

nie przeocz codziennych nowości na Elportal.pl

ELEKTRONIKA

dla wszystkich

nr 6/2022 (317) • czerwiec • www.elportal.pl

Czy potrafisz opanować THEREMIN?

DIY PLUS
tylko dla prenumeratorów

PROJEKTY dla elektroników

- ▶ Alarm do sejfów hotelowych dla podróżnych
- ▶ Frezarka CNC, część 6
- ▶ Wzmacniacz mocy Ultra-LD 200 W RMS, część 2

DIY dla wszystkich

- ▶ Sterownik pulsującej taśmy LED
- ▶ Czterokanałowa bezdotykowa tablica rozdzielcza dla świata po Covidzie
- ▶ Automatyczny dozownik wody dla zwierząt domowych
- ▶ Zapisywanie danych przy użyciu microSD i Raspberry Pi Pico
- ▶ Arduino programowane „ręcznie”
- ▶ IoT na oku: sterowanie urządzeniami domowymi

TUTORIALE

- ▶ Szkoła Konstruktorów
- ▶ Jak to działa
- ▶ Pomiar pH wody i gleby, część 1
- ▶ Projektowanie mini monitorów Hi-Fi, część 1
- ▶ Elektrozwory nie tylko do nawadniania, część 2
- ▶ Współczesne akumulatory – 5. Akumulatory litowe – historia
- ▶ Zasilanie do twojego projektu, część 3. Stabilizatory liniowe
- ▶ Projektowanie układów porównujących



**Miejsca dla
specjalistów**

PORTAL BRANŻOWY
ElektronikaB2B
PORTAL BRANŻOWY
AutomatykaB2B

EP.com.pl

Największy
portal dla
elektroników
konstruktorów

FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

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

www.piekarz.pl



16,90 zł (w tym 8% VAT)



epasa.pl 0d581f01a1

NOWOŚĆ NA POLSKIM RYNKU



Czy długie i szczęśliwe życie to przywilej nielicznych?

Poradnik „Żyj lepiej, żyj dłużej” to światowy bestseller, który pokazuje, że potrzeba niewiele, aby zadbać o swoje zdrowie.

Na ponad 120 stronach przedstawiamy aktualne wiadomości na temat profilaktyki oraz leczenia najczęstszych schorzeń.

To prawdziwe kompendium praktycznej i sprawdzonej wiedzy.

Przejrzyj i zamów z bezpłatną przesyłką
www.UlubionyKiosk.pl



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

**Zaprenumeruj
„Elektronikę
dla Wszystkich”,
a zawsze dostaniesz
najnowszy numer wprost
do Twojej skrzynki!**

**na start
do 6* wydań gratis**

**po 5 latach
nieprzerwanej
prenumeraty
do 12* wydań gratis**

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

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

Oferta dotyczy prenumeraty drukowanej.

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

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

prenumerata@avt.pl

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

eprasa.pl 0d581f01a1



8 Projekty dla elektroników:

Czy potrafisz opanować THEREMIN?.....	8
Alarm do sejfów hotelowych dla podróżnych	18
Frezarka CNC, część 6.....	24
Wzmacniacz mocy Ultra-LD 200 W RMS, część 2	26

Tutoriale:

Szkoła Konstruktorów	36
Jak to działa	47
Pomiary pH wody i gleby, część 1.....	51
Projektowanie mini monitorów Hi-Fi, część 1.....	54
Elektrozawory nie tylko do nawadniania, część 2	64
Współczesne akumulatory – 5. Akumulatory litowe – historia.....	66
Zasilanie do twojego projektu, część 3. Stabilizatory liniowe	68
Projektowanie układów porównujących	76

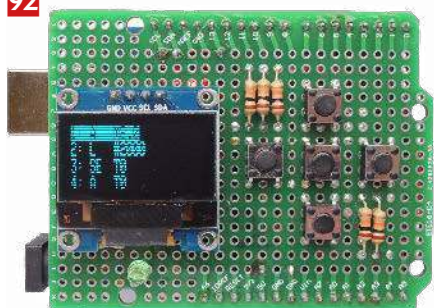
18



DIY dla wszystkich:

Sterownik pulsującej taśmy LED	81
Czterokanałowa bezdotykowa tablica rozdzielcza dla świata po Covidzie	84
Automatyczny dozownik wody dla zwierząt domowych	86
Zapisywanie danych przy użyciu microSD i Raspberry Pi Pico.....	88
Arduino programowane „ręcznie”	92
IoT na oku: sterowanie urządzeniami domowymi.....	94

92



DIY PLUS

98

Rubryki stałe:

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

A za miesiąc w lipcowym EdW



✱ **Generator sygnałów RF skanujący AM/FM/CW**
Marzenie każdego hobbysty – mieć w swojej pracowni elektronicznej generator sygnałów RF z możliwością wykorzystania do 150 MHz. Ta konstrukcja jest wynikiem długiego okresu rozwoju i testowania w ciągu ostatnich kilku lat.

✱ **Wzmacniacz 200 W**
Ten wzmacniacz o doskonałych parametrach wzbudził już ogromne zainteresowanie. W części trzeciej zajmujemy się budową zasilacza i prezentujemy tańszą wersję wzmacniacza o niższej mocy – 110 W przy obciążeniu 4-ohmowym.

✱ **Wabik na samce komarów**
Trochę to nie fair – samce komarów nie gryzą, a poniosą konsekwencje za paskudnie kłuszące samice. Dźwiękiem 484 Hz zwabimy samce komarów do pułapki, zanim wezmą udział w procesie rozmnażania. W ten sposób zmniejszymy populację komarów na naszym terenie.

✱ **Plus 6 intrygujących projektów DIY**

✱ **Plus 8 artykułów w stałych cyklach Tutoriali**

W kioskach od 28 czerwca

Theremin

Historia zna wiele odkryć przypadkowych, jednak zdarzają się one osobom nieprzypadkowym. Jabłko nie raz spadło komuś na głowę, ale trzeba było Izaaka Newtona, żeby z tego zdarzenia zrodziło się prawo powszechnego ciążenia.

Projektem miesiąca w tym wydaniu EdW jest **elektroniczny instrument muzyczny**, którym świat zachwycił się już 100 lat. To **theremin**, wynaleziony w październiku 1919 roku przez 23-letniego fizyka rosyjskiego Lwa Termena, znanego później pod nazwiskiem Leon Theremin (1896–1993). Nie było przypadkiem, że 17-letni Lew Termen w domowym laboratorium eksperymentował z falami radiowymi, zbudował cewkę Tesli, generującą milion woltów, a w 1919 roku prowadził już poważne badania dielektryków w petersburskim instytucie naukowym, znanym obecnie jako Instytut Joffe. Budując aparaturę do pomiaru stałej dielektrycznej skonstruował **oscylator przebiegów wysokiej częstotliwości z antenką** i zaczęły się przypadkowe odkrycia.

Najpierw zauważył, że jego układ może służyć jako **detektor ruchu** (patent „radio watchman”), a przy próbie wykorzystania dźwięku jako wskaźnika dostrojenia usłyszał niesamowicie brzmiące dźwięki, modulowane ruchami dłoni wokół antenki. I znów nie było przypadkiem, że Lew Termen był nie tylko fizykiem, ale miał też wykształcenie muzyczne i w odgłosach kojarzących się laikowi z wyciem, piskiem, skowytom zwierzęcia, itp., usłyszał brzmienia wiolonczeli i od razu potrafił „zagrać” konkretne melodie. Otaczający go profesorowie, współpracownicy i studenci byli zachwyceni. Już w roku 1920 Lew Termen dał pierwsze koncerty publiczne na wynalezionym instrumencie, nazywanym wówczas Termenvox.

Na tym powinienem zakończyć to wprowadzenie do projektu flagowego w tym wydaniu EdW, ale późniejsze życie i dokonania wynalazcy są tak fascynujące, że warto poświęcić tej biografii jeszcze kilka zdań. W roku 1927 Lew Termen wyruszył ze swoim wynalazkiem na podbój świata, przez Londyn, Paryż, Niemcy do Ameryki. W USA według jego patentu firma RCA uruchomiła seryjną produkcję thereminów, które wynalazca ciągle udoskonalał. Już pod nazwiskiem Leon Theremin (było to nazwisko rodowe jego ojca) w latach trzydziestych osiągnął w USA sukcesy biznesowe i towarzyskie. Stał się osobą bardzo popularną, a muzyka elektroniczna cieszyła się wielkim zainteresowaniem. Organizowano konkursy wirtuozów theremina, w których triumfy święciła wielka miłość wynalazcy – emigrantka z Rosji Clara Reisenberg. Nie ma tu miejsca na opis burzliwego życia osobistego Theremina – trzykrotnie żonaty, w tym w USA z tancerką afroamerykańską, ale Clara nie przyjęła jego wielokrotnych oświadczeń. W roku 1938, po jedenastu latach pobytu w USA nagle zniknął. Pogłoski o jego śmierci po wielu latach okazały się nieco przesadzone. W roku 1939 znalazł się w łagrze gułagu z bagatelny wyrokiem 8 lat. Było mu tam bardzo dobrze – w świetnie wyposażonym laboratorium NKWD mógł się całkowicie oddać pracy badawczo-rozwojowej, w grupie z takimi więźniami jak Korolow i Tupolew (w systemie gułag działały biura konstrukcyjno-rozwojowe, tzw. szarazki, w których pracowali wybitni konstruktorzy – więźniowie). W tym okresie opracował słynne **urządzenie do podsłuchu** zastosowane w ambasadzie amerykańskiej. Amerykanie przez 7 lat (1945–1952) nie wiedzieli, że rozmowy ambasadora są podsłuchiwane dzięki podarkowi od pionierów – płaskorzeźbie wielkiej pieczęci USA, powieszzonej przez ambasadora na ścianie w jego gabinecie. Wewnątrz tej pięknej pamiątki sojuszu wojennego USA – ZSRR znajdował się swojego rodzaju mikrofon pojemnościowy z drgającą lekką membraną jako okładką kondensatora. Rosjanie „naświetlali” wiązką 330 MHz mikrofon membranowy i odbierali sygnał odbity zmodulowany dźwiękiem wprawiającym w drgania membranę. Był to podsłuch trudny do wykrycia, bo działał pasywnie, bez zasilania i jakichkolwiek elementów elektronicznych. Współcześnie ta metoda podsłuchu jest stosowana z użyciem promieni lasera. Za ten wynalazek Lew Termen otrzymał w 1947 roku Nagrodę Stalinowską. Był też autorem wielu innych wynalazków – m.in. w roku 1927 zademonstrował pierwszy w ZSRR **telewizor**, działający podobnie jak modne dziś projektory (obraz miał ogromne rozmiary 1,5 m × 1,5 m), a podczas pobytu w USA skonstruował **wykrywacz metali** dla więzienia Alcatraz. Po wojnie wiele lat pracował nad instrumentami muzyki elektronicznej, uczył gry na thereminie w Konserwatorium Moskiewskim i był profesorem fizyki w Moskiewskim Uniwersytecie Państwowym. Żył 97 lat. Po rozpadzie ZSRR podróżował do USA, gdzie spotkał się po wielu latach z Clarą. Nakręcony o nim tuż przed śmiercią film dokumentalny *An Electronic Odyssey* nie wyjaśnia wszystkich zagadek jego życia. Może więcej wie jego prawnuk Piotr Theremin, 30-letni wirtuoz gry na thereminie.

Wykonanie własnego instrumentu theremin to wyzwanie podejmowane przez wielu elektroników na świecie. Obecna technologia pozwala zrealizować ten projekt na wiele różnych sposobów. Dlatego do tematu theremin wrócimy na łamach EdW jeszcze nie raz.

Wiesław Marciniak

W rubryce „Począta” zamieszczamy fragmenty Waszych listów oraz nasze odpowiedzi i komentarze. Prosimy o listy dotyczące bieżących wydań EdW, a także o listy z Waszymi komentarzami, propozycjami, problemami, pytaniami, oczekiwaniami względem nas, z propozycjami tematów do opracowania, itp. Piszcie do nas, bardzo cenimy Wasze listy, choć nie wszystkie prośby możemy zrealizować. Tym razem rubrykę począta zdominował temat rozwiązania zadania postawionego przez Czytelnika w liście opublikowanym w EdW 04. Spośród wielu nadesłanych rozwiązań zadania wybraliśmy do publikacji trzy listy – najdłuższy, średnio długi i najkrótszy.

Przypominamy zadanie

Wyobraźmy sobie obwód magnetyczny (jak na rysunku 1), przez który przenika zmienny strumień magnetyczny Φ , taki że:

$$\frac{E}{z} = \frac{d\Phi}{dt} = \frac{3V}{z}$$

przyjmijmy że na obwodzie zamontowano uzwojenie w postaci jednego zwoju zwartego. Zwój ten, jednak złożony jest z dwóch różnych materiałów, o różnych opornościach R_1 i R_2 . Dla uproszczenia przyjmijmy, że: $R_1=1[\text{Ohm}]$ i $R_2=2[\text{Ohm}]$.

Ponieważ $z=1$, to $E=3\text{ V}$.

W miejscach połączenia materiałów zwoju zwartego dołączono dwa woltomierze V_1 i V_2 , jak pokazano na rysunku 1.

Wyindukowana SEM wyniesie 3 V , a oporność całkowita pierścienia wynosi 3 Ohm ,

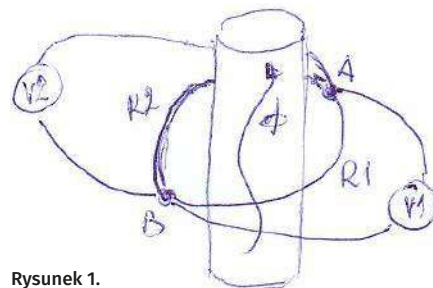
a więc wzbudzona wartość natężenia prądu wyniesie 1 A .

Pytania:

1. Co pokażą oba woltomierze?
2. Czy wskazania będą jednakowe, czy nie? Jeżeli tak, to dlaczego, a jeżeli nie, również dlaczego? (...)

Odpowiedzi:

1. Woltomierze będą miały różne wskazania, V_1 wskaże 1 V , a V_2 wskaże 2 V .
2. W przypadku silnych pól rozproszonych, konieczne jest uwzględnienie wpływu składowej Lorenza. Elektrycy silnoprądowi – rutyniarze, zapominają o tym zjawisku i może okazać się,



Rysunek 1.

że mierzymy nie to co chcemy. Obecnie mamy do czynienia z ładowarkami bezprzewodowymi, czy też transponderami przekazującymi zasilanie bezprzewodowo, a poza tym mamy na co dzień do czynienia z częstotliwościami GHz! Uwaga na pomiary!

Andrzej Lipczyk

Komentarz do listu Andrzeja Lipczyka z EdW 2022/04

Odpowiedź jest prawidłowa, mimo to rozwiązanie należy uznać za błędne.

Trzeba być bardzo ostrożnym ze stosowaniem drugiego prawa Kirchhoffa w obecności zmiennego pola magnetycznego. Pętla obwodu elektrycznego staje się źródłem SEM i prawo to nie obowiązuje. Podstawą staje się prawo Faradaya, a w rozwiązaniu autora, o nim ani słowa. Prawo to będące też treścią trzeciego równania Maxwella ma postać: $\text{rot } E = -dB/dt$. To najbardziej poprawna postać różniczkowa, czyli sprowadzona do punktu przestrzeni. Ze stosowaniem postaci całkowitej też trzeba być ostrożnym, ale wtedy gdy w grę wchodzi opóźnienia wynikające z przestrzennych wymiarów obwodu elektrycznego i czasu propagacji fali elektromagnetycznej w tym obwodzie. Taka sytuacja ma miejsce przy dużych częstotliwościach, gdy obwodu elektrycznego nie możemy traktować jako ostatek skupionych.

Tutaj jednak możemy spokojnie zapisać prawo Faradaya w postaci, którą wypowiem słownie: całka okrężna wektora natężenia pola elektrycznego E po zamkniętym konturze, jest równa pochodnej czasowej strumienia indukcji magnetycznej przenikającej ten kontur. Znak minus wynika z kierunku indukowanego napięcia będącego sumą (całką) wektora E po fragmentach przewodu (choć to samo obowiązuje w wolnej przestrzeni). Nie można zapomnieć, że nie tylko E i B są wektorami, ale także jako wektory trzeba traktować liniowe odcinki (fragmenty) przewodu, a także fragmenty powierzchni rozpiętej na konturze z którym mamy tu do czynienia. Inaczej mówiąc, całka okrężna wektora E (to napięcie) jest równa całce powierzchniowej

z pochodnej czasowej wektora B (to strumień). Ta zależność wraz z prawem Ampera uzupełnionym przez Maxwella jednoznacznie wiąże pole elektryczne i magnetyczne. A ponieważ ich postrzeganie może zależeć też od obserwatora (jego ruchu) określenie pola elektrycznego i magnetycznego może nie być jednoznaczne, dlatego pole to należy nazywać elektromagnetycznym. Prawo Faradaya obowiązuje dla każdej chwili czasu, i dlatego tylko w postaci różniczkowej (zerowy rozmiar obwodu) jest poprawne zawsze i w stu procentach. W tym zapisie rolę całki okrężnej przejęła rotacja.

Może wybiegłem zbyt daleko jak na potrzeby naszego zadania. Ale prawo Faradaya jest tu podstawą, i nie odwołanie się do niego w przedstawionym przez autora rozwiązaniu należy uznać... niestety – za niewybaczalne. Wypowiedziane (no bo nie zapisane) wyżej równanie może wyglądać na skomplikowane. Ale można wyrazić je o wiele prościej: zmiana pola magnetycznego indukuje pole elektryczne, a ilościowo to: napięcie $U = d\Phi/dt$.

Trzeba być ostrożnym w twierdzeniu – co mierzy który woltomierz nie zapominając, że on też wyposażony jest w przewody. W nich też indukuje się napięcie (będące wynikiem zmiennego strumienia pola magnetycznego), a więc wskazanie woltomierza i pomiar woltomierzem to nie to samo! Czy zatem wskazanie woltomierzy jest tylko zależne od rezystancji R_1 i R_2 , czy też od położenia w przestrzeni punktów (A i B) łączących te rezystancje? Takie pytanie powinno paść, mimo że odpowiedź jest negatywna (tj. – nie zależy).

Czy, jeśli punkty te będą bardzo blisko siebie, jest sytuacją równoważną z mniej więcej symetrycznym rozłożeniem tych punktów? Sedno odpowiedzi jest w tym,

że w przewodach woltomierza V_1 indukuje się dokładnie to samo napięcie co w odcinku pętli o rezystancji R_1 . Te napięcia się znoszą, i dla wskazania woltomierza V_1 pozostaje jedynie wpływ prądu „wpychanego” do odcinka pętli o rezystancji R_1 przez SEM w całej zamkniętej pętli. Ten składnik to faktycznie $1\text{ A} \times 1\text{ ohm}$ czyli 1 V . Takie będzie wskazanie V_1 , i należy jeszcze raz zdecydowanie rozróżnić określenie „wskazanie” od – pomiar. To samo obowiązuje dla woltomierza V_2 , który wskaże napięcie 2 V .

Wynik jest zgodny z rozwiązaniem podanym przez autora zadania. Jednak wobec braku poprawnej interpretacji, rozwiązanie autora należy uznać za błędne.

