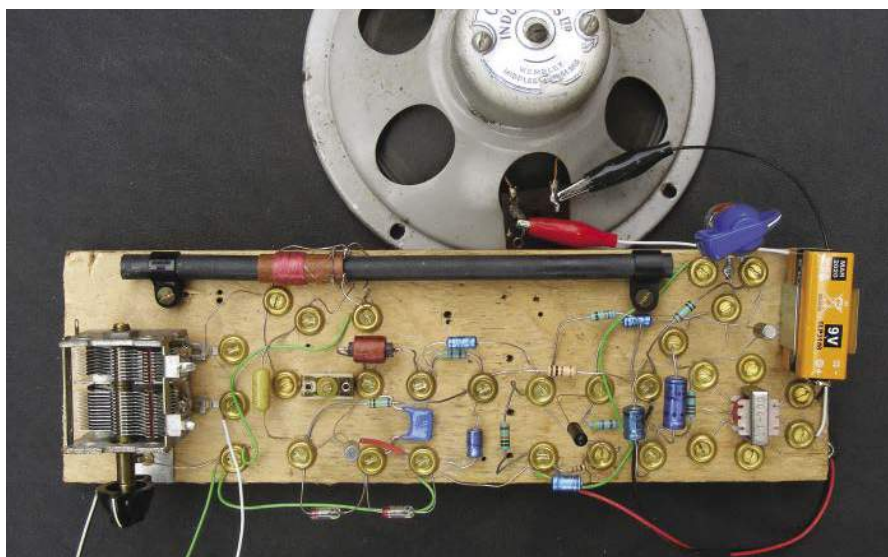
1 metr drutu d.n.e 0.2 mm plus taśma klejąca

Cewki MW lub LW z uzwojeniem sprzęgającym.

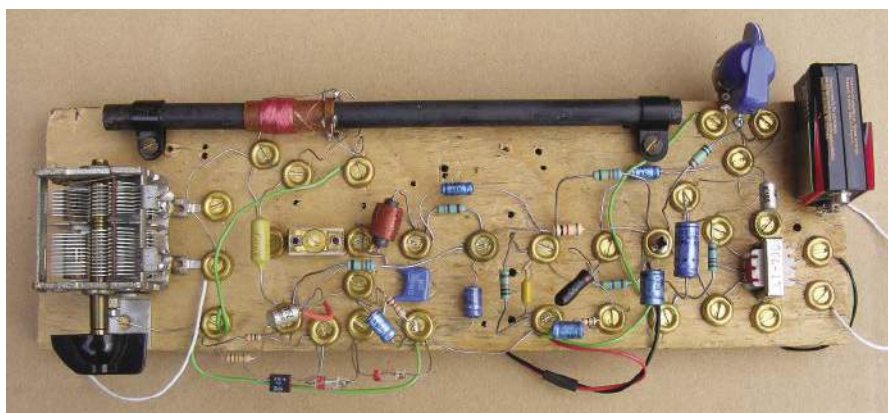
L2 Dławik Repanco 5 mH typ CH2 dla MW lub * 7,5 mH CH3 dla LW.

Inne

Wkręty mosiężne i podkładki



Rysunek 33. Radio kompletne, wersja długofalowa, z podstawowymi modyfikacjami układów polaryzacji



Rysunek 34. Radio kompletne, wersja LW ze wszystkimi modyfikacjami. Zauważ jak L2 jest skrócony w celu zmniejszenia sprzężenia z długim prętem ferrytowym

90°, aby leżeć płasko w celu zaciśnięcia ich podkładkami śrub.

Dławik częstotliwości radiowej (RFC)

W oryginalnym projekcie zastosowano dławik Repanco 2,5 mH CH1. Również Bourns robi takie pasujące typ, 2,4 mH 4666-RC z uzwojeniem koszykowym dostępne w Mouser, pokazany na **rysunku 29**. Dławiki 5 mH typu CH-2 Repanco są nadal dostępne (**rysunek 30**). Zmierzone wykazywały 3,5 mH, co jest w wystarczające dla radia „Biedronka”.

Pręt ferrytowy

Użyłem 75 milimetrowego standardowego pręta ferrytowego, o średnicy 10 mm, z nawiniętymi 50 zwojami emaliowanego drutu o średnicy 0,2 mm, z odczepem na 5 zwoju i zabezpieczonymi taśmą klejącą. Dało to całkowitą indukcyjność uzwojenia wynoszącą około 190 μ H. To może się różnić w zależności od typu użytego ferrytu. Nie

miałem pojęcia, jaki jest mój, więc jest to kolejny obszar, w którym konieczne jest doskonalenie, aby uzyskać pełny zakres strojenia. Jedną z wad konstrukcyjnych „Biedronki” jest pręt ferrytowy z uzwojeniem bez izolacji, co spowodowało, zwarcie 2 uzwojeń. Użyłem pary plastikowych klipsów „P” o średnicy 9,5 mm. Nawijanie ręczne cewek było normą w latach 70. Dla tych, którzy boją się to dziś zrobić, przygotowałem kilka gotowych cewek. Polecam też zastosowanie dłuższego, 20 cm pręta ferrytowego dla lepszego odbioru jeśli nie używasz anteny zewnętrznej.

Smutna historia

To smutne, ale teraz wiem, dlaczego radio mojego dzieciństwa nie działało. W układzie znalazło się wiele drobnych błędów, w czasach, gdy nie wiedziałem, jak je „namierzyć”. Ponadto, jak wielu początkujących, budowałem całość za jednym razem, bez testowania etapowego. Gdybym miał przy sobie starszego mentora z multimetrem, prawdopodobnie odnalazłby błędy w ustawieniu punktów pracy i radio by zadziało. Wydaje mi

się, że George Dobbs używał nadwyżki tranzystorów OC71, która miała bardzo niskie wartości h_{fe} i były najtańszym tranzystorem w Wielkiej Brytanii w tamtym czasie. Namówiłem mamę, aby kupiła urządzenie o pełnej specyfikacji od Electro Value Ltd. Należało zastosować rezystory emiterowe, aby projekt był mniej zależny od h_{fe} , lub dopisać dodatkową sekcję dotyczącą dobierania rezystorów polaryzacji, aby uzyskać odpowiednie napięcia.

Mechaniczna konstrukcja bez lutowania działała jednak zaskakująco dobrze. Byłem zdumiony, jak wrażliwy jest prosty obwód refleksowy, gdy trymer dodatkowego sprzężenia zwrotnego jest ustawiony tuż przed wystąpieniem oscylacji. Nawet zadziało bez anteny i uziemienia. Tylko z samym prętem ferrytowym, program **BBC Radio 5 Live** wypełnił warsztat. Jednakże poza dziesięcioma gwiazdami i brzęczeniem od zasilania, tylko tyle udało mi się uzyskać. Tam była duża czułość, ale mała selektywność. Gotowe radio pokazano na **rysunku 31**.

Pomoc domowa

Podszewam, że większość czytelników wołałaby wersję długofalową, aby uzyskać Polskie Radio PR1 225 kHz.

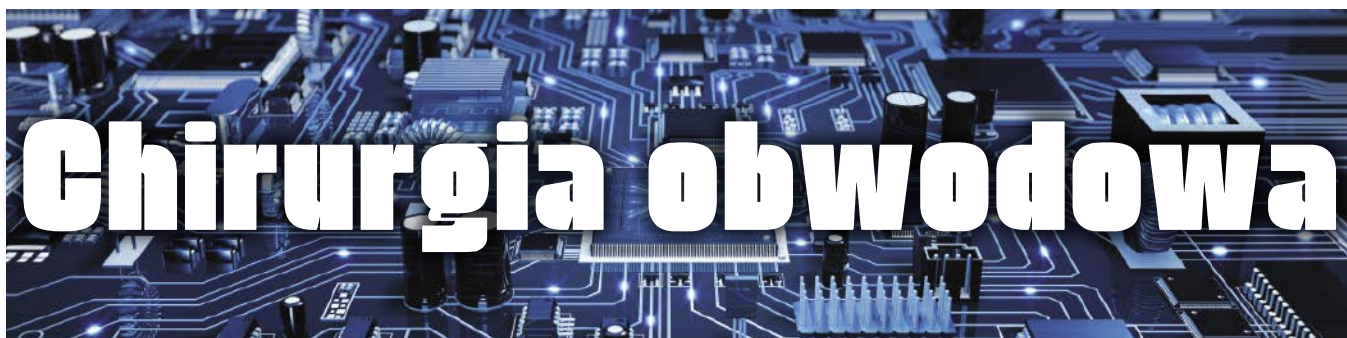
Aby uzyskać wersję LW, układ wymaga kilku zmian w obwodzie strojenia. Użyłem gotowej cewki 4,5 mH LW, z 20 cm prętem ferrytowym. Składa się ona z dwóch oddzielnych uzwojeń, które trzeba połączyć, tak jak pokazano na **rysunku 32**. Mniejsza sekcja kondensatora strojeniowego została odłączona, używana jest tylko sekcja 360 pF. Redukcja pojemności strojącej potrzebna jest również dla gotowych cewek średniofalowych. Dla fal długich, C1 zwiększono do 47 nF a dławik L2 do 7,5 mH (typ CH3). Tak zmodyfikowane radio pokazano na **rysunku 33**, ze starym czułym głośnikiem Goodmansa 3 Ω .

Wersja ostateczna z modyfikacjami polaryzacji, dla wersji releksowej, pokazano na **rysunku 34**. Jest nieco bardziej skomplikowany niż oryginalny projekt, ale nie mogę się oprzeć chęci by wszystko poprawić. Zatrzymałem się tylko na granicy przerobienia całości na elementy krzemowe.

Zużywając skąpe 7 mA, mógłbym słuchać wiadomości przez cały dzień za pomocą jednej tylko baterii PP3. To był także fantastyczny elektromagnetyczny wykrywacz zakłóceń od moich lamp i tester urządzeń „niewidocznych” dla normalnego radia. ■

Jake Rothman

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



Chirurgia obwodowa

Transformatory i LTspice, część 1

W tym miesiącu przyjrzymy się podstawom transformatorów i niektórym aspektom symulowania obwodów z ich użyciem w programie LTspice.

Transformator, to pasywny element elektroniczny, który przenosi energię elektryczną w postaci prądu przemiennego z jednego obwodu do drugiego bez użycia połączenia galwanicznego. Transformator składa się z dwóch lub więcej cewek (uzwojeń) znajdujących się blisko siebie, z których jedna służy do wprowadzania energii/sygnału (zwanego uzwojeniem pierwotnym). Cewka ta wytwarza zmienne pole magnetyczne, które ze względu na bliskie usytuowanie uzwojeń przechodzi przez inną cewkę (uzwojenie bądź uzwojenia wtórne), wytwarzając napięcie na tych uzwojeniach, które spowoduje przepływ prądu, jeśli są one podłączone do obciążenia.

Transformatory są dostępne w bardzo szerokiej gamie typów, od małych elementów do montażu powierzchniowego po ogromne (wielkości domu) transformatory mocy stosowane w krajowych sieciach dystrybucji energii elektrycznej. Wśród nich znajdują się transformatory stosowane w zasilaczach liniowych i impulsowych, transformatory impulsowe stosowane w transmisji danych, transformatory audio i wiele innych specjalistycznych typów. Kluczowe właściwości transformatorów to: izolacja galwaniczna pomiędzy obwodami, możliwość zmiany poziomów napięć

i możliwość zmiany efektywnej impedancji obciążenia. podłączonego za pośrednictwem transformatora.

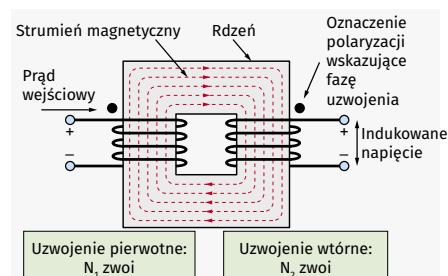
Indukcja elektromagnetyczna

Prąd elektryczny (DC lub AC) wytwarza pole magnetyczne wokół przewodu, w którym płynie. Jeśli drut jest zwinięty w pętlę, pole w środku pętli jest bardziej skoncentrowane. Podstawowe zjawisko fizyczne zachodzące w obrębie transformatora nazywa się „indukcją elektromagnetyczną” i zostało odkryte przez Michaela Faradaya i Josepha Henry’ego w latach trzydziestych XIX wieku (za co zostali wyróżnieni oznaczeniem ich nazwiskami ważnych jednostek układu SI – Farada i Henra.).

Indukcja elektromagnetyczna to wytwarzanie siły elektromotorycznej („emf” mierzonej w woltach) w przewodniku elektrycznym pod wpływem zmiennego pola magnetycznego. W transformatorze zmieniający się prąd w jednym przewodniku wytwarza zmienne pole magnetyczne, które indukuje emf w innym przewodniku. Do indukcji elektromagnetycznej wymagane jest zmienne pole magnetyczne, więc chociaż stały prąd (DC) również wytwarza pole magnetyczne, do działania transformatora wymagany jest prąd przemienny.

Siła elektromotoryczna

Siła elektromotoryczna nie jest siłą mechaniczną – jest to działanie elektryczne wytwarzane przez nieelektryczne źródło energii (np. energia chemiczna z akumulatora lub indukcja elektromagnetyczna w transformatorze). Sprawia, że prąd płynie w obwodzie zamkniętym lub wytwarza napięcie w obwodzie otwartym w wyniku oddzielenia ładunków. W przypadku obwodu otwartego separacja ładunków wytwarza pole

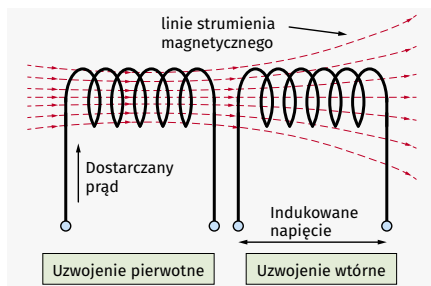


Rysunek 2. Transformator z rdzeniem – jest znacznie bardziej wydajny niż transformator z rysunku 1

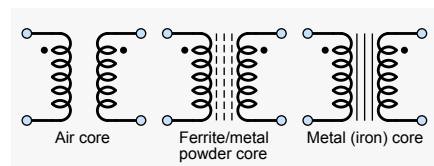
elektryczne, które przeciwdziała separacji ładunku w równowadze z siłą emf napędzającą separację. Napięcie w obwodzie otwartym jest równe emf.

Działanie transformatora

Rysunek 1 przedstawia dwie cewki znajdujące się blisko siebie. Przyłożenie zmiennego prądu do jednej cewki spowoduje wytworzenie pola magnetycznego, które jest wizualizowane jako linie strumienia magnetycznego. Część strumienia magnetycznego przejdzie przez drugą cewkę, powodując indukowanie emf (a tym samym napięcie na cewce). Przez drugą cewkę przechodzi tylko część strumienia magnetycznego, zatem taki układ stanowi transformator o dosyć słabych parametrach – nie będzie on zbyt efektywny w przenoszeniu energii pomiędzy uzwojeniami pierwotnym i wtórnym. Sytuację



Rysunek 1. Najprostszy transformator: dwie cewki połączone strumieniem magnetycznym



Rysunek 3. Symbole transformatorów – zwróć uwagę na rdzenie i oznaczenia polaryzacji

można poprawić dodając rdzeń przecho-
dzący przez uzwojenia transformatora, jak
pokazano na **rysunku 2**. Jeśli materiał i kon-
strukcja rdzenia zostaną starannie dobrane
(w szczególności jego właściwości magne-
tyczne), wówczas większość strumienia będzie
skupiona w rdzeniu i dlatego przejdzie przez
obie cewki, tworząc wydajny transformator.
W idealnym przypadku 100% strumienia zo-
stanie dostarczane do drugiej cewki.

Symbole obwodów transformatora
zwykle składają się z symboli cewek/in-
duktorów ustawionych tyłem do siebie, od-
powiadających uzwojeniom. Linie pomiędzy
cewkami można wykorzystać do wskazania
rodzaju materiału rdzenia. Niektóre przykłady
pokazano na **rysunku 3**.

Rysunek 2 pokazuje kierunek prądu wej-
ściowego i indukowanego napięcia. Aby
prawidłowo używać transformatora, czę-
sto konieczna jest wiedza o tym, w którą stronę
są wykonane połączenia. Jest to oznaczane
za pomocą kropek na elemencie i symbolach
schematycznych, zwanych znacznikami po-
laryzacji bądź kropkami fazowymi. Zaciski
oznaczone kropką mają tę samą polaryzację
napięcia chwilowego – gdy oznaczony kropką
przewód pierwotny jest zasilany przez do-
datnią połowę cyklu prądu przemiennego,
oznaczony kropką przewód wtórny będzie
miał dodatnią polaryzację w stosunku do za-
cisku nieoznaczonego kropką. Podczas cyklu
dodatniego prąd będzie wpływał do oznaczo-
nego kropką zacisku pierwotnego i wypływał
z oznaczonego kropką zacisku wtórnego.

Zwoje cewki pierwotnej i wtórnej

Zależność między napięciem pierwotnym
i wtórnym dla idealnej transformacji jest okre-
ślona przez względną liczbę zwojów w każ-
dym uzwojeniu. W szczególności dla napięcia
pierwotnego (u_p) przyłożonego do uzwojenia
ze zwojami N_p i wtórnego ze zwojami N_s na-
pięcie wtórne będzie wynosić:

$$u_s = \frac{N_s}{N_p} u_p$$

Napięcie wtórne to napięcie pierwotne
pomnożone przez współczynnik zwojów.
Zależność ta dotyczy zarówno wartości sku-
tecznej, jak i szczytowej napięć. Jeśli współ-
czynnik zwojów jest większy niż 1, wówczas
napięcie wtórne będzie większe niż pierwotne
i mamy transformator podwyższający napięcie.
Odwrotnie jest w transformatorze obniżają-
cym napięcie. Ten ostatni jest powszechnie
stosowany w obwodach liniowych zasilaczy
sieciowych w celu uzyskania niskiego napięcia
z sieci elektrycznej 230 V AC.

Transformator przenosi moc z uzwoje-
nia pierwotnego na wtórne. Idealny trans-
formator ma 100% sprawności, więc moc
wejściowa będzie równa mocy wyjściowej,
dla prądów pierwotnych i wtórnych (i_p i i_s)
mamy moc $u_s \cdot i_s = u_p \cdot i_p$. Z powyższego równania
transformatora wynika, że prąd w uzwojeniu
wtórnym wynosi:

$$i_s = \frac{N_p}{N_s} i_p$$

Zatem transformator obniżający napię-
cie będzie pobierał mniej prądu z uzwo-
jenia pierwotnego, niż jest dostarczane
przez uzwojenie wtórne. Na przykład, je-
śli $v_p = 240$ V i $v_s = 12$ V (przełożenie zwo-
jów 20:1) i 1 A zostanie pobrany z uzwojenia
wtórnego, prąd w uzwojeniu pierwotnym wy-
niesie 50 mA (1/20 A).

Jeśli podłączymy rezystor (R_L) do ob-
wodu wtórnego, prawo Ohma mówi nam, że
 $u_s/i_s = R_L$. Korzystając z zależności na przeło-
żenie transformatora:

$$R_L = \frac{u_p}{i_s} = \frac{\left(\frac{N_s}{N_p}\right)u_p}{\left(\frac{N_p}{N_s}\right)i_p} =$$

$$\left(\frac{N_s^2}{N_p^2}\right)\left(\frac{u_p}{i_p}\right) = \left(\frac{N_s^2}{N_p^2}\right)R'_L$$

gdzie $R'_L = u_p/i_p$ jest efektywną rezystan-
cją uzwojenia pierwotnego widzianą przez
źródło go zasilające. W ten sposób trans-
formator zmienił wartość skuteczną rezystora
obciążenia na $(N_p^2/N_s^2)R_L$. Jeśli rezystor
 R_L jest podłączony do strony wtórnej ideal-
nego transformatora, wówczas transformator
będzie wyglądał jak rezystor o wartości R'_L
w źródle zasilającym uzwojenie pierwotne.
Argument ten można zastosować bardziej
ogólnie do impedancji (obwodów o pojemności
i indukcyjności oraz rezystancji) podłączonych
do uzwojenia wtórnego. Ta właściwość trans-
formatora ma zastosowanie w dopasowywaniu
impedancji.

Cewki indukcyjne i transformatory

Cewka sama w sobie jest cewką indukcyjną.
W rzeczywistości prosty drut ma indukcyj-
ność, ale zwykle element cewki indukcyjnej jest
utworzony z jednej lub więcej pętli (zwojów)
izolowanego drutu (czasami bardzo wielu zwo-
jów). Indukcyjność (L) jest powiązana z liczbą
kwadratów zwojów (N^2), ale konkretna za-
leżność zależy od rodzaju, konstrukcji i wy-
miarów cewki indukcyjnej. Ogólnie możemy
zapisać $L = kN^2$, więc $N = \sqrt{L/k}$ dla danego typu
i wielkości konstrukcji, gdzie k jest stałą.
W przypadku idealnego transformatora mo-
żemy założyć, że k jest takie samo dla uzwojenia

pierwotnego i wtórnego, więc podstawiając
powyższą zależność napięcia, otrzymujemy:

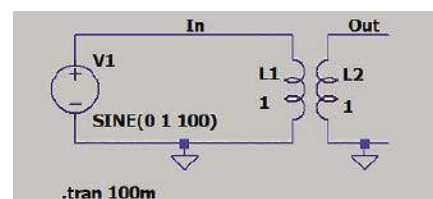
$$u_s = \frac{\sqrt{\frac{L_p}{k}}}{\sqrt{\frac{L_s}{k}}} u_p = \sqrt{\frac{L_p}{L_s}} u_p$$

Stosunek zwojów jest równy pierwiastkowi
kwadratowemu stosunku indukcyjności uzwo-
jeń (rozważanych jako pojedyncze cewki in-
dukcyjne). Ten stosunek indukcyjności jest
ważny przy ustawianiu transformatorów w sy-
mulacjach SPICE.

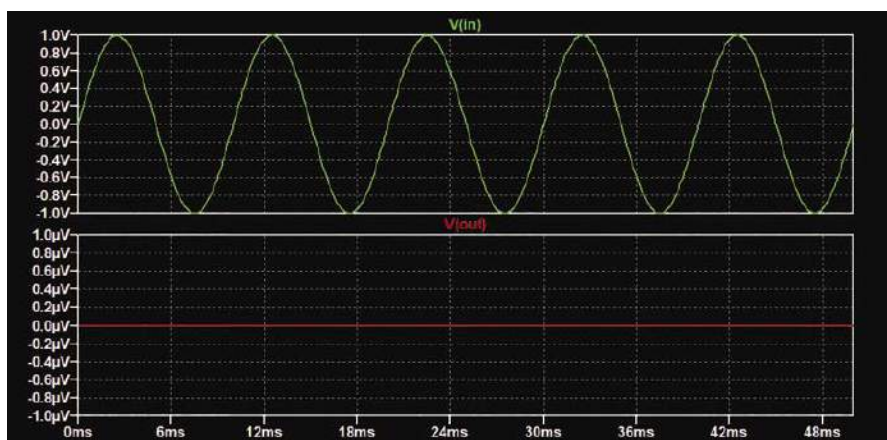
Transformatory to nie tylko cewki induk-
cyjne – powyższa dyskusja wskazuje, że ide-
alny transformator z rezystorem podłączonym
do uzwojenia wtórnego wygląda jak rezystor
dla źródła. Indukcyjność transformatorów ma
jednak znaczenie, ponieważ wpływa na zach-
owanie i wydajność obwodu. Jednak w niektó-
rych zastosowaniach, takich jak transformatory
sieciowe, ma to mniejsze znaczenie. Jeżeli sto-
pień sprzężenia uzwojenia pierwotnego i wtór-
nego nie jest idealny, transformator będzie
zachowywał się częściowo jak transformator,
a częściowo jak cewka indukcyjna – nazywa
się to indukcyjnością rozproszenia. Kolejną
ważną, nieidealną cechą jest rezystancja uzwo-
jeń – w idealnym przypadku jest to zero, ale
przewody w prawdziwych transformatorach
będą miały pewien opór. Wiele właściwości
rdzenia jest również ważnych w rzeczywistych
transformatorach – jednym ważnym, ale złożo-
nym czynnikiem jest nasycenie rdzenia, które
stanowi ograniczenie maksymalnego strumie-
nia magnetycznego prowadzące do często nie-
pożądanych skutków, takich jak zniekształcenie
(mimo słabej wydajności transformatora bez
rdzenia ma on jedną zaletę, brak efektu nasy-
cania się rdzenia).

Praca z LTspice

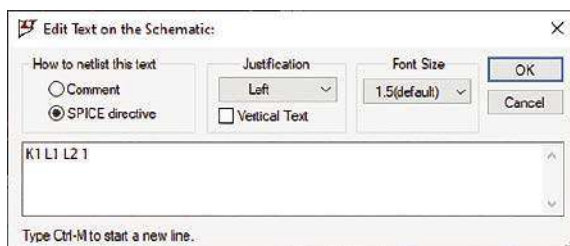
Symulacja transformatorów w LTspice
nie jest tak prosta, jak w przypadku innych
elementów pasywnych, takich jak rezystory
i kondensatory – nie możemy po prostu opu-
ścić symbolu transformatora na schemacie.
Transformator składa się z dwóch lub więcej
cewek, dlatego możemy narysować symbol
transformatora, odpowiednio umieszczając
dwie cewki na schemacie, jak pokazano
na **rysunku 4**. Symbole można odbijać



Rysunek 4. Dwie cewki w LTspice – to jeszcze nie transformator!



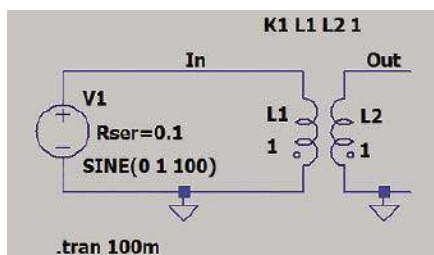
Rysunek 5. Wyniki symulacji z obwodu na rysunku 4



Rysunek 6. Dodanie zestawienia indukcyjności wzajemnych cewek

i obracać za pomocą klawiszy Ctrl-E i Ctrl-R, ale może być również konieczne użycie narzędzia przesuwania w celu zmiany położenia etykiety i wartości, aby transformator wyglądał prawidłowo. Na rysunku 4 przedstawiono napięcie sinusoidalne o częstotliwości 100 Hz i wartości 1 V zasilające „transformator” składający się z dwóch cewek 1H. Transformator powinien mieć przełożenie 1:1, więc powinien dawać napięcie 1 V, ale jeśli zasymulujemy ten obwód, nie otrzymamy żadnego sygnału wyjściowego, jak pokazano na **rysunku 5**.

Problem polega na tym, że jeśli chodzi o LTSpice, mamy po prostu dwie całkowicie oddzielne cewki bez żadnego związku między nimi. Musimy powiedzieć LTSpice, że tworzą transformator, co wymaga umieszczenia deklaracji na schemacie. Jak zapewne wiesz, rzeczywistym wejściem do symulatora nie jest schemat, ale uzyskana z niego lista sieci. Lista sieci to tekstowy opis obwodu wraz z poleceniami instrującymi symulator, co ma robić. Jest generowany automatycznie, ale



Rysunek 7. Obwód transformatora w LTSpice

można go wyświetlić z menu, wybierając opcję Widok -> Lista sieci. Na rysunku 4 lista sieci wygląda następująco:

```

L1 In 0 1
L2 Out 0 1
V1 In 0 SINUS(0 1 100)
.tran 100m
.backanno
.end

```

Wszystko to można skonfigurować za pomocą rysunku lub operacji w menu (polecenie **.tran** Simulation jest generowane poprzez pozycję menu Edit Simulation Cmd). Polecenia **.backanno** i **.end** są automatycznie dodawane do każdej listy sieci. Powinno być oczywiste do czego służy **.end**. Komenda **.backanno** powoduje zapisanie danych umożliwiających sondowanie prądów poprzez kliknięcie na symbole pinów schematu.

Aby utworzyć transformator, musimy dodać element cewki wzajemnej do listy sieci, ale nie można tego zrobić bezpośrednio za pomocą menu. Musimy dodać linię listy sieci dla wzajemnej cewki indukcyjnej jako tekst do schematu. Wszystkie komponenty zaczynają się od określonej litery początkowej

(**L** dla cewki indukcyjnej, **V** dla źródła napięcia itd.). Cewki wzajemne mają początkową literę **K** – być może bardziej oczywista **M** lub **T** dla transformatora jest już używana w tranzystorach MOSFET i liniach transmisyjnych. Składnia instrukcji wzajemnej indukcyjności jest następująca:

Kxxx L1 L2 [L3 ...] <współczynnik>

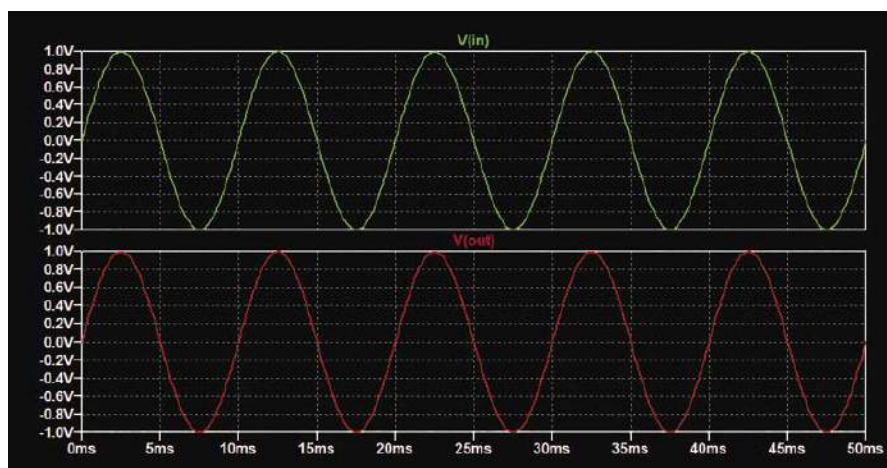
Jest to nazwa rozpoczynająca się na literę **K**, po której następuje lista sprzęganych ze sobą cewek (uzwojenia transformatora) i współczynnik sprzężenia. Dla transformatorów idealnych współczynnik sprzężenia wynosi 1, ale parametr ten można ustawić w zakresie od 0 do 1, aby modelować transformatory, w których nie cały strumień magnetyczny idealnie łączy cewki (część niesprężona tworzy indukcyjność rozproszenia). Zachowanie obwodu może być złożone w przypadku współczynników sprzężenia innego niż 1, może powodować powolne symulacje i często nie jest konieczne. Zaleca się, aby najpierw przeprowadzić symulacje ze współczynnikiem sprzężenia wynoszącym 1, nawet jeśli później mają zostać zbadane inne wartości. Dla obwodu na rysunku 4 potrzebujemy:

K1 L1 L2 1

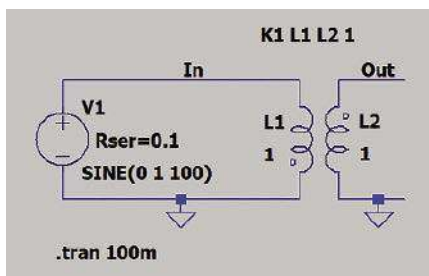
Aby dodać do schematu zestawienie indukcyjności wzajemnej, należy skorzystać z pozycji menu Dyrektywa Spice (**.op**) (patrz **rysunek 6**). Upewnij się, że wybrana jest opcja dyrektywy SPICE (nie komentarz), w przeciwnym razie nie będzie działać. Umieść tekst klikając na schemat obok transformatora (patrz **rysunek 7**). Zauważ, że po dodaniu zestawienia indukcyjności wzajemnej LTSpice automatycznie doda kropki fazowe do schematu – zmień orientację cewek, jeśli nie są one odpowiednie do tego, jak chcesz narysować schemat.

Rdzenie i sprzężenie w LTSpice

Linie rdzenia (patrz rysunek 3) nie są częścią symbolu induktora LTSpice, ale można je dodać jako dodatkowe elementy graficzne.



Rysunek 8. Wyniki symulacji dla obwodu z rysunku 7



Rysunek 9. Faza wyjściowa zmieniona w stosunku do obwodu z rysunku 7

Należy to zrobić klikając prawym przyciskiem myszy → **draw** → **line**, a nie rysując przewód. Rysowanie ich jako części symbolu transformatora nie będzie miało wpływu na symulację.

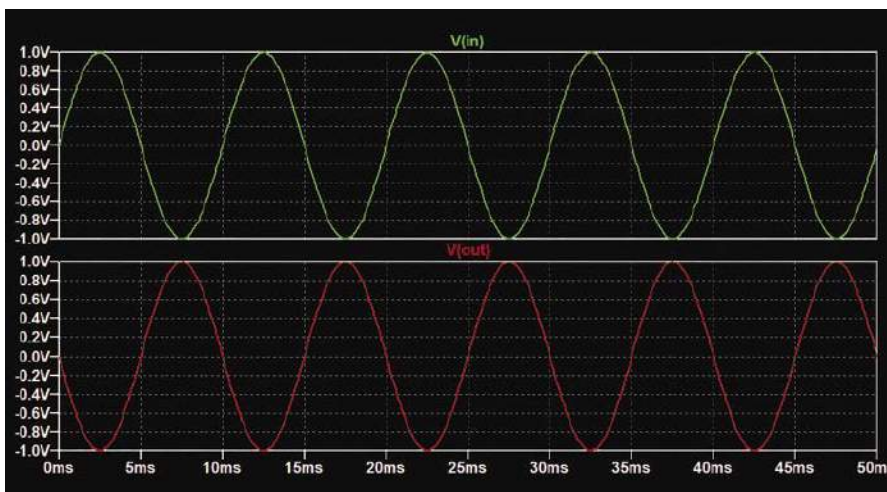
Chociaż możemy zastosować idealny współczynnik sprzężenia, zwykle dobrym pomysłem jest uwzględnienie w symulacji rezystancji szeregowej uzwojenia (w idealnym przypadku wynosi ona zero, ale nie będzie tak w przypadku prawdziwego transformatora). Można ją zmierzyć lub uzyskać ze specyfikacji transformatora. W przypadku obwodu na rysunku 7 dodaliśmy ją do źródła napięcia, co oznacza, że można ją wyświetlić na schemacie, ale działa to tylko dlatego, że mamy źródło napięcia podłączone bezpośrednio do uzwojenia. Rezystancję szeregową można również dodać bezpośrednio do cewki indukcyjnej, klikając prawym przyciskiem myszy bezpośrednio na symbolu cewki indukcyjnej. Wartość ta nie jest jednak wyświetlana, co może wprowadzać w błąd. Można go również dodać jako oddzielny rezystor, co ogólnie może być najlepszą opcją.

Symulacja obwodu z rysunku 7 daje sygnał wyjściowy pokazany na **rysunku 8** – oczekiwany sygnał 1 V.

Rysunek 9 jest taki sam jak rysunek 7, z tą różnicą, że przeciwny koniec uzwojenia został uziemiony (na co wskazuje zmienione położenie kropki fazowej dla L2). Wynikowy sygnał wyjściowy jest przesunięty w fazie o 180° (pół fali sinusoidalnej) z wejście (patrz **rysunek 10** i porównaj z rysunkiem 8).

Wartości indukcyjności dla transformatorów LTspice

Dotychczasowe przykłady wykorzystywały dwie równe cewki indukcyjne, więc cewki wejściowe i wyjściowe są równe na schemacie LTspice. Ponieważ transformator w LTspice jest skonfigurowany z cewek, nie ma bezpośredniego sposobu ustawienia współczynnika zwojów – musimy ustawić stosunek wartości cewki do kwadratu współczynnika zwojów transformatora. Na przykład, jeśli chcemy transformatora podwyższającego napięcie

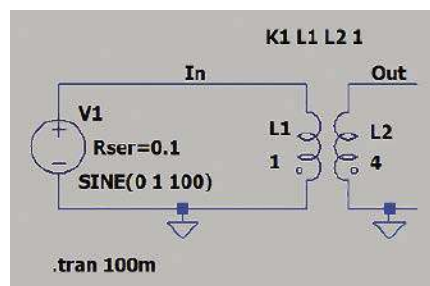


Rysunek 10. Wyniki symulacji dla obwodu z rysunku 9

o przełożeniu 1:2, potrzebujemy przełożenia cewki indukcyjnej 1:4. Pokazano to w obwodzie na **rysunku 11**. Powstałe przebiegi (**rysunek 12**) pokazują napięcie wyjściowe 2 V dla wejścia 1 V.

Nasuwa się pytanie, jakie wartości indukcyjności zastosować w prawdziwym projekcie – w tych przykładach dla wygody użyto po prostu okrągłych liczb. Oczywiście, jeśli wartości indukcyjności są określone dla prawdziwego elementu, należy ich użyć. W przeciwnym razie, jeśli posiadasz odpowiedni miernik (np. miernik RLC dla prądu stałego), to wartość można zmierzyć (przy rozwartym obwodzie pozostałych uzwojeń). W przeciwnym razie indukcyjność należy dobrać tak, aby dawała prąd jawny (I) w kontekście obwodu, stosując $I=U/Z_L$, gdzie U to oczekiwane napięcie uzwojenia, a Z_L to impedancja cewki indukcyjnej przy częstotliwości roboczej (f), znalezione przy użyciu $Z_L=2\pi fL$.

Prawdziwe transformatory charakteryzują się złożonym zachowaniem i znacznymi, nieidealnymi charakterystykami, co może utrudniać symulację (i projektowanie obwodów). Kluczową cechą symulacji jest wspomniana już rezystancja uzwojenia, którą należy zawsze uwzględnić, aby zapobiec nadmiernym prądom stałym. Prądy stałe występują wyraźnie, jeśli w określonym napięciu

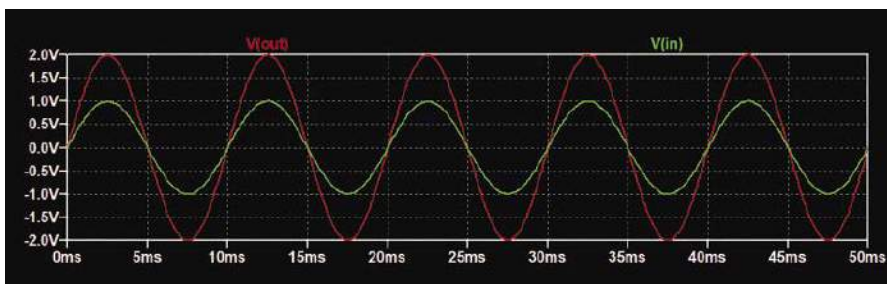


Rysunek 11. Transformator podwyższający o przełożeniu 1:2

wejściowym występuje prąd stały lub są skutkiem zachowania obwodu, ale mogą też pojawiać się w mniej oczywisty sposób ze względu na warunki początkowe stosowane przez symulator podczas rozruchu. Zaznaczenie opcji Simulate → Edit Simulation Cmd → Transient → Skip initial operating point solution może pomóc uniknąć powodowanych tym problemów. Indukcyjność rozproszenia można modelować za pomocą cewek dodatkowych o współczynniku sprzężenia cewek wzajemnych. Modelowanie efektów nieliniowych spowodowanych nasyceniem rdzenia wymaga bardziej złożonych obwodów zastępczych. ■

Ian Bell

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



Rysunek 12. Wyniki symulacji dla obwodu z rysunku 11

Diagnostyka uszkodzeń aktywnych urządzeń elektroakustycznych przy pomocy programu komputerowego audioTester

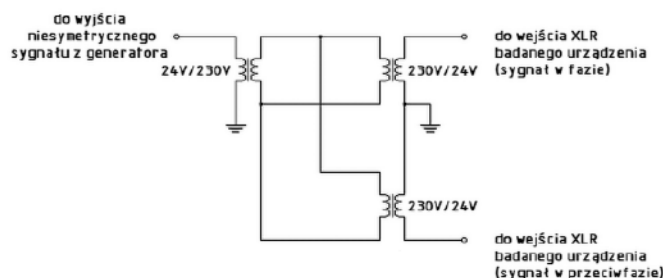
W artykule opisano prostą metodę diagnostyki uszkodzeń aktywnych urządzeń elektroakustycznych, takich jak subwoofery estradowe i monitory studyjne przy pomocy programu komputerowego audioTester. Publikacja zawiera m.in. opis techniczny symetryzatora sygnału XLR a także opis przystawki pozwalającej na przeprowadzanie wspomaganych komputerowo pomiarów charakterystyk amplitudowo-częstotliwościowych oraz fazowo-częstotliwościowych układów filtrów aktywnych oraz wzmacniaczy mocy metodą wobulacji sygnału sinusoidalnego. Urządzenie to może być szczególnie przydatne pasjonatom elektroniki i elektroakustyki pragnącym samodzielnie zajmować się diagnostyką i naprawą urządzeń elektroakustycznych.

Amatorzy elektroniki i elektroakustyki często spotykają się z problemami związanymi z diagnostyką i naprawą aktywnych urządzeń elektroakustycznych. Uszkodzenia w tego typu urządzeniach wynikają w głównej mierze z przeciążenia stopnia mocy podczas pracy w stanie przesterowania lub też są to drobne usterki powstałe pod wpływem drgań urządzenia, w wyniku czego dochodzi m.in. do przerw w obwodach elektronicznych spowodowanych głównie przez słabą jakość PCB oraz zastosowanie do ich montażu spoiw bezołowiowych. Podstawową przeszkodę stanowi konieczność doprowadzenia do wejść urządzenia symetrycznego sygnału XLR. Dostępne w handlu generatory sygnału sinusoidalnego posiadają przeważnie wyjścia niesymetryczne uniemożliwiające podanie dwóch sygnałów w fazie zgodnej i przeciwnej na wejścia urządzenia. W dalszym etapie potrzebna jest także możliwość pomiaru charakterystyk amplitudowo-częstotliwościowych oraz fazowo-częstotliwościowych badanego układu. Tego typu funkcjonalność możemy zrealizować wykorzystując program komputerowy audioTester wraz ze specjalną przystawką, której opis techniczny zostanie przedstawiony w treści publikacji.

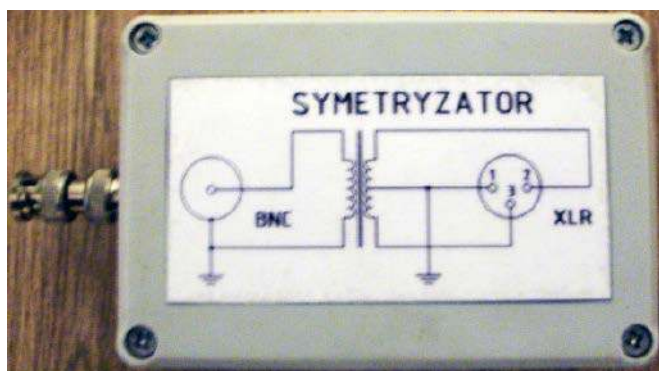
Pasywny symetryzator sygnału XLR

Dostępne w handlu pasywne symetryzatory sygnału XLR są obecnie stosunkowo drogie, chociaż ich niewątpliwą zaletą jest fakt zastosowania specjalnie nawiniętych transformatorów sygnałowych, które umożliwiają pracę w pełnym zakresie pasma akustycznego (tzn. od 20 Hz do 20 kHz). Są to urządzenia profesjonalne, stosowane powszechnie w technice nagłośnieniowej. W sytuacji, kiedy nie dysponujemy odpowiednią kwotą pieniędzy, do zastosowań amatorskich możemy wykonać samodzielnie własny pasywny symetryzator sygnału XLR o dużo gorszych parametrach ale za to przy bardzo niskich kosztach wykonania.

Symetryzator do zastosowań amatorskich możemy wykonać przy pomocy trzech transformatorów sieciowych małej mocy o napięciu



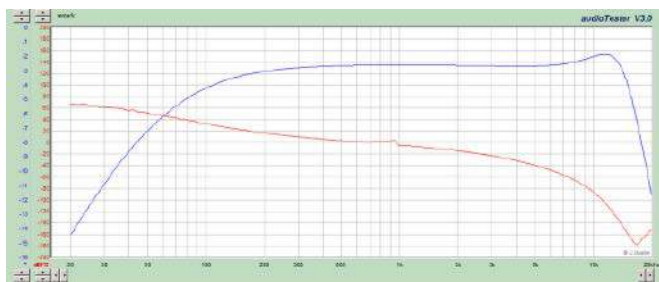
Rysunek 1. Schemat ideowy pasywnego symetryzatora sygnału XLR wykonanego przy pomocy trzech transformatorów sieciowych



Rysunek 2. Wygląd zewnętrzny pasywnego symetryzatora sygnału XLR umieszczonego w obudowie uniwersalnej

pierwotnym równym 230 V i napięciu wtórnym równym 24 V. Warto zabrać ze sobą do sklepu omomierz i wybrać trzy takie same transformatory, których rezystancja uzwojenia wtórnego jest wyższa niż 50 omów. Jest to szczególnie istotne ze względu na fakt ograniczonej obciążalności niesymetrycznego wyjścia generatora sygnału sinusoidalnego. Układ połączeń przedstawiono na rysunku 1. Transformatory łączymy ze sobą uzwojeniami pierwotnymi, natomiast uzwojenia wtórne posłużą nam odpowiednio do podania niesymetrycznego sygnału sinusoidalnego z generatora celem uzyskania dwóch sygnałów w fazie zgodnej i przeciwnej, które podamy na wejście XLR badanego urządzenia elektroakustycznego.

Wygląd zewnętrzny gotowego symetryzatora przedstawia rysunek 2, natomiast na rysunku 3 przedstawione zostały charakterystyki amplitudowo-częstotliwościowa oraz fazowo-częstotliwościowa symetryzatora. Głównymi zaletami takiego rozwiązania są niski koszt wykonania oraz brak konieczności wykorzystywania zewnętrznego źródła zasilania, ponieważ układ jest pasywny. Znaczącą wadę



Rysunek 3. Charakterystyki amplitudowo-częstotliwościowa oraz fazowo-częstotliwościowa pasywnego symetryzatora sygnału XLR

stanowi natomiast wąskie pasmo przenoszenia, ograniczone zarówno w dolnym jak i w górnym zakresie, co w przypadku diagnostyki uszkodzeń, np. dwudrożnych aktywnych monitorów studyjnych nie stanowi jednak problemu, ponieważ częstotliwość sygnału sinusoidalnego z generatora ustawia się przeważnie w okolicach częstotliwości podziału, aby mieć możliwość jednoczesnego badania toru wysokotonowego oraz nisko-średniotonowego. Metodyka wykonywania pomiarów polega na podaniu na wejścia XLR za pośrednictwem symetryzatora sygnału sinusoidalnego z generatora o częstotliwości z zakresu 2...4 kHz i przesłедzeniu całej drogi sygnału od wejścia do wyjścia układu przy pomocy oscyloskopu, celem lokalizacji miejsca, w którym sygnał zanika, a tym samym określenia przyczyny uszkodzenia.

Przystawka do systemu pomiarowego wykorzystującego program komputerowy audioTester

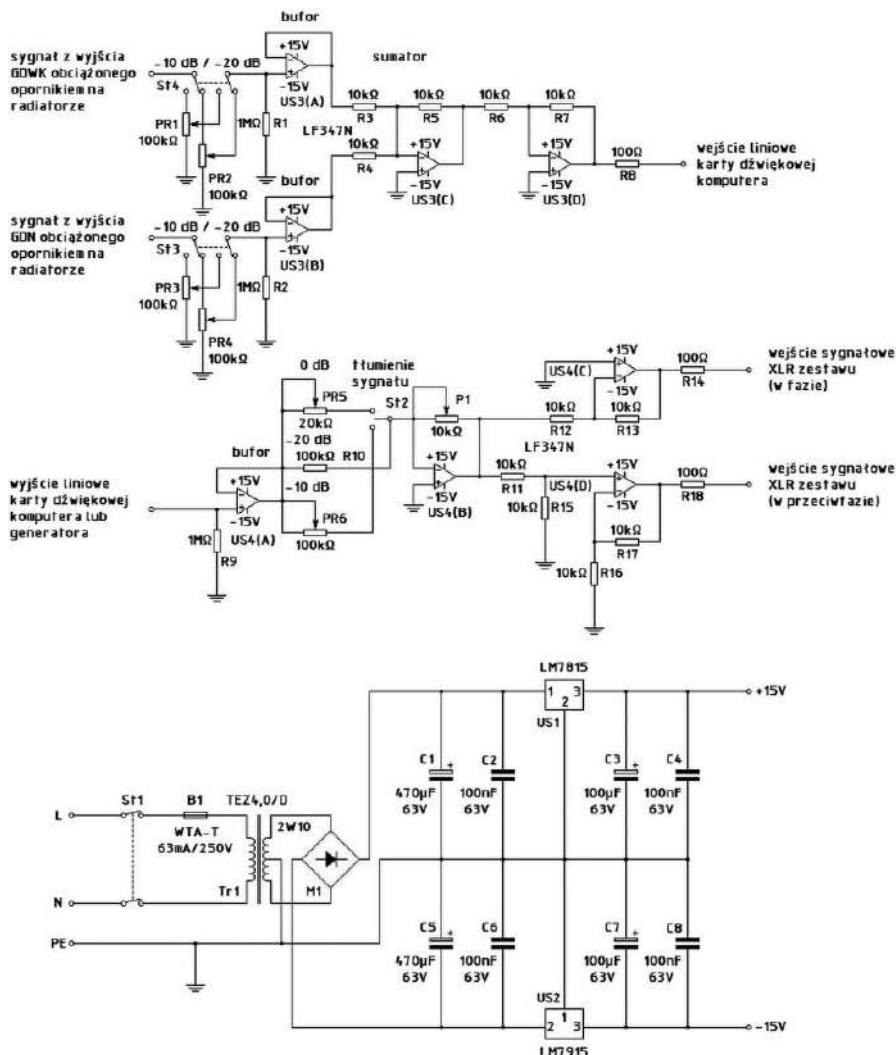
Schemat ideowy układu przedstawiono na rysunku 4. Układ zasilany jest z sieci prądu przemiennego o napięciu skutecznym równym 230 V. Załączenie układu następuje na skutek zwarcia włącznika St1 zasilającego transformator Tr1 za pośrednictwem bezpiecznika B1. W układzie zastosowano popularny czterowatowy transformator typu TEZ4,0/D. Transformator ten posiada odczep ze środka uzwojenia wtórnego umożliwiający wykonanie zasilacza symetrycznego. Za prostownikiem M1 znajduje się zestaw kondensatorów: C1, C2, C3, C4, C5, C6, C7 oraz C8 odpowiadający za filtrację napięć zasilających. Za stabilizację napięć zasilających odpowiadają układy scalone US1 oraz US2. Do pomiarów charakterystyk amplitudowo-częstotliwościowych oraz fazowo-częstotliwościowych przy pomocy programu komputerowego audioTester wykorzystuje się kartę dźwiękową komputera. Sygnał z wyjścia liniowego tej karty trafia na bufor zbudowany w oparciu o wzmacniacz operacyjny US4(A). Kolejny stopień zbudowany w oparciu o wzmacniacz operacyjny US4(B) służy do skokowej i płynnej regulacji wzmocnienia sygnału zadanego. Celem ustawienia odpowiedniego tłumienia, wieloobrotowy potencjometr P1 należy ustawić na maksymalną rezystancję, na wejście wzmacniacza operacyjnego należy podać sygnał z zasilacza stabilizowanego o napięciu +1 V względem masy, przełącznik St2 należy ustawić na pozycję „0 dB” i tak długo ustawiać wieloobrotowy potencjometr montażowy PR5 typu „helitrim” aż na wyjściu wzmacniacza operacyjnego US4(C) uzyska się napięcie o wartości +1 V względem masy, natomiast na wyjściu wzmacniacza operacyjnego US4(D) uzyska się napięcie o wartości -1 V względem masy. W następnej kolejności przełącznik St2 należy ustawić na pozycję „-10 dB” i tak długo ustawiać wieloobrotowy potencjometr montażowy PR6 typu „helitrim” aż na wyjściu wzmacniacza operacyjnego US4(C) uzyska się napięcie o wartości około +316 mV względem masy, natomiast na wyjściu wzmacniacza operacyjnego US4(D) uzyska się napięcie o wartości około -316 mV względem masy.

Korzystamy tutaj z następującej formuły określającej wzmocnienie (lub tłumienie) napięciowe wyrażone w decybelach, stanowiące dwadzieścia logarytmów przy podstawie 10 ze stosunku napięcia wyjściowego do napięcia wejściowego:

$$K_u[dB] = 20 \log_{10} \left(\frac{U_{wy}}{U_{we}} \right)$$

Położenie środkowe przełącznika St2, w którym styk środkowy jest rozwarty ustala tłumienie na stałą wartość równą „-20 dB”. Sygnał na wyjściu wzmacniacza operacyjnego US4(B) jest oczywiście odwrócony w fazie, w wyniku czego wzmacniacz odwracający zbudowany na wzmacniaczu operacyjnym US4(C) przywraca fazę zgodną, natomiast wzmacniacz nieodwracający zbudowany na wzmacniaczu operacyjnym US4(D) podaje na swoje wyjście sygnał w przeciwfazie. Oporniki R14 oraz R18 zabezpieczają wyjścia wzmacniaczy operacyjnych przed przeciążeniem w przypadku wystąpienia zwarcia. Uzyskany w ten sposób sygnał symetryczny służy do diagnostyki wejść XLR badanego urządzenia elektroakustycznego.

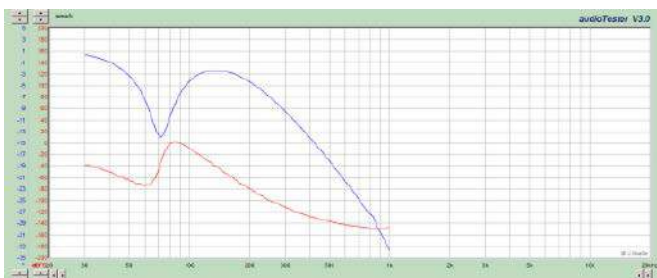
Wyjścia wzmacniacza mocy należy obciążyć opornikami o odpowiedniej mocy i rezystancji, przykręconymi do radiatora. Układ jest dwukanałowy i umożliwia pomiar charakterystyk amplitudowo-częstotliwościowych a także fazowo-częstotliwościowych sumy sygnałów pochodzących z kanału wysokotonowego oraz



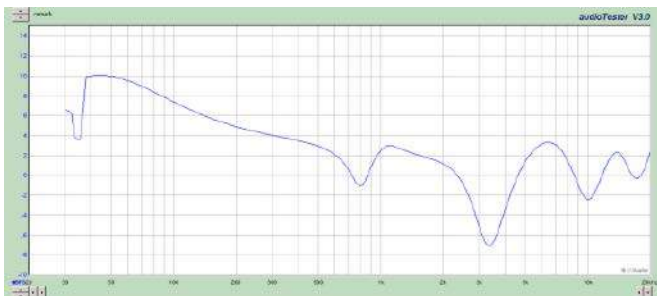
Rysunek 4. Schemat ideowy przystawki do systemu pomiarowego wykorzystującego program komputerowy audioTester



Rysunek 5. Wygląd zewnętrzny gotowej przystawki do systemu pomiarowego wykorzystującego program komputerowy audioTester



Rysunek 6. Przykładowe charakterystyki amplitudowo-częstotliwościowa oraz fazowo-częstotliwościowa subwoofera aktywnego ze wzmacniaczem transkonduktancyjnym zmierzona przy pomocy programu komputerowego audioTester



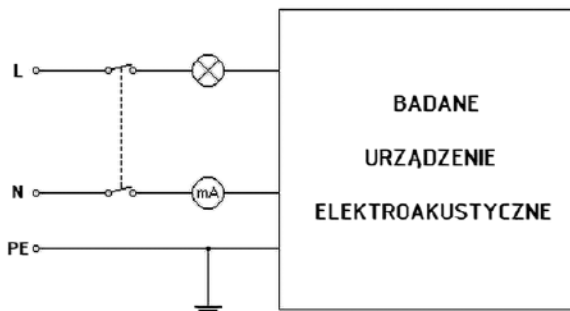
Rysunek 7. Przykładowa charakterystyka amplitudowo-częstotliwościowa sumy kanałów wysokotonowego oraz nisko-średniotonowego badanego aktywnego monitora studyjnego zmierzona przy pomocy programu komputerowego audioTester



Rysunek 8. Kompletny system pomiarowy składający się ze sztucznego obciążenia, przystawki oraz komputera wykorzystującego program audioTester

Tabela 1. Prąd zwarcia odpowiadający mocy elektrycznej zastosowanej żarówki

Moc żarówki	Prąd zwarcia
15 W	65 mA
25 W	109 mA
40 W	174 mA
60 W	261 mA
75 W	326 mA
100 W	435 mA
150 W	652 mA
200 W	870 mA



Rysunek 9. Schemat układu zabezpieczającego badane urządzenie elektroakustyczne przed uszkodzeniem

nisko-średniotonowego dwudrożnego aktywnego monitora studyjnego. W przypadku subwooferów wykorzystujemy tylko jeden z kanałów. Możemy także mierzyć charakterystyki kanału wysokotonowego oraz nisko-średniotonowego osobno, podając sygnał tylko na jedno z wejść układu. Przełączniki St3 oraz St4 ustalają tłumienie sygnałów pochodzących z wyjść wzmacniacza mocy. Za tłumienie sygnału odpowiedzialne są oporowe dzielniki napięcia zbudowane w oparciu o wieloobrotowe potencjometry montażowe PR1, PR2, PR3 oraz PR4 typu „helitrim”. Tłumienie ustawia się w bardzo podobny sposób, wykorzystując źródło napięcia o wartości +1 V (np. zasilacz stabilizowany) oraz woltomierz podłączony do wyjścia wzmacniacza operacyjnego US3(D). Przy prawidłowym ustawieniu wieloobrotowych potencjometrów montażowych PR1, PR2, PR3 oraz PR4, w pozycji „-10 dB” na wyjściu tego układu powinniśmy uzyskać napięcie o wartości wynoszącej około +316 mV, natomiast w pozycji „-20 dB” na wyjściu tego układu powinniśmy uzyskać napięcie o wartości wynoszącej około +100 mV. Wzmacniacze operacyjne US3(A) oraz US3(B) pracują jako buforzy natomiast wzmacniacz operacyjny US3(C) pracuje jako sumator sygnałów pochodzących z obydwu kanałów. Odwraca on jednocześnie fazę sygnału, dlatego potrzebujemy dołączyć na wyjściu jeszcze jeden wzmacniacz operacyjny US3(D) pracujący w układzie odwracającym, który zapewni nam zgodność fazy sygnału wejściowego z fazą sygnału wyjściowego. Opornik R8



Rysunek 10. Okładka książki pt. „Wprowadzenie do projektowania układów elektronicznych subwooferów aktywnych. Poradnik praktyczny”

zabezpiecza wyjście tego wzmacniacza operacyjnego przed przeciążeniem w przypadku wystąpienia zwarcia. Wyjście wzmacniacza operacyjnego US3(D) łączymy bezpośrednio z wejściem liniowym karty dźwiękowej komputera.