Na koniec zauważmy, że dla każdej zamkniętej pętli, przez którą nie przenika (zmienny!) strumień pola magnetycznego, możemy stosować drugie prawo Kirchhoffa. Takimi pętlami są: część obwodu o rezystancji R_1 i przewody woltomierza V_1 . I podobnie „po drugiej stronie”: R_2 i V_2 (przewody V_2). Dla poprawnej interpretacji postawionego tu zadania jest jednak istotne spostrzeżenie, że „nie indukuje się w obwodzie A- R_1 -B- V_1 i A- R_2 -B- V_2 żadne napięcie”. Ale, że indukowane napięcia w poszczególnych fragmentach obwodu (każdym „milimetrze” przewodu) sumarycznie znoszą się do zera.

Autor założył także milcząco, że rezystancje obu woltomierzy są nieskończenie duże, czyli że prąd w obwodach woltomierzy nie płynie. To założenie konieczne, aby odpowiedź (nie rozwiązanie) uznać za poprawną. Ale, jakie będzie wskazanie V_1 i V_2 jeśli rezystancje tych woltomierzy wyniosą odpowiednio R_3 i R_4 ? To pytanie kieruję do autora listu-zadania i do dociekliwych Czytelników.

Problem jest czysto teoretyczny, ale może być także sensowny ze względu

na błąd pomiaru w podobnych sytuacjach. A można też dodać, że skoro zagadnienie było postawione w latach 60. ubiegłego wieku, to i woltomierze były kiepskie, analogowe o niezbyt dużej rezystancji wewnętrznej (w tamtych czasach standardem było 20 kΩ/V).

Wskazówka:

W punktach A i B spełnione jest pierwsze prawo Kirchhoffa. W obwodach A-R1-B-V1 i A-R2-B-V2 także drugie (prawo Kirchhoffa). A każda pętla obejmująca Φ jest źródłem SEM.

Po trochę żmudnych, ale prostych przekształceniach dostaniemy odpowiedź:

$$V_1 = U_0 \times R_1 \times R_3 \times (R_2 + R_4) / ((R_2 \times R_4 \times (R_1 + R_3) + R_1 \times R_3 \times (R_2 + R_4)))$$

i symetryczny wzór dla V2:

$$V_2 = U_0 \times R_2 \times R_4 \times (R_1 + R_3) / ((R_1 \times R_3 \times (R_2 + R_4) + R_2 \times R_4 \times (R_1 + R_3)))$$

gdzie U_0 to napięcie $d\Phi/dt$, R_1 i R_2 jak w pierwotnym zadaniu, a R_3 i R_4 rezystancje mierników V_1 i V_2 .

Te wzory oczywiście redukują się do: $U_0 \times R_1 / (R_1 + R_2)$ i $U_0 \times R_2 / (R_1 + R_2)$ gdy rezystancja wewnętrzna woltomierzy jest zaniedbywalna, i co jest zgodne z zadaniem postawionym przez Pana Andrzeja Lipczyka.

Jeśli Ktoś lubi taką zabawę, można kontynuować. Prawdopodobnie we wnętrzu pierwotnej pętli jest jakiś rdzeń transformatora. A jak zmieniać się pomiary (wskazania) jeśli uwzględnimy też pole rozproszone poza tym rdzeniem? Czy wskazania V_1 i V_2 wzrosną, zmaleją, a może nie ulegną zmianie?

Karol Świerc

Po liście Andrzeja Lipczyka w EdW04 wystąpiła redakcyjna zachęta do skomentowania tego listu. Zatem uwzględniam tę zachętę i pragnę zauważyć, że jest niemożliwe, aby para woltomierzy przyłączonych do pary tych samych punktów pokazywała różne napięcia. Oczywiście możliwe jest, aby na tej samej parze przewodów panowały różne napięcia. Tak np. było w przypadku sieci telefonów stacjonarnych CB (czyli „Centralnej Baterii”). Na parze przewodów telefonicznych istniała stała składowa o napięciu 60 V w stanie spoczynkowym (po podniesieniu „słuchawki”, czyli tzw. mikrotelefonu, napięcie to spadało do około 20 V wskutek zamknięcia obwodu mikrofonu węglowego).

Istniała na tej samej parze przewodów składowa zmienna o częstotliwości 400 Hz (tzw. „sygnał zewowy” centrali), odpowiedzialna za charakterystyczne „buczenie” po podniesieniu słuchawki, czyli mikrotelefonu. W przypadku dzwonienia do kogoś na tej samej parze przewodów pojawiało się napięcie zmienne 110 V–120 V o częstotliwości 25 Hz (prąd o tych parametrach uruchamiał dzwonek polaryzowany w telefonie odbiorcy). Zatem woltomierze podłączone do wspólnej pary przewodów telefonicznych mogły pokazać różne napięcia pod warunkiem, że były to różne woltomierze (np. woltomierz prądu stałego i jakiś woltomierz „selektywny” mierzący napięcie w określonym i zadanym przedziale częstotliwości). Jednak w tym liście nie mamy do czynienia z tego rodzaju sytuacją. W połączonych ze sobą metalowych „półpięścieniach” pojawia się wyłącznie siła elektromotoryczna indukcji; zatem oba woltomierze powinny pokazać to samo. Błąd autora tego listu polegał na tym, że założył, że przez cały pierścień (złożony z dwóch różnych półpięści metalicznych) płynie prąd o tym samym natężeniu (co widać w przedłożonych przez niego rachunkach). Moim zdaniem ten układ półpięści można potraktować jak dwa różne oporniki połączone ze sobą równolegle. Na zaciskach obu rezystorów będzie panowało to samo napięcie, ale popłyną przez nie prądy o różnych natężeniach. (...)

Jacek Konieczny

Jestem zdziwiony treścią listu Pana Andrzeja Lipczyka w skrzynce POCZTA nr 4/2022r.

To taka trochę uwspółcześniona wersja paradoksów Zenona z Elei.

Podstawowy błąd, to mylenie SEM z napięciem na końcach przewodu=baterii.

Drugi błąd: kierunki SEM w gałęziach R_1 i R_2 skierowane są przeciwnie (dla R_1 od A do B, dla R_2 od B do A)

Trzeci błąd, rezystancje dobrych woltomierzy to 10 MΩ na zakres i prądy płynące w gałęziach woltomierzy nie mają praktycznej żadnego wpływu na omawiany układ.

Biorąc pod uwagę prawa Kirchhoffa, nie trudno obliczyć, nawet w pamięci, że

– wskazania obu woltomierzy będą IDENTYCZNE,

– wskazania woltomierzy wyniosą 0,000...V,

Pozdrawiam całą Redakcję EdW.

Andrzej Nowicki Chemik emeryt

PS Co do uwagi korespondenta, że trzeba umieć mierzyć, to się w pełni z nią zgadzam

Mam sprawę w związku z przełomem w Waszym czasopiśmie.

Na wstępie jestem czytelnikiem od pierwszego numeru w 1996 roku. Cieszę się ze współpracy zagranicznych czasopism, ale nie zapominajmy też o naszej elektronice. Oprócz Waszego czasopisma czytałem „Praktyczny Elektronik”, „Elektor Elektronik”, żałuję że ich już nie ma.

Sprawa dotyczy nowej wizji i przedstawiania artykułów, przedstawiając ją na przykładzie w celu zobrazowania.

Zainteresowałem mnie wzmacniacz z numeru 5/2022 (316), opis i wszystko dobrze, fajnie się czyta, czekam na kolejny opis, ale chcę go zbudować i tu się pojawia problem.

Jest ładny dymek z linkiem do pcb. Ale po wpisaniu linku, nie ma dostępu do całego oryginalnego artykułu. Zalogowałem się no i okazało się że płytke trzeba kupić, ale ja nie chcę bo lubię sam zrobić w domu, ale za szablon też trzeba zapłacić.

Skoro macie już współpracę handlowo-biznesową to zróbcie Państwo to sumiennie i w 100% do końca. Skoro publikujecie i chcecie zachęcić do dalszych numerów to pozwólcie zrealizować projekt a nie tylko o nim czytać. Dlatego powinny być do każdego projektu opisanego w czasopiśmie płytki pcb, programy, kody źródłowe i cała reszta materiałów do realizacji. Pcb może być w gazecie a jak jest mało miejsca to na serwerze FTP.

Artykuły zasilacza są też dobrze opisane i ładnie zobrazowane graficznie.

Z uwagi na zapracowany tryb życia rzadko się udzielam, ale sumiennie cyklicznie czytam. Ale ta sytuacja wymagała interwencji :-)

Pozdrawiam

Wierny czytelnik Michał

Red. Odpowiedzią na ten list jest „Ważny komunikat” w ramce poniżej.

WAŻNY KOMUNIKAT

Uzgodniliśmy z Redakcją „Silicon Chip”, że kody, pliki PCB i wszelki software do ich projektów opublikowanych w EdW będą dostępne bezpłatnie. Zatem Czytelnicy EdW korzystający z naszych linków do strony internetowej „Silicon Chip” mają zapewnioną możliwość darmowego pobrania materiałów powiązanych z danym artykułem.

Równocześnie działa też druga droga dostępu do wspomnianych powyżej plików. Jest nią link do naszej strony edw.elportal.pl. Jeśli więc chcesz ściągnąć kody, pliki PCB lub inny software powiązany z danym artykułem, to znajdziesz je za darmo na <https://edw.elportal.pl/materialy-dodatkowe/>.

Natomiast korzystając z linku do strony „Silicon Chip” znajdziesz tam zarówno kody, PCB i inny software do pobrania za darmo, jak również ofertę płatną płytek drukowanych i innych podzespołów do danego projektu.

Poniżej publikujemy skrócone linki do stron z materiałami dodatkowymi do artykułów z EdW04 i 05, aby dostęp do niezbędnych plików był jeszcze szybszy i łatwiejszy.

1. Wyjątkowo czuły Magnetometr – <https://bit.ly/3M0cdkd>
2. Jednoczipowy mini wzmacniacz stereo 2x5 W – <https://bit.ly/3Fz6NKE>
3. Stroboskop improwizowany – <https://bit.ly/3w0e6YC>
4. Killer szumów usznych i bezsenności – <https://bit.ly/3wgkmdF>
5. Wzmacniacz mocy Ultra-LD 200 W RMS, część 1 – <https://bit.ly/38i3BGX>
6. Czasomierz przemówień dla konkursów i debat – <https://bit.ly/3wgxqQ5>
7. Pięcizakresowy miernik panelowy LCD z funkcją wyświetlacza USB – <https://bit.ly/3F5JwH>



Czy potrafisz opanować THEREMIN?

Osoby grające w gry wideo wiedzą, że w niektóre z nich można grać za pomocą gestów dłoni – nie dotyka się wtedy sprzętu. Podobnie niektóre telefony i tablety mogą być sterowane za pomocą gestów. Istnieje jednak instrument muzyczny, na którym również się gra za pomocą ruchów rąk – i jest on sto lat starszy od takich gier i telefonów. Nazywa się Theremin (wymawia się tere-min) i wydaje naprawdę niesamowite, wręcz upiorne dźwięki. I można go zbudować samemu. Czy uda nam się go opanować... cóż, to już inna historia!

Niesamowite dźwięki tego niemal mistycznego instrumentu pojawiają się w wielu nagraniach i ścieżkach dźwiękowych filmów aż po dzień dzisiejszy – mimo że Theremin został skonstruowany przez Lwa Termena w 1919 roku!

Magazyn Silicon Chip opublikował w ciągu ostatnich lat sześć projektów Theremina, ale ten jest pierwszym, w którym zastosowano tranzystory, a nie układy scalone. Nie ma też żadnych elementów do montażu powierzchniowego, więc jego budowa jest naprawdę

łatwa, a uruchomienie go to po prostu kwestia regulacji kilku pokręteł.

Chociaż wszystkie nasze poprzednie konstrukcje cieszyły się sporą popularnością, niektórzy z naszych czytelników tęsknili za prostą, dyskretną konstrukcją i poprosili nas o poprawienie Theremina opublikowanego w Electronics Australia w czerwcu 1969 roku, w artykule Leo Simpsona.

Było to proste... ale ta konstrukcja nie miała płytki drukowanej, a jej złożenie wymagało stolarskich oraz innych umiejętności. W związku

z tym, choć w niewielkim stopniu zmieniliśmy pięćdziesięcioletni układ, prezentuje się on teraz bardziej nowocześnie.

W poprawionej konstrukcji zastosowano nieco inne tranzystory, ponieważ niektóre z pierwotnie użytych są obecnie niedostępne. Ponadto urządzenie można zasilac z małego transformatora 9 VAC w obudowie z wtyczką sieciową lub nawet z akumulatora 12 V.

W odróżnieniu od niektórych komercyjnych Thereminów, które posiadają udziwnioną gamę

elementów sterujących, w naszym Thereminie są tylko dwa, tak jak w oryginalnym wynalazku.

Jednym z nich jest pionowa „antena”, która służy do sterowania wysokością tonu. Wysokość tonu można zmieniać, przesuwając rękę w pobliżu anteny.

Poza zwykłą zmianą wysokości tonu można dodać efekt vibrato, trzepcząc dłonią lub palcami w pobliżu anteny. Przesunięcie ręki z jednej pozycji do drugiej o bardzo małą wartość spowoduje powstanie tonu ślizgającego się lub efektu glissando – nie można łatwo zagrać pojedynczych nut.

Nawiasem mówiąc, zachowaliśmy tradycyjną nazwę „antena” dla elementu sterującego wysokością tonu w Thereminie, mimo że tak naprawdę nic nie nadaje ani nie odbiera. Ponadto przypomina on antenę teleskopową w przenośnym radiu.

Drugi element sterujący to pozioma płytką, która służy do zmiany głośności. Można również dodać efekt tremolo (podobny, ale nie taki sam jak vibrato), poruszając dłonią lub palcami nad płytką głośności. Całe to machanie i trzepotanie rękami przy elementach sterujących to tylko wykorzystanie efektów pojemnościowych do zmiany parametrów obwodu, ale fakt, że niczego tak naprawdę nie dotykamy, sprawia, że proces ten wydaje się jeszcze bardziej intrygujący dla publiczności.

Gra na Thereminie nie jest szczególnie łatwa, ale jeśli mamy dobre „ucho” muzyczne i potrafimy grać na instrumencie, takim jak skrzypce lub wiolonczela, a może puzon, będziemy mieć łatwiejszy start w tworzeniu muzyki.

Mieszanie częstotliwości

Dźwięk lub nuta muzyczna są wytwarzane przez mieszanie sygnałów z dwóch oscylatorów o częstotliwości radiowej w celu wytworzenia słyszalnej sumy lub różnicy częstotliwości. Niektórzy czytelnicy



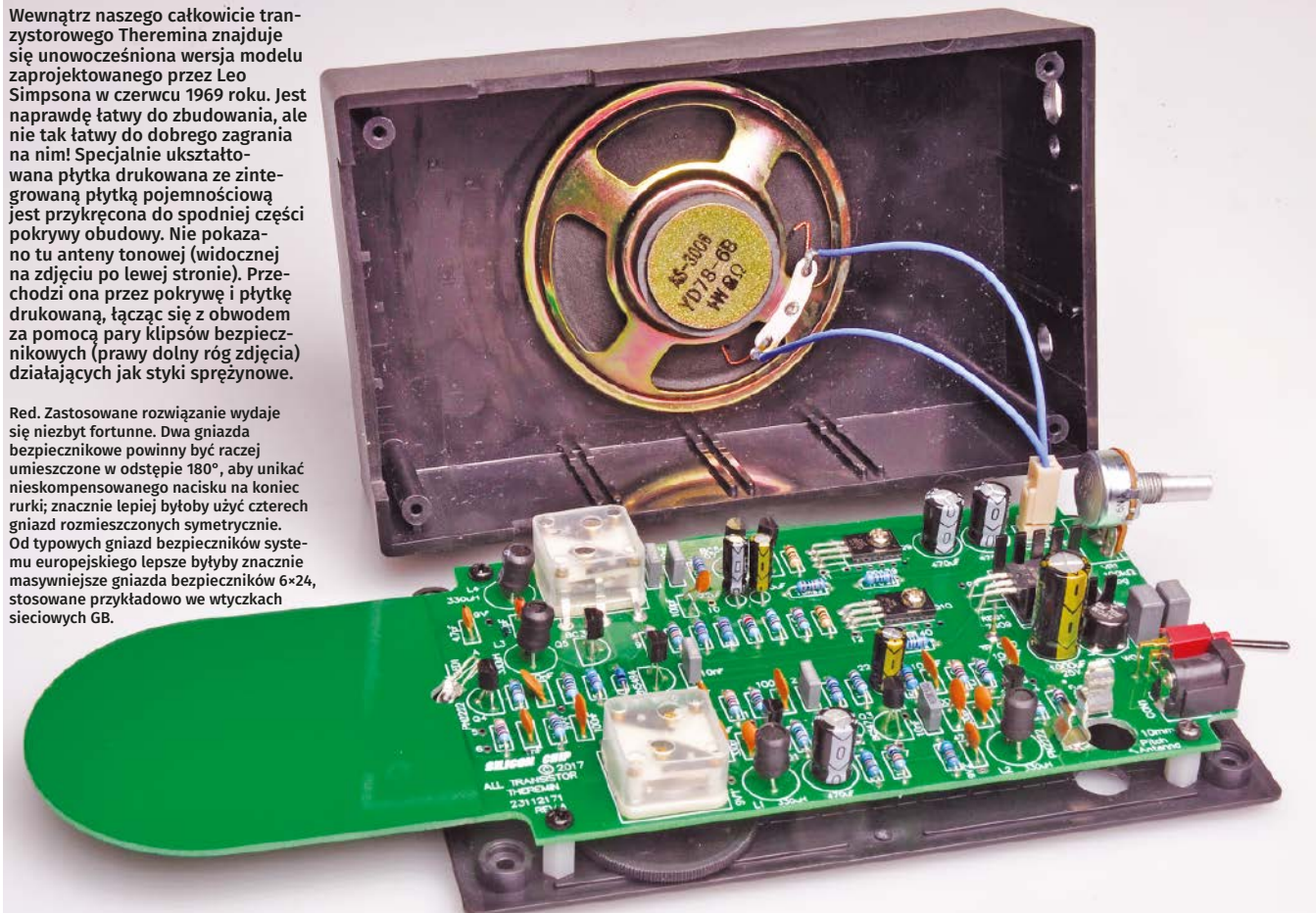
Lew Termen we własnej osobie, gra na instrumencie, który wymyślił w 1919 roku. Theremin zachwycił publiczność na trzech kontynentach

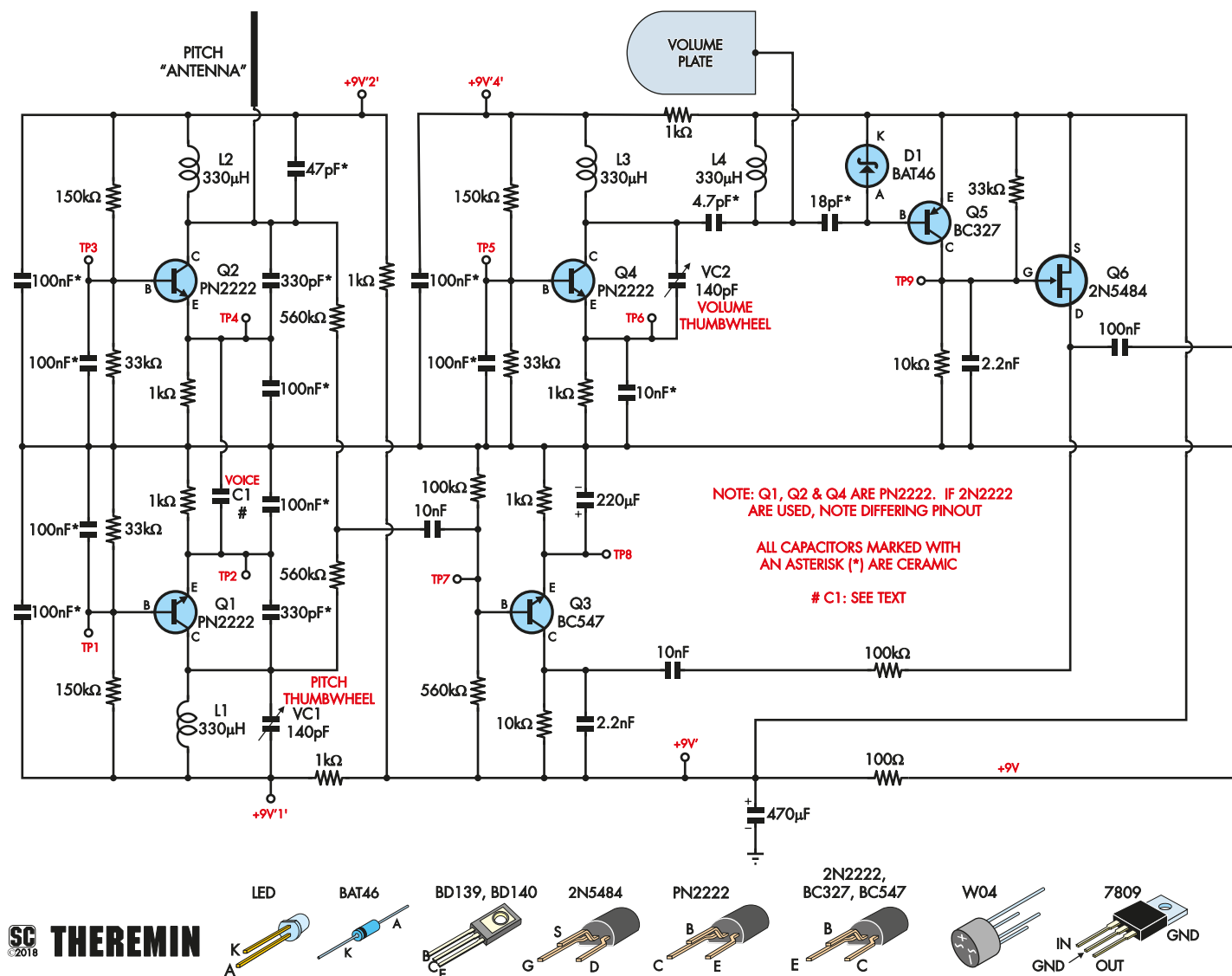
mogli słyszeć podobny rodzaj gwizdania, wydawanego przez odbiornik radiowy kiedy szukamy stacji na falach krótkich. Ostrożnie manipulując pokrętełłem odbiornika, można wytwarzać dźwięki od gwizdu o wysokiej częstotliwości do warczenia o niskiej częstotliwości.

Dwa oscylatory w Thereminie, używane do wytwarzania słyszalnego sygnału, muszą być tak zestrojone, aby mogły pracować na blisko siebie leżących częstotliwościach i bez zbyt wyraźnej tendencji do wzajemnego blokowania się na tej samej częstotliwości. Jeden oscylator musi być tak zaprojektowany, aby jego częstotliwość zmieniała się łatwo po zbliżeniu ręki do anteny tonowej. Drugi oscylator ma stałą częstotliwość.

Wewnątrz naszego całkowicie tranzystorowego Theremina znajduje się unowocześniona wersja modelu zaprojektowanego przez Leo Simpsona w czerwcu 1969 roku. Jest naprawdę łatwy do zbudowania, ale nie tak łatwy do dobrego zagrania na nim! Specjalnie ukształtowana płytką drukowaną ze zintegrowaną płytką pojemnościową jest przykręcona do spodniej części pokrywki obudowy. Nie pokazano tu anteny tonowej (widocznej na zdjęciu po lewej stronie). Przechodzi ona przez pokrywkę i płytkę drukowaną, łącząc się z obwodem za pomocą pary klipsów bezpiecznikowych (prawy dolny róg zdjęcia) działających jak styki sprężynowe.

Red. Zastosowane rozwiązanie wydaje się niezbyt fortunny. Dwa gniazda bezpiecznikowe powinny być raczej umieszczone w odstępnie 180°, aby unikać nieskompensowanego nacisku na koniec rurki; znacznie lepiej byłoby użyć czterech gniazd rozmieszczonych symetrycznie. Od typowych gniazd bezpieczników systemu europejskiego lepsze byłoby znacznie masywniejsze gniazda bezpieczników 6x24, stosowane przykładowo we wtyczkach sieciowych GB.





Gdy oba oscylatory są ustawione na tej samej częstotliwości, sygnał różnicowy ma bardzo niską częstotliwość zbliżoną do 0 Hz i z głośnika nie słychać żadnego dźwięku. Po zbliżeniu ręki do anteny zmienia się częstotliwość oscylatora zmiennego i powstaje słyszalny dźwięk (ton dudnienia).

Działanie układu

Dwa oscylatory sterujące wysokością dźwięku wykorzystują tranzystory NPN PN2222 (Q1 i Q2). Są one połączone w układ Colpittsa o częstotliwości roboczej około 470 kHz każdy.

Generator Colpittsa to rodzaj oscylatora LC, który bardzo dobrze nadaje się do tego typu układów. W Internecie można znaleźć wiele innych informacji na ten temat.

Antena tonowa jest podłączona do kolektora Q2, więc zbliżenie ręki do anteny spowoduje zmianę jej pojemności, a tym samym zmianę częstotliwości oscylacji.

Drugi oscylator wysokości tonu z udziałem Q1 jest dostrajany za pomocą

regulowanego trymera VC1 o pojemności 140 pF. Ten trymer jest standardowym kondensatorem strojeniowym w plastikowej kwadratowej obudowie, zwykle stosowanym w przenośnych radioodbiornikach AM, wykorzystywana jest jedna jego sekcja.

Podobny układ jest wykorzystywany do regulacji głośności. Oba trymery można łatwo regulować.

Sygnały z obu oscylatorów są praktycznie czystymi sinusoidami i w rezultacie ich różnica też miałaby taką samą postać. Gdyby oba oscylatory były zasilane z tego samego źródła, miałyby tendencję do ustawiania się na tej samej częstotliwości, gdy zbliżałyby się do siebie na odległość kilkuset herców. Oznaczałoby to, że częstotliwość różnicowa nie sięgałaby płynnie w rejon niskich tonów. Z tego powodu szyna zasilania każdego oscylatora jest odsprężona za pomocą rezystora 1 kΩ i kondensatora ceramicznego 100 nF.

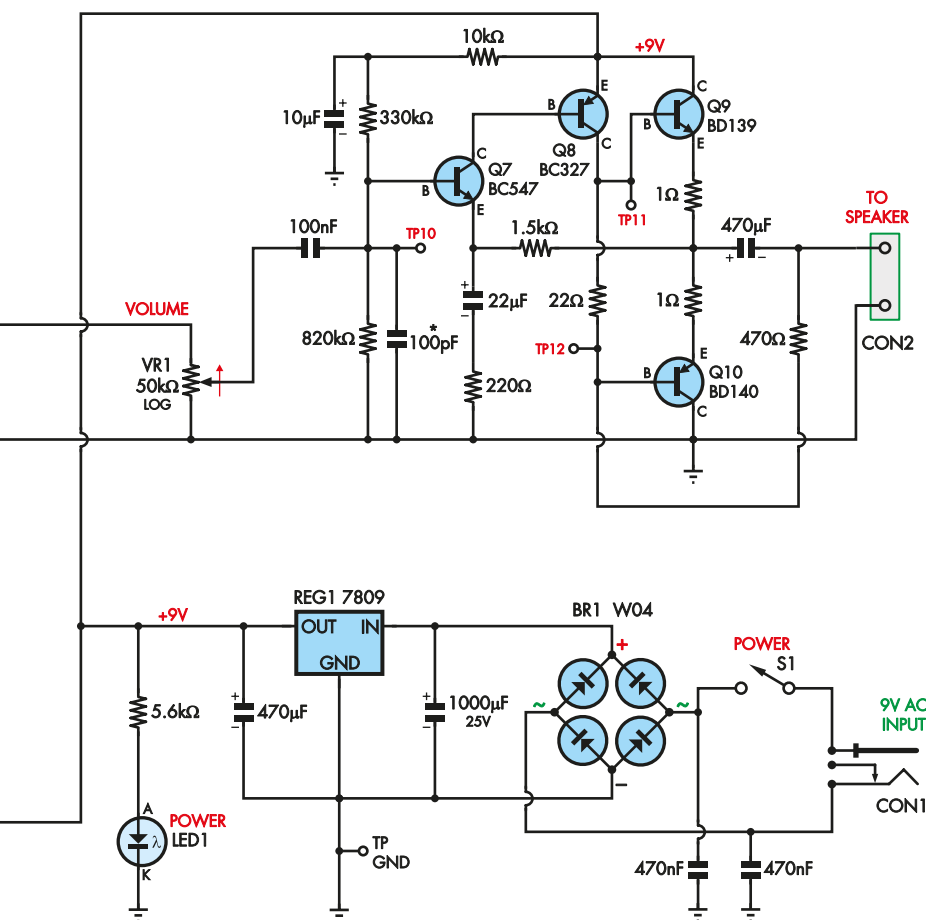
W rezultacie oba oscylatory nie blokują się, dopóki częstotliwość różnicowa nie osiągnie

zaledwie kilku Hz – jest to bardzo niskie buczanie. Pożądane jest jednak, aby oscylatory w końcu się wzajemnie zablokowały, w przeciwnym razie zbyt trudno byłoby wyregulować pokrętkę VC1 w celu uzyskania sygnału różnicowego o częstotliwości 0 Hz.

W oryginalnym układzie Theremina skonstruowanego przez Leo Simpsona tranzystorami oscylatora były tranzystory Philipsa BF115 RF, ale obecnie są one niedostępne. Używamy więc tanich tranzystorów ogólnego przeznaczenia PN2222, które mają bardzo dobrą częstotliwość graniczną (fT) równą 250 MHz, więc nie mają problemów z oscylacją przy 470 kHz.

Wyjście z każdego oscylatora przez rezystory 560 Ω jest doprowadzone do stopnia mieszacza na tranzystorze BC547 NPN ogólnego przeznaczenia, Q3, w konfiguracji wspólnego emitera.

Na wyjściu mieszacza otrzymujemy cztery częstotliwości: dwie częstotliwości oscylatorów wynoszące po około 470 kHz, sumę tych dwóch częstotliwości (około 940 kHz) oraz



Rysunek 1. Dwa sygnały o częstotliwości radiowej, generowane przez oscylatory Q1 i Q2, są miksowane w celu wytworzenia sygnału dźwiękowego, którego częstotliwość można zmieniać poprzez odległość dłoni od anteny. Nieco inny układ, ale również oparty na pojemności dłoni/płytki, zmienia głośność podawaną do konwencjonalnego wzmacniacza audio i małego głośnika

różnicę między tymi dwiema częstotliwościami, która jest słyszalnym sygnałem wyjściowym.

BC547 nie ma zbyt dużego wzmocnienia w zakresie RF, a kondensator 2,2 nF bocznikujący rezystor kolektorowy dodatkowo tłumi składowe RF, pozostawiając na wyjściu jedynie pożądane sygnały akustyczne. Stopień miesza jest nieco przesterowany, aby dodać składowe harmoniczne, dzięki czemu dźwięk będzie subiektywnie bardziej interesujący.

Niewielką zmianą, jaką wprowadziliśmy w oryginalnym układzie, jest możliwość sprzężenia oscylatorów: przestrajanego dłonią i referencyjnego za pomocą C1, co pozwala na uzyskanie „udźwięcznienia”.

Gdy częstotliwość oscylatora tonu różni się od częstotliwości oscylatora referencyjnego, a więc otrzymujemy dźwięk wyjściowy, różnica częstotliwości pomiędzy tymi dwoma oscylatorami moduluje kształt fali o częstotliwości akustycznej, tak że nie jest to idealna sinusoida.

Zazwyczaj w przypadku Theremina chcemy uzyskać dźwięk przypominający wiolonczelę

przy niskich częstotliwościach, który wraz ze wzrostem częstotliwości zmienia się w coś bardziej zbliżonego do fletu. Dodanie kondensatora C1 pozwala na eksperymentowanie w celu uzyskania innego brzmienia – można wypróbować wartości od około 220 pF do 470 pF.

Dzielnik rezystancyjny sterowany napięciem

Wyjście z kolektora Q3 jest doprowadzane do dzielnika napięcia składającego się z rezystora 100 kΩ oraz rezystancji dren-źródło N-kanalowego tranzystora JFET, Q6.

Rezystancja dren-źródło Q6 zależy od napięcia bramka-źródło, podawanego przez układ regulacji głośności, obejmujący oscylator na Q4, płytkę pojemnościową i wzmacniacz prądu stałego, Q5. Q4 to kolejny tranzystor PN2222 NPN, a oscylator głośności to również generator Colpittsa, pracujący z częstotliwością około 900 kHz. Oscylator głośności ma również odsprężone zasilanie przez rezystor 1 kΩ i kondensator 100 nF.

Wyjście oscylatora głośności jest podawane za pośrednictwem kondensatora ceramicznego o pojemności 4,7 pF do równoległego obwodu przestrajanego, składającego się z dławika radiowego o pojemności 330 μH i pojemności płytki regulacji głośności. Część sygnału z obwodu przestrajanego jest sprzężona z diodą Schottky'ego D1 przez kondensator 18 pF.

Powstały w ten sposób pulsujący prąd stały jest wzmacniany przez tranzystor PNP Q5, a napięcie na rezystorze 10 kΩ podawane na bramkę FET-a po odfiltrowaniu składowej RF za pomocą kondensatora 2,2 nF. Poziom odtwarzanego dźwięku powinien się zmniejszać po zbliżeniu ręki do płytki głośności.

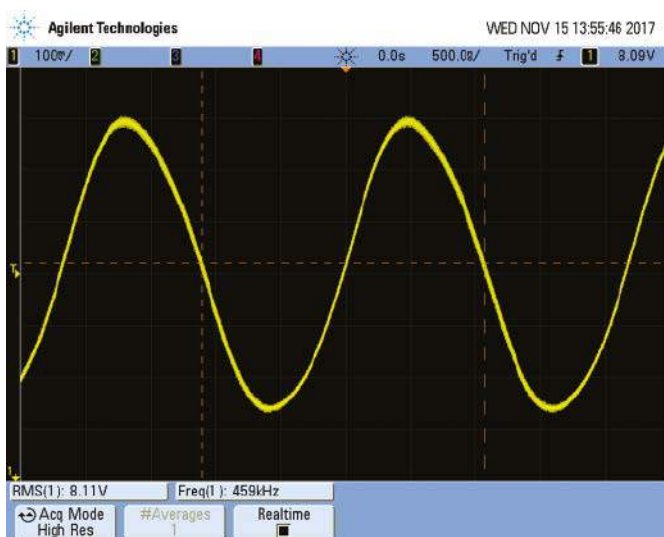
Początkowo oscylator głośności jest regulowany za pomocą trymera VC2 o pojemności 140 pF, aby uzyskać minimalną głośność sygnału dźwiękowego, gdy dłoń znajduje się w pobliżu płytki głośności. Polega to na dostrojeniu oscylatora w taki sposób, aby jego częstotliwość pokrywała się z częstotliwością rezonansową przestrajanego układu.

W rezultacie napięcie pochodzące od diody osiągnie maksimum, co spowoduje, że Q5 znajdzie się w stanie przewodzenia, a w konsekwencji zostanie włączony. Bramka FET jest polaryzowana w stronę dodatniej szyny zasilania, a rezystancja dren-źródło jest na niskim poziomie. Powoduje to zwarcie dużej części sygnału akustycznego do szyny zasilania.

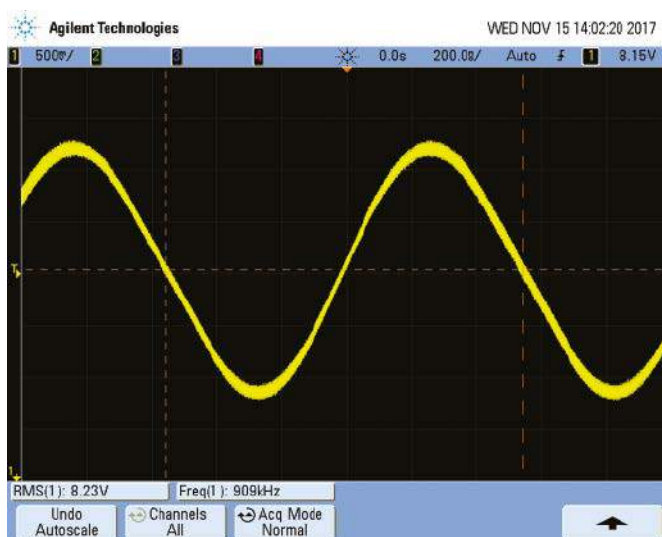
Po odsunięciu ręki od płytki głośności pojemność w obwodzie strojonym zmienia częstotliwość rezonansową, przez co zmniejsza się napięcie stałe na diodzie. Powoduje to stopniowe przesuwanie punktu pracy Q5 w kierunku odcięcia, co powoduje wzrost rezystancji dren-źródło FET-a. W ten sposób większa część sygnału dźwiękowego jest podawana do kolejnego wzmacniacza. W tym momencie uwiadcza się szczególna cecha FET-a. Dla małych napięć o dowolnej polaryzacji (lub przemiennej) przyłożonych między drenem a źródłem, FET zachowuje się jak rezystor, który można zmieniać liniowo za pomocą napięcia przyłożonego między źródłem a bramką.

Przy napięciu bramki zmieniającym się w zakresie od zera do około 4 V poniżej potencjału źródła zależność między napięciem bramka-źródło a rezystancją dren-źródło może być względnie liniowa, ale przestaje to być prawdą, gdy napięcie bramka-źródło zbliża się do napięcia wyłączania FET.

W tym obszarze zależność staje się bardzo nieliniowa – niewielki wzrost napięcia bramka-źródło powoduje bardzo dużą zmianę rezystancji dren-źródło, a więc FET wyłącza się w niewielkim zakresie napięcia. Oznacza to, że w pewnym obszarze w pobliżu płytki głośności, niewielki ruch ręką powoduje dużą zmianę głośności, tak że działa on niemal jak przełącznik. Aby



Przebieg 1. Sygnał ten jest wyjściem oscylatora referencyjnego „tonu” (opartego na Q1), którego częstotliwość jest regulowana przez kondensator pokrętki tonu, VC1. Należy zauważyć, że sygnał wyjściowy jest czystą sinusoidą



Przebieg 2. Podobnie, wyjście oscylatora „głośności” na podstawie Q4. Jest ono regulowane przez VC2. Oba te pomiary są trudne do wykonania z powodu obciążenia przez sondę oscyloskopu

ograniczyć ten efekt, między kolektor i emiter Q5 podłączono rezystor 33 kΩ. Gdy tranzystor jest wyłączony, rezystor 33 kΩ i rezystor 10 kΩ w obciążeniu kolektora tworzą dzielnik napięcia, który ogranicza napięcie między bramką a źródłem FET do około minus sześciu woltów. Powoduje to, że regulacja głośności działa bardziej płynnie, ale zmniejsza też dostępny zakres regulacji.