Zabezpieczenie badanego urządzenia elektroakustycznego przed uszkodzeniem

Przed przystąpieniem do pomiarów warto jest podłączyć najpierw badane urządzenie elektroakustyczne do sieci zasilającej za pośrednictwem żarówki z włóknem wolframowym (żarówki halogenowe oraz LED się do tego nie nadają) a także miliamperomierza prądu przemiennego. Żarówka o mocy dobranej do mocy badanego urządzenia elektroakustycznego zabezpieczy nam jego układ elektroniczny przed dalszymi uszkodzeniami, jakie mogą nastąpić np. na skutek zwarcia w jednym z uzwojeń transformatora zasilającego, przebicia w którymś z kondensatorów elektrolitycznych filtra zasilacza lub

w którymś z tranzystorów stopnia mocy. Prawidłowo działające urządzenie elektroakustyczne podłączone za pośrednictwem żarówki powinno spowodować krótki rozbłysk żarnika tej żarówki a następnie jego wygaszenie lub lekkie żarzenie. Jeśli żarówka świeci się niemal z pełną mocą przez cały czas a miliamperomierz wskazuje wartość prądu zbliżoną do tej, którą podano w poniższej tabeli, to świadczy to o uszkodzeniu urządzenia i konieczności jego naprawy przed podłączeniem go bezpośrednio do sieci zasilającej.

Książka o układach elektronicznych do subwooferów aktywnych

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

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

mgr inż. Tomasz Łysek

REKLAMA

UWAGA! Tylko prenumeratorzy czasopism Elektronika dla Wszystkich, Elektronika Praktyczna, Świat Radio oraz Elektronik mogą korzystać z atrakcyjnych rabatów w Sklepie AVT:

- ✓ do 50% na wydania specjalne czasopism Wydawnictwa AVT
- ✓ 20% na kity w wersji A (płytki drukowane do projektów AVT)
- ✓ 10% na pozostałe wersje kitów: (A+, B, C, D)
- ✓ 10% na książki
- ✓ 5% na pozostałe produkty z oferty sklepu

K L U B
AVT
ELEKTRONIKA

Ponadto każdy prenumerator ww. czasopism korzysta z rabatów od 30% do 50% na zakup czasopism z oferty www.UlubionyKiosk.pl

Jak uzyskać rabat? Podczas zamówienia powołaj się na swój numer prenumeraty – otrzymasz go mailowo po zakupie prenumeraty wraz z kartą członkowską Klubu AVT-Elektronika.

Regulamin Klubu AVT-Elektronika znajdziesz na stronie <https://sklep.avt.pl/klub-avt-elektronika>

Patronat EdW nad szkołami i uczelnianymi Kołami Naukowymi rozkwiata i daje redakcji EdW impulsy zachęcające do wspierania edukacji szkolnej i uczelnianej. Działa sprzężenie zwrotne. Dostajemy mnóstwo wiadomości od uczniów, nauczycieli i studentów. Dla nich jest ta rubryka.

Wykład 16

Wzmacniacze jedno- i dwutranzystorowe

Wzmacniacze jednotranzystorowe to najprostsze wzmacniacze tranzystorowe, które, jak sama nazwa wskazuje, mają tylko jeden tranzystor jako element wzmacniający.

Idealny obwód dla małych wzmocnień

Podstawowy obwód

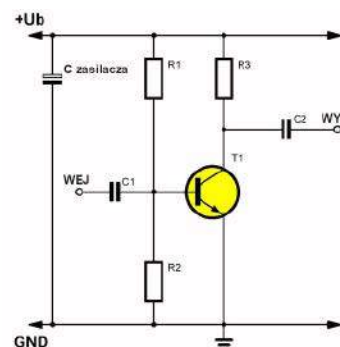
Jeśli zachodzi potrzeba wzmocnienia sygnału maksymalnie dziesięciokrotnie, można to bez problemu zrobić za pomocą obwodu z tylko jednym tranzystorem. Najprostszy schemat jednostopniowego wzmacniacza tranzystorowego przedstawiono na rysunku obok. Baza tranzystora jest spolaryzowana za pomocą dzielnika napięcia $R1/R2$. Emiter jest podłączony bezpośrednio do masy, a wzmocniony sygnał jest pobierany z kolektora. Na bazie zatem ustawione jest napięcie około 0,65 V. Sygnał do wzmocnienia moduluje to napięcie polaryzacji, powodując mniejszy lub większy przepływ prądu bazy w tranzystorze, a to z kolei zmienia prąd kolektora.

Wady

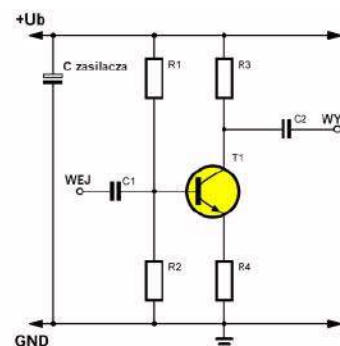
Obwód ten ma duże wzmocnienie prądowe, a tym samym wysokie wzmocnienie sygnału, ale ma też wady. Największą wadą jest to, że obwód jest bardzo niestabilny. Jeśli temperatura tranzystora nieznacznie się zmieni, prąd kolektora wzrośnie lub spadnie, zmieniając ustalony punkt pracy stopnia. Co więcej, nie ma ujemnego sprzężenia zwrotnego, co oznacza, że zniekształcenia sygnału będą bardzo wysokie.

Ujemne sprzężenie zwrotne emitera

Pierwszym usprawnieniem jest dołożenie małego rezystora $R4$ w emiterze, jak pokazano na rysunku obok. Rezystor ten tworzy ujemne sprzężenie zwrotne, które drastycznie zmniejsza wzmocnienie, ale zwiększa stabilność i zmniejsza zniekształcenia. Prąd kolektora przepływa teraz również przez rezystor $R4$ i wytwarza na nim pewne napięcie. To napięcie emitera stabilizuje działanie obwodu. Załóżmy na przykład, że prąd kolektora chciałby wzrosnąć pod wpływem temperatury. W rezultacie większe napięcie odkłada się na rezystorze emitera, powodując spadek napięcia baza/emiter. Baza otrzymuje mniej prądu, przeciwdziałając niepożądanemu wzrostowi prądu kolektora.



Najprostszy schemat wzmacniacza jednotranzystorowego (© 2018 Jos Verstraten)



Stabilizacja wzmacniacza poprzez dołożenie rezystora emiterowego (© 2018 Jos Verstraten)

Mniejsze wzmocnienie

Prąd sygnału wejściowego też będzie generował niewielkie napięcie sygnału na rezystorze emitera R4. Napięcie to powoduje mniejsze otwarcie tranzystora, również dla tego sygnału, a prąd sygnału w kolektorze staje się znacznie mniejszy niż w przypadku bez rezystora emitera. Powoduje to zmniejszenie wzmocnienia sygnału. Zazwyczaj wartość rezystora emitera R4 jest w przybliżeniu dziesięć razy mniejsza od rezystancji kolektora. Wzmocnienie sygnału jest w przybliżeniu określone przez stosunek wartości rezystorów kolektora i emitera.

Odsprężony rezystor emiterowy

Po dodaniu tego jednego małego rezystora emiterowego, wzmocnienie sygnału stopnia spada z kilkuset razy do maksymalnie 10...20. Jeśli ten współczynnik wzmocnienia jest zbyt niski i nadal chcemy pracować z pojedynczym tranzystorem, możemy użyć obwodu przedstawionego na rysunku obok. Duży kondensator C3 jest podłączony równolegle z małym rezystorem emitera. Jak wiadomo, kondensator ma bardzo małą impedancję (czyli rezystancję dla prądu przemiennego). W rezultacie rezystor emitera jest obecny dla prądu stałego, który stabilizuje punkt pracy tranzystora, ale jest praktycznie zwarty dla prądów przemiennych, które przepływają przez stopień. W ten sposób obwód pozostaje stabilny termicznie, ale nadal można osiągnąć wzmocnienie sygnału rzędu kilkuset razy.

Duże rozrzuty

Wydaje się to idealne, ale tak nie jest. Istnieją duże różnice wartości współczynnika wzmocnienia między tranzystorami tego samego typu (nawet w obrębie pojedynczej partii produkcyjnej – przyp. tłum.). W rezultacie wzmocnienie AC obwodu może wahać się między 200 a 800 (każdy typ tranzystora ma inny zakres wzmocnień, zależny też od prądu kolektora – przyp. tłum.), co jest sytuacją niepożądaną. Trzeba więc w jakiś sposób ustabilizować również wzmocnienie AC stopnia. Można to zrobić, rozbudowując obwód do schematu pokazanego na poniższym rysunku. Teraz mały rezystor R5 jest połączony szeregowo z kondensatorem emitera C3. W rezultacie impedancja emitera dla napięcia AC jest teraz określona przez wartość tego rezystora, a wzmocnienie napięcia AC jest określone przez stosunek R3 do R5.

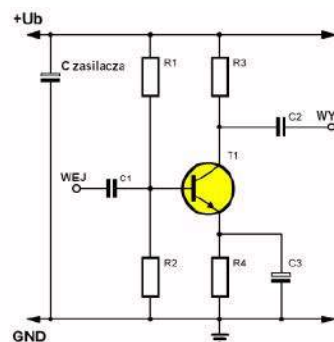
Wniosek: punkt pracy dla napięcia DC stopnia jest stabilizowany przez rezystor R4, a wzmocnienie napięcia AC przez rezystor R5. Dla sygnałów zmiennych R4 i R5 są połączone równolegle, co trzeba uwzględnić przy obliczeniach wzmocnienia, jeśli R4 nie jest znacząco większy od R5 – przyp. tłum.

Przesunięcie fazowe

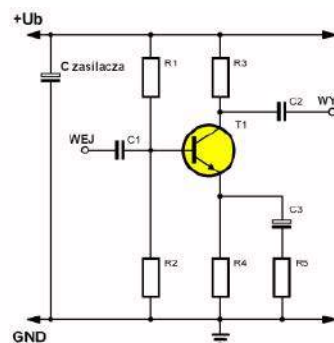
Wzmacniacz jednotranzystorowy z wejściem w bazie i wyjściem w kolektorze jest z definicji obwodem, który powoduje przesunięcie fazowe o 180°. Jeśli sygnał na bazie wzrośnie, wzmocniony sygnał na kolektorze spadnie.

Bootstrapping ze wzmacniaczem jednotranzystorowym

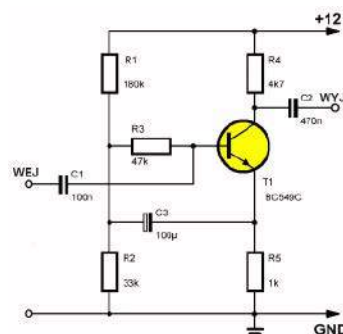
Wzmacniacz jednotranzystorowy ma dość niską impedancję wejściową. Jest ona określana przez rezystory R1 i R2, które są równoległe do sygnału wejściowego. W końcu między napięciem zasilania +Ub a masą znajduje się bardzo duży kondensator elektrolityczny C zasilacza, a ten element ma bardzo małą impedancję dla sygnałów AC. Sygnał „widzi” zatem, że rezystor R2 jest podłączony do masy poprzez impedancję kondensatora zasilacza, czyli jest równoległy do R1 (dodatkowo trzeba doliczyć równolegle połączone rezystory emitera, których wypadkowa rezystancja jest pomnożona przez wzmocnienie prądowe tranzystora – przyp. tłum.). Jednym z możliwych rozwiązań jest dołożenie wtórnika emiterowego przed stopniem wzmacniającym. Można również zastosować zasadę bootstrapu. Jak to działa, pokazano na poniższym rysunku. Duży kondensator C3 jest kondensatorem bootstrapu, który przekazuje sygnał z emitera będący w fazie z powrotem do bazy. Dla napięć zmiennych dzielnik R1/R2 będzie miał znacznie większą impedancję, niż dla napięć stałych. Napięcie na styku rezystorów R1, R2 i R3 podąża teraz za zmianami sygnału wejściowego. Na rezystorze R3 nie ma teraz prawie żadnego spadku napięcia, więc wpływ niskiej rezystancji dzielnika napięcia R1/R2 jest kompensowany. Sygnał wejściowy „widzi” teraz znacznie wyższą rezystancję dzięki sprzężeniu zwrotnemu przez C3. Oczywiście teraz nie da się połączyć dodatkowego kondensatora równolegle z rezystorem emiterowym, ponieważ wtedy żaden sygnał nie może być doprowadzony z powrotem z emitera do bazy. Dlatego wzmocnienie napięcia AC obwodu jest niewielkie, przy pokazanych wartościach komponentów uzyskuje się



Stabilny punkt pracy DC, a jednocześnie wysokie wzmocnienie sygnału (© 2018 Jos Verstraten)



Schemat najbardziej idealnego wzmacniacza jednotranzystorowego (© 2018 Jos Verstraten)

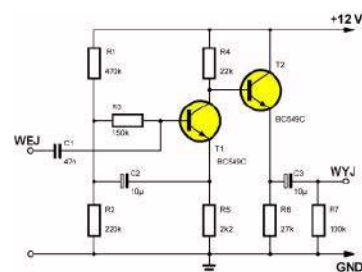


Zastosowanie bootstrapu do wzmacniacza jednotranzystorowego (© 2018 Jos Verstraten)

wzmocnienie około 4,2. Jednak ze względu na efekt działania bootstrapu, obwód ma impedancję wejściową około 500 kΩ, co jest więcej niż wystarczające dla większości zastosowań.

Obniżenie impedancji wyjściowej

Impedancja wyjściowa wzmacniacza z pojedynczym tranzystorem jest dość wysoka. Problem ten można rozwiązać jedynie poprzez dołączenie do wyjścia wzmacniacza wtórnika emiterowego. Oczywiście nie jest konieczne oddzielne ustalenie punktu pracy wtórnika emiterowego. Bazę tego stopnia można podłączyć bezpośrednio do kolektora wzmacniacza. Rezultatem jest układ na rysunku obok, praktycznie bardzo użyteczny schemat wzmacniacza o impedancji wejściowej ponad 0,5 MΩ, impedancji wyjściowej około 47 Ω i wzmocnieniu sygnału około 10.



Wzmacniacz tranzystorowy z bootstrappem na wejściu i wtórnikiem emiterowym na wyjściu (© 2018 Jos Verstraten)

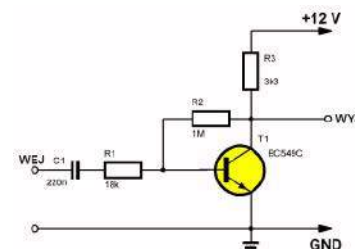
Alternatywna metoda polaryzacji tranzystora

Od kolektora do bazy

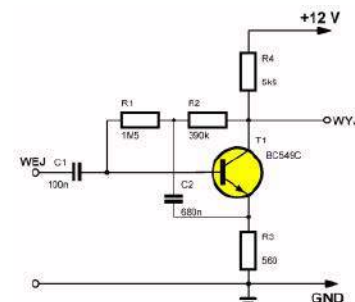
Schematy przedstawione na dwóch poprzednich rysunkach są standardowymi schematami wzmacniaczy jednorozmiarowych. Czasami jednak można znaleźć alternatywny sposób stabilizacji punktu pracy i wzmocnienia wzmacniacza jednorozmiarowego. Na poniższym schemacie emiter jest ponownie podłączony bezpośrednio do masy. Baza nie jest teraz polaryzowana z zasilacza przez dzielnik napięcia, ale z kolektora. To bezpośrednie sprzężenie zwrotne poprzez R2 z wyjścia do wejścia tworzy efekt stabilizujący. Załóżmy, że tranzystor chce przewodzić więcej pod wpływem rosnącej temperatury. Prąd kolektora wzrasta, napięcie na kolektorze spada. W rezultacie mniejszy prąd popłynie do bazy przez R2, co z kolei spowoduje spadek prądu kolektora. Rezystor sprzężenia zwrotnego R2 zapewnia również zmniejszenie wzmocnienia sygnału. Wzmocnienie napięcia AC jest określane przez stosunek wartości R2 i R1, więc w pokazanym przykładzie można oczekiwać wzmocnienia sygnału około 50.

A teraz znowu bootstrapping!

Niska impedancja wejściowa obwodu może zostać skompensowana poprzez zastosowanie zasady bootstrap. Ponieważ zasada bootstrapu zawsze wymaga sprzężenia zwrotnego w fazie, w emiterze musi znajdować się rezystor. Praktyczny przykład technik omówionych w tym artykule jest przedstawiony na poniższym rysunku, wzmacniacz jednorozmiarowy z dobrą stabilizacją, impedancją wejściową około 500 kΩ i wzmocnieniem sygnału około 10.



Alternatywna polaryzacja i stabilizacja wzmacniacza jednorozmiarowego (© 2018 Jos Verstraten)



Zastosowanie bootstrapu do układu z poprzedniego rysunku (© 2018 Jos Verstraten)



Wzmacniacze jednorozmiarowe

Rozwiązanie znajdziesz na www.elportal.pl/quizy

1. Podstawową wadą najprostszego wzmacniacza jednorozmiarowego, w którym emiter dotychczas jest do masy jest:

- ☐ a. zbyt duże wzmocnienie;
- ☐ b. zbyt duże zniekształcenia;
- ☐ c. zbyt duża wrażliwość na zmiany temperatury.

2. Nieduży rezystor połączony szeregowo z emiterem:

- ☐ a. poprawia stabilność kosztem wzmocnienia;
- ☐ b. obniża zniekształcenia kosztem stabilności;
- ☐ c. obniża wzmocnienie chroniąc układ przed nasyceniem się.

3. Kondensator równoległy z rezystorem emitera:

- ☐ a. jeszcze bardziej poprawia stabilność układu;
- ☐ b. stabilizuje napięcie na rezystorze emitera;
- ☐ c. zwiększa wzmocnienie.

4. Rezystor włączony szeregowo z kondensatorem w emiterze:

- ☐ a. ogranicza prąd emitera dla sygnałów zmiennych;
- ☐ b. określa wzmocnienie dla sygnałów zmiennych;
- ☐ c. zmniejsza zniekształcenia.

5. Impedancja wejściowa tego układu:

- ☐ a. jest bardzo niska;
- ☐ b. jest równa rezystancji wypadkowej rezystorów bazy połączonych równoległo;
- ☐ c. jest zależna od rezystora kolektora.

6. Bootstrap:

- ☐ a. zwiększa liniowość wzmacniacza;
- ☐ b. zwiększa wzmocnienie;
- ☐ c. zwiększa impedancję wejściową.

7. Jedynym sposobem na obniżenie impedancji wyjściowej wzmacniacza jest:

- ☐ a. dodanie wtórnika emiterowego;
- ☐ b. dodanie buforu na wzmacniaczu operacyjnym;
- ☐ c. zastosowanie tranzystora o mniejszym wzmocnieniu, i mniejszego rezystora kolektora.

8. Alternatywną metodą polaryzacji bazy tranzystora jest:

- ☐ a. połączenie bazy z kolektorem przez diodę;
- ☐ b. połączenie bazy z emiterem przez rezystor;
- ☐ c. połączenie bazy z kolektorem przez rezystor.

9. By zwiększyć impedancję wejściową alternatywnego układu potrzeba:

- ☐ a. dodać kondensator bootstrapu;
- ☐ b. dodać kondensator bootstrapu i rezystor emiterowy;
- ☐ c. dodać kondensator bootstrapu oraz rezystor i kondensator emiterowy;

10. Kondensator bootstrapu działa, gdyż:

- ☐ a. przewodzi część sygnału wyjściowego w fazie na rezystor polaryzujący zmniejszając spadek napięcia na nim;
- ☐ b. przewodzi część sygnału wejściowego na emiter, przez co układ pracuje częściowo w konfiguracji wspólnej bazy, który to układ ma dużą impedancję wejściową;
- ☐ c. zwiększa spadek napięcia Vbe dla sygnałów zmiennych tym samym znacząco redukując prąd bazy.

Jeśli chcesz uzyskać większe wzmocnienie, niż oferuje wzmacniacz jednotranzystorowy, powinieneś zdecydować się na układy dwutranzystorowe. Obwody te są również znacznie bardziej stabilne i mają mniejsze zniekształcenia.

Obwody z 2 tranzystorami NPN

Czy łączyć dwa identyczne stopnie jeden po drugim?

Niestety, za pomocą wzmacniacza z jednym tranzystorem można osiągnąć dość niskie wzmocnienie sygnału, przynajmniej zakładając, że cenisz sobie stabilizację punktu pracy i niskie zniekształcenia. Jeśli chcesz większe wzmocnienie, musisz użyć więcej niż jednego tranzystora. Zasadniczo można połączyć dwa obwody jednorozmiarowe jeden po drugim, jak pokazano na rysunku obok, tak aby powstał bezpośredni sprzężony obwód dwutranzystorowy. Wzmocnienie pierwszego stopnia wynosi $390\text{ k}\Omega$ podzielone przez $3,9\text{ k}\Omega$ czyli sto razy. Wzmocnienie dla sygnałów zmiennych ustalone jest stosunkiem między rezystancją kolektora R_3 i rezystancją R_5 w szeregu z kondensatorem emitera. Drugi stopień ma wzmocnienie $4,7\text{ k}\Omega$ podzielone przez $470\text{ }\Omega$, czyli dziesięć razy. Całkowite wzmocnienie tego obwodu jest zatem w przybliżeniu równe 1000. W przybliżeniu, ponieważ wzmocnienie pierwszego stopnia będzie nieco niższe z powodu obciążenia tego stopnia przez prąd bazy drugiego tranzystora.

Właściwości obwodu

Wzmocnienie napięciowe tego obwodu nie pozostawia wiele do życzenia. Ale jeśli chodzi o stabilność, należy obawiać się najgorszego. Stabilizacja za pomocą rezystora emiterowego jest przydatna dla pojedynczego stopnia. Niewielkie zmiany polaryzacji nie powodują dużych przesunięć w napięciu kolektora. Ale jeśli użyjesz napięcia kolektora pierwszego stopnia do polaryzacji bazy drugiego stopnia, te małe wahania zostaną dodatkowo wzmocnione przez drugi stopień, co sprawi, że drugi stopień może zostać całkowicie nasycony.

Ujemne sprzężenie zwrotne jest absolutnie konieczne

Jeśli chcesz bezpośrednio połączyć dwa stopnie, musisz użyć jakiejś formy ujemnego sprzężenia zwrotnego, tak aby napięcie wyjściowe drugiego stopnia również określało polaryzację pierwszego stopnia. Układ sam się ustabilizuje, dzięki czemu odchylenia od ustawionych napięć, na przykład pod wpływem temperatury, będą automatycznie kompensowane.

Wzmacniacz dwutranzystorowy z ujemnym sprzężeniem zwrotnym

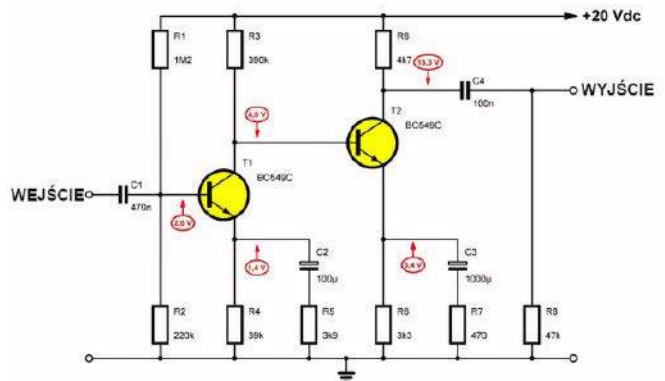
Schemat praktycznie użytecznego wzmacniacza dwutranzystorowego z ujemnym sprzężeniem zwrotnym pokazano na poniższym rysunku. Napięcie na emiterze drugiego tranzystora jest teraz używane do polaryzacji bazy pierwszego stopnia. Zasada działania jest łatwa do zrozumienia. Po włączeniu zasilania wszystkie punkty mają oczywiście potencjał masy. W rezultacie baza pierwszego tranzystora nie jest wystawiona i półprzewodnik ten jest całkowicie zatkany. Napięcie kolektora jest zatem równe napięciu zasilania. Drugi tranzystor będzie teraz pobierał duży prąd bazy przez rezystor R1. W rezultacie T2 jest całkowicie otwarty i przepuszcza duży prąd kolektora przez R5 i R4. Na R5 powstaje duże napięcie, które steruje bazą pierwszego tranzystora poprzez rezystor R3. Ten tranzystor będzie teraz również przewodził, powodując przepływ prądu kolektora przez R1. Spadek napięcia na tym rezystorze powoduje spadek napięcia kolektora i zmniejszenie prądu bazy drugiego tranzystora.

Im bardziej pierwszy tranzystor przewodzi, tym mniej przewodzi drugi tranzystor. Wartość rezystorów R_5 i R_3 określa teraz punkt równowagi, w której oba tranzystory utrzymują wzajemne przewodzenie na stałym poziomie.

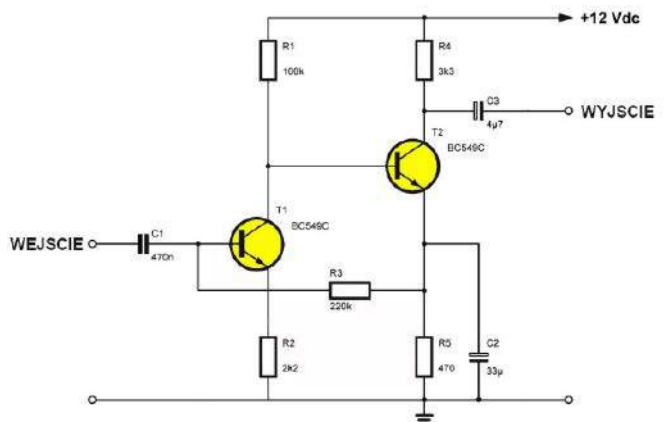
Duża stabilność

W oczywisty sposób taki układ jest bardzo stabilny. Każda zmiana, na przykład, współczynnika wzmocnienia jednego z tranzystorów jest natychmiast korygowana, ponieważ drugi tranzystor przewodzi mniej lub bardziej, a tym samym kontroluje pierwszy tranzystor w taki sposób, by przeciwdziałać tej zmianie.

Stabilizacja obwodu w zakresie napięcia stalego wpływa również na wzmocnienie napięcia przemiennego. W końcu na emiterze T2 występuje również niewielki sygnał, który jest również przesyłany z powrotem do bazy przez rezystor R3. W rezultacie pierwszy tranzystor jest



Najprostszy układ wzmacniacza dwutranzystorowego (© 2018 Jos Verstraten)



Dwutranzystorowy wzmacniacz z ujemnym sprzężeniem zwrotnym
(© 2018 Jos Verstraten)

sterowany mniejszym sygnałem (w końcu sygnał na emiterze T2 jest w przeciwfazie), co powoduje nieznaczne zmniejszenie wzmocnienia sygnału. Jeśli jednak wartość R3 jest bardzo duża w porównaniu do innych rezystorów, ta utrata wzmocnienia jest pomijalna.