Należy zauważyć, że nie jest możliwe, aby układ regulacji głośności dawał zerowy sygnał wyjściowy, ponieważ minimalna rezystancja FET-u wynosi zwykle 100 Ω i nie może on zewrzeć całego sygnału do szyny zasilania.

Podsumowując, wysokość dźwięku Theremina jest kontrolowana przez różnicę częstotliwości dwóch oscylatorów RF pracujących z częstotliwością około 470 kHz, z których jeden jest wrażliwy na pojemność dłoń-antena. Wysokość uzyskanego w ten sposób tonu można zmieniać w całym zakresie słyszalności.

Głośność dźwięku jest sterowana przez trzeci oscylator pracujący z częstotliwością około 900 kHz, którego obwód strojony zmienia częstotliwość rezonansową pod wpływem zmian pojemności dłoni i płytki regulacji głośności. Napięcie stałe, pochodzące ze strojonego obwodu, jest wykorzystywane do zmiany rezystancji dren-źródło FET-a, stanowiącego część rezystancyjnego dzielnika napięcia, do którego jest doprowadzany sygnał akustyczny.

Zrozumiały to, łatwo można prześledzić pozostałe aspekty działania Theremina. Sygnał z dzielnika FET jest podawany do potencjometru 50 kΩ, a następnie do wzmacniacza audio i głośnika.

Czterotranzystorowy wzmacniacz jest konwencjonalną konstrukcją z bezpośrednim sprzężeniem, w której dwa tranzystory wyjściowe są połączone w symetrycznym układzie

komplementarnym i pracują w czystej klasie B, tzn. bez prądu spoczynkowego redukującego zniekształcenia skrośne.

W tym projekcie nie zwracamy uwagi na zniekształcenia skrośne, częściowo dlatego, że ustawienie prądu spoczynkowego zwiększyłoby całkowity pobór prądu, co nie jest pożądane w przypadku zasilania Theremina z baterii.

Jak się okazuje, co widać na przebiegu 4, zniekształcenia skrośne nie są zauważalne na wyjściu. Całkowity pobór prądu jest głównie spowodowany prądem kolektora Q8, stopnia wzmocnienia napięciowego wzmacniacza klasy A. Maksymalna moc wyjściowa wynosi około 400 mW przy 8-omowym głośniku.

Ciekawostką dotyczącą wzmacniacza jest to, że stosujemy standardowy układ, w którym obciążenie kolektora Q8 jest zmniejszane przez „boot-strapping” (podciąganie) z wyjścia. Zamiast podłączać obciążenie kolektora Q8, tj. rezystor 470 Ω do masy, podłączyliśmy je do aktywnego zacisku głośnika, czyli do ujemnej elektrody wyjściowego kondensatora separującego 470 μF.

Ze względu na działanie tranzystorów wyjściowych Q9 i Q10 w układzie wtórnika emiterowego, impedancja obciążenia dla prądu przemiennego „widziana” przez kolektor Q8 jest znacznie wyższa niż 470 Ω.

W efekcie, ze względu na działanie tego wtórnika emiterowego, napięcie zmienne (tj. napięcie sygnału audio) jest praktycznie takie samo na obu końcach rezystora 470 Ω i dlatego prąd przemienny jest znacznie zmniejszony.

Należy zauważyć, że niewielki prąd stały obciążenia Q8 przepływa przez cewkę głośnika do masy.

Rysunek 2. Schemat montażowy płytki drukowanej Theremina. Wszystkie elementy, z wyjątkiem głośnika, są montowane na tej płytce drukowanej. Bezpośrednio poniżej znajduje się zdjęcie płytki drukowanej w tym samym rozmiarze, tym razem zamontowanej na pokrywie używanej przez nas obudowy UB1 Jiffy box. Jeśli wykazalibyśmy dużo chęci, moglibyśmy wykonać obudowę z drewna, tak jak w przypadku oryginalnego Theremina i większości wczesnych komercyjnych Thereminów. Nawiasem mówiąc, są dwie drobne różnice między zdjęciem prototypu po prawej stronie (PCB Rev „A”) a ostateczną płytką drukowaną / schemat montażowy powyżej (Rev „B”). Wartość VR1 została zmieniona na 50 kΩ (wcześniej było 100 kΩ), a w pobliżu Q7 dodano kondensator 2,2 nF. Podczas montażu należy zawsze postępować zgodnie z instrukcją montażu elementów

Poprawia to wzmocnienie i liniowość napięcia wyjściowego Q8. Jedyną potencjalną wadą tego układu jest to, że jeżeli głośnik jest odłączony, Q8 nie ma połączenia do masy, a zatem wzmacniacz jest zablokowany, pobierając znikomy prąd.

Zasilanie

Układ jest zasilany z transformatora 9 VAC w obudowie z wtyczką sieciową. Można również użyć akumulatora 12 V, ale parametry układu mogą być gorsze niż w przypadku zasilania prądem przemiennym.

Należy pamiętać, że zasilacz impulsowy 12 V prądu stałego nie nadaje się

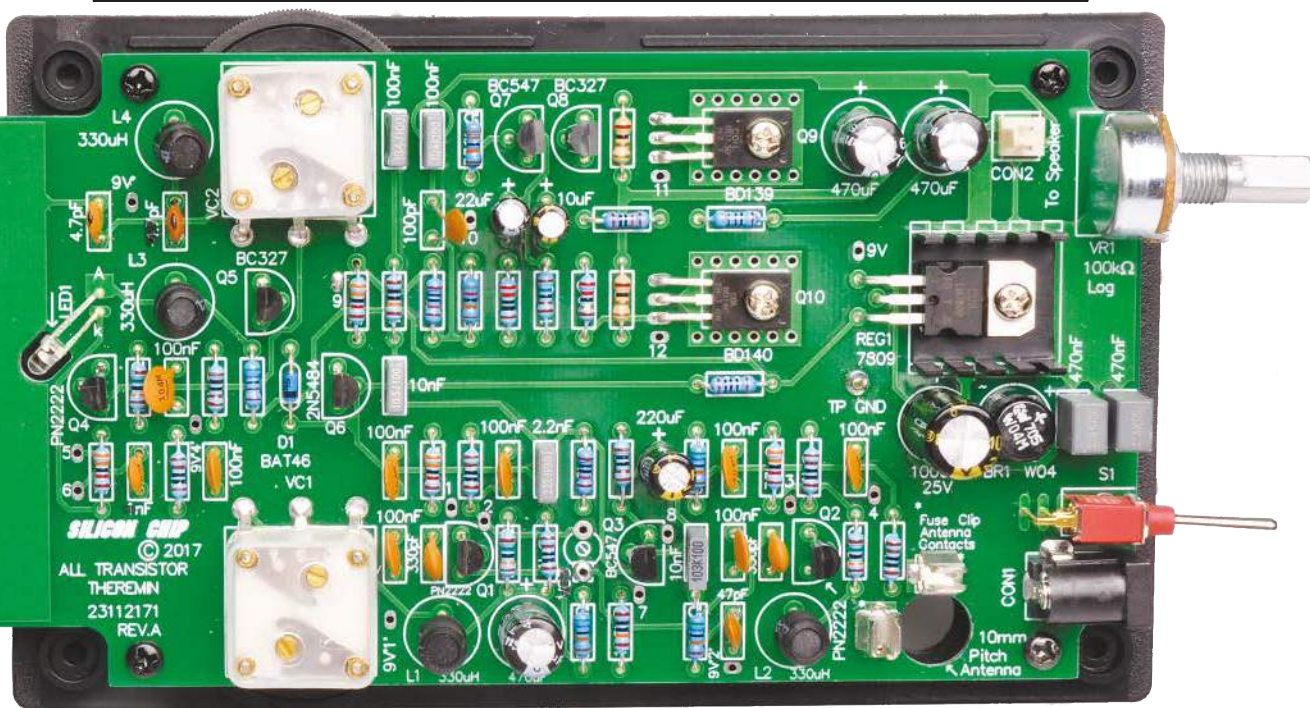
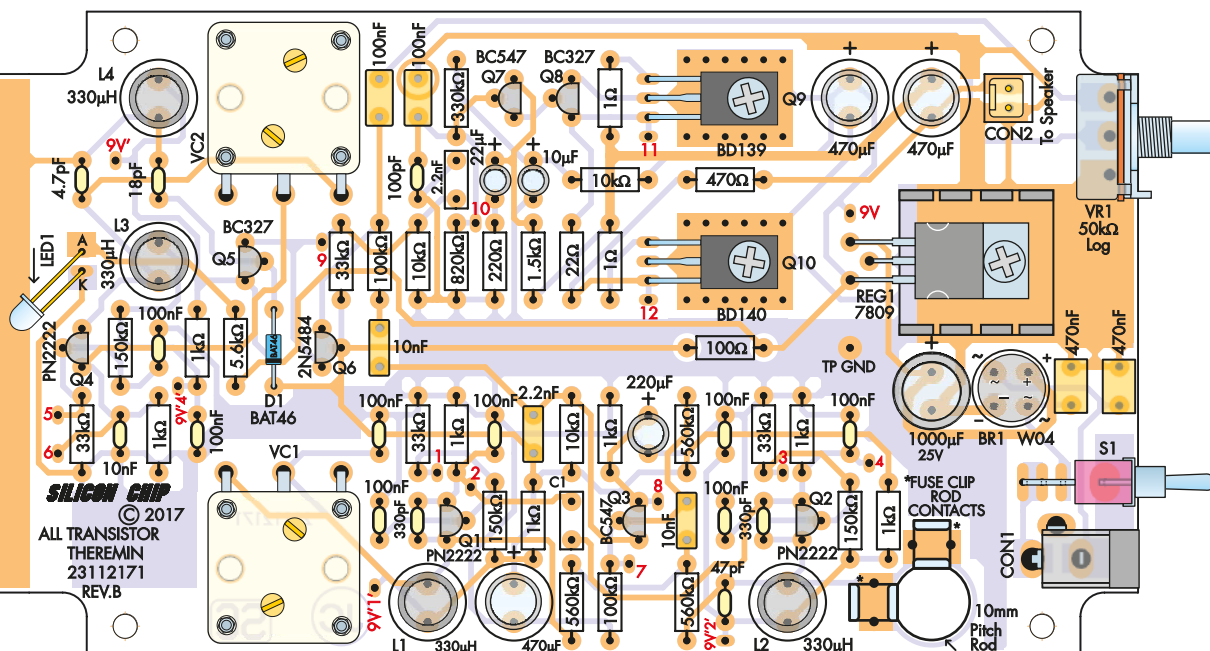
do tego projektu ze względu na dużą ilość harmonicznych i szumów, które zwykle emituje – istnieje duże prawdopodobieństwo, że zakłóci to pracę oscylatorów, przedostanie się ich do wzmacniacza audio... albo jedno i drugie.

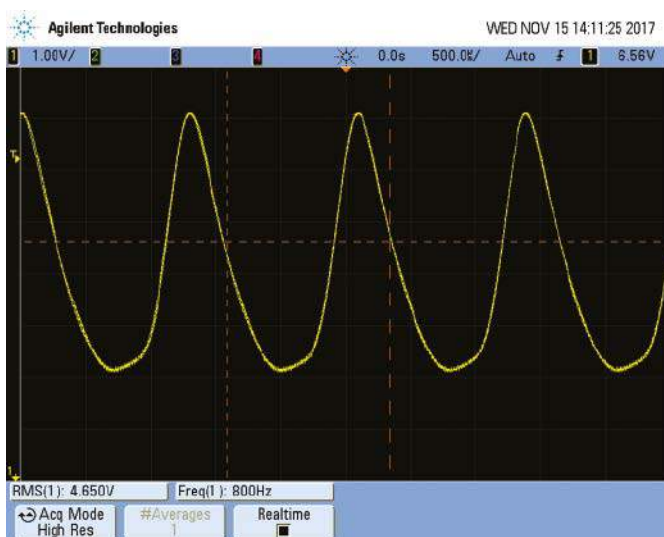
Przełącznik S1 włącza zasilanie układu. Kondensatory 470 nF podłączone do każdej linii zasilania wejściowego zwierają do masy wyższe harmoniczne sieci przedostające się do transformatora. W ten sposób można uniknąć niepożądanych dźwięków z Theremina spowodowanych przez zasilanie z sieci. Dodatkową korzyścią jest to, że kondensatory filtrujące 470 nF zapewniają efekt wirtualnego uziemienia linii zasilającej, dzięki czemu pojemność

dłoni jest bardziej efektywna w przypadku regulatorów tonu i głośności.

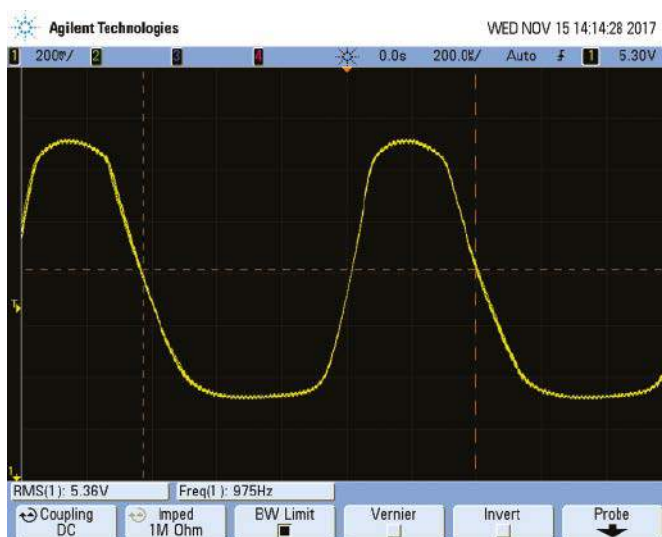
Napięcie 9 VAC jest prostowane przez mostek prostowniczy BR1, a następnie filtrowane za pomocą kondensatora 1000 μ F, aby zapewnić stosunkowo niskie tętnienia napięcia wynoszące ok. ~12 VDC dla REG1, scalonego stabilizatora 9 V, który dostarcza stabilne napięcie stałe 9 V do układu. Kondensator 470 μ F na wyjściu stabilizatora poprawia pracę zasilacza i może zapewnić krótkotrwałą prąd szczytowy dla wzmacniacza. Dioda LED1 sygnalizuje, że zasilanie jest włączone.

Przy okazji należy zauważyć, że zasilanie Theremina z baterii ma pewną wadę,





Przebieg 3. Ten fragment przebiegu pokazuje sygnał na wyjściu miksera, Q3, mierzony na jego kolektorze. Jego amplituda jest w następnym stopniu zmieniana przez tranzystor JFET Q6, a następnie podawana do regulatora głośności VR1 i wzmacniacza audio



Przebieg 4. Wyjście ze wzmacniacza audio na głośnik. Należy zauważyć, że nie ma żadnych widocznych zniekształceń skrośnych, mimo że w tranzystorach wyjściowych nie występuje prąd spoczynkowy: układ pracuje w czystej klasie B. Red. Ciekawe byłoby zbadanie tego sygnału analizatorem widma lub analizatorem harmonicznym.

ponieważ efekt wirtualnego uziemienia jest zapewniony przez dwa kondensatory 470 nF, które w przypadku zasilania prądem stałym nie odgrywają żadnej roli. Może to w pewnym stopniu ograniczyć czułość obu regulacji.

Montaż

Wszystkie elementy układu są umieszczone na stosunkowo niewielkiej płytce drukowanej, na której znajduje się również płytka regulacji głośności, co sprawia, że urządzenie jest łatwe w budowie. Antena tonowa to rurka aluminiowa o długości 400 mm i średnicy 10 mm, włożona w otwór w panelu przednim i płytce drukowanej, stykająca się z układem za pomocą dwóch sprężystych styków, które są w rzeczywistości gniazdami standardowego bezpiecznika na płytce drukowanej.

Dwa kondensatory strojeniowe są zamontowane bezpośrednio na płytce drukowanej, a ich pokrętła wystają nieco z każdej strony obudowy. Rozpocznijmy montaż od zamontowania rezystorów. Niezależnie od tego, czy znamy kody kolorów rezystorów, czy nie, zachęcamy sprawdzić wartości każdego z nich multimetrem cyfrowym przed włożeniem i wlutowaniem na miejsce (niektóre paski kolorów są bardzo podobne do innych).

Rezystory nie są spolaryzowane i można je montować do płytki w dowolny sposób. Dobrą praktyką jest jednak montowanie ich w taki sposób, aby wszystkie ich kody kolorystyczne były ułożone w tym samym kierunku (np. pasek tolerancji na dole lub po prawej stronie). Dzięki temu o wiele łatwiej jest później sprawdzić ich wartości.

Teraz można umieścić cztery cewki 330 μ H. Następnie należy zamontować kondensatory. W tym układzie są stosowane ich trzy typy.

Jednym z typów są poliestrowe MKT, które można rozpoznać po ich małym kształcie „klocka”. Drugi typ to ceramiczne w kształcie dysku.

Ani kondensatory poliestrowe, ani ceramiczne nie są spolaryzowane – można je wkładać dowolną stroną. Czasami są one oznaczone kodem, aby wskazać ich pojemność, ale najczęściej mają nadrukowaną wartość pojemności.

Trzecim rodzajem kondensatorów są kondensatory elektrolityczne. Mają one (zazwyczaj) kształt cylindryczny i z rzadkimi wyjątkami (w tym obwodzie nie ma ich wcale) są spolaryzowane – muszą być włożone konkretną stroną, jak pokazano na schemacie montażowym płytki drukowanej. Po jednej stronie mają oznaczenie polaryzacji w postaci symboli „-”, które wskazują przewód ujemny.

Następne w kolejności są półprzewodniki, z których wszystkie są spolaryzowane. Zamontujemy diodę D1 i mostek prostowniczy BR1, a następnie tranzystory.

Upewnijmy się, że w każdej pozycji umieściliśmy właściwy tranzystor – niektóre wyglądają identycznie.

Zwróćmy uwagę, że płytka drukowana jest zaprojektowana dla tranzystorów PN2222A dla Q1, Q2 i Q4. W przypadku stosowania tranzystorów 2N2222A należy je wstawić po obroceniu w pionie pod kątem 180° w stosunku do pokazanego na schemacie montażowym płytki drukowanej, przy czym wyprowadzenie bazy należy wygiąć do przodu, aby dopasować je do położenia otworów w płycie.

Tranzystory Q9 i Q10 są zamontowane poziomo, metalową stroną w kierunku płytki drukowanej. Ich przewody są zgięte pod kątem 90°, aby można je było wsunąć w otwory płytki drukowanej.

Oprócz lutowania, tranzystory są mocowane do płytki drukowanej za pomocą śrub M3×10 mm i nakrętek, przy czym śruba jest umieszczona od strony lutowania na płytce drukowanej, a nakrętka na tranzystorze. Przed przylutowaniem zamocujemy śrubę i nakrętkę, aby upewnić się, że znajdują się we właściwym położeniu.