Podwójne ujemne sprzężenie zwrotne

Opracowano niezliczone warianty podstawowego schematu z poprzedniego rysunku. Pierwsze sprzężenie zwrotne od emitera do bazy jest często uzupełniane drugim sprzężeniem zwrotnym od kolektora T2 do emitera T1. W końcu te punkty są również w opozycji! Celem jest uczynienie obwodu jeszcze bardziej stabilnym i zminimalizowanie zniekształceń sygnału. Rysunek obok przedstawia praktyczny przykład takiego obwodu. Wzmocnienie sygnału jest równe 46, impedancja wyjściowa jest niska, a impedancja wejściowa wysoka. Sprzężenie zwrotne przez R3 jest rezystancyjne i dlatego wpływa na punkt pracy dla napięcia stałego obwodu. W końcu emiter T2 jest odsprężniony przez C2 i nie zawiera napięcia sygnału zmiennego. Drugie sprzężenie zwrotne przez R6 jest pojemnościowe (w końcu za kondensatorem C3 nie ma napięcia stałego kolektora) i dlatego wpływa tylko na wzmocnienie sygnału zmiennego. Wartość rezystora R6 określa zarówno impedancję wejściową, jak i wyjściową. Wynika to z ujemnego sprzężenia zwrotnego sygnału wprowadzanego przez ten rezystor. W miarę zmniejszania wartości R6 impedancja wyjściowa będzie maleć, ale impedancja wejściowa będzie rosnąć.

Alternatywny obwód

Rysunek obok alternatywny obwód, w którym dwa ujemne sprzężenia zwrotne są rezystancyjne i dlatego oba stabilizują punkty pracy. Ponadto, ponieważ emiter drugiego stopnia nie jest odsprężniony pojemnościowo, ujemne sprzężenie zwrotne poprzez rezystor R4 również wpływa na wzmocnienie sygnału zmiennego. Wzmocnienie sygnału zależy głównie od stosunku wartości rezystorów R3 i R1 i przy pokazanych wartościach wynosi około 300 razy. Kondensatory C2 i C3 służą do nieznacznego ograniczenia szerokości pasma przenoszenia obwodu. Nie ma sensu w budowaniu typowych wzmacniaczy małej częstotliwości z pasmem sięgającym megaherców. Pasma przenoszenia ograniczone do 500 kHz jest więcej niż wystarczające dla tego typu zastosowań. Kondensatory o małych pojemnościach tworzą zwarcie dla wysokich częstotliwości, powodując spadek wzmocnienia dla tych częstotliwości.

Obwody z tranzystorami PNP/NPN

Pierwszy przykład

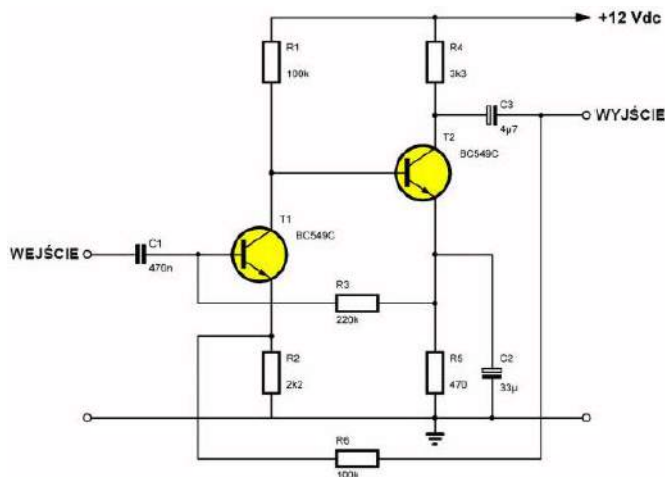
Oprócz wzmacniaczy dwutranzystorowych z dwoma tranzystorami NPN, opracowano różne układy wykorzystujące kombinację tranzystora NPN i PNP. Dużą zaletą jest to, że potrzeba znacznie mniej rezystorów, aby uzyskać identyczne współczynniki wzmocnienia. Poniższy rysunek przedstawia typowy przykład takiego obwodu. Wzmacniacz ten ma bardzo silne ujemne sprzężenie zwrotne. Wadą jest to, że wzmocnienie sygnału spada do około 5, ale z drugiej strony zniekształcenia sygnału spadają znacznie poniżej 0,01%!

Obliczanie wzmocnienia sygnału

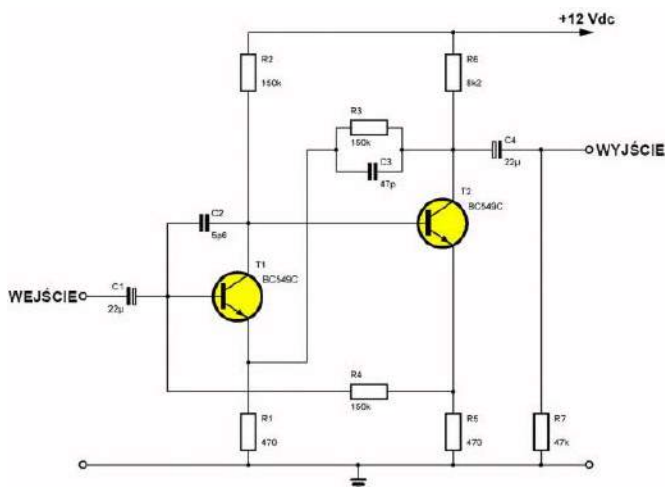
Wzmocnienie A_v można określić na podstawie stosunku wartości między rezystorami R5 i R4 zgodnie ze wzorem:

$$A_v = (R_4 + R_5) / R_4$$

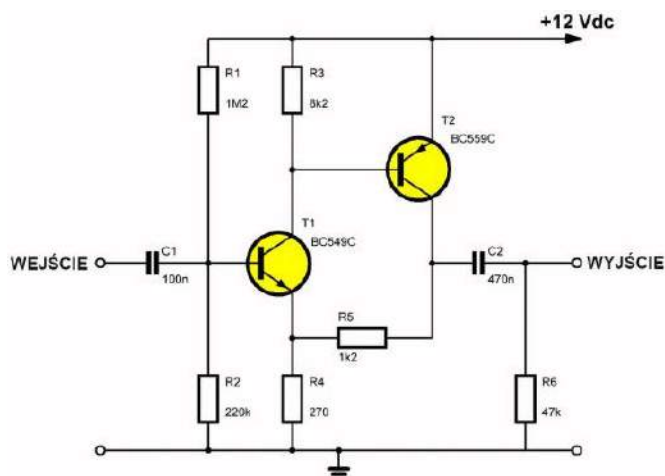
Następnie można ustawić napięcie na wyjściu obwodu na połowę napięcia zasilania, zmieniając stosunek wartości między rezystorami R1 i R2. Ze względu na silne ujemne



Dwutranzystorowy wzmacniacz z podwójnym sprzężeniem zwrotnym (© 2018 Jos Verstraten)



Alternatywny obwód z podwójnym sprzężeniem zwrotnym (© 2018 Jos Verstraten)



Przykład wzmacniacza dwutranzystorowego z tranzystorami PNP i NPN (© 2018 Jos Verstraten)

sprężenie zwrotne przez R5, obwód ma bardzo wysoką impedancję wejściową i bardzo niską impedancję wyjściową. Za pomocą obliczeń matematycznych ustalono, że impedancja wyjściowa pokazanego obwodu wynosi tylko około 10 Ω. Dołożenie wtórnika emiterowego na wyjściu jest zatem całkowicie niepotrzebne.

Wzmacniacz bardzo wysokiej jakości

Obwód z poprzedniego rysunku ma doskonałe właściwości. Takie obwody można stosować w wysokiej jakości sprzęcie audio, na przykład w celu zwiększenia napięcia wyjściowego tunera do standardowej wartości 0,775 V (0 dB). Kolejną zaletą tego układu jest brak różnicy faz między sygnałem wejściowym i wyjściowym.

Jedyną wadą jest to, że poprzez regulację wzmocnienia sygnału (R4, R5) zmienia się również polaryzacja stopnia, w wyniku czego trzeba ponownie majstrować przy R1 i R2. Dlatego też nie jest możliwa regulacja wzmocnienia w dużym zakresie za pomocą potencjometru regulacyjnego.

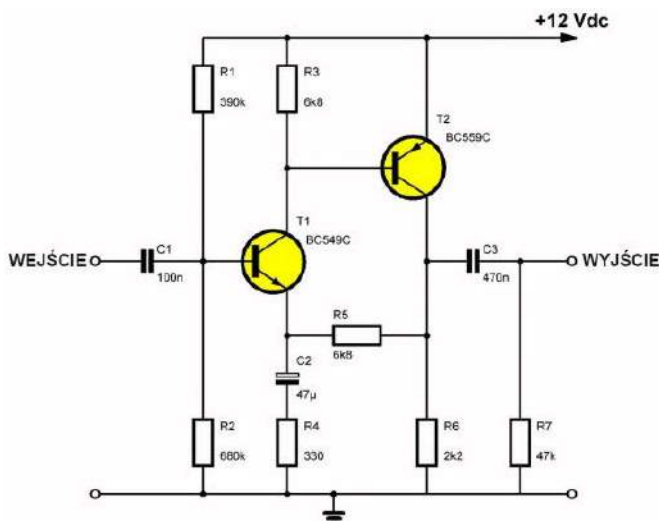
Regulowane wzmocnienie

Jest to możliwe w przypadku obwodu przedstawionego na rysunku

obok. Jedyną różnicą jest to, że ujemne sprzężenie zwrotne sygnału AC jest oddzielone od ujemnego sprzężenia zwrotnego polaryzacji poprzez wprowadzenie kondensatora C2. Oczywiście należy teraz dodać dodatkowy rezystor R6, w przeciwnym razie tranzystor T2 nie będzie mógł przewodzić prądu.

Ponownie, wzmocnienie sygnału jest określane zgodnie ze wzorem już podanym przez stosunek między R4 i R5. Można jednak teraz zmienić ten stosunek bez zmieniania napięć polaryzacji stopnia. Na przykład, projektując R4 w postaci potencjometru regulacyjnego, można dostosować wzmocnienie sygnału AC tego obwodu do własnych potrzeb. Ponieważ sygnał AC i polaryzacja są podawane z odwrotnością w inny sposób, należy wziąć pod uwagę nieco większe zniekształcenia niż w przypadku obwodu na poprzednim rysunku. ■

Jos Verstraten



Obwód ten umożliwia zmianę wzmocnienia sygnału bez wpływu na ustawienie (© 2018 Jos Verstraten)



Wzmacniacze dwutranzystorowe

Rozwiązanie znajdziesz na www.elportal.pl/quizy

1. Wadą wzmacniacza jednotranzystorowego jest:

- ☐ a. ograniczone wzmocnienie pojedynczego stopnia;
- ☐ b. ograniczone pasmo przenoszenia;
- ☐ c. odwracanie fazy sygnału.

2. Czy można bezpośrednio połączyć dwa stopnie wzmacniające tak, by wyjście pierwszego polaryzowało od razu bazę drugiego?

- ☐ a. tak;
- ☐ b. nie;
- ☐ c. tak, ale układ będzie niestabilny.

3. By układ z bezpośrednim połączeniem stopni działał, należy:

- ☐ a. dodać ujemne sprzężenie zwrotne;
- ☐ b. dodać bootstrap do każdej bazy;
- ☐ c. dodać wtórniki emiterowe między stopniami.

4. W prezentowanym dwutranzystorowym układzie z ujemnym sprzężeniem zwrotnym wzmocnienie zależy od:

- ☐ a. wartości rezystorów emiterowych;
- ☐ b. stosunku wartości rezystorów R5 i R3 na schemacie;
- ☐ c. wartości rezystorów emiterowych i stosunku R5 i R3.

5. Układ z podwójnym sprzężeniem zwrotnym oferuje:

- ☐ a. większe wzmocnienie i większą impedancję wejściową;
- ☐ b. większą stabilność;
- ☐ c. większą stabilność i wzmocnienie;

6. W układzie z artykułu R6 ustala:

- ☐ a. impedancję wejściową;
- ☐ b. impedancję wyjściową;
- ☐ c. impedancję wejściową i wyjściową.

7. W układzie alternatywnym z podwójnym sprzężeniem wzmocnienie zależy głównie od:

- ☐ a. wartości rezystorów emiterowych;
- ☐ b. stosunku wartości rezystora emiterowego T1 i rezystora sprzęgającego emiter T1 i kolektor T2;
- ☐ c. rezystora łączącego emiter T2 i bazę T1.

8. Używając w układzie dwutranzystorowym jednego tranzystora PNP zamiast NPN można:

- ☐ a. uprościć układ bez pogorszenia parametrów;
- ☐ b. uzyskać głębsze sprzężenie zwrotne;
- ☐ c. zaoszczędzić pieniądze na rezystorach.

9. Pokazany w artykule układ z tranzystorem NPN i PNP ma:

- ☐ a. niskie wzmocnienie, ale większą stabilność;
- ☐ b. wysokie wzmocnienie i większą stabilność;
- ☐ c. niskie wzmocnienie, ale zniekształcenia poniżej 0,01%.

10. Ostatni z prezentowanych w artykule układów pozwala na:

- ☐ a. uzyskanie znacznych wzmocnień rzędu tysięcy razy;
- ☐ b. uzyskanie jeszcze mniejszych zniekształceń;
- ☐ c. regulację wzmocnienia jednym rezystorem.

REKLAMA

<http://www.ep.com.pl/EPwtoku>

czytaj artykuły „Elektroniki Praktycznej” zanim zostaną wydane w formie papierowej

Piszemy prosty program dla ARM Cortex M4,

czyli jak skomplikować prostą rzecz w przykładach

Każdy programista mikrokontrolerów zaczynał swoją przygodę od prostych programów, na przykład błyskających diodą LED. Nieważne, czy ktoś zaczynał od PIC16F628 i Assemblera, czy od ATMegi i Bascoma, czy w końcu od Arduino (jak większość ludzi w ostatniej dekadzie), pierwszym programem jest prosta pętla zapalająca i gasząca diodę LED. Takie „Hello world!” świata mikrokontrolerów. Sam zaczynałem od programowania PIC16F628A i PIC18F45K50 za pomocą XC8, mikroPascala czy PICBasica, i zawsze pierwszym programem było proste „machanie nóżką” w nieskończonej pętli.

Ostatnio postanowiłem wyjść poza swoją strefę komfortu i spróbować swych sił w programowaniu dużo potężniejszych układów – padło na płytkę WeAct Studio BlackPill z układem STM32F401CCU6. Jest to przedstawiciel rodziny ARM Cortex M4, oferujący duży przyrost wydajności względem ośmiobitowych braci mniejszych, funkcje DSP znane z szesnastobitowych braci mniejszych i moduł FPU do obliczeń zmiennoprzecinkowych. Wszystko to na płytce za około 30 złotych. Ba, domyślnie jest wgrany program przyciemniający i rozjaśniający diodę LED, który robi to prawie płynnie. Zróbmy to jednak po swojemu.

Witaj, okrutny świecie!

Na początek przydałoby się jakieś IDE. Trzy główne opcje to STM32Cube IDE, Arduino IDE oraz PlatformIO. Arduino IDE nie znoszę – przez lata cierpiało na chroniczne braki użytecznych funkcji, a i wygląd też mi się nie widzi. STM32Cube IDE wydaje się dobrą opcją, ale taka „kobyła” do jednego układu mi nie jest potrzebna. Padło więc na PlatformIO, które to środowisko jest rozszerzeniem do Microsoft VisualStudio Code. Nie będę przedstawiał procesu instalacji, ani konfiguracji pierwszego projektu z BlackPill. Jako framework wybrałem Arduino, bo jestem leniwy, jak każdy programista. Jedna uwaga,

Listing 1. Zawartość pliku platformio.ini

```
[env:blackpill_f401cc]
platform = ststm32
board = blackpill_f401cc
upload_protocol = dfu
framework = arduino
```

Listing 2. „Hello World!” w wersji podstawowej

```
#include <Arduino.h>

void setup() {
    // put your setup code here, to run once:
    pinMode(PC13, OUTPUT);
}

void loop() {
    // put your main code here, to run repeatedly:

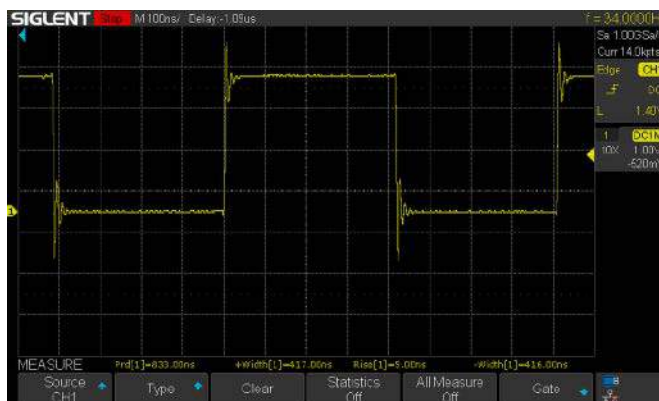
    digitalWrite(PC13, HIGH);
    delay(500);
    digitalWrite(PC13, LOW);
    delay(500);
}
```

Listing 3. „Hello World!” po kilku iteracjach

```
#include <Arduino.h>

void setup() {
    // put your setup code here, to run once:
    pinMode(PC13, OUTPUT);
}

void loop() {
    unsigned char k = 0;
    // put your main code here, to run repeatedly:
    for(unsigned int a = 0; a < 1000; a++){
        digitalWrite(PC13, HIGH);
        digitalWrite(PC13, LOW);
    }
    delay(1000);
    for(unsigned char i = 0; i < 10; i++){
        digitalWrite(PC13, HIGH);
        delay(500);
        digitalWrite(PC13, LOW);
        delay(500);
    }
    unsigned char l = 0;
    for(unsigned int j = 0; j <= 1000; j++){
        for(unsigned char k = 0; k < 100; k++){
            if(l <= k) digitalWrite(PC13, HIGH);
            if(l > k) digitalWrite(PC13, LOW);
            delayMicroseconds(250);
        }
        delayMicroseconds(750);
        l++;
        if(l == 101) l = 0;
    }
    delay(1000);
}
```



Rysunek 1. Test prędkości GPIO z użyciem standardowych funkcji dostępu

do pliku PlatformIO.ini trzeba dopisać jedną linię zmieniającą tryb programowania na DFU. **Listing 1** pokazuje, jak to powinno wyglądać.

Dzięki temu wystarczy mieć kabel USB typu C by zaprogramować układ. Domyślnie potrzebny byłby programator ST-LINK V2, który przy okazji udostępnia nam opcje debugowania, ale moim zdaniem nie są one nam potrzebne. Ale w praktyce nie jest jednak tak łatwo. Z jakiegoś powodu tryb DFU nie zawsze startuje poprawnie, system widzi wtedy urządzenie USB bez deskryptorów opisujących, czym ono jest. Nawet jak połączenie zostanie nawiązane poprawnie, to urządzenie jest widziane jako STM32_BOOTLOADER, ale VisualStudio Code nie rozpoznaje go. Potrzebny jest program instalujący sterowniki urządzeń USB, np. Zadig, by nasza „czarna pigułka” dawała się zaprogramować. No, ale tak działa USB, zwłaszcza pod Windowsem – miało być „od ręki” i bez instalowania sterowników, ale to działa tylko dla urządzeń z prostymi deskryptorami. Wszystko, czego nie objęła standardowa specyfikacja USB wymaga sterowników, i nie każde sterowniki system może automatycznie pobrać.

Podstawowy program „Hello, World!” przedstawia **listing 2**. Kod ten sprawia, że dioda LED podłączona do pinu oznaczonego jako PC13 zacznie migać z częstotliwością 1 Hz. No ale przedtem dioda prawie płynnie zmieniała jasność. A tu tylko sobie mruga. Czy da się lepiej? Tak! Na początek umieścimy ten kod w pętli, by nie działał w kółko. Potem dopiszemy programowe PWM. Przy okazji dodamy jeszcze jeden kawałek kodu na początek, by sprawdzić, jak szybko STM32F401CC6U pod kontrolą Arduino „macha nóżką”. Kompletny kod po kilku iteracjach (dawno nie pisałem w C++) i trochę pozapominałem) przedstawia **listing 3**. Warto wspomnieć, iż cały ten kod zabiera 1152 bajty pamięci RAM i 11508 bajtów pamięci Flash. Stanowi to odpowiednio 1,8% dostępnego RAMu i 4,4% dostępnej pamięci programu. Większość tego „zjada” sam framework potrzebny do wykonania aplikacji użytkownika. Dla STM32 jest trochę większy od wersji dla ośmiobitowców, bo bardziej złożony układ wymaga więcej do uruchomienia.

Na **rysunku 1** możemy zaobserwować wynik naszego małego eksperymentu z prędkością. Częstotliwość wynosi 1,2 MHz. Czas narastania zboczy wynosi 5 ns, czasy stanu wysokiego i niskiego to odpowiednio 417 ns i 416 ns. Prawie idealne 50%. Dzwonienie wynika z podłączenia do płytki stykowej, która to dodaje swoją impedancję pasożytniczą. Jak to się ma do tradycyjnego Arduino z układem ATmega328? Nie mam takiej płytki do testów, ale wg dokumentacji funkcja digitalWrite() pozwala przełączyć pin w ciągu około czterech mikrosekund. Czyli maksymalna częstotliwość wyniesie ledwo 125 kHz. W praktyce jednak jest to bliżej 100 kHz, co oznacza że Arduino na STM32 jest ponad 12 razy szybsze. Ale wartość ta dotyczy tylko dostępu do GPIO i innych peryferiów. W teście użyłem funkcji digitalWrite(), które nie są najszybsze, a do tego nie wiadomo, czy zegar peryferiów jest odpowiednio ustawiony. STMicroelectronics stosuje niezależny od taktowania CPU zegar dla peryferiów. To z kolei oznacza bardziej deterministyczne zachowanie tychże. Ba, dzięki DMA ciężar kontroli nad peryferiami może być zdjęty z CPU – program przetwarza informacje zapisane w pamięci RAM, a peryferia same

Listing 4. Dodajemy obsługę przycisku

```
#include <Arduino.h>

// put function declarations here:
void buttonPressed(void);

struct ef {
    unsigned : 30;
    unsigned bInt : 1;
    unsigned buttonToggle : 1;
};
volatile ef exec_flags;

void setup() {
    // put your setup code here, to run once:
    pinMode(PC13, OUTPUT);
    pinMode(PA0, INPUT_PULLUP);
    attachInterrupt(digitalPinToInterrupt(PA0), buttonPressed, CHANGE);
}

void loop() {
    unsigned char k = 0;
    // put your main code here, to run repeatedly:
    for(unsigned int a = 0; a < 1000; a++){
        digitalWrite(PC13, HIGH);
        digitalWrite(PC13, LOW);
    }
    delay(1000);

    while(exec_flags.buttonToggle == 0){
        for(unsigned char i = 0; i < 10; i++){
            digitalWrite(PC13, HIGH);
            delay(500);
            digitalWrite(PC13, LOW);
            delay(500);
            if (exec_flags.bInt || exec_flags.buttonToggle) {
                exec_flags.buttonToggle = 1;
                exec_flags.bInt = 0;
                break;
            }
        }
        unsigned char l = 0;
        for(unsigned int j = 0; j <= 1000; j++){
            for(unsigned char k = 0; k < 100; k++){
                if(1 <= k) digitalWrite(PC13, HIGH);
                if(1 > k) digitalWrite(PC13, LOW);
                delayMicroseconds(250);
            }
            delayMicroseconds(750);
            l++;
            if(l == 101) l = 0;
            if (exec_flags.bInt || exec_flags.buttonToggle) {
                exec_flags.buttonToggle = 1;
                exec_flags.bInt = 0;
                break;
            }
        }
    }

    delay(1000);
    if (exec_flags.bInt && exec_flags.buttonToggle) {
        exec_flags.buttonToggle = 0;
        exec_flags.bInt = 0;
    }
}

// put function definitions here:
void buttonPressed(){
    exec_flags.bInt = 1;
}
```

zajmują się ich zapisaniem do pamięci lub pobraniem ich. Ale to zagadnienie jest zbyt zaawansowane dla nas, początkujących pisarzy kodu.

Wracając do naszego programu, to nadal nie działa on tak, jak powinien – nie możemy wyłączyć diody przyciskiem, jak w oryginalnym programie. Musimy zatem dodać tę funkcję. Zrobimy to poprawnie, czyli z użyciem przerwania. Ponieważ Arduino jest zbudowane wokół języka C++, kod musi być napisany inaczej, niż to robiłem w XC8 dla

mikrokontrolerów PIC. W XC8 w razie jakiegokolwiek przerwania program przechodzi do globalnej funkcji jego obsługi, w której trzeba sprawdzić flagę przerwania, a potem podjąć jakąś akcję, przy czym ten kod musi być krótki. W C++, a zatem i w przypadku Arduino, przy okazji konfigurowania przerwania wskazujemy funkcję, która zostanie wykonana w razie jego wystąpienia. Ułatwia to tworzenie rozbudowanych programów z wieloma przerwaniami. Dokumentacja wskazuje jednak, iż w czasie działania funkcji obsługującej przerwanie nie działa na przykład `delay()`. Dlaczego? Bo ta funkcja wykorzystuje w swoim działaniu przerwanie od jednego z timerów. Generalnie to nie problem, bo i tak nie powinniśmy tworzyć funkcji obsługi przerw, które są długie lub/i powolne. Nawet jeśli mikrokontroler ARM oferuje znaczny przyrost prędkości wykonywania kodu. Przy okazji przyda się też rozwiązanie stosowane przeze mnie w programach dla PICów, czyli „lotna” struktura, w którą spakowane są różne flagi. W przypadku ośmiobitowców oszczędza to pamięć RAM, bo jeden bajt może obsłużyć do ośmiu flag, ale dodatkowo instrukcje warunkowe używające dwóch lub więcej flag naraz są optymalizowane do porównania całej struktury z maską bitową. Oczywiście układ STM32F401CCU6 ma bardzo dużo pamięci RAM, więc nie musimy każdego bitu oszczędzać, ale po co nabierać złych nawyków? Aha, nasza struktura musi być „ulotna”, czyli `volatile`, bo będą do niej odwoływać się różne funkcje, co oznacza iż nie tylko musi być zmienną globalną, ale też zmienną, której wartość nie jest stale zdeterminowana i w każdej chwili może ulec zmianie. W przeciwnym razie program może nie działać poprawnie.

Popatrzmy zatem na **listing 4**. Ojoj, ale się porobiło! Na początku mamy deklarację funkcji, która obsługuje nasze przerwanie. To po to, by kompilator wiedział, że ta funkcja gdzieś jest zdefiniowana. Zamiast deklaracji możemy w tym miejscu umieścić definicję, ale z reguły definicje trafiają na koniec programu, a nie na sam początek. W ten sposób jednak zachowamy większy porządek w kodzie. Następnie definiujemy strukturę, a potem ją deklarujemy jako „ulotną”. Ponieważ deklaracja i definicja są oddzielone, możemy użyć tej samej struktury pod różnymi nazwami w różnych miejscach programu, a nawet zadeklarować całą tablicę złożoną z tych struktur. **Listing 5** pokazuje alternatywną formę deklaracji połączonej z definicją. Obie formy są poprawne. Sama struktura zawiera trzy zmienne, pierwsza, niezdefiniowanego typu `unsigned` zawiera 30 bitów i nie ma nazwy. Kolejne dwie zmienne też nie mają zdefiniowanego typu, choć mógłbym użyć typu `boolean`, gdyż mają rozmiar jednego bitu. Obie mają też nazwy: `bInt` i `buttonToggle`.

W sekcji `setup()` też zaszły zmiany. Po pierwsze, pin PA0, do którego przyłączony jest przycisk użytkownika, został skonfigurowany jako wejście z rezystorem podciągającym. W następnej linii deklarujemy przerwanie dla tego pinu w razie pojawienia się zmiany. Tutaj też deklarujemy, która funkcja ma być wykonana w razie wystąpienia przerwania. Przeskoczmy na moment do tej funkcji – jedyną, co ona robi, to zmienia wartość bitu `bInt` w naszej strukturze `exec_flags` na 1. Właściwa magia dzieje się dopiero w pętli głównej. Pierwsza sekcja, która przetłacza pin PC13, pozostaje bez zmian. Dwie kolejne sekcje, odpowiadające za właściwe błyskanie diodą zostały umieszczone w pętli `while()`. Warunkiem wykonania pętli jest wartość 0 flagi `buttonToggle`. We fragmencie zapalającym diodę LED na pół sekundy, oraz w pętli odpowiedzialnej za powolne rozjaśnianie diody dodane są instrukcje warunkowe z ciekawymi warunkami i jeszcze ciekawszymi operacjami. Otóż jeśli którakolwiek z naszych flag ma wartość 1, wartość `buttonToggle` jest zmieniana na 1, a `bInt` na 0. Ostatnie polecenie opuszcza pętlę, w której znajduje się nasza instrukcja warunkowa, by przejść do kolejnego fragmentu pętli `while()`. Dzięki temu, że ten fragment powtórzony jest też w funkcji emulującej efekt PWM, wciśnięcie przycisku w dowolnym momencie, niezależnie od wartości `buttonToggle`,

Listing 5. Alternatywna forma struktury

```
volatile struct {
    unsigned : 30;
    unsigned bInt : 1;
    unsigned buttonToggle : 1;
} exec_flags;
```

spowoduje opuszczenie pętli `while()`. Na końcu pętli głównej dodana jest inna funkcja warunkowa, która zresetuje obie flagi, ale tylko wtedy gdy obie flagi są ustawione.