Stabilizator REG1 jest zamontowany poziomo, podobnie jak Q9 i Q10, ale jest umieszczony na małym radiatorze, który znajduje się pomiędzy nim a płytka drukowaną. Wygnijmy wyprowadzenia pod kątem 90° przed włożeniem ich do płytki drukowanej, przymocujemy wypustkę do radiatora i płytki drukowanej za pomocą śruby M3×10 mm i nakrętki, a następnie przylutujemy wyprowadzenia.

Teraz można przystąpić do montażu złączy CON1 i S1. Przed przylutowaniem upewnijmy się, że te dwie części są mocno dociśnięte do płytki drukowanej. Następnie można wlutować dwa klipsy bezpiecznikowe, które stykają się z anteną tonową. Klipsy mogą wymagać niewielkiego rozchylenia, aby zapewnić dobry styk z anteną z rurki aluminiowej 10 mm (red. patrz uwagę przy spisie elementów).

CON2 służy do połączenia z głośnikiem. Zamontujemy dwustykowe gniazdo szpilkowe na płytce drukowanej. Do dwustykowego gniazda podłącza się przewody o długości ~100 mm, zaciskając najpierw końcówki przewodów w złączach zaciskowych wtyku (można je również przylutować, aby uzyskać pewne połączenie), a następnie wkładając je do gniazda. Drugie końce przewodu są przylutowane do zacisków głośnikowych.

Dioda LED1 jest montowana poziomo w wycięciu w płytce drukowanej, a jej



Obudowa UB1 Jiffy box z wklejonym głośnikiem oraz trzema szczelinami i trzema otworami wymaganymi na ściankach obudowy. Należy również wywiercić w podstawie obudowy okrągły wzór otworów, aby umożliwić wydobywanie się dźwięku

wyprowadzenia są wygięte tak, aby wchodziły w otwory w płytce. Upewnijmy się, że polaryzacja jest prawidłowa – dłuższy przewód jest anodą.

Dwa kondensatory dostrajające (VC1 i VC2) w plastikowych obudowach są przymocowane do płytki drukowanej za pomocą dwóch krótkich wkrętów M3 (powinny być dostarczone razem z kondensatorami). Trzy końcówki muszą być zgięte pod kątem prostym, aby można je było wsunąć w otwory w płytce drukowanej. Następnie są one lutowane na miejscu.

Potencjometr regulacji głośności instrumentu lutujemy do płytki po uprzednim dopasowaniu długości jego osi do posiadanej gałki.

Testowanie

Sprawdźmy dokładnie swoją konstrukcję, aby upewnić się, że nie ma w niej żadnych błędów – zwłaszcza jeśli chodzi o orientację wszystkich spolaryzowanych elementów (kondensatorów elektrolitycznych, diod, tranzystorów i stabilizatora), a właściwe elementy znajdują się we właściwych miejscach. Po upewnieniu się, że wszystko jest w porządku, podłączmy transformator 9 VAC (lub akumulator 12 V prądu stałego – plus do środkowego bolca) i włączmy urządzenie. Dioda LED1 powinna się zaświecić. Na płytce

drukowanej umieściliśmy wiele punktów testowych. Są one oznaczone od jednego do dwunastu, a dodatkowe punkty testowe są oznaczone jako TP GND, 9V, 9V', 9V1', 9V2 i 9V4.

Podłączmy ujemny przewód multimetru do masy TP. Napięcie TP 9V powinno być zbliżone do 9 V (ale może wynosić od 8,85 do 9,15 V). Punkt testowy 9V' powinien wynosić około 8,6 V, a punkty testowe 9V1', 9V2 i 9V4 powinny mieć wartości od około 8 V do 8,6 V.

W punktach testowych 1, 3 i 5 napięcie powinno wynosić około 1,0 V, choć w punkcie TP5 może być nieco niższe i wynosić około 0,8 V. TP2, 4 i 6 powinny mieć napięcie 0,4 V, przy czym TP6 może mieć napięcie nawet 0,22 V. TP7 powinien mieć napięcie około 1,1 V, a TP8 +0,6 V. Punkt testowy 9 zależy od ustawienia VC2, ale powinien mieścić się w zakresie od 2 V do 8,6 V i być regulowany za pomocą VC2.

W punkcie pomiarowym TP10 powinno być 6,2 V. Podłączmy głośnik w celu uzyskania kolejnych odczytów. W punkcie TP11 powinniśmy zmierzyć napięcie 5,5 V, a w punkcie TP12 około 5,3 V.

Jeśli wszystkie napięcia są prawidłowe, odłączmy zasilanie i umieścimy urządzenie w obudowie.

Obudowa

My umieściliśmy nasz Theremin w pudełku UB1 Jiffy box, ale dla zachowania autentyczności możemy wykonać własne pudełko z drewna lub lepiej z cienkiej sklejki, tak jak w oryginalnym projekcie Lwa Termena (i większości wczesnych modeli). To zależy od nas.

Płytkę drukowaną jest zamontowana do góry nogami na pokrywie obudowy (tak, aby strona z elementami była skierowana w dół). Jeśli wykonamy obudowę z drewna, powinna ona mieć taki sam lub podobny układ. Cztery otwory o średnicy 3 mm w pokrywie utrzymują płytkę drukowaną na miejscu.

Należy wyciąć w obudowie trzy szczeliny. Jedna z nich umożliwia wysuwanie płytki głośności (część płytki drukowanej) z lewej strony, a dwie inne umożliwiają wysuwanie pokręteł dołączonych do VC1 i VC2 z przodu i z tyłu.

Pozostałe wymagane otwory znajdują się w prawej części obudowy (7 mm na potencjometr regulacji głośności, 10 mm na gniazdo zasilania, 5 mm na wyłącznik zasilania) oraz

Wykaz elementów:

- 1× płytka drukowana oznaczona kodem 23112171, 226×85 mm (zawiera zintegrowaną płytkę głośności)
- 1× UB1 obudowa Jiffy box 158×95×53 mm
- 1× zestaw transformatorów 9 VAC 350 mA w obudowie z wtyczką sieciową
- 1× rurka aluminiowa o średnicy 10 mm i długości ~400 mm (na antenę wysokości tonu)
- 2× mini kondensatory typu radiowego w obudowach plastikowych do strojenia (z pokrętłem i śrubami montażowymi) (VC1, VC2)
- 4× dławiki 330 μH (L1-L4) pionowe
- 1× głośnik 3" (4 Ω lub 8 Ω)
- 1× pokrętło pasujące do potencjometru
- 1× przełącznik SPDT (S1) montowany na płytce drukowanej
- 1× gniazdo zasilania prądu stałego do montażu na płytce drukowanej (5,5/2,1 lub 5,5/2,5 mm) (CON1)
- 1× miniaturowy radiator 19×19×9,5 mm lub podobny
- 2× gniazda bezpieczników M20x5 do płytek drukowanych
- 4× gwintowane kołki dystansowe M3×9 mm
- 11× wkrętów M3×10 mm (4 o długości 5 mm są opcjonalne – patrz tekst)
- 3× nakrętki M3
- 1× dwuszlankowe gniazdo do druku
- 1× dwustykowa wtyczka do gniazda szpilkowego (CON2)
- 4× przyklejane gumowe nożyki (im wyższe, tym lepiej!)
- 1× opaska PC dla masy TP GND
- 1× opaska termokurczliwa o długości 15 mm i średnicy 10 mm

Półprzewodniki:

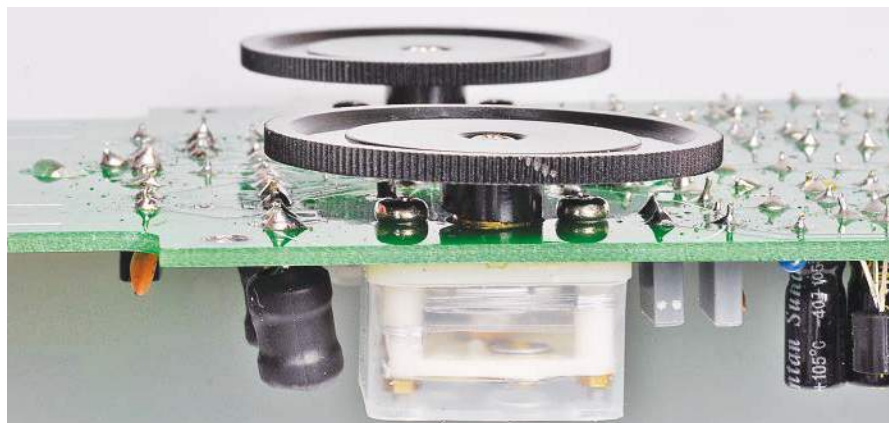
- 3× tranzystory PN2222 NPN (Q1, Q2, Q4) lub 2N2222A (patrz tekst)
- 2× tranzystory BC547 NPN (Q3, Q7)
- 2× tranzystory PNP BC327 (Q5, Q8)
- 1× 2N5484 JFET (Q6)
- 1× tranzystor NPN BD139 (Q9)
- 1× tranzystor PNP BD140 (Q10)
- 1× dioda Schottky'ego BAT46 (D1)
- 1× regulator/stabilizator napięcia 7809 9 V (REG1)
- 1× prostownik mostkowy W04 1 A (BR1) lub podobny
- 1× niebieski LED 3 mm o wysokiej intensywności światła (LED1)

Kondensatory:

- 1× 1000 μF 25 V PC elektrolityczny
 - 3× 470 μF 16 V PC elektrolityczny
 - 1× 220 μF 16 V PC elektrolityczny
 - 1× 22 μF 16 V PC elektrolityczny
 - 1× 10 μF PC elektrolityczny
 - 2× 470 nF MKT poliestrowy
 - 2× 100 nF MKT poliestrowy
 - 2× 10 nF MKT poliestrowy
 - 2× 2,2 nF MKT poliestrowy
 - 8× 100 nF ceramiczny
 - 1× 10 nF NPO (COG) ceramiczny
 - 2× 330 pF NPO (COG) ceramiczny
 - 1× 100 pF ceramiczny
 - 1× 47 pF NPO (COG) ceramiczny
 - 1× 18 pF NPO (COG) ceramiczny
 - 1× 4,7 pF NPO (COG) ceramiczny
- Red. kondensatory ceramiczne o pojemności do 10 nF mogą być z bardzo dobrym skutkiem zastąpione kondensatorami styrofleksowymi KSF, dostępnymi na wyprzedziałach lub z demontażu

Rezystory: (0,25 W, 1%)

- 1× 820 kΩ
- 3× 150 kΩ
- 3× 10 kΩ
- 7× 1 kΩ
- 1× 100 Ω
- 3× 560 kΩ
- 2× 100 kΩ
- 1× 5,6 kΩ
- 1× 470 Ω
- 1× 22 Ω
- 1× 330 kΩ
- 4× 33 kΩ
- 1× 1,5 kΩ
- 1× 220 Ω
- 2× 1 Ω 5%
- 1× 50 kΩ potencjometr logarytmiczny 16 mm (VR1)



To zbliżenie płytki drukowanej pokazuje, jak zamocowane są dwa kondensatory zmienne (w rzeczywistości są to miniaturowe kondensatory do strojenia radia AM). Jeśli okaże się, że pokrętło zaczepia się lub blokuje na płytce lub obudowie, być może trzeba będzie wyregulować jego położenie lub pogłębić szczelinę

Dowiedz się więcej o Thereminie (a nawet naucz się na nim grać!)

W internecie można znaleźć tysiące przykładów propagatorów Theremina (wystarczy wpisać w Google hasło „Theremin”). Wielu z nich to znakomici muzycy, którzy naprawdę wiedzą, jak sprawić, by ten instrument dosłownie „śpiewał”.

Jedną z najlepszych jest w zasadzie wnuczka Lwa Termena – demonstracja Lydii Kaviny na stronie www.bbc.com/news/magazine-17340257, która trwa tylko kilka minut, ale warto ją obejrzeć. Na tej samej stronie znajduje się ciekawy artykuł Martina Vennarda z BBC World Service o Lwie Termenie i wynalezionym przez niego instrumencie.



Lidia Kavina demonstruje instrument, który wynalazł jej wuj

Thereminowa interpretacja utworu Clair de Lune Debussy'ego jest po prostu urzekająca. Warto poszukać w internecie jej innych utworów.

Innym mistrzowskim przykładem gry na Thereminie jest prawie 17-minutowy film zamieszczony na stronie <https://youtu.be/MJACNHuGp0>, na którym Carolina Eyck, niemiecka kompozytorka grająca na Thereminie (uważana za jedną z najlepszych na świecie), nie tylko demonstruje swoją sprawność na tym instrumencie, ale także szczegółowo wyjaśnia, jak na nim gra.

Co prawda Theremin, na którym gra, jest znacznie bardziej skomplikowany (i droższy!) niż nasz prosty model, a ponadto oferuje szereg elementów sterujących, które odstraszyłyby wszystkich, z wyjątkiem najbardziej doświadczonych graczy. Ale ten film pomoże naprawdę zrozumieć zawiłości Theremina – zwłaszcza jeśli chcemy wydobyć z niego coś więcej niż tylko zwykłe wycie i piski początkującego muzyka!



Carolina Eyck wyjaśnia, co potrafi Theremin!

jeden otwór 10 mm w pokrywie obudowy, przez który ma przechodzić antena tonowa (plus cztery wspomniane już otwory do mocowania płytki drukowanej). W podstawie obudowy należy również wykonać szereg otworów, przez które będzie wydostawał się dźwięk z głośnika. Przedstawiliśmy schematy wszystkich tych otworów. Można zmierzyć i zaznaczyć położenie otworów albo skserować schematy ze strony siliconchip.com.au, wydrukować je i użyć jako szablonów).

Przymocujemy dwa pokręta do VC1 i VC2 za pomocą dostarczonych wkrętów M3. Upewnijmy się, że podczas obracania pokręta nie blokują się na płytce drukowanej.

Jeśli tak jest, być może trzeba będzie nieco spiłować tuleję pokręta, aby zapewnić dodatkowy luz nad płytką drukowaną.

Głośnik przyklejamy do podstawy obudowy klejem kontaktowym, uszczelniaczem silikonowym lub podobnym środkiem. Gumowe nóżki są przymocowane do spodu obudowy, aby ją podnieść i umożliwić wydobywanie się dźwięku.

Zamknięcie w obudowie

Płytkę drukowaną może być przymocowana do obudowy nakrętką potencjometru, ale postanowiliśmy przymocować ją również do pokrywy obudowy czterema tulejami dystansowymi z gwintem M3 10 mm, z których każdy ma na górze i na dole śrubę M3 5 mm (alternatywnie można użyć 10-milimetrowych, niegwintowanych prętów dystansowych z gwintem M3 o długości ok. 5 mm na każdym końcu o łącznej długości 20 mm i nakrętkami M3, bezpośrednio na panelu przednim).

Takie podejście nieco utrudnia instalację płytki drukowanej w obudowie, ale można to zrobić, co potwierdzają nasze zdjęcia!

Wsuńmy pokrywę wraz z płytką drukowaną do obudowy tak, aby dzwignia przełącznika i wałek potencjometru wychodziły przez otwory w prawej części. Następnie przymocujemy potencjometr za pomocą nakrętki.

Przed przystąpieniem do regulacji należy zainstalować antenę. Antena jest umieszczona w pokrywie górnej na głębokości 24 mm. Umieściliśmy odcinek opaski termokurczliwej o średnicy 10 mm na dolnym końcu aluminiowej rury, aby zaznaczyć, kiedy należy przerwać dalsze wkładanie rury do obudowy.

Upewnijmy się, że dwa pokręta VC1 i VC2 mogą swobodnie poruszać się w obudowie, gdy pokrywa jest założona. Jeśli się blokują, być może trzeba pogłębić szczeliny, w których



Na tym zdjęciu widać płytkę drukowaną zamontowaną na pokrywie obudowy, gotową do montażu. Płytkę drukowaną „zwisa” z pokrywy obudowy, a elementy znajdują się pod nią. Antena tonowa przechodzi przez pokrywę, przez odpowiedni otwór w płytce drukowanej i jest utrzymywana na miejscu za pomocą sprężynowych klipsów bezpiecznikowych widocznych w pobliżu wyłącznika zasilania (po lewej stronie)

są osadzone. Jeśli wszystko jest w porządku, przy-
mocujemy pokrywę do obudowy za pomocą śrub.

Regulacja wysokości dźwięku i głośności

Ustawmy VR1 w pozycji środkowej, podłączmy zasilacz i włączmy go. Wyregulujemy pokrętkę głośności i pokrętkę tonu do momentu usłyszenia dźwięku, a następnie ustawmy pokrętkę głośności tak, aby dźwięk był słyszalny nawet wtedy, gdy ręka znajduje się blisko płytki. Ustawmy pokrętkę wysokości dźwięku palcem wskazującym lewej ręki i dłonią nad płytką głośności. W ten sposób ręka jest trzymana z dala od anteny tonowej.

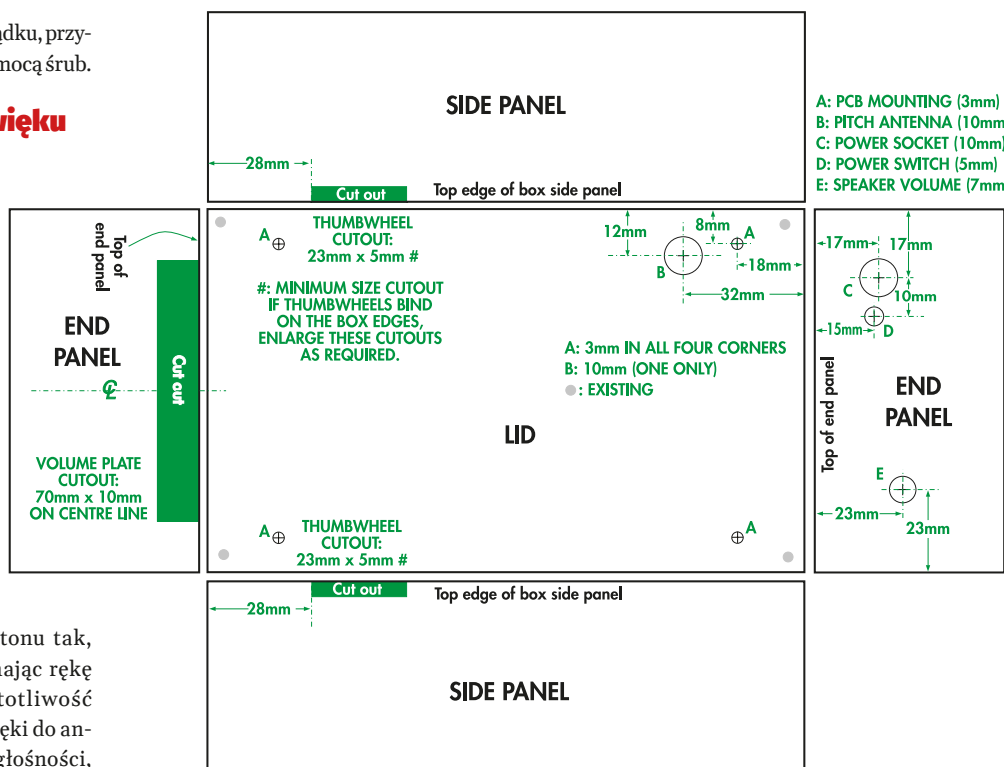
Wyregulujemy trymer pokrętki tonu tak, aby uzyskać zero tętnienia, trzymając rękę z dala od anteny tonowej. Częstotliwość powinna rosnąć w miarę zbliżania ręki do anteny. Trzymając rękę blisko płytki głośności, wyregulujemy trymer regulacji głośności, aby uzyskać minimalną głośność.