Jak to działa w praktyce? Działa prawie że dobrze. Problemem jest sam przycisk – brak tu tzw. debouncingu, czyli eliminacji drgań styków. Przez to po pojawieniu się przerwania program może opuścić pętlę `while()`, ale pojawi się nowe przerwanie zanim program dojdzie do ostatniej instrukcji warunkowej w pętli głównej, co zresetuje flagę `buttonToggle`.

Zróbmy to z timerem

Ta wersja programu nie podoba mi się z wielu względów. Po pierwsze, używamy za dużo funkcji `delay()` i `delayMicroseconds()`. Przez większość czasu program zajęty jest czekaniem. Ponadto przydałaby się eliminacja drgań styków. Po trzecie, używanie kopii/wklej jest mało eleganckie – w końcu po to są funkcje, by kompilator robił to za nas. Po czwarte, fajnie by było, gdyby dioda płynnie się rozjaśniała, a potem przygasła, czyli zachowywała się lepiej, niż nawet w oryginalnym programie. Ktoś się zapyta: na co to komu? To tylko „Hello, World!”, po co to komplikować? Nie lepiej zrobić czegoś innego, bardziej użytecznego? Ja mam prostą odpowiedź: bo możemy to zrobić lepiej. Więc czemużby nie zrobić tego lepiej? Przy okazji czegoś się może nauczymy.

Zanim jednak zaczniemy pisać kod, trzeba się zastanowić, jak ma taki program działać. Chcemy by program animował PWM w sposób prawidłowy i płynny, na zmianę rozjaśniając i przyciemniając diodę. Ale nie możemy użyć sprzętowego PWM, bo jest podłączone do innego pinu układu, więc będziemy to realizować programowo. Po drugie, musimy obsłużyć przycisk tak, by wyeliminować drgania styków. Użycie przerwania nie działa za dobrze, bo potem mamy całą serię przerw z powodu drgań tych styków. Zależnie od przycisku drgania mogą trwać od 400 µs do nawet 2–3 ms, generując od trzech do ponad 10 fałszywych impulsów. Nie możemy też spędzać każdej wolnej chwili na odpytывanie stanu pinu. Nie dość, że program by nic innego nie robił, to jeszcze każde wciśnięcie przycisku mógłby odczytać jako wiele wciśnieć. Po trzecie, powinniśmy zostawić więcej czasu dla innych funkcji programu, które moglibyśmy dodać w przyszłości. Jak to zrobić?

Zacznijmy od PWM. Jeśli rzucimy okiem do listingu 3 lub 4, widać iż funkcja PWM składa się z trzech prostych zadań: inkrementacji jednej zmiennej od zera do stu, sprawdzania, czy zmienna osiągnęła wartość wypełnienia, i jeśli tego nie robiła, to pin ma stan wysoki, w przeciwnym razie ma stan niski. Na końcu jest dodana pauza, by ta dioda trochę świeciła i by pętla nie wykonała się za szybko. Gdyby tylko istniał jakiś sposób, by wszystkie operacje związane z PWM dałoby się wykonywać w równych odstępach czasu bez konieczności dodawania jakichkolwiek opóźnień, a w międzyczasie robić coś innego. Otóż jest na to sposób, i to niemal tak stary, jak komputerowe systemy kontrolne. Czyli sięgający lat 50. ubiegłego wieku. Tym sposobem jest użycie timera, który w określonych odstępach czasu wywołuje przerwanie, które to zrealizuje naszą funkcję PWM. Częstotliwość PWM będzie równa częstotliwości przerwania pomnożonej przez liczbę inkrementacji na jeden okres, czyli przez rozdzielczość. Jeśli ustawimy przerwanie, by było co 5 µs, i użyjemy 1% rozdzielczości, to częstotliwość PWM

wyniesie 2 kHz, bo pełny okres powinien trwać 500 μ s. Ponieważ po doliczeniu do stu musimy inkrementowaną wartość okresu wyzerować, możemy przy okazji sprawdzić stan naszego przycisku. Ba, możemy dodać kolejny, programowy licznik do tego, gdyż w praktyce czas wciśnięcia przycisku przez człowieka trwać może dziesiątki milisekund, niezależnie od tego, jak szybkim palcem się dysponuje. Sprawdzając stan przycisku co, powiedzmy 25 ms nie tylko pozwoli nam zapomnieć o drganiach, ale dodając jeszcze jeden licznik możemy odróżnić krótkie wciśnięcie od długiego wciśnięcia, a nawet rozpoznać dwuklik. Zatem jak możemy użyć taki timer w naszym ARMie?

Praca z mikrokontrolerem ARM różni się od pracy ze skromnym ośmiobitowcem w sposobie dostępu do sprzętu. Ośmiobitowe PICi przyzwyczyły mnie do bezpośredniego zwracania się do rejestrów, a nawet poszczególnych bitów w tych rejestrach, ale w świecie 16-bitowym i 32-bitowym (zarówno ARM, jak i MIPS czy x86) jest to praktyka raczej niespotykana. Producenci mikroprocesorów i mikrokontrolerów od lat oferują narzędzia do automatycznej konfiguracji peryferiów i HAL (Hardware Abstraction Layer), czyli sposób by te same funkcje dostępu do sprzętu działały na każdym układzie w obrębie konkretnej rodziny, albo i całej klasy. W świecie PICów o skromnych ośmiu bitach rozwiązano to w ten sposób, że wszystkie rejestry i bity w tych rejestrach mają te same nazwy. W przypadku Arduino abstrakcja sprzętu jest realizowana w kodzie – po wybraniu układu/platformy z listy odpowiednie pliki są dodawane do programu po to, by kompilator mógł zamienić wysokopoziomą funkcję, jak konfiguracja portu szeregowego czy timera, na zestaw instrukcji odwołujących się do peryferiów przez rejestry. To samo dzieje się, gdy używamy środowiska STM32CubeIDE czy MPLAB-X IDE. W ostatniej dekadzie Microchip dodał automatyczną konfigurację i generowanie funkcji dla swoich ośmiobitowców PIC i Atmel. Przejdźmy jednak do Arduino. Tu mamy bowiem ciekawą sytuację – na pierwszy rzut oka lista wbudowanych funkcji zawiera obsługę I²C, SPI, UART, czy programową obsługę 1-Wire. Ale nie widać nic do obsługi timerów. W przypadku ośmiobitowych płytek Arduino mamy sytuację podobną, do normalnego IDE – przykłady pokazują, jak sterować bezpośrednio rejestrami, by ustawić timer. W przypadku ARM STM32 jest inaczej – potrzebujemy dedykowanej biblioteki, która jest częścią STM32Duino. Biblioteka załatwia za nas większość konfiguracji i dostarcza wszystko, co jest potrzebne by użyć timera zarówno do odmierzania równych odstępów czasu, jak i do generowania sygnałów PWM na potrzeby m.in. przetwornic i silników elektrycznych, a także mierzyć częstotliwość czy odliczać impulsy. My jednak użyjemy timera tylko do odmierzania równych odcinków czasu. Będziemy używać biblioteki `HardwareTimer`, która jest częścią `STM32Duino`.

Listing 6 prezentuje zawartość funkcji `setup()`. To tutaj następuje konfiguracja timera `TIM2`. Odchodzimy tu od C, do którego jestem przyzwyczajony, i zaczynamy wykorzystywać w pełni składnię i możliwości C++. Pierwsze dwie linijki to już znana nam konfiguracja pinów. Od razu też ustawmy flagę od przycisku, by właściwy program miał bardziej logiczną budowę. W następnej jednak tworzymy instancję obiektu timera `myTim` klasy `HardwareTimer`, dodatkowo użycie słowa kluczowego `new` tworzy instancję, która będzie istnieć nawet po opuszczeniu `setup()`. Przy okazji przekazujemy też parametr, który timer sprzętowy ma być użyty (`TIM2`). W następnej linijce używamy funkcji będącej częścią klasy `HardwareTimer` aby skonfigurować `TIM2`, by przepełniał się co 5 mikrosekund. Kolejna linijka dołącza przerwanie do naszego timera, i gdy ten będzie się przepełniał co 5 mikrosekund, wzywana będzie funkcja `T2_task()`. Funkcja ta została zadeklarowana przed funkcją `setup()`. Ostatnia linijka wznawia timer `TIM2` przypisany do obiektu `myTim`. To moje pierwsze poważne spotkanie

Listing 6. Funkcja `setup()` w wariacie z timerem

```
void setup() {
    // put your setup code here, to run once:
    pinMode(PC13, OUTPUT);
    pinMode(PA0, INPUT_PULLUP);
    softPWM.bToggle = 1;
    HardwareTimer *myTim = new HardwareTimer(TIM2);
    myTim->setOverflow(5, MICROSEC_FORMAT);
    myTim->attachInterrupt(T2_task);
    myTim->resume();
}
```

Listing 7. Nowa struktura ze wszystkimi niezbędnymi zmiennymi do wariantu z timerem

```
volatile struct {
    unsigned : 5;
    unsigned bCount : 4;
    unsigned bPressed : 1;
    unsigned bIntCount : 6;
    unsigned sPWMset : 7;
    unsigned sPWMdir : 1;
    unsigned sPWMper : 7;
    unsigned bToggle : 1;
} softPWM;
```

z programowaniem obiektywnym od dwóch dekad, więc mam nadzieję, iż to wyjaśnienie jest zrozumiałe.

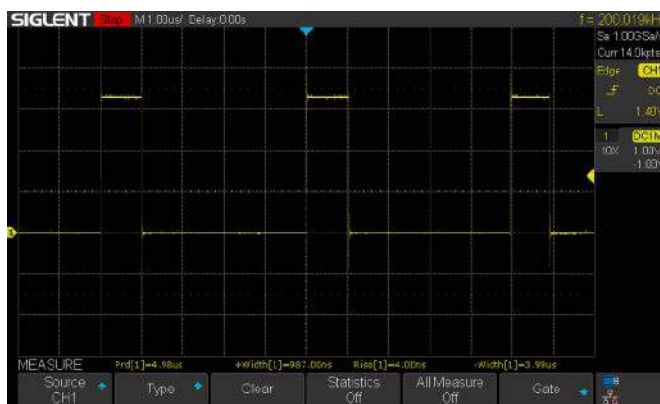
Reszta obsługi PWM znajdzie się w funkcji `T2_task()`, jednak wcześniej wypadałoby mieć kilka zmiennych globalnych. Ponownie użyjemy struktury, by wszystko ładnie spakować. Potrzebujemy

Listing 8. Funkcja `T2_task()` wykonywana w przerwaniu timera

```
void T2_task(){
    softPWM.sPWMper++;
    if (softPWM.sPWMper == 100){
        softPWM.sPWMper = 0;
        softPWM.bIntCount++;
    }
    if ((softPWM.bToggle == 1) && (softPWM.sPWMper <=
    softPWM.sPWMset)) digitalWrite(PC13, 0);
    else digitalWrite(PC13, 1);

    if (softPWM.bIntCount == 49){
        softPWM.bIntCount = 0;
        if (digitalRead(PA0) == 0) softPWM.bCount++;
        if (softPWM.sPWMdir == 0){
            softPWM.sPWMset++;
            if (softPWM.sPWMset == 100) softPWM.sPWMdir = 1;
        }
        else {
            softPWM.sPWMset--;
            if (softPWM.sPWMset == 0) softPWM.sPWMdir = 0;
        }
    }

    if (softPWM.bCount > 10){
        softPWM.bCount = 0;
        softPWM.bPressed = 1;
    }
    if (softPWM.bPressed == 1){
        if (softPWM.bToggle == 1){
            softPWM.bToggle = 0;
            digitalWrite(PC13, 1);
        }
        else {
            softPWM.bToggle = 1;
            softPWM.sPWMper = 0;
            softPWM.sPWMset = 0;
            softPWM.sPWMdir = 0;
        }
        softPWM.bPressed = 0;
    }
}
```



Rysunek 2. Czas trwania przerwania w programie z listingów 5–8

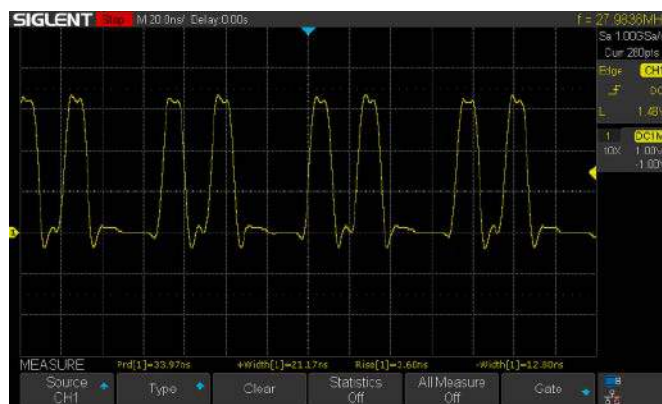
mieć zmienną odliczającą do stu, żeby jeden okres trwał 500 μ s, zmienną przechowującą zadane wypełnienie, kolejną zmienną zliczającą do 50 dla programowej eliminacji drgań styków, i zmiany wypełnienia. Do tego jeszcze przyda się flaga by zmieniać kierunek zliczania wypełnienia, i druga by włączać i wyłączać naszą diodę po wciśnięciu przycisku. Nasza oszczędna struktura przedstawiona jest na **listingu 7**. Zostało nam jeszcze 5 bitów do wykorzystania w przyszłości.

Listing 8 pokazuje kod funkcji T2_task(). Na początek zwiększamy licznik okresu, a gdy ten osiągnie wartość 100, zerujemy go i zwiększamy licznik dla przycisku. Następnie sprawdzamy, czy PWM jest włączone, i czy wartość okresu jest mniejsza lub równa wartości ustawionego wypełnienia. Gdy oba te warunki są spełnione, dioda jest zapalona. W przeciwnym razie jest zgaszona. Następna sekcja sprawdza stan licznika dla przycisku, gdy licznik doliczy do 49, jest resetowany, a potem sprawdzamy, czy przycisk został wciśnięty, i w razie czego inkrementujemy licznik wciśnięcia. W tym miejscu testowany jest kierunek naliczania wypełnienia, jak bit softPWM.sPWMdir ma wartość 0, to następuje inkrementacja do stu, gdzie kierunek zostaje odwrócony, aż wartość osiągnie znów wartość 0. Ostatnia część kodu sprawdza wartość licznika dla przycisku, gdy osiągnie zadaną wartość, jest zerowana, i ustawiana jest flaga bPressed. Dalej sprawdzamy obecność tej flagi i stan flagi bToggle. Jak bToggle ma wartość 1, to zerujemy ją i wygaszamy diodę. Jak ma wartość zero, to zmieniamy ją na 1 i zerujemy wybrane zmienne by PWM zaczęło od zera. Na koniec zerowana jest flaga bPressed.

I to tyle. Program jest kompletny. Sekcja loop() jest kompletnie pusta. Wszystko dzieje się w przerwaniu. A to oznacza, że możemy w pętli głównej oddać się innym czynnościom programowym. No ale pisałem wcześniej, że przerwania powinny być małe i lekkie, a to na takie nie wygląda. W rzeczywistości tylko część kodu odpowiedzialna za emulację PWM jest wykonywana za każdym razem. Pozostałe mikrokontroler zwyczajnie przeskakuje, gdy warunki nie są spełnione. No ale co, jeśli chcielibyśmy wykonać bardziej rozbudowany program w cyklicznych odstępach czasu, i to by „zjadło” nam za dużo czasu w przerwaniu? Czy możemy to jeszcze bardziej poprawić? Ależ oczywiście, że tak.

Komplikacje i debugowanie

Najprostszą metodą redukcji czasu trwania przerwania byłoby przeniesienie całego kodu do pętli głównej i wykonanie go, gdy przerwanie daje znak. Będzie to działać, ale jeśli program wykonuje inne zadania, to nasz kod PWM może wykonywać się rzadziej, niż byśmy tego chcieli. Nie jest to zatem idealne rozwiązanie. Z drugiej strony nie wiemy nawet, jak długo trwa wykonywanie naszego kodu – nigdy nie zrobiliśmy



Rysunek 3. Test prędkości GPIO z użyciem funkcji digitalWriteFast(). Skok do początku pętli powoduje dłuższy stan niski po drugim impulsie

żadnego testu. A zrobić test jest dość łatwo – wystarczy dodać zmianę stanu innego pinu, na przykład PC14 ze stanu niskiego na wysoki, gdy zaczyna się przerwanie, i ze stanu wysokiego na niski, gdy się kończy.

Rysunek 2 przedstawia wyniki tego testu. Jak widać, przerwanie trwa niecałą mikrosekundę. Użyłem szybszej funkcji dostępu do GPIO, by wynik był bardziej miarodajny. **Rysunek 3** pokazuje test podobny do tego z rysunku 1, ale z użyciem tych szybszych funkcji dostępu, ale zdublowanych, dzięki czemu widać zarówno czas trwania każdego stanu, jak i czas trwania stanu niskiego, gdy program wykonuje skok do początku pętli.

Przy okazji tych eksperymentów możemy przenieść część kodu do pętli głównej, co pokazuje **listing 9**. Jeśli jednak umieścimy w pętli głównej jakiś dodatkowy, rozbudowany kod, którego wykonanie zajmie za dużo czasu, to nie dość, że zaburzymy generowanie przebiegu PWM, to jeszcze w ogóle możemy go nie wytworzyć z powodu jednego błędu. Błąd polega na tym, iż sprawdzamy, czy zmienne softPWM.bIntCount oraz softPWM.sPWMset równe stosownym wartościom. W wersji z listingu 9. nie ma to znaczenia, i nic złego się nie stanie. Ale jeśli w pętli głównej będzie dodatkowy kod, i w trakcie jego wykonywania pojawi się przerwanie timera, to wartości tych zmiennych mogą z łatwością być większe niż wskazane wartości, gdy program wróci na początek pętli głównej. By tego uniknąć wystarczy zamienić „=” na „>=” w stosownych warunkach. Kod jednak przestanie być czasowo dokładny. Dlatego w tym przypadku lepiej zostawić kod w przerwaniu, nawet jak nie jest ono tak szybkie, jak byśmy chcieli.

Warto na koniec poruszyć kwestię debugowania naszego kodu. Te proste przykłady generalnie tego nie wymagały, ale przy okazji pisania innego, bardziej skomplikowanego programu (maszyny stanów skończonych) pojawiła się potrzeba podejrzenia wartości kilku zmiennych. Najprościej wysłać więc te wartości przez UART do komputera, gdzie aplikacja w rodzaju CoolTerm albo RealTerm pozwoli nam na ich odczytanie w jednym z kilku formatów. Standardowo STM32Duino do realizacji komunikacji szeregowo używa sprzętowego portu UART zawartego w każdym kompatybilnym mikrokontrolerze. Jest to o tyle niewygodne, że potrzebny jest port szeregowy w komputerze. I to nie zwykły port RS-232, jakiego używało się jeszcze ponad dwie dekady temu, lecz port kompatybilny z napięciem 3,3 V, jakie jest używane we współczesnych ARMach, jak na przykład płytka BlackPill tutaj używana. Skoro zatem programujemy płytkę przez port USB, to może by ten port wykorzystać jako wirtualny port szeregowy? Nie stanowi to żadnego problemu, musimy jedynie wprowadzić pewne zmiany w pliku platformio.ini, co pokazuje **listing 10**. Ma to pewien wpływ na wielkość wynikowego programu: kod, gdzie wszystko dzieje się

w przerwaniu zajmuje 1156 bajtów pamięci RAM i 13784 bajty pamięci Flash, co stanowi odpowiednio 1,8% RAM i 5,3% Flash. W wersji z obsługą USB-CDC kod zajmuje odpowiednio 4976 bajtów RAM (7,6%) i 26364 bajty pamięci Flash (10,1%). Ten dodatkowy kod odpowiada za to, by mikrokontroler „przedstawił się” jako urządzenie USB-CDC (Communication Device Class – klasa urządzenia komunikacyjnego). Jest to standardowa metoda tworzenia urządzeń szeregowych, interfejsów sieciowych i tym podobnych. Na listingu 10

widać też, jak urządzenie jest opisane za pomocą zmiennych VID i PID oraz nazwy. Numery VID (Vendor ID) są sprzedawane przez USB Implementation Forum. Za pięć tysięcy dolarów rocznie można zostać członkiem USB IF, ale tańszą opcją jest nabycie licencji na logo USB i własny numer VID, co kosztuje 3500 dolarów co dwa lata. My używamy VID należącego do STMicroelectronics w celach edukacyjnych. Zakup własnego VID jest wymagany, gdy chcemy sprzedawać nasz produkt i mieć logo USB Certified.

Listing 9. Kompletny program, ale część kodu z przerwania została umieszczona w pętli głównej

```
#include <Arduino.h>

// put function declarations here:
void T2_task(void);

volatile struct {
    unsigned : 5;
    unsigned bCount : 4;
    unsigned bPressed : 1;
    unsigned bIntCount : 6;
    unsigned sPWMset : 7;
    unsigned sPWMdir : 1;
    unsigned sPWMper : 7;
    unsigned bToggle : 1;
} softPWM;

void setup() {
    // put your setup code here, to run once:
    pinMode(PC13, OUTPUT);
    pinMode(PA0, INPUT_PULLUP);
    softPWM.bToggle = 1;
    HardwareTimer *myTim = new HardwareTimer(TIM2);
    myTim->setOverflow(5, MICROSEC_FORMAT);
    myTim->attachInterrupt(T2_task);
    myTim->resume();
}

void loop() {
    if (softPWM.bIntCount == 49){
        softPWM.bIntCount = 0;
        if (digitalRead(PA0) == 0) softPWM.bCount++;
        if (softPWM.sPWMdir == 0){
            softPWM.sPWMset++;
            if (softPWM.sPWMset == 100) softPWM.sPWMdir = 1;
        }
        else {
            softPWM.sPWMset--;
            if (softPWM.sPWMset == 0) softPWM.sPWMdir = 0;
        }
    }

    if (softPWM.bCount > 10){
        softPWM.bCount = 0;
        softPWM.bPressed = 1;
    }
    if (softPWM.bPressed == 1){
        if (softPWM.bToggle == 1){
            softPWM.bToggle = 0;
            digitalWrite(PC13, 1);
        }
        else {
            softPWM.bToggle = 1;
            softPWM.sPWMper = 0;
            softPWM.sPWMset = 0;
            softPWM.sPWMdir = 0;
        }
        softPWM.bPressed = 0;
    }
}

void T2_task(){
    softPWM.sPWMper++;
    if (softPWM.sPWMper == 100){
        softPWM.sPWMper = 0;
        softPWM.bIntCount++;
    }
    if ((softPWM.bToggle == 1) && (softPWM.sPWMper <= softPWM.sPWMset)) digitalWrite(PC13, 0);
    else digitalWrite(PC13, 1);
}
```

Listing 10. Plik platformio.ini z dodaną obsługą USB-CDC oraz ustawieniami dla terminalu wbudowanego w IDE

```
[env:blackpill_f401cc]
platform = ststm32
board = blackpill_f401cc
upload_protocol = dfu
framework = arduino
build_flags =
    -D PIO_FRAMEWORK_ARDUINO_ENABLE_CDC
    -D USBCON
    -D USB_D_VID=0x0483
    -D USB_D_PID=0x0100
    -D USB_MANUFACTURER="Unknown"
    -D USB_PRODUCT="\BlackPill_F401CC\"

monitor_port = COM[3]
monitor_speed = 115200
```

Listing 11. Podstawowa implementacja funkcji debugera za pomocą dyrektyw preprocesora

```
#define DEBUG 0

#if DEBUG == 1
    #define debugStart(x) Serial.begin(x)
    #define debug(x, ...) Serial.print(x, ##__VA_ARGS__)
    #define debugln(x, ...) Serial.println(x, ##__VA_ARGS__)
#else
    #define debugStart(x)
    #define debug(x, ...)
    #define debugln(x, ...)
#endif
```

Listing 12. Ostateczna wersja programu z obsługą diod WS2812B

```
#include <Arduino.h>

#define DEBUG 0

#if DEBUG == 1
#define debugStart(x) Serial.begin(x)
#define debug(x, ...) Serial.print(x, ##__VA_ARGS__)
#define debugln(x, ...) Serial.println(x, ##__VA_ARGS__)
#else
#define debugStart(x)
#define debug(x, ...)
#define debugln(x, ...)
#endif

#define WSTh 4
#define WSTl 14
#define LEDMax 24

// put function declarations here:
void T2_task(void);
void WSWrite(unsigned int L);
void WSClear(void);

volatile struct {
    unsigned : 5;
    unsigned bCount : 4;
    unsigned bPressed : 1;
    unsigned bIntCount : 6;
    unsigned sPWMset : 7;
    unsigned sPWMDir : 1;
    unsigned sPWMper : 7;
    unsigned bToggle : 1;
} softPWM;

static unsigned int randVal;
static union
{
    unsigned int colors;
    struct {
        unsigned char G : 8;
        unsigned char R : 8;
        unsigned char B : 8;
        unsigned : 8;
    } color;
} /* data */;

static unsigned char LEDCount;
static unsigned char LEDNum;

HardwareTimer *myTim = new HardwareTimer(TIM2);

void setup() {
    // put your setup code here, to run once:
    pinMode(PC13, OUTPUT);
    pinMode(PC14, OUTPUT);
    pinMode(PC15, OUTPUT);
    pinMode(PA0, INPUT_PULLUP);
    debugStart(115200);
    digitalWrite(PC13, 0);
    delay(500);
    if (digitalRead(PA0) == 0){
        unsigned char g = 0;
        unsigned char f = 0;
        while(1){
            digitalWriteFast(PC_14, 1);
            digitalWriteFast(PC_14, 0);
            digitalWriteFast(PC_14, 1);
            digitalWriteFast(PC_14, 0);
        }
    }
    digitalWrite(PC13, 1);
    for(unsigned int k = 0; k < LEDMax; k++){
        WSClear();
    }
    delay(250);
    softPWM.bToggle = 1;
    //HardwareTimer *myTim = new HardwareTimer(TIM2);
    myTim->setOverflow(5, MICROSEC_FORMAT);
    myTim->attachInterrupt(T2_task);
    myTim->resume();
}

void loop(){
    randVal = random();
    debug("RandVal, dec: ");
    debug(randVal, DEC);
    debug(" , hes: 0x");
    debugln(randVal, HEX);
    //colors = 0x00AA55AA;
    //colors = 0xFF3F0F;
    //colors++;
    colors = randVal >> 8;
    debug("colors, dec: ");
    debug(colors, DEC);
    debug(" , hex: 0x");
    debug(colors, HEX);
    colors = colors & 0x00FCFCFC;
    debug(" , limit: 0x");
    debugln(colors, HEX);
    debug(" , B: 0x");
    debug(color.B, HEX);
    debug(" , R: 0x");
    debug(color.R, HEX);
    debug(" , G: 0x");
    debugln(color.G, HEX);
    if(LEDCount < LEDNum){
        myTim->pause();
        WSWrite(colors);
        myTim->resume();
    }
    else{
        myTim->pause();
        WSClear();
        myTim->resume();
    }
    LEDCount++;
    if(LEDCount > LEDMax){
        LEDCount = 0;
        LEDNum++;
        if(LEDNum > LEDMax) LEDNum = 0;
        delay(50);
    }
}

void T2_task(){
    digitalWriteFast(PC_14, HIGH);
    softPWM.sPWMper++;
    if (softPWM.sPWMper == 100){
        softPWM.sPWMper = 0;
        softPWM.bIntCount++;
    }
    if ((softPWM.bToggle == 1) && (softPWM.sPWMper <=
    softPWM.sPWMset)) digitalWrite(PC13, 0);
    else digitalWrite(PC13, 1);

    if (softPWM.bIntCount == 49){
        softPWM.bIntCount = 0;
        if (digitalRead(PA0) == 0) softPWM.bCount++;
        if (softPWM.sPWMDir == 0){
            softPWM.sPWMset++;
            if (softPWM.sPWMset == 100) softPWM.sPWMDir = 1;
        }
        else {
            softPWM.sPWMset--;
            if (softPWM.sPWMset == 0) softPWM.sPWMDir = 0;
        }
    }

    if (softPWM.bCount > 10){
        softPWM.bCount = 0;
        softPWM.bPressed = 1;
    }
    if (softPWM.bPressed == 1){
        if (softPWM.bToggle == 1){
            softPWM.bToggle = 0;
            digitalWrite(PC13, 1);
        }
        else {
            softPWM.bToggle = 1;
            softPWM.sPWMper = 0;
            softPWM.sPWMset = 0;
            softPWM.sPWMDir = 0;
        }
    }
}
```

Listing 12. Ostateczna wersja programu z obsługą diod WS2812B

```

    }
    softPWM.bPressed = 0;
}
digitalWriteFast(PC_14, LOW);
}

void WSWrite(unsigned int L){
//WS2812B timing:
//0 = Ton < 560ns, Ttot 1250ns
//1 = Ton >560ns, Ttot 1250ns

    unsigned char i= 0;
    for(unsigned char k = 0; k < 24; k++){

        if(L & 0b1){
            while(i< WSTl) {
                digitalWriteFast(PC_15, 1);
                i++;
            }
            i = 0;
            while(i< WSTh) {
                digitalWriteFast(PC_15, 0);
                i++;
            }
            i = 0;
        }
        else {
            while(i< WSTh) {
                digitalWriteFast(PC_15, 1);
                i++;
            }
        }
    }
}

    i = 0;
    while(i< WSTl) {
        digitalWriteFast(PC_15, 0);
        i++;
    }
    i = 0;
    L = L >> 1;
}

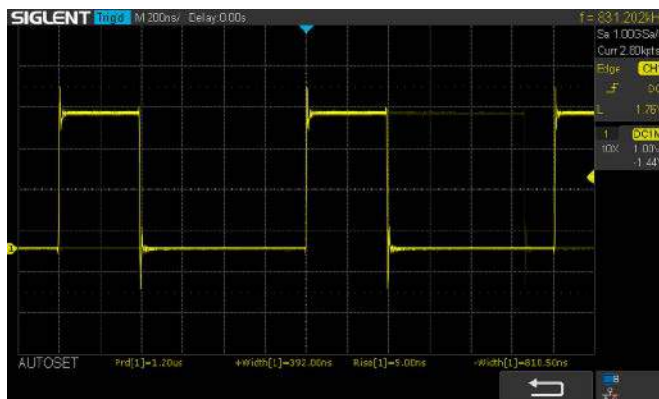
void WSClear(){
//WS2812B timing:
//0 = Ton < 560ns, Ttot 1250ns
//1 = Ton >560ns, Ttot 1250ns

    unsigned char i= 0;
    for(unsigned char k = 0; k < 24; k++){
        while(i< WSTh){
            digitalWriteFast(PC_15, 1);
            i++;
        }
        i = 0;
        while(i< WSTl) {
            digitalWriteFast(PC_15, 0);
            i++;
        }
        i = 0;
    }
}

```

Przejdźmy zatem do rzeczy przyjemniejszych, czyli poprawnej implementacji funkcji debugowania. Implementacja ta zakłada, iż nasz projekt rozrósł się do niebagatelnych rozmiarów, i ma wiele zmiennych, które chcemy podejrzeć w różnych miejscach, ale po doprowadzeniu do pełnej sprawności nie chcemy przypadkowo zostawić żadnej funkcji print() czy println(). Huawei kiedyś popełniło taki błąd i mało nie straciło ogromnej inwestycji w Wielkiej Brytanii z powodu pomyłki jednego programisty, który zostawił taką „furtkę” do debugowania w swoim kodzie. Zatem jak to się robi? Dość prosto, używając dyrektyw preprocesora. **Listing 11** pokazuje, jak to się robi. Taka implementacja ma to do siebie, że zmiana jednej linijki powoduje usunięcie wszystkich funkcji debugujących z naszego kodu. Jeśli DEBUG ma wartość 1, to dyrektywa #ifdef dokonuje prostej substytucji poleceń na odpowiadające im funkcje łączności szeregowej. Jeśli DEBUG ma wartość zero, to w miejscu tych funkcji pojawiają się puste linijki – kompilator je pomija. Warto zwrócić uwagę na budowę definicji dla debug() i debugln(). Pokazany sposób pozwala użyć jednego lub więcej argumentów, a preprocesor automatycznie dodaje lub usuwa argumenty i przecinki, dzięki czemu kompilator poprawnie kompiluje kod, niezależnie od tego, czy te funkcje mają jeden czy więcej argumentów.