Należy pamiętać, że nie jest możliwe całkowite wyłączenie dźwięku za pomocą regulatora głośności z powodów, które zostały już wyjaśnione. Regulacje te trzeba będzie powtarzać za każdym razem, gdy urządzenie zostanie ustawione w innym położeniu.

Przekonamy się, że Theremin potrafi wydobywać nieskończenie wiele różnych dźwięków. Szybkie, zamazyste ruchy rąk mogą wywoływać niskie pomruki, buczenie i warknięcia. Podobnie można uzyskać zawrozenia i piski w zakresie wysokich tonów.

Aby uzyskać efekt vibrato, należy trzymać rękę od głośności w stałej pozycji, a ręką od wysokości tonu trzepotać blisko anteny z żądaną szybkością.

Drobniejsze zmiany można wprowadzać, poruszając palcami, podczas gdy ręka pozostaje nieruchoma. Podobnie, aby uzyskać efekt tremolo, należy trzymać rękę od tonu



Ten schemat wiercenia/wycinania dla obudowy UB1 Jiffy box jest przedstawiony w skali 1:2, więc aby użyć go jako szablonu, należy go dwukrotnie powiększyć



Panel przedni został przyklejony do pokrywy obudowy, aby uzyskać naprawdę profesjonalne wykończenie. Można go również pobrać ze strony siliconchip.com.au, jeśli chcemy go wydrukować na papierze o większej gramaturze lub błyszczącym

w stałej pozycji i poruszać ręką z głośnością (w naszych przykładach (patrz ramka) wiadać, że dwie panie grające na Thereminie w dużym stopniu używają palców).

Jak wspomnieliśmy wcześniej, jeśli chcemy zmienić stopień wzajemnej modulacji oscylatorów dźwięku, możemy dodać pojemność pomiędzy emiterami Q1 i Q2, pokazaną na schemacie i na płytce drukowanej jako C1.

Dobrym punktem wyjścia do eksperymentów są wartości od 220 pF do 470 pF, ale można je zwiększyć lub zmniejszyć bez żadnego ryzyka. ■

John Clarke

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



Moduł alarmu znajduje się wewnątrz sejfów i uruchamia alarm w przypadku otwarcia drzwi sejfów, jeśli przed upływem zadanego czasu nie zostanie wprowadzony prawidłowy kod

Alarm do sejfów hotelowych dla podróżnych

Jeżeli turystyka nie jest nam obca, to pewnie znamy się już z małymi sejfami znajdującymi się w każdym pokoju w większości hoteli i w kabinach na statkach wycieczkowych. Opisany tu alarm do sejfów hotelowych informuje, że sejf został otwarty pod nieobecność użytkownika, a także daje sprawcy bardzo złe przeczucie, że został wykryty. Ich naturalną reakcją będzie natychmiastowe zamknięcie sejfów i ucieczka.

Szablony do produkcji PCB i gotowe płytki PCB dostępne są na stronie:
<https://bit.ly/3snhDxC>
Materiały dodatkowe są również dostępne na stronie edw.elportal.pl:
<https://bit.ly/3wuGemZ>

Każdy, kto regularnie podróżuje statkami wycieczkowymi lub zatrzymuje się w hotelach, zna wszechobecny sejf pokojowy, który zazwyczaj znajduje się w szafie. Sejf jest wyposażony w 4-cyfrowy wyświetlacz LED oraz klawiaturę numeryczną, która umożliwia wprowadzenie 4-cyfrowego kodu przed zamknięciem sejfów i ponownie przed jego otwarciem. Są one bardzo poręczne, ale naiwnością byłoby sądzić, że zapewniają wysoki stopień bezpieczeństwa naszym kosztownościom. W końcu, jeśli zapomnimy kodu lub sejf ulegnie awarii, otwarcie go przez personel hotelowy jest dość prostym zadaniem. Oznacza to, że w korytarzach hoteli lub statków mogą czaić się osoby, którym nie leży na sercu dobro pasażera. A skoro będą próbować swoich podstępnych działań pod naszą nieobecność, jak możemy ich do tego zniechęcić? Odpowiedź brzmi: skorzystać z naszego Alarmu sejfów hotelowych.

Oczywiście alarm ten można także umieścić w sejfie domowym, w szafce na dokumenty lub

w szufladzie biurka, które mają być monitorowane. Można go też użyć do monitorowania szafki z narzędziami, spiżarni (przed głodnymi nastolatkami grasującymi nocą?) lub innych miejsc.

Alarm sejfów hotelowych to małe plastikowe pudełko z dwoma LEDami (czerwonym i zielonym) oraz dwoma przyciskami. Fotorezystor (LDR) wykrywa otwarcie sejfów i reaguje na światło w pomieszczeniu lub latarkę. LED zaczyna natychmiast migać i jeśli w ciągu 15 sekund nie wprowadzi się kodu za pomocą dwóch przycisków, wbudowany przetwornik piezoelektryczny zacznie krzyczeć na nas (lub sprawcę).

Czas trwania alarmu wynosi 60 sekund (ustawienie domyślne), ale można go ustawić w zakresie od 10 sekund do 120 sekund, w dziesięciosiekundowych odstępach.

Jeśli sejf został otwarty podczas nieobecności użytkownika, alarm zasygnalizuje to miganiem na przemian czerwonego i zielonego LEDa. Aby

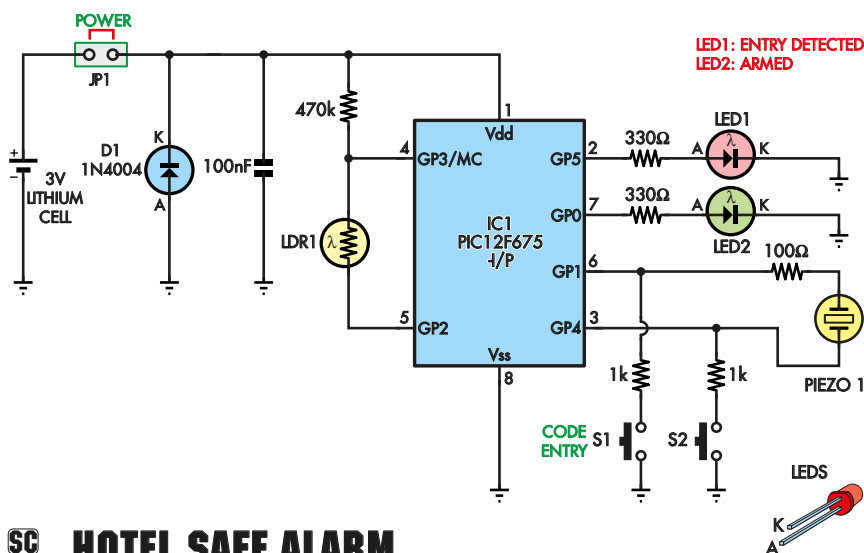
anulować ten stan alarmowy, wystarczy wprowadzić kod, naciskając, w normalny sposób, dwa przyciski.

Sposób wprowadzania kodu i różne ustawienia czasu zostaną opisane w dalszej części artykułu.

Opis układu

Układ jest bardzo prosty; wystarczy 8-pinowy mikrokontroler PIC, dwa LEDy, dwa przyciski i kilka innych elementów – patrz rysunek 1. Jest on zasilany z baterii litowej 3 V i włączany za pomocą zworki JP1. Można ją usunąć, gdy alarm nie jest używany, aby oszczędzać baterię.

Układ scalony IC1 to mikrokontroler PIC12F675-I/P, z wgranym programem, które umożliwia wprowadzanie niezbędnych ustawień za pomocą tylko dwóch przycisków. W normalnych warunkach układ scalony IC1 znajduje się w trybie uśpienia, a jego zegar sterujący budzi go co 2,3 s i przez krótki czas



SC 2016 HOTEL SAFE ALARM

Rysunek 1. Układ składa się z mikrokontrolera PIC IC1, fotorezystora LDR1, kilku LEDów i przetwornika piezoelektrycznego. Jeśli sejf zostanie otwarty, rezystancja LDR1 obniża się i powoduje przełączenie pinu 4 układu scalonego IC1 w stan niski, co uruchamia timer alarmu. Następnie należy wprowadzić prawidłowy kod w ciągu 15 sekund za pomocą przełączników przyciskowych S1 i S2 w celu wyłączenia alarmu

Specyfikacje

Zasilanie: 3 V przy typowym prądzie 2,5 μ A

Prąd alarmu: 0,5 mA

Opóźnienie alarmu (czas na wpisanie kodu): regulowane w zakresie 1–60 s w krokach co 1s; wartość początkowa wynosi 15 s

Czas trwania alarmu: 10 s do 120 s, w 10 s interwałach. Wartość domyślna to 60 s

Kod wyłączenia alarmu: dowolna sekwencja kodów od jednego do ośmiu naciśnień przełączników

Sygnal alarmowy: serie tonów trwające 280 ms o częstotliwości 4–6 kHz, z przerwą 220 ms między seriami

sprawdza natężenie oświetlenia otoczenia za pomocą fotorezystora (LDR). Dzieje się to w następujący sposób.

W normalnych warunkach wyjście GP2 układu IC1 na styku 5 jest ustawione w stanie wysokim (przy napięciu 3 V), więc przez rezystor 470 k Ω i LDR nie płynie żaden prąd. Ma to na celu zminimalizowanie prądu pobieranego z ogniwa 3 V. Po przebudzeniu IC1 ustawia GP2 w stan niski (0 V), a następnie monitoruje napięcie na wejściu GP3 (styk 4).

W ciemności, rezystancja LDR jest duża (znacznie powyżej 1 M Ω), więc napięcie na styku 4 będzie wysokie, bliskie 3 V, więc układ scalony IC1 ziewnie i ponownie przejdzie w stan uśpienia. Po przebudzeniu i wystawieniu LDR na działanie światła otoczenia jego rezystancja będzie znacznie niższa, być może nawet 10 k Ω w jasnym świetle. W związku z tym napięcie na styku 4 będzie niskie i układ scalony IC1 zacznie się wzbudzać. Cóż, może nie, ale zaczyna migać zielona dioda LED2, sygnalizując, że alarm zaraz się włączy.

Jeśli w ciągu 15-sekundowego opóźnienia za pomocą obu przycisków zostanie

wprowadzony prawidłowy kod, alarm zostanie wyłączony. Jeśli nie zostanie wprowadzony żaden kod lub zostanie wprowadzony kod nieprawidłowy, przetwornik piezoelektryczny wydaje dźwięk, ponieważ styki 6 i 3 (GP1 i GP4) na przemian ustawiają się w stan wysoki i niski, co powoduje emisję impulsów o częstotliwości 4 kHz. W zamkniętej przestrzeni sejfu hotelowego i w bliskim sąsiedztwie może to być dość głośne.

Z pewnością nie ma wątpliwości, że sprawca wykroczenia został „wykryty”. Jak już wspomniano, alarm będzie emitował dźwięk przez domyślny okres 60 sekund (chyba że zostanie zaprogramowany inaczej).

Zrzuty ekranu z rysunków 2 i 3 pokazują komplementarne sygnały sterujące przyłożone do przetwornika piezoelektrycznego. Na rysunku 2 oba sygnały mają częstotliwość 3,99 kHz i amplitudę bardzo zbliżoną do 3 V międzyszczytowo, bez uwzględnienia skoków nadmiarowych. Dlatego całkowity sygnał przyłożony do przetwornika będzie bardzo bliski 6 V międzyszczytowego lub około 3 V RMS, jak pokazano na czerwonym przebiegu na rysunku 2.

Wykaz elementów:

- 1× dwustronna płytka drukowana, kod 03106161, 61×47 mm
- 1× etykieta na panel przedni, 74×47 mm
- 1× przezroczysta lub niebieska obudowa UB5, 83×54×31 mm
- 1× koszyk baterii CR2032 do druku
- 1× ogniwo litowe CR2032 (3 V)
- 2× przełączniki SPST do montażu na płytce drukowanej (S1, S2)
- 1× przetwornik piezoelektryczny o średnicy 30 mm
- 1× fotorezystor 10 k Ω
- 1× podstawa DIL8
- 2× odstępniaki M3 z gwintem 12 mm
- 2× odstępniaki M3 z gwintem 6 mm
- 6× wkrętów maszynowych M3×6 mm
- 2× wkręty maszynowe M3×6 mm
- 2× wkręty z łbem stożkowym M3×6 mm
- 1× złącze dwustykowe (rozstaw pinów 2,54 mm) (JP1)
- 1× boczniak zworkowy
- 2× kołki PC
- 1× opaska termokurczliwa o długości 25 mm i średnicy 2 mm

Półprzewodniki:

- 1× PIC12F675-I/P zaprogramowany kodem 0310616A. hex (IC1)
- 1× dioda 1N4004 (D1)
- 1× czerwony LED 3 mm o wysokiej jasności (LED1)
- 1× zielony LED 3 mm o wysokiej jasności (LED2)

Kondensator:

- 1× 100 nF 63 V lub 100 V MKT, poliestrowy lub ceramiczny

Rezystory:

- 1× 470 k Ω
- 2× 330 Ω
- 2× 1 k Ω
- 1× 100 Ω

Na rysunku 3 przedstawiono te same komplementarne sygnały sterujące, ale przy znacznie mniejszej szybkości przebiegania 100 ms/div. Widać na nim wybuchy sygnału o długości około 280 ms, oddzielone przerwami o długości około 220 ms.

Jeśli drzwi sejfu zostaną ponownie pociągnięte, alarm będzie emitował dźwięk przez pozostałą część 60-sekundowego okresu, a następnie przejdzie w stan uśpienia. Po ponownym otwarciu drzwi sejfu czerwony i zielony LED będą migać naprzemiennie przez 15 sekund, chyba że zostanie wprowadzony prawidłowy kod. Jeśli nie, alarm piezoelektryczny zacznie ponownie emitować sygnał dźwiękowy. I tak przez pełen cykl... W ten sposób nie tylko odstrasza, uruchamiając alarm, jeśli nie zostanie wprowadzony prawidłowy kod, ale także informuje o otwarciu sejfu pod nieobecność użytkownika, nawet jeśli został on zamknięty po wykryciu.

Funkcje

- Zasilanie 3 V z baterii CR2032
- Sygnalizacja uzbrojenia (zielona dioda LED2)
- Wskaźnik wejścia (czerwona dioda LED1)
- Alarm piezoelektryczny
- Mały pobór prądu: w stanie czuwania 2,5 μ A
- Regulowany czas opóźnienia alarmu
- Regulowany czas trwania alarmu



Rysunek 2. Żółte i zielone przebiegi pokazują komplementarne sygnały sterujące przyłożone do przetwornika piezoelektrycznego. Oba sygnały mają częstotliwość 3,99 kHz i amplitudę międzyszczytową zbliżoną do 3 V, bez uwzględnienia skoków przesterowania. Czerwony sygnał przyłożony do przetwornika jest pokazany na czerwonym przebiegu i wynosi międzyszczytowo 6 V



Rysunek 3. Ten oscylogram pokazuje te same sygnały komplementarne, ale przy znacznie mniejszej szybkości przemiatania (100 ms / div). Sygnały mają długość ok. 280 ms i są oddzielone przerwami o długości ok. 220 ms. Czerwony przebieg pokazuje całkowity sygnał przyłożony do przetwornika i wynosi międzyszczytowo 6 V

Wykrywanie stanu przycisków

Oprócz dostarczania sygnału sterującego dla przetwornika piezoelektrycznego, styki 6 i 3 (GP1 i GP4) monitorują stan dwóch przycisków z zestykiem chwilowym S1 i S2. W tym celu GP1 i GP4 są ustawione jako wejścia, które są w stanie normalnie wysokim, ale można je ściągnąć w stan niski za pomocą rezystorów 1 kΩ połączonych szeregowo z przełącznikami. Jeśli więc S1 jest zamknięty, styk 6 (GP1) zostanie ściągnięty do stanu niskiego.

Rezystory 1 kΩ są dołączone, aby naciśnięcie przełączników w czasie trwania alarmu nie spowodowało zwarcia sygnału alarmowego podawanego do przetwornika piezoelektrycznego.

Zasilanie z akumulatora

Jak już wspomniano, układ jest zasilany z baterii 3 V przez złącze JP1. Gdy IC1 znajduje się w trybie uśpienia, natężenie prądu jest dość niskie i wynosi około 2,5 µA.

Pobór prądu podczas włączania się alarmu piezoelektrycznego wynosi 0,5 mA. Gdy dioda LED2 miga, prąd wynosi 1,5 mA (przy napięciu ogniwa 3 V).

Dioda D1 jest dołączona jako zabezpieczenie przed uszkodzeniem układu IC1 w przypadku nieprawidłowego podłączenia ogniwa. Jeśli biegunowość jest niewłaściwa, D1 przewodzi, a więc „zwiera” ogniwo.

Odwrotna polaryzacja ogniwa może wystąpić, jeśli mocowanie ogniwa zostanie zamontowane odwrotnie. Alternatywnie, jeśli mocowanie ogniwa jest zainstalowane prawidłowo, dioda chroni obwód w przypadku nieprawidłowego zamontowania samego ogniwa. Zwróćmy uwagę, że w przypadku używanego przez nas uchwytu na ogniwa nie ma możliwości, aby ogniwo zostało włożone niepoprawnie i nawiązało połączenie z obwodem.

Zasilanie układu IC1 jest bocznikowane kondensatorem 100 nF, a układ IC1 pracuje

z wykorzystaniem wewnętrznego oscylatora 4 MHz, który jest wyłączany w trybie uśpienia.

Jasność diody LED2 informuje o napięciu ogniwa. Przy napięciu zasilania 3 V dioda LED2 świeci dość jasno, ale gdy napięcie ogniwa spadnie do 2 V, będzie przygasać, co oznacza, że należy ogniwo wymienić.

Sztuczki programistyczne

Należy pamiętać, że wejście GP3 układu PIC12F675 jest zwykle skonfigurowane jako wejście MCLR (master clear), co pozwala mikrokontrolerowi na zewnętrzny reset po włączeniu zasilania. Jednak w naszym układzie musimy użyć tego wejścia jako wejścia ogólnego przeznaczenia do monitorowania fotorezystora.

Gdy MCLR jest ustawiony jako wejście, operacja MCLR jest przełączana na wewnętrzne połączenie w mikrokontrolerze, dzięki czemu nie jest tracona funkcja zerowania zasilania master clear.

Jedną z wad używania pinu MCLR jako wejścia ogólnego przeznaczenia jest to, że może wystąpić problem podczas programowania mikrokontrolera. Dzieje się tak, gdy wewnętrzny oscylator jest używany również do zasilania mikrokontrolera (co ma miejsce w naszym przypadku). Problem ten rozwiązaliśmy w oprogramowaniu, co zostało omówione w ramce dot. programowania.

Montaż płytki drukowanej

Wszystkie elementy są zamontowane na małej dwustronnej płytce drukowanej oznaczonej kodem 03106161 (61×47 mm). Mieści się ona w małej plastikowej obudowie (UB5). Zwróćmy uwagę, że LEDy, przełączniki, fotorezystor i przetwornik piezoelektryczny są zamontowane po jednej stronie płytki drukowanej, a pozostałe elementy, po drugiej.