Napiszmy zatem przykładowy kod by przetestować funkcje debugowania. Posłużymy się przedstawionym wcześniej kodem z pustą pętlą główną, i ją zapełnimy czymś z lekka losowym. By pokazać efekty działania naszego kodu użyjemy drobnego dodatku sprzętowego w formie pierścienia z 24. programowalnymi diodami LED WS2812B. Kompletny kod jest przedstawiony na **listingu 12**. Mamy więc dyrektywy preprocesora, w tym dodatkowe dyrektywy #define potrzebne do realizacji obsługi WS2812B: WSTh i WSTl ustalają czasy trwania stanów pinów dla wartości 0 i wartości 1. Czasy przy użytych wartościach pokazuje **rysunek 4**. Sumaryczny okres jednego znaku wynosi 1,2 µs, wartość WSTh daje nam 392 ns, a WSTl 810,5 ns. LEDMax określa maksymalną liczbę diod do zapalenia. W teorii moglibyśmy obsługiwać dowolnie długi łańcuch diod, ale w praktyce ogranicza nas zarówno zasilanie, jak i czas potrzebny na zaprogramowanie wszystkich diod. Dla zachowania płynności maksymalna liczba diod na jeden łańcuch



Rysunek 4. Czasy trwania sygnałów dla poprawnej transmisji danych do WS2812B. W tym wypadku wysyłany jest bit „0”

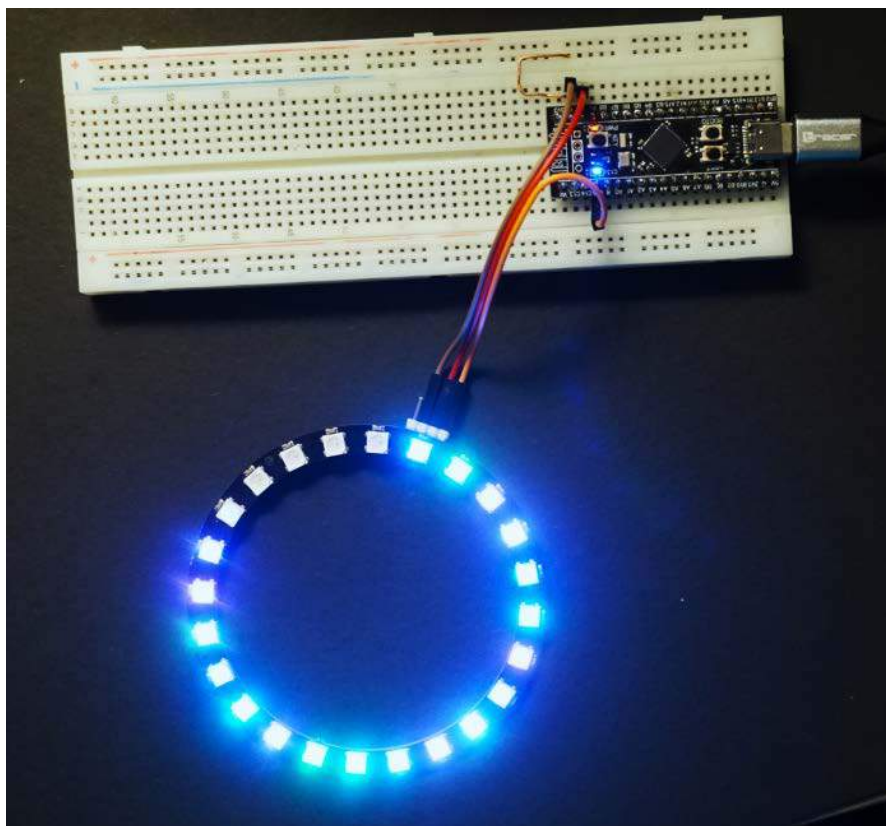
nie powinna być większa, niż 1440. Wracając do kodu, to w tej sekcji znajdują się też deklaracje trzech funkcji, znanej już nam funkcji obsługi przerwania i dwóch funkcji potrzebnych do obsługi diod oraz czterech zmiennych pomocniczych. Jedna z nich jest w formie unii zmiennej colors i struktury color, dzięki czemu można uzyskać dostęp do poszczególnych kolorów, jak i do całej zmiennej. Deklaracja obiektu myTim też znajduje się tutaj, by obiekt był globalny, a nie ograniczony do funkcji setup().

Funkcja setup() zawiera kilka zmian. Przede wszystkim dodałem kod pozwalający przeprowadzić test z rysunku 3. Poza konfiguracją pinu PC15, ustawiony jest też wirtualny port szeregowy do debugowania oraz wzywana jest też funkcja zerująca wartości kolorów diod WS2812B. Na koniec ustawiony jest timer używając wcześniej zadeklarowanego obiektu.

Pętla główna zawiera całą „magię”. Do zmiennej randVal przypisywana jest wartość wygenerowana za pomocą funkcji random() używającej programowego rejestru przesuwającego z liniowym sprzężeniem zwrotnym – STM32F401CCU8 nie posiada sprzętowego generatora liczb losowych, choć moglibyśmy taki zrealizować używając garści

Listing 13. Wynik działania funkcji debugowania programu z listingu 12 w programie CoolTerm

```
21:12:51.402 --> RandVal, dec: 1663234727, hes: 0x6322F2A7
21:12:51.402 --> colors, dec: 6497010, hex: 0x6322F2, limit: 0x6020F0
21:12:51.402 --> , B: 0x60, R: 0x20, G: 0xF0
21:12:51.402 --> RandVal, dec: 795202630, hes: 0x2F65D446
21:12:51.402 --> colors, dec: 3106260, hex: 0x2F65D4, limit: 0x2C64D4
21:12:51.402 --> , B: 0x2C, R: 0x64, G: 0xD4
21:12:51.402 --> RandVal, dec: 1643691567, hes: 0x61F8BE2F
21:12:51.402 --> colors, dec: 6420670, hex: 0x61F8BE, limit: 0x60F8BC
21:12:51.402 --> , B: 0x60, R: 0xF8, G: 0xBC
21:12:51.433 --> RandVal, dec: 100351560, hes: 0x5FB3E48
21:12:51.433 --> colors, dec: 391998, hex: 0x5FB3E, limit: 0x4F83C
21:12:51.433 --> , B: 0x4, R: 0xF8, G: 0x3C
21:12:51.433 --> RandVal, dec: 482588153, hes: 0x1CC3B5F9
21:12:51.433 --> colors, dec: 1885109, hex: 0x1CC3B5, limit: 0x1CC0B4
21:12:51.433 --> , B: 0x1C, R: 0xC0, G: 0xB4
```



Fotografia 1. Efekt działania programu z listingu 12

elementów. W ramach debugowania wartości tej zmiennej, oraz zmiennej colors są wysyłane do komputera. Zmienna colors zawiera wartość randVal przesuniętą w prawo o 8 pozycji, a następnie ograniczoną do maksymalnych wartości dla poszczególnych kolorów wynoszącej 63, zamiast 255. W ten sposób ograniczamy maksymalny średni pobór prądu przez diody do ~5 mA na diodę na kolor – bez tego okolice gniazda USB-C zaczynały się mocno rozgrzewać. Ponadto diody te są wyjątkowo jasne. Warto tu zwrócić uwagę na fakt, iż kolejność znaczenia bitów jest odwrotna, i najbardziej znaczący bit będzie pierwszym wysłanym. Następna część porównuje wartości LEDCount i LEDNum,

REKLAMA

by zapalić tylko część diod funkcją WSWrite(), resztę zaś wygasza funkcją WSClear(). Obie funkcje są poprzedzone funkcją wyłączającą przerwanie timera, a po nich timer jest ponownie włączany. To dlatego, że przesłanie 24 bitów zajmuje więcej czasu, niż wynosi odstęp między przerwaniami, a czas przerwania dodany do transmisji bitu zaburzy długość któregoś ze stanów. Reszta pętli głównej odpowiada za inkrementację LEDCount i LEDNum, zerowanie ich w odpowiednich momentach i dodawanie opóźnienia 50 ms, dzięki czemu efekt animuje się ładnie.

Funkcje WSWrite() i WSClear() są dość podobne do siebie. Pętla for przesyła 24 bity, przy czym w WSClear są to same zera, podczas gdy WSWrite używa instrukcji warunkowej by odczytać ostatni bit zmiennej L, i wybrać odpowiedni wariant, gdzie stan wysoki jest albo dłuższy, albo krótszy od stanu niskiego. Pętle while() są używane do realizacji opóźnienia wynoszącego tyle, ile zadeklarowaliśmy dyrektywami #define. Funkcja T2_task() nie była modyfikowana.

Na koniec porównajmy objętość programu z włączonym i wyłączonym debugowaniem. Z debugowaniem włączonym kod potrzebuje 5128 bajtów pamięci RAM i 32720 bajtów pamięci programu. **Listing 13** pokazuje fragment transmisji z programu CoolTerm. Funkcje debugowania zaburzają jednak łączność z WS2812B, bo dodają zbyt wiele opóźnień, przez co diody się resetują, zanim dostaną kolejny pakiet danych. Z debugowaniem wyłączonym kod potrzebuje 5124 bajty pamięci RAM i 32224 bajty pamięci programu. **Fotografia 1** pokazuje wynik działania programu. Funkcja PWM działa nadal, ale przez jej wstrzymanie na czas programowania diod zmienia częstotliwość pracy. Moglibyśmy to poprawić redukując częstotliwość PWM i modyfikując niektóre wartości, ale to zadanie zostawiam jako ćwiczenie dla Czytelników.

Co dalej?

Moglibyśmy kontynuować rozbudowę tego programu, dodać inne tryby świecenia pojedynczą diodą albo kółkiem z diodami WS2812B, pokusić się o tęczę, gonące się światełka, efekt ognia lub zorzy polarnej. Napisać od zera całą bibliotekę NeoPixel lub podobną. Myślę jednak, iż taki obszerny wstęp jest wystarczający, ale zachęcam Czytelników do własnych eksperymentów, zwłaszcza że koszt rozpoczęcia przygody z PlatformIO i STM32Duino jest relatywnie niski. ■

Paweł Kowalczyk

numery archiwalne • prenumerata • książki
www.UlubionyKiosk.pl

Elektrony

Według fizyki klasycznej elektrony to bardzo małe, twarde jak skała kule, które krążą po kołowych i eliptycznych orbitach wokół jądra atomu. Są one absolutną podstawą elektroniki.

Przez fale do wolnych elektronów

Fale elektromagnetyczne

Ten artykuł zaczyna się od czegoś, co wydaje się być tematem pobocznym: fal elektromagnetycznych. Jednak omówienie tego zjawiska jest niezbędne do zrozumienia szczególnych właściwości elektronów.

Promieniujące elektrony

Następnie omówimy fakt, że w pewnych warunkach elektrony mogą emitować światło. Właściwość ta jest bardzo ważna w elektronice, ponieważ bez niej nie istniałyby na przykład diody LED.

Wolne elektrony

Na zakończenie skupimy się na zjawisku odłączania się elektronów od swoich atomów, co jest niezbędnym warunkiem dla przepływu prądu elektrycznego. Takie elektrony nazywane są elektronami swobodnymi, a to, jak działają procesy fizyczne, które mogą uwolnić elektrony z ich atomów, jest ostatnim tematem tego artykułu.

Fale elektromagnetyczne

Pola elektryczne i magnetyczne

Przez cały XIX wiek uzyskano dobry wgląd w strukturę światła. Wiadomo było, że światło jest częścią spektrum promieniowania elektromagnetycznego. Promieniowanie, które powstaje w wyniku wzajemnego oddziaływania pola magnetycznego i elektrycznego. Zmienne pole elektryczne powoduje powstanie zmiennego pola magnetycznego. To pole z kolei powoduje powstanie zmiennego pola elektrycznego. Te połączone pola tworzą zjawisko falowe w przestrzeni, falę elektromagnetyczną. Fala ta charakteryzuje się długością fali, oznaczaną literą λ (lambda), która jest wyrażona w metrach.

Drgania fal elektrycznych (E) i magnetycznych (H) powodują powstawanie linii pola w przestrzeni. Na rysunku poniżej pola te są reprezentowane za pomocą linii kolistych. Linie pola elektrycznego E1 indukują prostopadłe do nich linie pola magnetycznego H1, te zaś indukują prostopadłe do nich linie pola E2, i tak dalej. W ten sposób powstaje fala elektromagnetyczna złożona z prostopadłych względem siebie linii pól elektrycznego i magnetycznego.

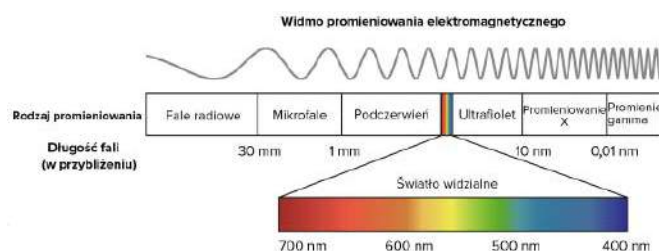
Prędkość fal elektromagnetycznych

Prędkość, z jaką to zjawisko rozchodzi się w przestrzeni jest równa 299 792,5 km/s! Oznacza to, że fala elektromagnetyczna może okrążyć Ziemię około 7,5 razy w ciągu jednej sekundy!

Widmo elektromagnetyczne

Światło widzialne stanowi tylko niewielką część widma elektromagnetycznego. Zakresy fal, od pasma radiowego, przez pasma telewizyjne, GSM, Wi-Fi, aż do sygnałów radarowych obejmują długości od 1000 m do około 1 mm. Następnie znajduje się obszar podczerwieni, czyli elektromagnetycznego promieniowania ciepłego, które

emitują słońce, grzejniki i ogień. Następną część spektrum to niezwykle wąski obszar światła widzialnego. Światło czerwone ma falę najdłuższą, a fioletowe najkrótszą. Za obszarem światła widzialnego znajduje się obszar promieniowania ultrafioletowego, niebezpiecznego dla ludzkiej skóry. Obszar jeszcze krótszych fal (10^{-8} m) to już promieniowanie rentgenowskie, znane ze szpitali, i promieniowanie gamma, które jest bardzo niebezpieczne dla zdrowia. Chociaż wszystkie te rodzaje promieniowania wydają się zupełnie inne, podstawa jest zawsze taka sama: interakcja linii pola elektrycznego i magnetycznego, które rozchodzą się w przestrzeni ze stałą prędkością.



Widmo promieniowania elektromagnetycznego (Khan Academy, CC BY-SA 3.0)

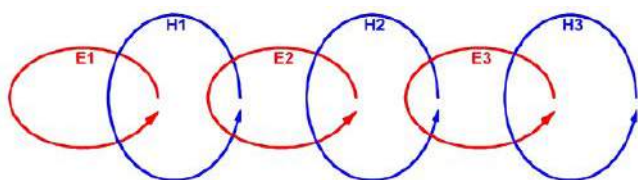
Promieniowanie złożone

Niektóre źródła promieniowania elektromagnetycznego emitują tylko jedną określoną długość fali, na przykład niemodulowany nadajnik radiowy lub laser. Jednak większość źródeł elektromagnetycznych emituje złożony sygnał, zawierający fale o różnych długościach. Typowym przykładem jest słońce, które emituje promieniowanie elektromagnetyczne od dalekiej podczerwieni do dalekiego ultrafioletu. Opracowano urządzenia, za pomocą których można badać skład długości fal takiego promieniowania. Urządzenia te nazywane są „spektroskopami”.

Najprostszym spektroskopem jest szklany pryzmat, który może być wykorzystywany do analizy światła słonecznego. Urządzenie składa się z nieprzepuszczającej światła płytki, w której wykonano niewielki otwór. Przez ten otwór światło słoneczne pada na pryzmat. Promieniowanie elektromagnetyczne ze słońca jest przez niego załamane. Jednak każda długość fali ma inny kąt załamania. W rezultacie wąska wiązka białego światła słonecznego jest przekształcana w szeroką wiązkę, w której widoczne są wszystkie kolory tęczy. Jeśli skierujesz tę wiązkę na biały ekran, zobaczysz piękne kolory poszczególnych długości fal światła słonecznego, od głębokiej czerwieni do fioletu. Przy użyciu bardziej skomplikowanego sprzętu możliwe jest również obserwowanie fal promieniowania elektromagnetycznego, których ludzkie oko nie widzi.



Najprostszym spektroskopem jest szklany pryzmat (Wikimedia Commons, edytowane przez Jos Verstraten)



Zjawisko fali elektromagnetycznej w przestrzeni (© 2017 Jos Verstraten)

Atomy i promieniowanie

Tuba wyładowcza

Jeśli gazowy wodór zostanie zamknięty w szklanej rurce, a do rurki zostanie przyłożone duże napięcie stałe, gazowy wodór wyemituje światło. Analizując to zjawisko elektromagnetyczne za pomocą spektroskopu, można zaobserwować coś niezwykłego. Widmo elektromagnetyczne emitowane przez wodór nie składa się z ciągłego widma, jak w przypadku światła słonecznego, ale z dobrze zdefiniowanych długości fal. Każda linia na rysunku reprezentuje określoną długość fali z zebranego i przeanalizowanego światła wodorowego. Okazuje się, że wszystkie substancje mają takie widma liniowe. Rozkład linii jest na tyle specyficzny dla danej substancji, że można go uznać za rodzaj odcisku palca danej substancji.



Widmo elektromagnetyczne emitowane przez wodór w tubie wyładowczej (Wikimedia Commons)

Pochodzenie linii

Pytanie, na które musieli odpowiedzieć fizycy, brzmiało: skąd wzięły się te charakterystyczne linie? Wiadomo już było, że do emisji promieniowania elektromagnetycznego potrzebna jest energia. Jedynymi obiektami w atomie, które posiadają taką energię są elektrony, które krążą wokół jądra z dużą prędkością (a więc i energią). Gdyby jednak elektrony emitowały promieniowanie elektromagnetyczne, straciłyby energię (a tym samym prędkość). W rezultacie elektrony krążyłyby bliżej jądra, a ostatecznie nawet zderzyłyby się z nim. Jest to w wyraźnej sprzeczności z zaobserwowanymi faktami!

Nowy model atomowy Bohra

Duński fizyk Niels Bohr opracował około 1914 roku teorię, która mogłaby wyjaśnić te dziwne linie widma. Bohr pracował z Thomsonem i Rutherfordem podczas swoich lat studenckich i znał istniejące modele atomowe jak własną kieszeń. Był jednak również świadomy problemów związanych z tymi modelami. Bohr opracował nowy model atomu, który w rzeczywistości dobrze pasował do modelu planetarnego Rutherforda, ale jednocześnie wyjaśniał stacjonarne orbity elektronów i emisję widm liniowych.

Tezy Bohra

Bohr stwierdził:

- Elektrony mogą krążyć wokół jądra jedynie po ściśle określonych orbitach (powłokach) (nic nowego!). Jednak wykorzystując pewne twierdzenia mechaniki kwantowej Maxa Plancka, Bohr był w stanie obliczyć promienie tych orbit. I to było nowe, ponieważ rozwiązało pytanie, dlaczego zaobserwowano istnienie tych powłok.
- W przeciwieństwie do wszystkiego, co twierdzi teoria elektromagnetyczna Maxwella, elektrony na tych orbitach nie emitują żadnej energii. Bohr rozwiązał tę sprzeczność, twierdząc, że tradycyjna teoria elektromagnetyczna nie ma zastosowania w małej skali struktury atomu. Elektrony mogą zatem krążyć po tych orbitach w nieskończoność, ponieważ są to stany podstawowe atomu. Te powłoki nazywane są „orbitami Bohra” elektronów. Elektrony wolą zatem znaleźć miejsce na jednej z orbit Bohra, ponieważ jest to najbardziej stabilny stan atomu.
- Każda orbita (powłoka) odpowiada określonej energii spoczynkowej elektronów. Energia ta zależy od odległości między jądrem a powłoką. Jeśli do atomu zostanie dodana dodatkowa energia, na przykład poprzez podgrzanie go lub wystawienie na działanie silnego pola elektrycznego, elektrony mogą zaabsorbować tę energię. Są wtedy w stanie przeskoczyć na orbitę lub powłokę,

która ma wyższą energię spoczynkową. Warunkiem jest dostarczenie wystarczającej ilości energii do atomu, aby pokonać różnicę w energii spoczynkowej między dwiema powłokami. Elektron, który opuścił swoją podstawową orbitę Bohra poprzez dodanie dodatkowej energii, nazywany jest elektronem „wzbudzonym”.

- W tym momencie atom znajduje się w stanie niestabilnym. Wzbudzony elektron powróci na podstawową orbitę Bohra tak szybko, jak to możliwe. Jednak wracając na tę orbitę elektron musi wyemitować pewną ilość energii w postaci promieniowania elektromagnetycznego. Ilość emitowanej energii jest równa różnicy energii spoczynkowej między dwiema powłokami. Przeście elektronu ze stanu wzbudzonego do podstawowego stanu Bohra nazywane jest „skokiem kwantowym”.

Wyjaśnienie widm liniowych

Dzięki tym stwierdzeniom Bohr był w stanie wyjaśnić emisję widm liniowych. W końcu każda długość fali w widmie odpowiada przeskokowi wzbudzonego elektronu z powłoki zewnętrznej do powłoki wewnętrznej. Ponieważ możliwych jest wiele przeskoków między siedmioma powłokami, można również oczekiwać wielu linii w widmie, którym odpowiadają różne długości fal.

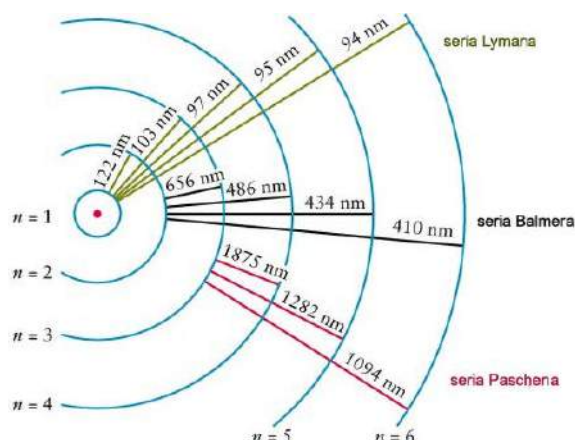
Teoria jest poprawna!

Możliwe było obliczenie energii odpowiadających wszystkim możliwym skokom kwantowym wzbudzonego elektronu. Znając energię uwalnianą podczas takiego skoku, można również obliczyć odpowiadającą mu długość fali promieniowania elektromagnetycznego. Im więcej energii zostaje uwolnione, tym mniejsza będzie długość fali generowanego promieniowania. Każdy skok kwantowy generuje zatem jedną linię o określonej długości fali, w widmie badanego atomu. Później niezliczeni fizycy eksperymentalni próbowali znaleźć wszystkie linie, które odpowiadają wszystkim możliwym skokom kwantowym zgodnie z teorią Bohra. Na poniższym rysunku wymieniono nazwiska naukowców, którzy całkowicie rozwikłali strukturę linii wodoru. Ponadto wskazano, które długości fal, których skoków kwantowych zaobserwowali.

Znaczenie dla elektroniki

Oczywiste jest, że opisane właściwości fizyczne elektronu mają bardzo duże znaczenie dla elektroniki. Fakt, że elektrony wytwarzają promieniowanie elektromagnetyczne, gdy powracają ze stanu wzbudzonego do stanu podstawowego, jest wykorzystywany w wielu komponentach elektronicznych. Rozważmy:

- lampy kineskopowe i oscyloskopowe,
- diody LED,
- panele elektroluminescencyjne,
- światłówki,
- „neonówki”.



Przegląd skoków kwantowych elektronu w atomie wodoru (Wikimedia Commons)

Podstawy optoelektroniki

W tych komponentach atomy są najpierw wprowadzane w stan wzbudzony poprzez dodanie energii z zewnątrz. Pod jej wpływem elektron przechodzi na wyższą orbitę, stając się elektronem wzbudzonym, po czym „spada” na orbitę podstawową. W tym momencie emitowane jest promieniowanie elektromagnetyczne o określonej długości fali. Obecnie te procesy są na tyle dokładnie poznane, iż możliwe jest teoretyczne określenie, jaka substancja wyemituje jaki kolor. Później pozostaje już tylko kwestia wyprodukowania niezbędnych surowców w powtarzalny i tani sposób na skalę przemysłową.



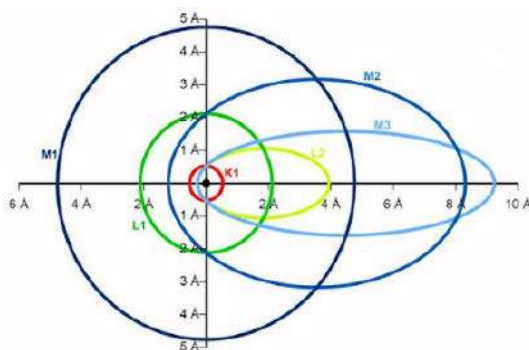
Wyjaśnienie emisji promieniowania elektromagnetycznego przez elektron (Wikimedia Commons, edytowane przez Jos Verstraten)

Model atomowy Sommerfelda

Gdy naukowcy opracowali dużo precyzyjniejsze spektroskopy i zaczęli dokładnie badać linie widmowe wodoru, okazało się, że pojedyncze linie składały się w rzeczywistości z kilku linii, które znajdowały się bardzo blisko siebie. Nazwano to „subtelną strukturą” widma. Teoria Bohra nie potrafiła tego wyjaśnić. Wyjaśnienia tego zjawiska dostarczył Sommerfeld. Postulował on, że elektrony mogą krążyć wokół jądra nie tylko po orbitach kołowych, ale także eliptycznych. Istnieje wówczas znacznie więcej możliwości energii spoczynkowych na powłokę, co skutkuje znacznie większą liczbą przejść z powłoki na powłokę. Jednak energie spoczynkowe na orbicie eliptycznej są bardzo zbliżone do siebie, a zatem również długości fal emitowanego promieniowania elektromagnetycznego też są zbliżone. Dzięki bardzo dokładnemu zbadaniu struktury subtelnej, Sommerfeld stworzył model przedstawiony na poniższym rysunku:

- powłoka K ma jedną orbitę kołową,
- powłoka L ma jedną orbitę kołową i jedną eliptyczną,
- powłoka M ma jedną orbitę kołową i dwie eliptyczne,
- itd.