Na rysunku 4 pokazano rozmieszczenie elementów po obu stronach płytki drukowanej.

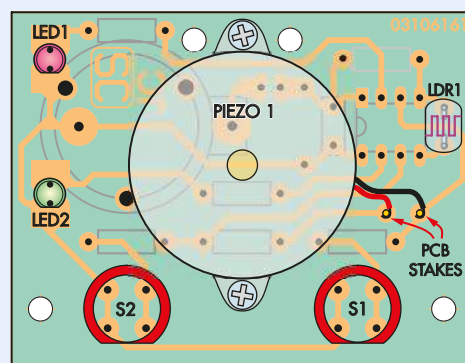
Rozpocznijmy konstrukcję od wlutowania rezystorów, sprawdzając multimetrem wartość każdego z nich przed włożeniem go na miejsce. Teraz można zamontować diodę D1, zwracając uwagę na jej prawidłowy kierunek, a następnie dopasować gniazdo układu scalonego, ustawiając wycięcie tak jak pokazano na rysunku 4. Następnie wlutowuje się kondensator 100 nF. Potem wlutowujemy dwustykowe złącze nagłówkowe (header) do zwory JP1 wraz z mocowaniem ogniwa. Upewnijmy się, że zacisk plusowy jest skierowany w stronę diody D1 na płytce drukowanej.

Diody LED1 (czerwona) i LED2 (zielona) są zamontowane w taki sposób, że górna część soczewki LEDa znajduje się 14 mm nad górną powierzchnią płytki drukowanej. Upewnijmy się, że dłuższy przewód każdego LEDa (anoda) jest umieszczony w pozycji „A” na płytce drukowanej. Fotorezystor jest również zamontowany 14 mm nad powierzchnią płytki drukowanej. Po zamontowaniu LEDów należy zainstalować przełączniki S1 i S2, zwracając uwagę na ich prawidłowe ułożenie (płaska strona w kierunku pokazanym na rysunku).

Montaż przetwornika piezoelektrycznego

Przetwornik piezoelektryczny jest zamontowany z zachowaniem odstępu od płytki drukowanej, wsparty na odstępnikach M3×6 mm i zamocowany za pomocą śrub M3. W otworach montażowych w uchwytach przetwornika piezoelektrycznego należy wywiercić otwory o średnicy 3 mm, aby można było użyć tych śrub. Przewody są przylutowane do kołków PC oznaczonych na płytce drukowanej jako „piezo”.

Do podłączenia przetwornika piezoelektrycznego użyliśmy kołków PC, ponieważ umożliwiają one nasunięcie opaski termokurczliwej na przewody, a kołki PC zapobiegają wyrwaniu przewodów.



eprasa.pl 0d581f01a1

Programowanie mikroprocesora PIC

Zaprogramowany do tego projektu układ PIC można kupić w sklepie internetowym Silicon Chip (www.siliconchip.com.au) lub zaprogramować go samodzielnie. Oprogramowanie jest również dostępne na stronie internetowej Silicon Chip.

Jeśli programujemy mikrokontroler samodzielnie, możemy spotkać się z ostrzeżeniem programatora, że programowanie nie jest obsługiwane, gdy pin MCLR jest ustawiony jako wejście ogólnego przeznaczenia i używany jest wewnętrzny oscylator. Zignorujmy jednak to ostrzeżenie. Dzieje się tak dlatego, że wszelkie problemy związane z taką konfiguracją są już rozwiązane w oprogramowaniu. Jeśli chcemy poznać więcej szczegółów, należy czytać dalej.

Jak wspomniano, ustawiamy MCLR jako wejście ogólnego przeznaczenia i wykorzystujemy wewnętrzny oscylator w układzie IC1. Może to stanowić problem dla programatora podczas procesu weryfikacji kodu oprogramowania po zakończeniu programowania.

Problem polega na tym, że gdy tylko mikrokontroler zostanie zaprogramowany, zacznie wykonywać swój program. Typowy program początkowo konfiguruje mikrokontroler z liniami

ogólnego przeznaczenia ustawionymi jako wejścia lub wyjścia (I/O). Jest to sprzeczne z potrzebą wykorzystania przez programator linii we/wy do programowania zegara i danych do weryfikacji programu.

Ten problem nie występuje, jeśli pin MCLR jest ustawiony jako zewnętrzne wejście MCLR, ponieważ programator ma wtedy kontrolę nad mikrokontrolerem, powstrzymując go przed wykonaniem zaprogramowanego kodu. Należy również pamiętać, że aby kod mógł zostać uruchomiony, mikrokontroler musi mieć skonfigurowany wewnętrzny oscylator, a nie zewnętrzny kwarc, RC lub zewnętrzny oscylator zegarowy.

Problem programowania został rozwiązany w dostarczonym oprogramowaniu poprzez włączenie trzysekundowego opóźnienia na początek programu. Opóźnienie to występuje przed ustawieniem linii we/wy jako wejść lub wyjść. Linie we/wy pozostają zatem jako wejścia o wysokiej impedancji, podczas gdy programator weryfikuje wewnętrznie zaprogramowany kod za pomocą linii programowania zegara i danych.

Ostrzeżenie od programatora nadal będzie emitowane, ale mikrokontroler

może zostać pomyślnie zaprogramowany i poprawnie zweryfikowany przez programator.

Należy pamiętać, że układ PIC12F675 również wymaga specjalnego programowania, ponieważ posiada wartość kalibracji oscylatora (OS-CAL), która jest przechowywana w pamięci układu PIC. Ta wartość kalibracji jest indywidualnie programowana w każdym układzie PIC przez producenta i zapewnia wartość, która ustawia układ PIC na pracę z dokładną częstotliwością 4 MHz przy użyciu wewnętrznego oscylatora.

Wartość ta musi być odczytana przed skasowaniem i zaprogramowaniem, aby podczas programowania mogła być dołączona do reszty kodu. Jeśli ta procedura nie zostanie wykonana, oscylator może być poza częstotliwością, co będzie miało wpływ na dzwiek hotelowego alarmu sejfowego.

Większość programatorów PIC automatycznie uwzględnia tę wartość OSCAL, ale warto sprawdzić, czy programator poprawnie obsługuje tę funkcję, zwłaszcza jeśli występują trudności. Na koniec należy pamiętać, że układ PIC12F675 wymaga zasilania 5 V do programowania, mimo że w układzie jest zasilany napięciem 3 V.

Stan baterii litowej można kontrolować, obserwując diodę LED2. Jeśli miga jasno, oznacza to, że z ogniwnem jest wszystko w porządku. W miarę rozładowywania się ogniwa, LED będzie świecić słabiej.

Zmiana ustawień

Ustawienia alarmu do sejfów hotelowych można zmienić w trzech aspektach: opóźnienie alarmu, czas trwania alarmu i kod dostępu. Można je zmienić tylko po wyłączeniu alarmu przez wyjęcie zworki JP1, a następnie naciśnięciu jednego lub obu przełączników przy ponownie włożonej zworki JP1 w celu podłączenia zasilania.

Zmiana opóźnienia wejścia i czasu trwania alarmu jest opcjonalna i można pozostawić

ustawienia domyślne, odpowiednio 15 i 60 sekund. Konieczne będzie jednak ustawienie kodu dostępu.

Opóźnienie alarmu

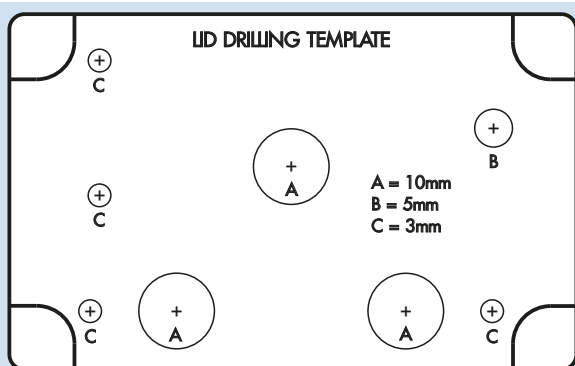
Aby ustawić opóźnienie (czas na wpisanie kodu), należy wyłączyć zasilanie urządzenia przez wyjęcie zworki JP1 i przytrzymać wciśnięty przełącznik S1, gdy zworka JP1 jest wkładana. Przytrzymajmy wciśnięty przycisk S1 do momentu usłyszenia krótkiego sygnału dźwiękowego z przetwornika piezoelektrycznego (po około trzech sekundach). Po zwolnieniu przycisku S1 usłyszymy kolejny sygnał dźwiękowy. Po naciśnięciu przełącznika S2 można wprowadzić czas opóźnienia. Zaczyna się od jednej sekundy (plus

początkowy czas budzenia – 2,3 sekundy), a każde naciśnięcie przycisku S2 powoduje bardzo krótki, podwójny sygnał dźwiękowy z modułu piezoelektrycznego, który oznacza, że opóźnienie wejścia zostało zwiększone o jedną sekundę.

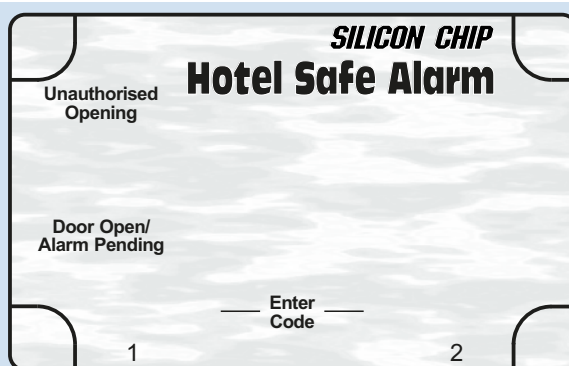
Opóźnienie można zwiększyć do 60 sekund, ale naszym zdaniem 15 sekund jest wystarczające. Następnie należy zapisać ustawienie opóźnienia alarmu, naciskając przycisk S1, co zostanie zasygnalizowane krótkim sygnałem dźwiękowym przetwornika.

Czas trwania alarmu

Proces ustawiania czasu trwania alarmu jest bardzo podobny do ustawiania czasu na wpisanie kodu, ale teraz robimy to za pomocą



Rysunek 5. Ten szablon wiercenia można pobrać jako plik PDF ze strony internetowej Silicon Chip



Rysunek 6. Ten projekt graficzny panelu przedniego jest również dostępny w formacie PDF na stronie internetowej silicon chip (patrz ramka)

przełącznika S2. Aby ustawić opóźnienie (czas na wpisanie kodu), należy wyłączyć zasilanie urządzenia przez wyjęcie zworki JP1 i przytrzymać wciśnięty przełącznik S2, gdy zworka JP1 jest wkładana. Przytrzymujemy wciśnięty przycisk S2 do momentu usłyszenia krótkiego sygnału dźwiękowego (po około trzech sekundach). Po zwolnieniu przycisku S2 usłyszymy kolejny sygnał dźwiękowy.

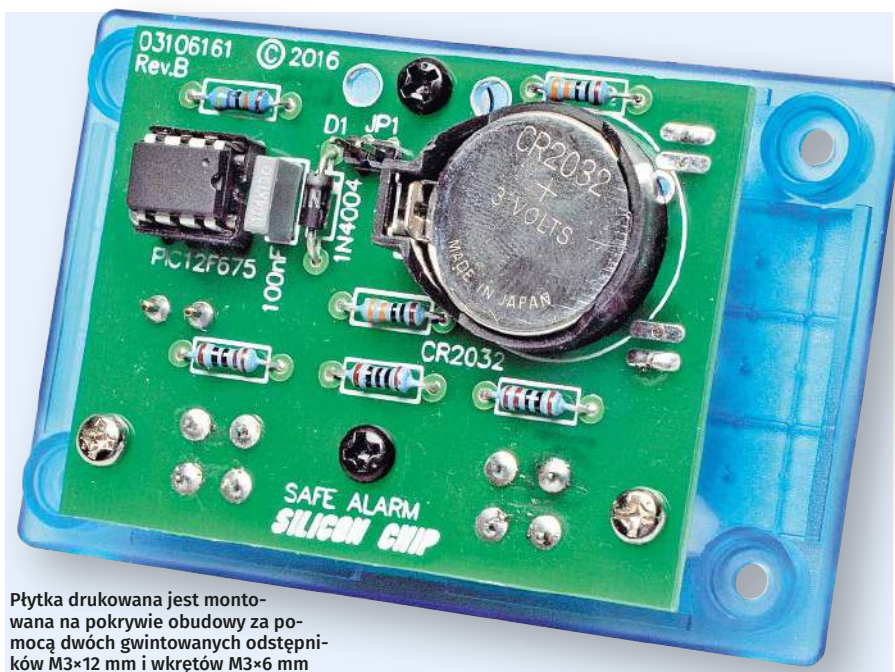
Czas trwania alarmu jest wprowadzany przez naciśnięcie przełącznika S1. Czas trwania alarmu zaczyna się od 10 sekund, a każde naciśnięcie przycisku S1 jest sygnalizowane krótkim sygnałem dźwiękowym z modułu piezoelektrycznego, informującym, że czas trwania alarmu został wydłużony o 10 sekund. Czas trwania alarmu można regulować w zakresie od 10 do 120 sekund w 10-sekundowych odstępach. Po naciśnięciu przycisku S2 wprowadzony czas trwania alarmu zostanie zapisany w pamięci i zasygnalizowany krótkim sygnałem dźwiękowym z przetwornika piezoelektrycznego.

Kod dostępu

Kod dostępu składa się z sekwencji naciśnieć przycisków S1 i S2. Może to być tak proste, jak 1, 2 lub 2, 1, albo może to być nawet osiem naciśnieć, np. 1 2 2 2 1 2 1 2. Większość osób będzie chciała, aby był on w miarę krótki, tak aby łatwo go zapamiętać, np. 1221. Może to być jednak dowolna sekwencja od 1 do 8 naciśnieć.

W celu ustawienia kodu wejścia należy wyłączyć zasilanie urządzenia poprzez wyjęcie zworki JP1 i przytrzymać oba przełączniki S1 i S2 wciśnięte, gdy zworka JP1 jest wkładana. Kontynuujemy trzymanie wciśniętych przycisków S1 i S2 aż do usłyszenia krótkiego sygnału dźwiękowego z przetwornika piezoelektrycznego (po około trzech sekundach). Po zwolnieniu przycisków S1 i S2 usłyszymy kolejny sygnał dźwiękowy.

Teraz wprowadzany jest kod wejściowy, a każde naciśnięcie przełącznika jest potwierdzane krótkim sygnałem dźwiękowym.



Płytkę drukowaną jest montowana na pokrywie obudowy za pomocą dwóch gwintowanych odstępników M3×12 mm i wkrętów M3×6 mm

Wprowadzony kod zostanie zapamiętany po pozostawieniu obu przełączników w stanie otwartym (tzn. po tym jak żaden z nich nie będzie naciśnięty) przez pięć sekund. Następnie rozlegnie się sygnał dźwiękowy potwierdzenia.

Korzystanie z alarmu

Prawidłowy kod musi zostać wprowadzony w czasie wyznaczonym na opóźnienie alarmu. Nie próbujemy wprowadzać kodu zbyt szybko. Po każdym naciśnięciu przycisku należy poczekać na krótki sygnał dźwiękowy, a następnie nacisnąć kolejny przycisk. Na przykład, jeśli nasz kod to 1221, wykonajmy go w następującej kolejności: 1 sygnał dźwiękowy, 2 sygnał dźwiękowy, 2 sygnał dźwiękowy, 1 sygnał dźwiękowy. Jeśli kod jest poprawny, alarm nie włączy się (zielony LED przestaje migać po naciśnięciu przełącznika). Jeśli popełnimy błąd podczas wprowadzania kodu lub wprowadzimy go zbyt szybko, włączy się alarm, a sejf można zamknąć, aby stłumić jego dźwięk.

Wprowadzenie prawidłowego kodu zapobiega uruchomieniu alarmu tylko wtedy, gdy nie zostanie naciśnięty żaden inny przełącznik. Każde kolejne naciśnięcie przycisku po wprowadzeniu prawidłowego kodu spowoduje uruchomienie alarmu.

W przypadku wykrycia włamania, oba LEDy będą migać. Przeszaną one migać po naciśnięciu jednego z przełączników w celu rozpoczęcia sekwencji kodu dostępu. Zgaśnięcie LEDów może nawet dać intruzowi złudną nadzieję, że wprowadzony kod był prawidłowy.

Alarm jest ponownie uzbrajany po przejściu w stan ciemności, tzn. po zamknięciu drzwi sejfu. Gdy tylko światło zaświeci na fotorezystor, należy wprowadzić kod, aby zatrzymać alarm. ■

John Clarke

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

REKLAMA

Frezarka CNC, część 6

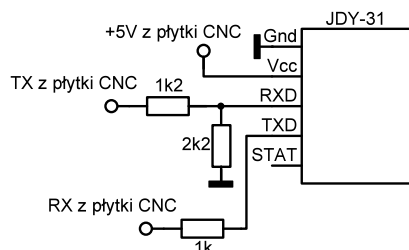
Materiały dodatkowe są dostępne na stronie [edw.elportal.pl](https://bit.ly/3MsiPb9):
<https://bit.ly/3MsiPb9>

Interesującą modyfikacją układu jest zastosowanie telefonu komórkowego do sterowania frezarką CNC. Oto dokładny opis jak to zrobić.

Możliwości modyfikacji

Jedną z możliwości modyfikacji układu jest zastosowanie telefonu komórkowego do sterowania frezarką CNC, zarówno w opcji ze sterowaniem przewodowym lub za pomocą modułu Bluetooth. Propozycję tę jeden z czytelników w listach kierowanych do redakcji uznał za śmieszną. Zdaniem autora, absolutnie tak nie jest! Współczesny telefon komórkowy to nic innego jak komputer o dużych możliwościach. Z punktu widzenia autora znacznie ważniejsze są dwa problemy: ograniczona ilość miejsca w warsztacie amatora oraz pyły powstające podczas frezowania, stanowiące zagrożenie dla komputera PC czy smartfona. Komputer PC, szczególnie w połączeniu z monitorem, zajmuje znacznie więcej miejsca niż telefon komórkowy. Znacznie ważniejszym problemem, jaki napotkamy stosując komputer PC, jest narażenie komputera na pył, który zawsze powstaje podczas frezowania. Pył ten, wciągany przez wentylatory komputera, osadza się na podzespołach w komputerze. W przypadku frezowania tworzy sztucznych pogarsza to chłodzenie podzespołów komputera, w przypadku frezowania metalu stanowi to wręcz zagrożenie uszkodzeniem, ze względu na możliwość wystąpienia zwarc. W przypadku telefonu komórkowego pyły osadzają się na mikrofonie i głośniku, pogarszając z czasem jakość komunikacji w takim telefonie (sprawdzone w praktyce...). Frezarki przemysłowe, które autor widział, mają zapewnione odpowiednie zabezpieczenia sterownika przed pyłami. Autor w celu zmniejszenia zapylenia w pomieszczeniu, w którym frezuje, stosuje miniszklarnię foliową, kupioną na jednym z portali aukcyjnych, nakładając ją na frezarkę podczas frezowania.

Moduł Bluetooth łączymy z płytą modułu CNC, jak pokazano na **rysunku 1**. Autor



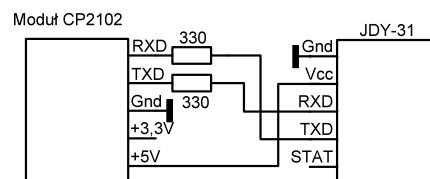
Rysunek 1.