Oczywiście, tylko pewna liczba elektronów może pojawić się na każdej orbicie. Jednak całkowita liczba elektronów na powłoce pozostaje taka sama, jak stwierdził Rutherford. To udoskonalenie modelu Bohra zapewnia również proste wyjaśnienie przewodnictwa elektrycznego niektórych substancji.



Orbity eliptyczne rozwiązują problem subtelnej struktury widma (Wikimedia Commons, edytowane przez Jos Verstraten)

Epilog

Model atomu Sommerfelda jest najbardziej dokładną reprezentacją struktury atomu, do jakiej jest zdolna mechanika klasyczna. Jednak model ten pozostawia również wiele otwartych pytań. Na wiele z tych pytań można odpowiedzieć, jeśli atom nie jest już postrzegany jako rodzaj miniaturowego układu słonecznego z kulkami twardej materii, ale jako spójność pól materii. Mechanika kwantowa zapewnia teoretyczną, matematyczną podstawę dla takiej reprezentacji struktury atomu. Ale ta skomplikowana teoria całkowicie wykracza poza zakres tego artykułu. Aby zrozumieć fizyczne podstawy elektroniki, można całkowicie polegać na modelu atomowym Sommerfelda.

Skład jądra atomowego

Skład jądra atomowego nie jest omawiany, ponieważ nie jest to ważne dla elektroniki. Ze względu na kompletność wspomnijmy jedynie, że jądro składa się z protonów i neutronów. Protony są nośnikami dodatniego ładunku elektrycznego i każdy atom zawiera tyle samo protonów, co elektronów. Neutrony nie są naładowane, ale mają znaczną masę. Ze względu na swoją elektryczną neutralność neutrony nie odgrywają żadnej roli w zjawiskach elektrycznych i elektronicznych.

Elektrony swobodne w przewodnikach

Emisja elektronów

Elektrony są nośnikami energii elektrycznej, a prąd elektryczny to inaczej przepływ elektronów w określonym kierunku w jednostce czasu. Dopóki jednak elektrony są związane z jądrem atomowym, nie może być mowy o przepływie prądu elektrycznego. Następnym pytaniem, na które należy odpowiedzieć, jest to, w jaki sposób elektrony mogą uciec z atomowego więzienia. Zwykle ma to również związek z dostarczeniem wystarczającej ilości energii do atomu. Tak jak elektron może zostać przeniesiony ze swojej orbity Bohra na bardziej odległą orbitę, tak samo elektron może zostać przyspieszony do takiego stopnia, że ucieknie z wiązania atomowego. Wymagana do tego ilość energii nazywana jest „energią wyjścia”, reprezentowaną przez symbol ϕ (phi). Zjawisko opuszczania wiązania przez elektrony nazywane jest „emisją”.

Energia wyjścia

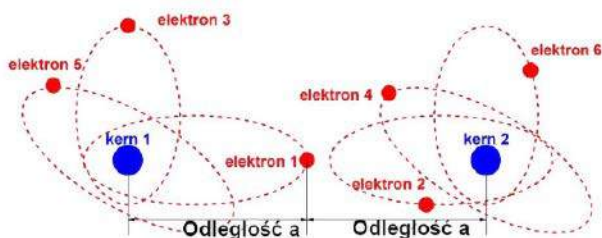
Energia wyjścia zależy od wielu czynników. Ogólnie rzecz biorąc, atomy, które mają całkowicie zapełnioną górną powłokę, muszą zaabsorbować znacznie więcej energii, aby umożliwić elektronowi ucieczkę niż atomy, w których ostatnia powłoka jest praktycznie pusta.

Zewnętrzne pole elektryczne

Drugim warunkiem przepływu prądu elektrycznego jest istnienie pola elektrycznego w poprzek przewodnika. W niektórych substancjach, takich jak metale, zewnętrzne elektrony są tak słabo związane z atomem, że mogą wędrować od atomu do atomu. Jednak te już wolne elektrony nie wywołują przepływu prądu elektrycznego! Nie wystarczy bowiem uwolnić elektronów z ich atomów. Bez zewnętrznego pola elektrycznego elektrony te wykonywałyby losowe ruchy od atomu do atomu, jak rodzaj chmury. Jednak po przyłożeniu zewnętrznego pola elektrycznego swobodne elektrony zostaną przyciągnięte do dodatniego ładunku pola, co zainicjuje przepływ elektronów.

Elektrony swobodne w przewodnikach

Poniższy rysunek pokazuje, jak można sobie wyobrazić powstawanie elektronów swobodnych w dobrym przewodniku zgodnie z modelem atomowym Sommerfelda. W metalach atomy znajdują się blisko siebie. Jeden z elektronów lewego atomu znajduje się na orbicie eliptycznej. W pewnych momentach elektron ten znajduje się w odległości a od swojego jądra. Jednak ze względu na gęstość atomów, elektron ten znajduje się w tym momencie w identycznej odległości a od sąsiedniego atomu (tego po prawej). W tym momencie elektron może przeskoczyć



Ruch swobodnych elektronów w przewodnikach (© 2017 Jos Verstraten)

z jednego atomu do drugiego bez większego wysiłku. Oczywiście zjawisko to nie występuje tylko raz, ale miliardy razy na sekundę.

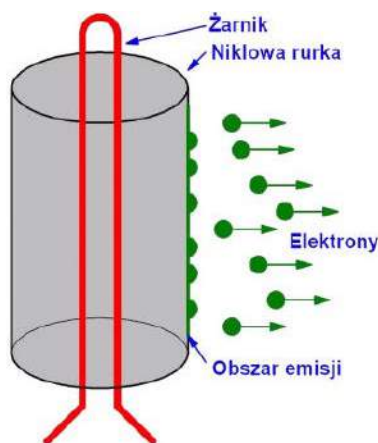
Prąd jest wytwarzany przez przyłożenie pola.

Miliardy elektronów przemieszczają się od atomu do atomu wewnątrz substancji. Bez wpływu czynników zewnętrznych ruch ten jest chaotyczny i dlatego nic się nie zmienia. Jeśli jednak do metalu zostanie przyłożone niewielkie napięcie elektryczne, w przewodniku powstanie pole elektryczne. Pole to wywiera siłę na swobodne elektrony, powodując ich ruch w kierunku określonym przez kierunek pola. Chaotyczny ruch staje się uporządkowany. Zjawisko to bardzo szybko rozprzestrzenia się w metalu, powodując przepływ elektronów od bieguna ujemnego do dodatniego. Ciężkie jądra utknęły w strukturze krystalicznej metalu i nie mogą się poruszać. W przewodniku tylko bardzo lekkie swobodne elektrony przemieszczają się tworząc przepływ prądu.

Emisja termiczna

Najczęstszym sposobem uwalniania elektronów z ich struktury atomowej jest emisja termiczna. W skali atomowej ciepło jest niczym innym, jak drganiami atomów. Intensywność tych drgań jest proporcjonalna do energii. Jeśli atomy mają dostatecznie niską energię wyjścia, ogrzewanie da elektronom wystarczającą ilość dodatkowej energii, aby uwolnić je z wiązań atomowych. Wystarczy wtedy zastosować zewnętrzne pole elektryczne, aby umożliwić przepływ prądu elektronowego. W tym kontekście gdy mowa o ogrzewaniu, nie chodzi o rozgrzewanie materii do wysokich temperatur. W fizyce ogrzewanie oznacza podniesienie temperatury substancji powyżej zera bezwzględnego, czyli $-273,16^{\circ}\text{C}$. Z fizycznego punktu widzenia, substancja, która ma temperaturę -200°C już znacznie się ogrzała!

Warto dodać, iż zjawisko przewodnictwa cieplnego polega na tym, iż atomy o większej energii drgając wchodzi w interakcje z sąsiednimi atomami przekazując im nadwyżkę energii i powodując ich większe drgania. Przepływ ten zawsze zachodzi od części cieplejszej, do chłodniejszej – wymuszenie przepływu w przeciwnym kierunku zawsze wymaga dostarczenia dodatkowej energii – przyp. tłum.



Emisja termiczna w katodzie kineskopu generuje chmurę swobodnych elektronów (© 2017 Jos Verstraten)

Prąd upływu zależy od temperatury

Emisja termiczna jest przyczyną zjawiska polegającego na tym, że prąd upływu półprzewodników wzrasta wraz ze wzrostem jego temperatury.

Lampy oscyloskopowe i kineskopy działają dzięki emisji termicznej

Emisja termiczna była na przykład wykorzystywana w staromodnych lampach elektronowych do generowania elektronów. W kineskopach lub lampach oscyloskopowych katoda jest podgrzewana pośrednio (za pomocą żarnika) lub bezpośrednio (jest jednocześnie żarnikiem). Chmura swobodnych elektronów emitowanych przez katodę jest następnie przyspieszana i skupiana za pomocą dodatkowych elektrod o dodatnim potencjale lub zewnętrznych pól magnetycznych. Tak skupiona i przyspieszona wiązka uderza w ekran pokryty luminoforem, który pochłaniając energię rozprędzonych elektronów emituje światło. Ruch wiązki po ekranie jest kontrolowany za pomocą elektrod lub cewek odchylających.

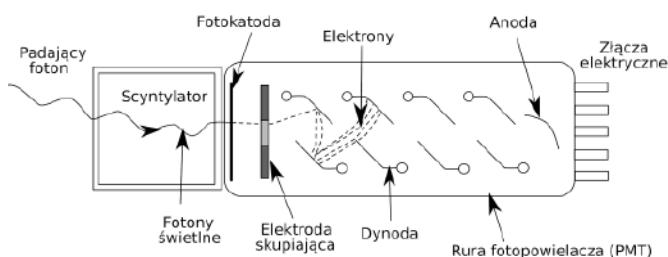
Fotoemisja

Niezbędna energia wyjściowa może być również dostarczona w postaci promieniowania elektromagnetycznego. Z ogólnych praw elektromagnetyzmu wiadomo, że energia promieniowania jest odwrotnie proporcjonalna do długości fali. Jeśli więc napromieniamy substancję promieniowaniem elektromagnetycznym, w pewnym momencie ilość energii dostarczonej będzie wystarczająca, by elektrony mogły się oderwać od swoich atomów i stać się elektronami swobodnymi. Zastosowanie zewnętrznego pola elektrycznego jest wystarczające, aby uwolnione elektrony mogły przepływać w jednym kierunku przez substancję. To zjawisko zwane fotoemisją jest fizyczną podstawą działania takich komponentów, jak fotodiody, fototranzystory czy fotorezystory (LDR).

Emisja wtórna

Elektrony mogą być również uwalniane z atomów poprzez bombardowanie substancji wysokoenergetycznymi elektronami. Elektrony te zderzają się z elektronami związanymi z atomami. Podczas zderzenia zawsze dochodzi do transferu energii. Powoduje to, że niektóre związane elektrony zyskują wystarczająco dużo energii, by móc uciec z wiązania atomowego. Jeśli energia padających elektronów jest tak duża, że są one w stanie uwolnić kilka elektronów ze swoich atomów, nazywa się to „wtórnym zwielokrotnieniem”. Efekt ten występuje na przykład w fotopowielaczach, gdzie niektóre elektrony są najpierw uwalniane z elektrody (zwanej dynodą – przyp. tłum.) poprzez fotoemisję. Następnie elektrony te są przyspieszane przez pole elektryczne. Przyspieszone elektrony mogą następnie zderzyć się z drugą elektrodą. Elektrony pierwotne dzięki przyspieszeniu zyskały tyle energii, że są w stanie zderzyć się z różnymi związanymi elektronami i uwolnić je z atomów. Powtarzając opisany proces kilka razy, można uzyskać wtórne zwielokrotnienie o współczynniku od 100 do 10 000. ■

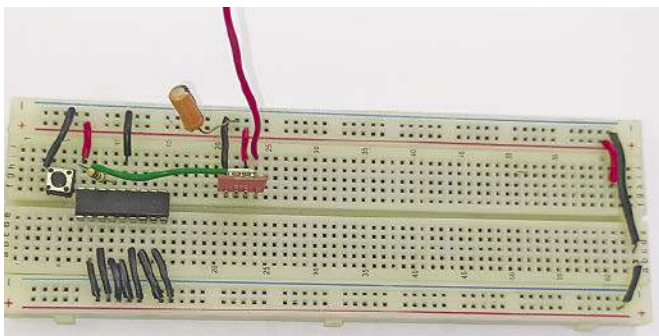
Jos Verstraten



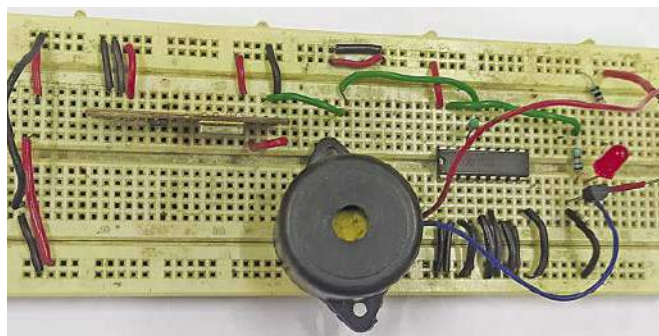
Wtórna emisja elektronów w fotopowielaczu. Scyntylator (specjalny kryształ) zamienia fotony promieniowania rentgenowskiego i gamma na fotony światła, które wybijają pojedyncze elektrony z fotokatody (Wikipedia Commons)

Dzwonek bezprzewodowy

Tradycyjne dzwonki montowane na furtce lub na drzwiach mieszkania bywają kłopotliwe w montażu. Zwykle występuje problem z ciągnięciem przewodów, z umocowaniem samego dzwonka, a jeśli system zasilany jest wprost z sieci 230 V AC, to montaż powinien być wykonany przez wykwalifikowanego elektryka. Alternatywą są dzwonki bezprzewodowe, bardzo proste i wygodne w montażu, a jedyną chyba ceną za te korzyści jest potrzeba zasilania bateryjnego w sekcji nadajnika takiego dzwonka. Nadajnik i odbiornik dzwonka bezprzewodowego można zamontować niemal gdziekolwiek, byle odbiornik znajdował się w zasięgu nadajnika. Bieżący artykuł pokazuje jak taki dzwonek wykonać.



Rysunek 1. Prototyp nadajnika



Rysunek 2. Prototyp odbiornika dzwonka bezprzewodowego

Konstrukcja jest prosta i jednocześnie tania dzięki dostępności układów scalonych kodera i dekodera, a także sparowanych ze sobą części w.cz. nadajnika i odbiornika. Konstrukcja dzwonka podąża za powszechną modą i tendencją aby wszystko było „bez drutu” i „bez kabla” (wireless i cordless), a w ten trend wpisują się także akumulatorene elektronarzędzia, a od pewnego czasu również szeroko rozumiana elektromobilność.

Rysunki 1 i 2 pokazują prototyp wykonany przez autora, odpowiednio nadajnika i odbiornika bezprzewodowego dzwonka.

Wykaz elementów:

Półprzewodniki:

IC1: Koder HT12E
IC2: Dekoder HT12D
T1: tranzystor NPN BC547
LED1: dioda LED 5 mm

Rezystory:

(wszystkie 0,25 W/±5%)
R1: 1 MΩ
R2: 33 kΩ
R3, R4: 1 kΩ

Kondensatory:

C1, C3: 100 µF/16 V elektrolityczny
C2: 0,01 µF ceramiczny

Inne:

S1: switch „push-to-on”
S2: przelącznik ON/OFF SPST
CON1, CON2: złączce 2-pinowe
TX1: Nadajnik RF 433 MHz
RX1: Odbiornik RF 433 MHz
PZ1: buzzer
ANT1, ANT2: odcinek przewodu 25–30 cm

Ponadto opcjonalnie:

RL1: przekaźnik 5 V SPDT
D1: dioda prostownicza 1N4007
CON3, CON4: złączce 2-pinowe
oraz dzwonek zasilany z sieci 230 VAC

Budowa układu i jego działanie

Z istoty dzwonka bezprzewodowego wynika, że układ składa się z dwóch oddzielnych części. System pracuje w paśmie 433 MHz wykorzystując gotowe moduły po stronie nadawczej jak i odbiorczej.

Konstrukcja nadajnika

W nadajniku wykorzystano specjalizowany układ scalony kodera (co zapewni, że nasz dzwonek nie zadzwoni do sąsiada – przypis red.). Switch włączający dzwonek ulokowany jest na linii zasilania, co jest rozsądne z uwagi na oszczędność energii w nadajniku. Na rysunku 3 pokazano schemat nadajnika.

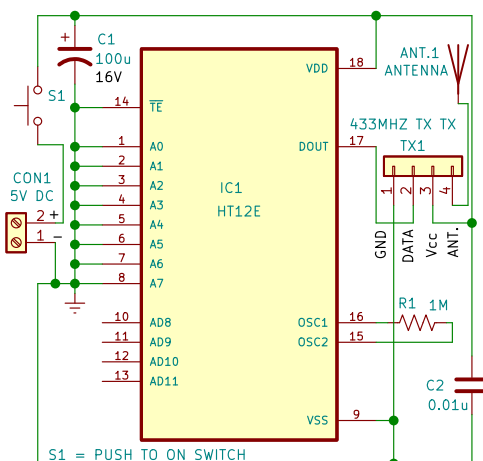
Enkoderem jest popularny układ scalony Holtek-a HT12E (IC1). Drugim istotnym podzespołem jest nadajnik RF 433 MHz,

a ponadto jest tu jedynie kilka drobnych dyskretnych elementów. HT12E koduje sygnał z wykorzystaniem linii adresowych i danych. Wyboru adresu dokonuje się poprzez przypisanie dedykowanym liniom stanu niskiego lub wysokiego. Tutaj, wszystkie linie A0 do A7 połączone są z potencjałem masy. Sparowanie nadajnika z odbiornikiem polega na wyborze tego samego adresu i danych. Chwilowe zwarcie styków S1 skutkuje sekwencją danych na linii Dout (pin 17). Linia ta podłączona jest na wejście Data TX1. Moduł ten moduluje nośną 433 MHz, która dociera do odbiornika.

Konstrukcja odbiornika

Tutaj wyróżniamy odbiornik w.cz., dekodek oraz obwód wykonawczy, którym jest dioda LED i buzzer lub głośniejszy dzwonek. Schemat ideowy części odbiorczej pokazano na rysunku 4.

Widzimy tu RF-receiver (RX1), dekodek HT12D (IC2), tranzystor npn BC547 (T1), buzzer (PZ1) i także kilka drobnych pasywnych elementów. Komunikacja w paśmie 433 MHz spoczywa na modułach TX1 i RX1. Sygnał po demodulacji w.cz. trafia na wejście DIN (pin 14) dekodera HT12D. Gdy dekodek stwierdzi zgodność adresu w strumieniu danych z adresem jaki mu przypisano na liniach A0 do A7, wówczas wystawia stan wysoki linii VT (pin 17). Linia ta włącza tranzystor T1 uruchamiając w ten sposób dźwięk buzzera oraz zaświecenie diody LED1. W ten sposób zwarcie switcha S1 w nadajniku



Rysunek 3. Schemat ideowy części nadawczej dzwonka



Korzystając z podzespołów wyszczególnionych w spisie elementów, wykonanie bezprzewodowego dzwonka jest bardzo proste. Najpierw zmontuj oddzielnie nadajnik

Problem w tym, że głośniejszy dzwonek jakiegokolwiek by on był konstrukcji, jest zapewne też bardziej energochłonny. W opcji z rysunku 5 wykorzystano przekątnik, który należy wpiąć w miejsce „wykropkowanej” części schematu z rysunku 4. Na obu schematach zaznaczono węzły A i B, które należy wzajemnie wymienić. Po stronie wykonawczej przekątnika można zastosować niemal

Konstrukcja i testowanie działania układu

Po uzbrojeniu PCB, obwód wraz z baterią zasilającą należy zamknąć w odpowiednio przygotowanej obudowie. W wygodnym miejscu należy zamontować switch S1 i całość nadajnika umieścić na furtce lub na zewnętrznej stronie drzwi mieszkania.

Odpowiednio, na **rysunkach 8 i 9** mamy projekt płytki drukowanej dla odbiornika oraz schemat montażowy pomocny przy rozmieszczaniu elementów na PCB. Płytkę PCB uwzględnia schemat z rysunku 4, nie przygotowano wersji dla odbiornika z przekaźnikiem.

Po zmontowaniu części odbiorczej, również należy przewidzieć dla niej odpowiednio przygotowaną obudowę. Stronę odbiorczą montujemy wewnątrz mieszkania. Jedynym ograniczeniem jest zasięg nadajnika oraz fakt, że po stronie odbiorczej prawdopodobnie będziemy chcieli zasilić układ z jakiegoś zasilacza sieciowego.

Oczywistym jest, że transmisja bezprzewodowa wymaga anteny nadawczej i odbiorczej. Tutaj, po obu stronach wystarczy użyć ok. 15 do 20 cm przewodu, który spełni zadanie takowej.

Pod adresem: <https://www.electronicshobby.com/videos-slideshows/live-diy-wireless-call-bell> można obejrzeć działanie układu dzwonka bezprzewodowego wykonanego przez autora. ■

S.C. Dwivedi

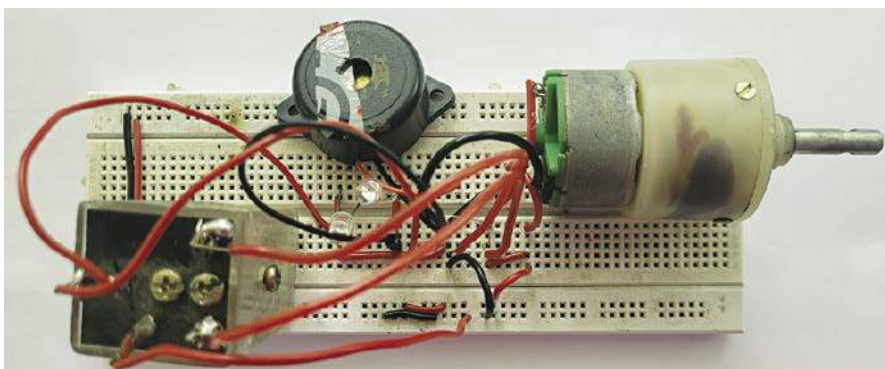


Ten artykuł był wcześniej opublikowany na łamach „EFY”, wrzesień 2023 (efymag.com)

Zabezpieczenie przeciwzwarceniowe w pojazdach elektrycznych

Być może najważniejszą sprawą w dzisiejszym świecie są źródła energii. Jak do tej pory źródła zasilania pojazdów bazują na pochodnych ropy naftowej, benzynie i oleju napędowym. Równocześnie mamy wszyscy świadomość, iż na dłuższą metę źródła kopalne są ograniczone. Z tego, i nie tylko z tego powodu rośnie świadomość wagi i nadzieje pokładane w energii elektrycznej. Zasilanie urządzeń stacjonarnych niemal wyłącznie bazuje na energii elektrycznej, natomiast od niedawna to samo ma się tyczyć urządzeń mobilnych. Największym wyzwaniem są pojazdy, dla których poszukuje się coraz nowszych form zasilania elektrycznego.

Obserwujemy ciągły wzrost zainteresowania „elektrykami” EV (Electric Vehicles). Z nową technologią wkraczają nowe problemy i wyzwania. Najistotniejsze zagrożenia dotyczą zwarć międzyzwojowych w silnikach elektrycznych, zjawisko wydzielania ciepła w samych akumulatorach, nagrzewanie się indukcyjności w silniku i w obwodach przetwornic zasilających i ogólnie problem sprawności energetycznej i związanych z tym strat energii. Sprawa jest kluczowej wagi, gdyż coraz częściej słyszymy o zdarzających się wypadkach pożarów oraz przegrzewania się i eksplozji baterii. Bieżący projekt wychodzi naprzeciw tym zagrożeniom. Za cenę prostego obwodu z dodatkowym przełącznikiem proponowany układ może uchronić przed tak groźnymi następstwami. Na **rysunku 1** widzimy prototyp, który przeszedł pozytywnie testy w naszym laboratorium (redakcji „Electronics For You”). Bardzo często przyczyną wypadków jest zjawisko zwarć prowadzące do przeciążenia prądowego akumulatora. Jeśli nawet skutek nie jest tak tragiczny jak pożar samochodu, to z pewnością ucierpi na tym akumulator, który jest być może najdroższym podzespołem w elektrycznym pojeździe. Różne są przyczyny zwarć, lecz generalnie oznacza to, że prąd płynie w niezamierzonym obwodzie gdzie impedancja widziana od źródła zasilania drastycznie maleje. Proponowane zabezpieczenie przeciwzwarceniowe powinno wykryć zjawisko przeciążenia akumulatora i natychmiast odłączyć obciążenie



Rysunek 1. Prototyp przetestowany w laboratorium EFY

od źródła zasilania. To uchroni i akumulator i silnik pojazdu elektrycznego. Zabezpieczy być może również przed tak groźnym zdarzeniem jakim jest samozapłon i w konsekwencji pożar pojazdu elektrycznego.

Praca układu i jego testowanie

Schemat obwodu zabezpieczenia przeciwzwarceniowego w pojazdach elektrycznych pokazano na **rysunku 2**. Układ wykonano z wykorzystaniem niewielu elementów, w tym:

silnik prądu stałego 12 VDC (M1), 12-to voltowy przełącznik (RL1), buzzer piezoelektryczny (PZ1).

Przystępując do wykonania projektu, należy zgromadzić podzespoły zgodnie z listą „spis elementów”. Następnie zalecamy kierowanie się wskazówkami wg poniższych punktów, które wyjaśnią sposób działania obwodu „protection”.

Krok 1

Po zmontowaniu układu i podłączeniu zasilania, styki przełącznika pozostają w pozycji NC, pozostawiając rozarty styk NO. Po zamknięciu switcha S1 dodatni biegun baterii zostaje podłączony do cewki przełącznika, który pozostaje „nieuzbrojony”. Ujemny biegun akumulatora jest połączony ze wspólnym wyprowadzeniem styku NO/NC. W tym stanie powinna zaświecić dioda LED2 (czerwona) jak również sygnał dźwiękowy buzzera oznacza gotowość obwodu zabezpieczenia baterii do włączenia.

Krok 2

Należy nacisnąć chwilowy przycisk S2, podczas gdy stabilny S1 pozostaje zwarty. S2 zamknie obwód zasilania cewki przełącznika, co spowoduje przełączenie styku z pozycji NC do NO. Styk NO jest włączony równoległe do S2, zatem mimo puszczenia

Wykaz elementów:

Półprzewodniki:

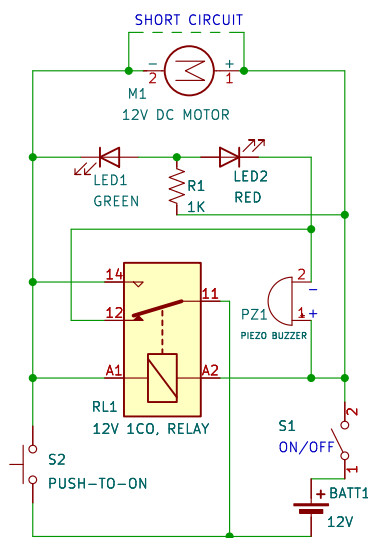
LED1: dioda LED zielona 5 mm
LED2: dioda LED czerwona 5 mm

Rezystory:

R1: 1 kΩ

Inne:

BATT1: bateria 12 VDC 7 Ah lub większy
S1: przełącznik ON/OFF 6 A
S2: przycisk niestabilny
PZ1: buzzer piezoelektryczny
RL1: przełącznik 12 V 1CO
M1: silnik 12 VDC



Rysunek 2. Schemat ideowy obwodu zabezpieczenia przeciwzwarceniowego

przycisku chwilowego, przekaźnik wejdzie w stan samopodtrzymania. Styk NO przekaźnika zamknie równocześnie obwód zasilania silnika. Dodatkowo zaświeci diodę LED1 (zieloną), a dioda czerwona powinna zgasnąć.

Krok 3

Aby sprawdzić działanie układu, należy zasymulować zwarcie na silniku. Zwarcie wyprowadzeń na silniku skutkuje zwarciem akumulatora. Stan taki powinien trwać bardzo krótko, gdyż zniknie równocześnie zasilanie cewki przekaźnika i styk NO powinien zostać natychmiast rozwarty. Stan taki jest stabilny, gdyż zniknie samopodtrzymanie przekaźnika. Ponowne wystartowanie układu wymaga chwilowego zwarcia styków niestabilnego przycisku S2.

Krok 4

Układ wrócił do stanu „spoczynku”, o czym powinna sygnalizować dioda czerwona (LED1) i dźwięk buzzera.

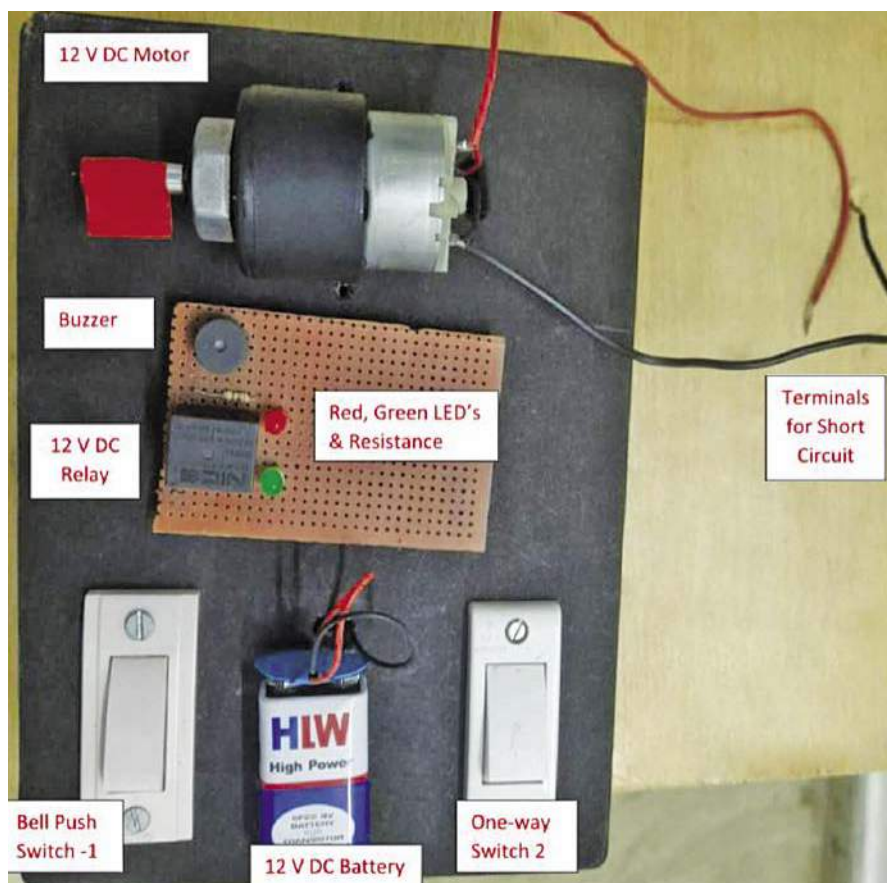
Ten prosty obwód nie rozwiąże problemu zwarcia akumulatora w pojeździe elektrycznym, ale zabezpieczy przez skutkami takiej sytuacji. Natychmiastowe odłączenie zasilania powinno uchronić silnik przed zniszczeniem, a także akumulator przed pełnym rozładowaniem. Uchroni więc nie tylko baterię przed uszkodzeniem, lecz także przed groźną w skutkach eksplozją akumulatora. Obwód „protection” wydłuży żywotność baterii, która jest kluczowym elementem w elektrycznym pojeździe. Mając na uwadze wysoki koszt akumulatora w pojazdach EV, można wysoko ocenić wartość bieżącego projektu, jako że przyczyni się do wydłużenia żywotności pojazdu i ograniczy ryzyko poważnych zdarzeń towarzyszących nowym technologiom elektromobilności. Na **rysunku 3** widzimy prototyp układu wykonany wg schematu z rysunku 2.

Abhishek Sharma

Ten artykuł był wcześniej opublikowany na łamach „EFY”, wrzesień 2023 (efymag.com)

Od Redakcji EdW: Jak widzimy na rysunku 1 i 3, autor przetestował układ na małym silniczku i zapewne jakiejś małej baterii 12 V. Sytuacja w elektrycznym samochodzie może być diametralnie odmienna. Są tu akumulatory o stosunkowo dużym napięciu wielu

REKLAMA



Rysunek 3. Prototyp zmontowany przez autora

szeregowych ogniw i o ogromnym prądzie zwarciovym. Nawet krótkotrwałe sztywne zwarcie takiej baterii może być groźne. Czy w aucie elektrycznym można włączyć cewkę przekaźnika bezpośrednio na akumulator i liczyć na to, że zwarcie w silniku spowoduje obniżenie napięcia do tego stopnia, że przekaźnik puści samopodtrzymanie? Ten projekt zapewne spełni zadanie w jakiejś zabawce. Ponadto autor sugeruje próbę rozłączenia stykiem mechanicznego przekaźnika olbrzymiego prądu zwarciovego przy stałym napięciu z akumulatora. Próba taka zakończy się z pewnością wygenerowaniem olbrzymiego łuku elektrycznego, który być może w ułamkach sekund „zespawa” trwale ów styk. W przypadku chęci rozłączania dużych prądów w obwodach DC należałoby pomyśleć o użyciu przekaźników półprzewodnikowych SSR (Solid-State Relay), które umożliwiają załączanie i rozłączanie dużych prądów bez efektu iskrzenia/łuków elektrycznych. To samo zjawisko bardzo utrudnia zasilanie np. grzałek do grzania wody wprost

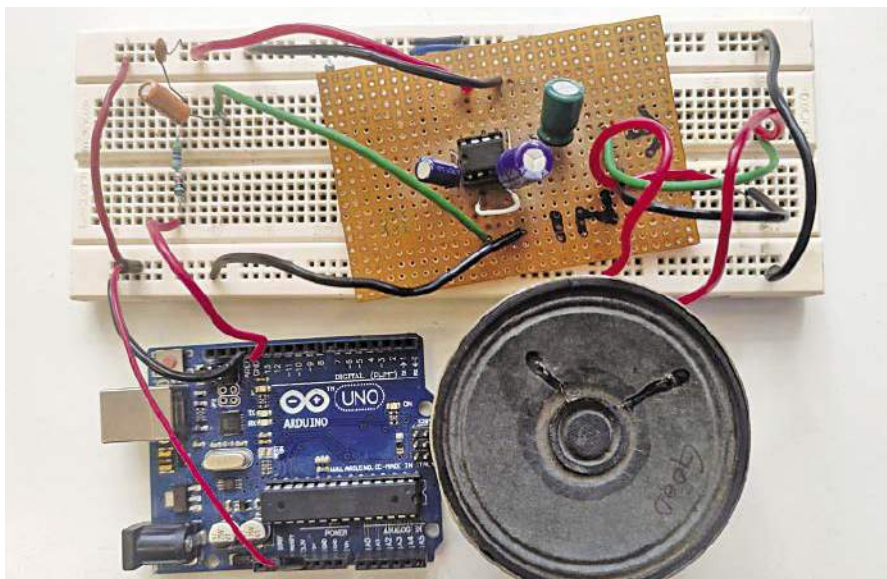
z ogniw fotowoltaicznych (a więc prądem stałym). Pomijając kwestie bezpieczeństwa i bardzo słabej efektywności takiego rozwiązania, główną trudnością będzie tu realizacja termostatu. Styki, prawdopodobnie po kilku przełączeniach spieką się ze sobą na skutek sporych łuków elektrycznych. To jeden z powodów, dla których w prostych systemach do grzania wody, mimo, że nie stosuje się drogich falowników, to jednak wykorzystuje się proste przetwornice z napięcia stałego na przemiennie, a dalej już tradycyjne termostaty pracujące przy napięciu przemiennym. W „prawdziwym Electric Vehicles” można skorzystać z pomysłu autora, lecz układ należałoby zdecydowanie przeprojektować. Prosiłoby się pewnie o realizację zabezpieczenia nadprądowego z prawdziwego zdarzenia (może na jakimś mili-ohmowym rezystorze), a nie liczenie na rezystancję wewnętrzną baterii, która miałaby pozwolić na zwarcie jej wyprowadzeń do tego stopnia, iż napięcie spadnie do zera.

ELPORTAL.pl

Generator tonów relaksujących na Arduino

Nie wszystkie częstotliwości w paśmie akustycznym są jednakowo miłe dla ucha. Już od czasów Starożytnych temat ten nurtował wielu badaczy. W szczególności, tzw. częstotliwości Solfeggio mają bardzo korzystny wpływ dla naszej psychiki i umysłu. W szczególności ton 528 Hz jest bardzo miły i korzystny w odsłuchu. Aczkolwiek do tonów „solfeżowych” zalicza się także częstotliwości 396, 417, 639, 741, 852 i 963 Hz.

Każdy zapewne doświadczył wrażenia jak muzyka może wpływać na nasz nastrój. Kiedy jesteśmy zestresowani i szukamy ukojenia i relaksu, spokojna i wolna muzyka może być bardzo pomocna. Wśród wielu tonów, właśnie ton o częstotliwości 528 Hz wydaje się, iż daje największe ukojenie. Można nawet wyczytać, że to „cudowny ton”, który regeneruje komórki DNA i ma ozdrowieńczy wpływ na naszą duszę i ciało. „Miłosny ton” 528 Hz redukuje stres i napięcie nerwowe. Zwiększa naszą witalną energię, poprawia naszą koncentrację i zdolność skupienia. Mówi się także, iż wpływa korzystnie na trawienie, redukuje uczucie bólu a także ma ozdrowieńczy wpływ na stany zapalne. Ton 528 Hz bywa nazywany „miłosną częstotliwością” ponieważ wydaje się, iż rezonuje z częstotliwością naszego serca i przyczynia się do wewnętrznej duchowej harmonii. Wsłuchiwanie się w ton 528 Hz jest istotną częścią medytacji, która przenosi umysł w stan duchowej jedności z Bogiem. Dr Masaru Emoto eksperymentując z częstotliwościami Solfeggio zauważył ich wpływ na krystalizację wody. Komórki naszego ciała składają się w większości z wody. Zatem nic dziwnego, że słuchanie wybranych częstotliwości może pomóc w aktywacji energii zawartej w naszych komórkach. Wymienione wyżej częstotliwości mają jednak różny wpływ i każda z nich koncentruje inne cechy i cele. I tak np. ton 396 Hz usuwa tzw. energię negatywną i może likwidować przykre uczucie winy. Ton o częstotliwości 417 Hz także usuwa negatywną energię skumulowaną w naszym ciele. 639 Hz posiada



Rysunek 1. Prototyp generatora przetestowany w laboratorium EFY

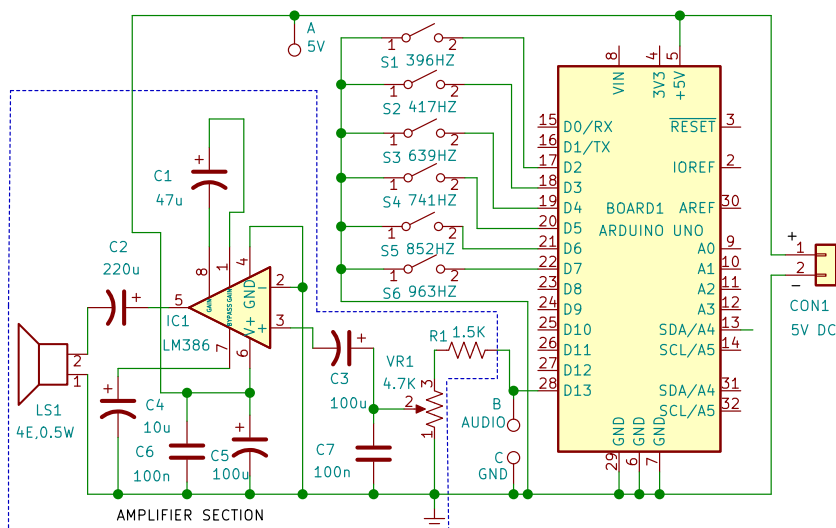
„tajemną moc” poprawy harmonijnych relacji między ludźmi. Mówi się, iż 741 Hz czyści komórki naszego ciała. 852 Hz potrafi wzmocnić energię w komórkach. Ton 943 Hz jest stosowany w terapii Sahasrara Nadi, aktywuje naszą wewnętrzną moc i przyczynia się do pozytywnego spojrzenia na otaczający Świat.

To wszystko teoria, którą trudno zweryfikować. Ale możesz eksperymentować na sobie. Zacznij od słuchania tonu o częstotliwości

528 Hz. Generator taki nie jest trudno wykonać. Zdjęcie na **rysunku 1** pokazuje prototyp generatora tonów Solfeggio, który został zmontowany i przetestowany w laboratorium Redakcji „Electronics For You”.

Budowa układu i jego działanie

Schemat ideowy generatora tych „magicznych częstotliwości” pokazano na **rysunku 2**. Układ wykonano korzystając z platformy Arduino.



Rysunek 2. Schemat ideowy układu

Wykaz elementów:

Półprzewodniki:

Płytką 1: Arduino Uno
IC1: LM386 – wzmacniacz akustyczny małej mocy

Rezystory:

(wszystkie 0,25 W/±5%)
R1: 1,5 kΩ
VR1: 4,7 kΩ – potencjometr

Kondensatory:

C1: 47 μF/16 V elektrolityczny
C2: 220 μF/16 V elektrolityczny
C3, C5: 100 μF/16 V elektrolityczny
C4: 10 μF/16 V elektrolityczny
C6, C7: 100 nF ceramiczny

Inne:

CON1: złącze 2-pinowe
S1-S6: switchy on/off
LS1: głośnik 4 Ω/0,5 W

Dołożyć trzeba tylko niewielki wzmacniacz mocy, może to być moduł z wykorzystaniem układu scalonego LM386 lub np. PAM8403.

Zadane częstotliwości generowane są z wykorzystaniem software-u mikroprocesora na płytce Arduino. Wyjściowy sygnał pozyskany jest z pinu 28 (D13) Arduino. Należy go jedynie wzmocnić, tak aby wysterować niewielki głośniczek. Arduino generuje sygnał prostokątny i wyjście D13 zawiera składową stałą o wartości połowy zasilania. Składowa DC usunięta jest dzięki zastosowaniu kondensatorów C3 i C7 pośredniczących między Arduino i wzmacniaczem akustycznym. Jeśli wykorzystamy gotowe moduły jak LM385 lub PAM8403, o składową stałą nie trzeba się martwić, gdyż kondensatory sprzęgające na tych mini-modułach już przewidziano. Gdy wykorzystamy gotowe moduły (płytki), wtedy dodatkowych elementów dyskretnych jest niewiele. Dobrej jakości powinien być głośnik, aczkolwiek nie jest wymagana duża moc. Autor wykorzystał „mini speaker” WS-887, który jest niedrogi a odtwarza dobrą jakość dźwięku.

Wybór częstotliwości wykonano wykorzystując 6 switchy S1 do S6. Te przełączniki są normalnie w pozycji otwartej. Wybór przydzielonej częstotliwości dokonuje się zwierając do masy jedno z wejść cyfrowych D2 do D7.

Generator na Arduino nie wytwarza selektywnej częstotliwości (fali sinusoidalnej), lecz falę prostokątną, której pierwsza harmoniczna jest zaprogramowaną częstotliwością Solfeggio. Wyższe harmoniczne można by odfiltrować z użyciem prostego filtra dolnoprzepustowego. Nie jest to potrzebne, ani nawet pożądanego. Ton z zawartością nieparzystych harmonicznych jest nawet korzystniejszy i przyjemniejszy dla ucha. I tak np. ton 528 Hz będzie zawierał częstotliwości 1584 Hz i 2640 Hz (3-cia i 5-ta harmoniczna). Aczkolwiek amplituda składowych harmonicznych w fali prostokątnej maleje odwrotnie proporcjonalnie do częstotliwości (tu odpowiednio 1/3 i 1/5 harmonicznej podstawowej 528 Hz). Zatem, w generatorze „magicznych” częstotliwości Solfeggio można wprost wykorzystać przebieg prostokątny bez dodatkowej filtracji. Zasada ta dotyczy wszystkich częstotliwości „solfeżowego” szeregu.

W programie mikrokontrolera na Arduino wykorzystano funkcję opóźnienia odliczaną w jednostkach – mikrosekundach. Na przykład: częstotliwość 528 Hz to okres 1894 mikrosekund (ponieważ $1/528 \times 1000000 = 1893,9$). To jest okres jednego pełnego cyklu. Jeśli chcemy wygenerować falę prostokątną z wypełnieniem 50%, to zarówno stan wysoki jak i niski powinien trwać połowę wyżej obliczonej wartości (w mikrosekundach). Zatem przez odcinek

czasu 947 μ s na wyjściu D13 powinien być stan bliski zasilaniu i tyle samo napięcie bliskie 0 V.

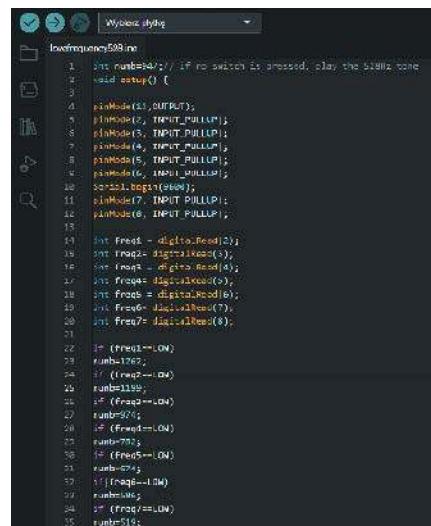
Program jest napisany tak, że dla wygenerowania „miłosnej częstotliwości” 528 Hz, wszystkie wejścia D2 do D7 powinny pozostać w stanie wysokim, czyli wszystkie switchy S1 do S6 powinny pozostać otwarte. Dla wyboru innych częstotliwości z szeregu Solfeggio należy zewrzeć tylko jeden ze switchy S1 do S6. Jednak mikroprocesor czyta stan wejść D2 do D7 tylko po włączeniu zasilania. Zatem procedura musi być taka, iż najpierw należy zewrzeć switch i następnie włączyć zasilanie. Odpowiedni ton o zadanej częstotliwości jest zawsze wygenerowany na pinie nr 28 (wyjście cyfrowe D13). Przyjęta logika programu nie pozwala na zmianę częstotliwości kiedy układ jest zasilany i pracuje. Trudno to uznać za wadę czy specjalne utrudnienie obsługi. Częstotliwości nie należy zmieniać zbyt często. Aby terapia była skuteczna, wybranemu tonowi należy się wsłuchiwać przez dłuższy czas. Zatem częste zmiany częstotliwości z szeregu Solfeggio są co najmniej nie zalecane.

Oprogramowanie

Program dla Arduino napisano z wykorzystaniem Arduino IDE. Gotowy „szkielet” można przegrać z komputera do płytki Arduino korzystając z tego samego software-u IDE. Po wykonaniu tych czynności, „Arduino board” można odłączyć od komputera i podłączyć do zewnętrznego zasilacza 5 VDC. Można także wykorzystać zasilanie bateryjne 9 V podłączając je do oddzielnego złącza, w które moduły Arduino są wyposażone. Na **rysunku 3** pokazano zrzut ekranu kodu źródłowego, który zarządza pracą mikrokontrolera w roli generatora częstotliwości Solfeggio.

Konstrukcja i testowanie pracy generatora

W pierwszej kolejności należy załadować kod źródłowy o nazwie lovefrequency528.ino. W tym celu pod programem Arduino IDE trzeba wybrać właściwą nazwę płytki oraz ustawić port, pod którym Arduino Uno jest widziany. Teraz odłączyć płytkę od komputera i podłączyć zasilanie 5 V do złącza CON1. Teraz należy zamontować całość układu zgodnie ze schematem ideowym na **rysunku 2**. Możesz użyć gotowej płytki wzmacniacza audio z np. układem scalonym LM386. Na **rysunku 2** zawartość tej płytki wyodrębniona jest linią kropkowaną. Można też zamontować wzmacniacz z elementów dyskretnych zgodnie ze schematem. Dla montażu można użyć płytki uniwersalnej. Jeśli dla wzmacniacza akustycznego użyjesz odrębnego zasilania, wtedy z modułem Arduino łączysz go jedynie w dwóch punktach: B i C. Jeśli chcesz wykorzystać zasilanie 5-cio voltowe



Rysunek 3. Zrzut ekranu kodu źródłowego

z Arduino, wystarczy że dodatkowo połączysz punkty na schemacie oznaczone A.

Po zamontowaniu elementów na PCB, umieść układ w odpowiednio przygotowanej obudowie. Prototyp wykonany w celu przetestowania w laboratorium redakcji zamontowano na płytce uniwersalnej, co widać na **rysunku 1**. Podłączono także miniaturowy głośnik typu WS-887.

Teraz można testować działanie układu, aczkolwiek ocena terapeutyczna jest indywidualna i bardzo subiektywna. Aby to ocenić, należy się wsłuchiwać w wybrany ton przez co najmniej 30 minut.

Uwaga Redakcji EFY

Jeśli dysponujesz miernikiem częstotliwości, można sprawdzić na ile wygenerowana częstotliwość jest bliska założeniom. W razie potrzeby należy się dostroić do 528 Hz na ile w granicach błędu jest to możliwe.

W laboratorium EFY przetestowano poprawność pracy układu z technicznego punktu widzenia, i zgodność częstotliwości z założonym szeregiem Solfeggio. Nie sprawdzono jednak wartości terapeutycznej urządzenia, czyli jaki jest faktyczny jego wpływ na ludzki organizm. ■

K. Padmanabhan

Ten artykuł był wcześniej opublikowany na łamach „EFY”, maj 2020 (efymag.com)

Od Redakcji EdW: Zachęcamy Czytelników do eksperymentowania z elektroniką, a także do sprawdzenia czy słuchanie wybranych tonów faktycznie „czyni cuda”. Niewątpliwie muzykoterapia jest już uznaną metodą terapeutyczną współczesnej medycyny. O dużym znaczeniu częstotliwości Solfeggio w muzykoterapii świadczy mnogość publikacji w Internecie na ten temat. Polecamy <https://mindeasy.com/solfeggio-frequency-science/>

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**.

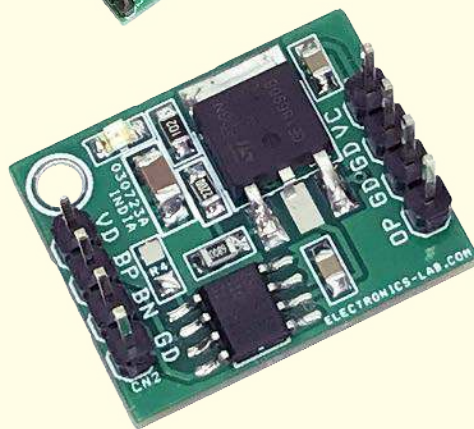
Sterownik silnika krokowego z joystickiem

To generator impulsów dla sterownika silnika krokowego wykorzystującego joystick. Jest to płytka zawierająca mikrokontroler kompatybilny z Arduino i obwody do sterowania maksymalnie 2-kanalowymi (2-osiowymi) silnikami krokowymi. Płytka ma wiele opcji do opracowania systemów sterowania związanych z silnikami krokowymi, od 2-osiowego joysticka analogowego do potencjometru trymera do 7-kanalowego wyjścia TTL z otwartym kolektorem do interfejsu sterowników krokowych z wejściami transoptorowymi itp.



Programowalny kondycjoner sygnału z czujnika rezystancyjnego mostkowego

Prezentowany moduł to kondycjoner sygnału czujnika oparty na układzie scalonym ZSC31010 CMOS, który umożliwia łatwą i precyzyjną kalibrację rezystancyjnych czujników mostkowych poprzez programowanie EEPROM. Po połączeniu z rezystancyjnym czujnikiem mostkowym, będzie on cyfrowo kalibrował offset i wzmacnienie z opcją kalibracji współczynników offsetu i wzmacnienia oraz liniowości w zależności od temperatury. Kompensacja drugiego rzędu może być włączona dla współczynników temperaturowych wzmacnienia, offsetu lub liniowości mostka. Moduł ZSC31010 komunikuje się za pośrednictwem interfejsu szeregowego ZACwire firmy IDT z komputerem hosta i można go łatwo skalibrować masowo w środowisku Windows. Po skalibrowaniu, pin sygnału wyjściowego może zapewnić wybierane absolutne wyjście analogowe 0 do 1 V; ratiometryczne wyjście analogowe rail-to-rail; lub cyfrowe wyjście szeregowo danych mostka z opcjonalnymi danymi temperatury. Układ U1 LM317 dostarcza napięcie 5 V DC do układu czujnika. D1 to dioda LED zasilania.



Niektóre projekty aktualnie dostępne tylko dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl:

- | | | |
|--|---|---|
| <ol style="list-style-type: none"> 1. 20-segmentowy wyświetlacz słupkowy w rozmiarze jumbo 2. Stacja pogodowa lilygo ttgo t5-4.7 z wyświetlaczem typu e-papier 3. Półprzewodnikowy przełącznik mocy DC z prądowym sprzężeniem zwrotnym 4. Wyłącznik nadprądowy – przełącznik wyłączający nadprądowy 5. Uniwersalny konwerter napięcia AC – wyjście 18 V DC z wejścia 85...265 V AC 6. Moduł procesora echa głosu – urządzenie opóźniające do efektów dźwiękowych, echo, reverb 7. Najlepszy sposób na próbkowanie dźwięku za pomocą ESP32 8. Choinka z Arduino i pikselowymi diodami | <ol style="list-style-type: none"> 9. RPi – stacja pogodowa IoT 10. Niskobudżetowy monitor jakości powietrza IoT oparty o RaspberryPi 4 11. Automatyczny system ogrodniczy z NodeMCU i Blynk, ArduFarmBot 2 12. TinyML – Rozpoznawanie ruchu przy pomocy Raspberry Pi Pico 13. Wzmacniacz piezoelektryczny do gitary i skrzypiec 14. Wysokowydajny i niezawodny sterownik bipolarnego silnika krokowego 15. Sterownik silnika prądu stałego z wykorzystaniem przełącznika i mosfetu – interfejs Arduino 16. Przedwzmacniacz do mikrofonu MEMS 17. Super prosty czuły wykrywacz metali 18. Stymulator czaszkowy Arduino (Bio-BrainTuner) | <ol style="list-style-type: none"> 19. Generator sygnałów AD9833 20. Obserwacja charakterystyk tranzystora 21. Wyświetlacz EKG z użyciem Arduino 22. Łatwy do zbudowania robot krocący 23. Sonarowy theremin MIDI 24. Zamek elektroniczny na kod 25. Prosty tester tranzystorów 26. Zegar binarny z użyciem Microbit 27. Przetwornik częstotliwości na napięcie (tachometr) – przetwornik częstotliwości na napięcie z czujnikiem magnetycznym o zmiennej reluktancji 28. Izolowany obwód wykrywania napięcia 250 V AC z pojedynczym wyjściem (wejście 250 V prądu przemiennego, wyjście 5 V) |
|--|---|---|

Miesięcznik „Elektronika dla Wszystkich” (12 numerów w roku) jest wydawany we współpracy z kilkoma redakcjami zagranicznymi



Wydawnictwo:
AVT-Korporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Adres redakcji:
03-197 Warszawa, ul. Leszczyńska 11
e-mail: edw@elportal.pl, www.elportal.pl

Redaktor merytoryczny:
Mariusz Ciszewski, Paweł Sujko

Dział Reklamy:
Katarzyna Gugała
katarzyna.gugala@elportal.pl, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobański
jakub.sobanski@elportal.pl

Copyright AVT-Korporacja Sp. z o.o., Warszawa, ul. Leszczyńska 11. Projekty publikowane w „Elektronice dla Wszystkich” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki dla Wszystkich”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice dla Wszystkich” jest dozwolone wyłącznie po uzyskaniu pisemnej zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice dla Wszystkich”.

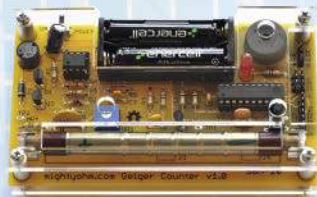
DTP, okładka, redakcja strony internetowej www.elportal.pl:
MAD Sp. z o.o.

Prenumerata:
W Wydawnictwie AVT, e-mail: prenumeratora@avt.pl
tel. 22 257 84 22, (godz. 10:00–14:00)

W RUCH S.A., e-mail: prenumeratora@ruch.com.pl
tel. 801 800 803, 22 717 59 59, www.prenumerata.ruch.com.pl

Elektor Bestsellers

SAVE UP TO
26% NOW!



www.elektor.com/sale/deals



eprasa.pl 9f4bfc67bd