Fotografia 2.

poleca użycie modułu Bluetooth typu JDY-31. Moduł ten może pełnić tylko funkcję Slave, tzn. nie może sam inicjować połączenia, jest jednak tańszy niż najczęściej wykorzystywany moduł HC-05 (konfigurowalny, może pełnić zarówno funkcję Master jak i Slave), czy moduł HC-06 (tylko Slave). Wspomniane moduły Bluetooth mogą być zasilane z napięcia w zakresie 3,3 do 6 V. W urzędzeniu autora moduł Bluetooth zasilany jest napięciem 5 V z płytki kontrolera CNC. Aby dopasować poziomy logiczne między torem RX Bluetooth a wyjściem TX modułu Arduino, zastosowano dzielnik rezystancyjny 1,2 kΩ/2,2 kΩ zapewniający dopasowanie poziomów napięć między oboma układami. W przypadku połączenia TX modułu Bluetooth i modułu RX modułu Arduino teoretycznie można nie stosować żadnych dodatkowych elementów, gdyż moduł Arduino zasilany z 5 V akceptuje sygnały logiczne o poziomach logicznych układu CMOS Bluetootha zasilanego z 3,3 V. Autor zastosował jednak pomiędzy pinem RX modułu Arduino a pinem TX modułu Bluetooth opornik ograniczający prąd, jaki może płynąć między oboma pinami. Ogranicza to możliwość uszkodzenia przez moduł Arduino modułu Bluetooth w wypadku gdyby wprowadzenie modułu Arduino zostało skonfigurowane jako wyjście, autor kilkakrotnie modyfikował program i zdarzało się, że nie wszystko poszło po jego myśli. Sygnały RX i TX modułu CNC wyprowadzone są na module CNC. W przypadku używania wersji GRBL skompilowanej dla prędkości transmisji 9600 modułu JDY-31, można użyć bezpośredniego połączenia do naszego urządzenia. Prędkość 9600 obsługują starsze wersje programu GRBL z wersji 0.9 oraz samodzielne kompilacje użytkowników, nowsze wersje programu z serii 1.1 pracują domyślnie z prędkością 115200.

Możemy również samodzielnie zmodyfikować źródła programu GRBL, tak aby uzyskać wymaganą prędkość transmisji. W tym celu pobieramy je z platformy GitHub (<https://github.com/grbl/grbl>), w pliku config.h znajdujemy frazę #define BAUD115200, ustawiamy wymaganą prędkość transmisji i kompilujemy program. Aby moduł Bluetooth współpracował z oprogramowaniem GRBL wykorzystującym prędkość transmisji 115200, musi być odpowiednio skonfigurowany. Moduł Bluetooth konfiguruje się za pomocą komend AT, wykorzystując konwerter USB na RS232 o poziomach logicznych CMOS. Autor poleca konwertery wykorzystujące układ CP2102, przykład takiego konwertera pokazany jest na **fotografii 2**. Aby skonfigurować moduł Bluetooth postępujemy następująco: łączymy moduł Bluetooth z modulem konwertera, jak pokazano na **rysunku 3**. Znajdujemy w systemie Windows **Menedżer urządzeń**, sprawdzamy czy w zakładce **Inne urządzenia** nie ma żadnych nieobsługiwanych urządzeń, jeśli takie wystąpią i znikną po włożeniu konwertera, ściągamy ze strony producenta układu scalonego, wykorzystanego w naszym konwerterze, odpowiednie sterowniki (w przypadku układu CP2102 jest to firma Silabs). Jeśli mamy prawidłowo zainstalowane sterowniki w zakładce **Porty (COM i LPT)**, po rozwinięciu listy widzimy zainstalowane sterowniki naszego konwertera. Numer portu COM przypisany do konwertera rozpoznamy



Rysunek 3.



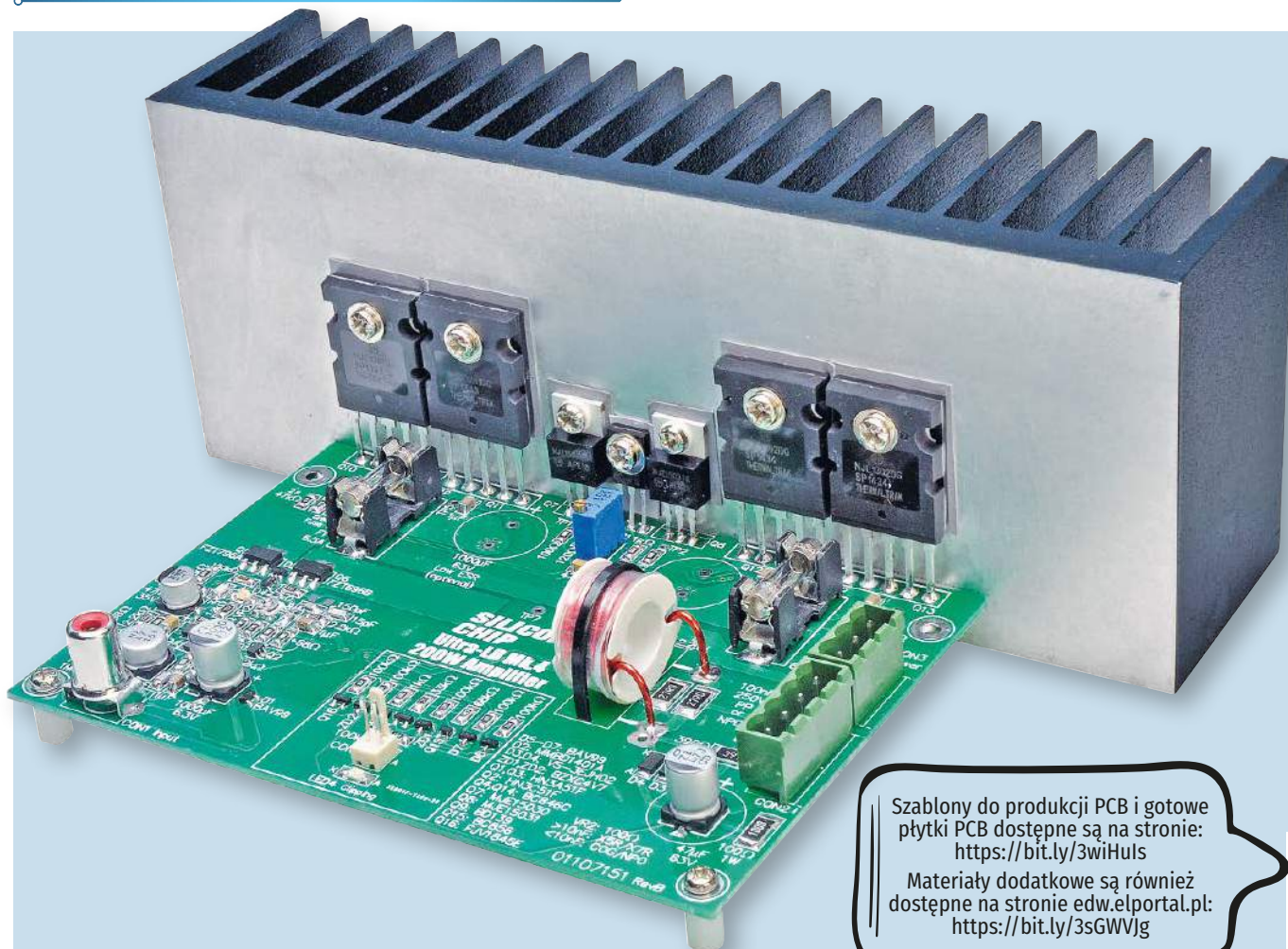
Fotografia 4.

po odłączeniu konwertera, sprawdzając o jakim numerze port COM zniknął. Uruchamiamy platformę Arduino, wybieramy opcję **Narzędzia, Monitor portu szeregowego**, zaznaczamy opcję **Zarówno NL, jak i CR** i ustawiamy prędkość transmisji równą **9600 baud**. Prędkość 9600 jest domyślną prędkością transmisji modułu Bluetooth JDY31. W górnym oknie wpisujemy ciąg znaków **AT+ Version** i naciskamy przycisk **Wyślij**. W dolnym oknie monitora portu szeregowego (terminala) pojawi się odpowiedź (jeśli nasz moduł Bluetooth jest sprawny) podobna do **+VERSION=JDY-31-V1.2,Bluetooth V3.0**. Teraz można ustawić prędkość 115200 w module Bluetooth, wysyłając komendę **AT+BAUD**. Moduł Bluetooth powinien odpowiedzieć **+BAUD=8**. Wartość 8 odpowiada prędkości 115200, inne wartości liczbowe odpowiadają innym prędkościom transmisji. Na zakończenie możemy sprawdzić czy moduł ma prawidłowo ustawiono prędkość transmisji. W tym celu przestawiamy w oknie terminala prędkość transmisji na 115200 i pytamy moduł Bluetooth np. o wersję modułu komendą **AT+Version**. W przypadku jakichkolwiek problemów z wysłaniem czy odczytaniem komend AT przez terminal należy zamknąć i ponownie otworzyć okno terminala. Niższe prędkości transmisji zapewniają większą niezawodność połączenia, ale kosztem wolniejszej reakcji frezarki CNC, co w przypadku opisanej konstrukcji nie ma praktycznie żadnego znaczenia. Największą niezawodność połączenia zapewnia połączenie kablowe. Autor do sterowania telefonem zastosował starszy smartfon z systemem Android ze złączem typu mikro USB. W przypadku połączenia frezarki ze smartfonem należy zastosować przewód wspierający technologię OTG (On-The-Go). W złączu mikro USB, oprócz pinów Vcc, Gnd, D+ i D- występuje dodatkowe wyprowadzenie Sense. W przewodach wspierających technologię OTG wyprowadzenie Sense połączone jest z wyprowadzeniem Gnd, natomiast w przewodach OTG niewspierających technologii OTG wyprowadzenie Sense jest

niepodłączone do innych wyprowadzeń. Stosując typowy przewód mikro USB-USB mamy na myśli przewód zakończony złączem USB typu A (jak np. w pendrivach). Moduł Arduino UNO wykorzystuje złącze typu B (drukarkowe), trzeba więc odciąć końcówkę USB typu A i zastąpić ją rozbierną końcówką typu B lub zastosować przejściówkę USB typ A na USB typ B pokazaną na **fotografii 4**. Autor do sterowania frezarką CNC za pomocą smartfona używa aplikacji **Grbl Controller**, aplikacja ta wymaga Androida w wersji 4.4 lub wyższej i współpracuje z oprogramowaniem GRBL od wersji 1.1f. Autor wykorzystuje program w wersji płatnej (koszt w chwili pisania artykułu 50 zł). Główną zaletą płatnej wersji w porównaniu do wersji darmowej, jest możliwość wznowiania zadań niemal od miejsca, w którym nastąpiło ich przerwanie i przeglądanie historii zadań. Sam interfejs aplikacji jest niezwykle prosty i intuicyjny. Strona domowa aplikacji Grbl Controller <https://zevy.github.io/grblcontroller/> zawiera dokładne opisy sposobów podłączenia modułu Bluetooth w wersji HC-05 i HC-06, a także bardzo dokładny opis interfejsu użytkownika. W przypadku frezowania z użyciem telefonu moduł nasz należy podłączyć pod ładowarkę gdyż sam proces frezowania może trwać dość długo, a starsze smartfony mają zwykle akumulator już w najgorszym stanie. W sklepie Google Play możemy znaleźć szereg innych aplikacji sterujących pracą frezarki CNC przez smartfona wpisując frazę **grbl controller**. Aby nasz moduł Arduino obsługiwał frezarkę CNC (a frezarka mogła w ogóle pracować) musimy wgrać do modułu Arduino UNO program GRBL. Pierwszą czynnością, jaką należy wykonać przed wgraniem programu, jest zapisanie wszystkich komórek EPROM mikrokontrolera wartościami 0. Odpowiednie sformatowanie EPROM-u jest ważne, gdyż w niej przetrzymywane są nastawy konfiguracyjne sterownika CNC, a domyślne wartości komórek EPROM równe 255 (FF w kodzie szesnastkowym) nie uprzyjemniają uruchamiania urządzenia... Aby sformatować EPROM, łączymy przewodem

USB komputer z płytą Arduino, otwieramy platformę Arduino, otwieramy plik znajdujący się w katalogu `eprom_cleaner` (<https://blog.protoner.co.nz/grbl-how-to-clear-eprom-settings>), (na Elportalu znajdują się wszystkie pliki i programy wykorzystane w tym artykule), otwieramy **Narzędzia, Płytkę: „x”**, gdzie x jest nazwą aktualnie wybranej płytki. Wybieramy płytkę **Arduino UNO** w opcji **Port**, ustawiamy numer portu przypisanego do płytki Arduino (numer portu przypisanego płytce Arduino UNO sprawdzamy analogicznie jak w przypadku wykrywania numeru portu przypisanego do konwertera USB-RS232), wybieramy zakładkę **Szkie i Wgraj**. Po wgraniu programu zobaczymy stosowny komunikat na ekranie platformy Arduino, a dioda podłączona do pinu PB5, pokazująca postęp wgrывania programu, zaświeci się na stałe. Następnie uruchamiamy program **Xloader**. Program Xloader umożliwia wgrывanie plików w formacie hex na modułach Arduino posiadających bootloader (moduły Arduino mają bootloader wgrany przez producenta modułu). Bootloader jest programem umożliwiającym wgrывanie programu bez posiadania programatora procesora. Aby wgrać plik hex programu do modułu Arduino, należy z prawej strony pola **Hex file** programu Xloader nacisnąć na przycisk z wielokropkiem, wskazać plik o nazwie `grbl_v1.1h.20190825.hex`, w polu **Device** wybrać płytkę **Uno(Atmega328)**, w polu **COM port** wybrać numer portu COM przypisany do naszej płytki Arduino, w polu **Baud rate** ustawić prędkość transmisji równą 57600, a następnie nacisnąć przycisk **Upload**. Po prawidłowym załadowaniu pliku hex program Xloader wyświetli komunikat potwierdzający wgrывanie programu. W przypadku, gdy wgrывanie programu zakończy się niepowodzeniem, należy wybrać inną prędkość transmisji. Płytki Arduino sprzedawane są z różnymi wersjami bootloadera, które wykorzystują różne prędkości transmisji. Najczęściej można kupić płytki Arduino UNO zawierające bootloader obsługujący prędkość transmisji 57600, najnowsze bootloadery wykorzystują transmisję o prędkości 115200, autor spotkał jednak płytkę z inną niż wspomniane prędkością transmisji. Mimo iż szybsze bootloadery są już dostępne od jakiegoś czasu, to i tak większość płytek posiada wolniejszy bootloader. W części siódmej tego artykułu (ostatniej) zajmiemy się konfiguracją i uruchomieniem frezarki CNC. Na zakończenie artykułu autor chciałby podziękować Waldkowi 3Z6AEF za uwagi do tego tekstu. ■

Jerzy Wilczewski
Rafał Orodziński
sq4avs@gmail.com



Szablony do produkcji PCB i gotowe płytki PCB dostępne są na stronie: <https://bit.ly/3wiHuls>
Materiały dodatkowe są również dostępne na stronie [edw.elportal.pl](https://bit.ly/3sGWVjg): <https://bit.ly/3sGWVjg>

Wzmacniacz mocy Ultra-LD 200 W RMS, część 2

W tej części przedstawiamy szczegóły konstrukcyjne modułu wzmacniacza o ultra niskich zniekształceniach. Większość elementów na płycie drukowanej to elementy do montażu powierzchniowego, co pozwoliło zachować zwartą konstrukcję i uzyskać niespotykane niski poziom zniekształceń. Uniknęliśmy elementów trudnych do lutowania.

W części pierwszej poznaliśmy funkcje, specyfikacje i schemat elektryczny wzmacniacza Ultra-LD Mk.4. W tej części omówimy niektóre aspekty projektu płytki drukowanej, opiszemy, w jaki sposób ją zmodyfikowaliśmy, aby poprawić jej działanie, a następnie przedstawimy przebieg montażu modułu.

Projekt płytki drukowanej

Jedną z zalet nowej płytki drukowanej w porównaniu z konstrukcjami Mk.2 i Mk.3 jest całkowite wyeliminowanie wszystkich wysokoprądowych przelotek, dzięki czemu nie ma już obaw, że przelotki ulegną uszkodzeniu w czasie awarii i nie są wymagane

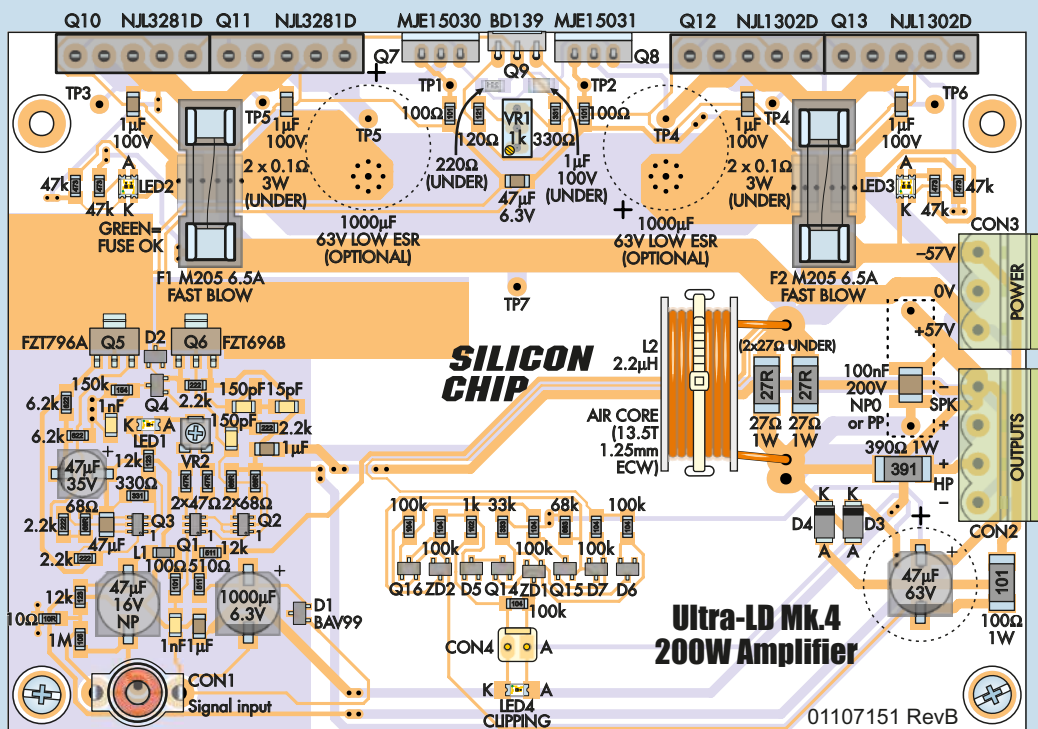
żadne przepusty dla połączeń na skroś płytki. Wszystkie ścieżki wysokoprądowe znajdują się po tej samej stronie płytki drukowanej.

Stopnie wejściowe są zmontowane w całości na górnej warstwie, z kilkoma przelotkami łączącymi elementy z analogową płaszczyzną masy pod spodem. Pozostałe przelotki są rozmieszczone parami (lub w większej ilości) w celu zapewnienia zapasu i są w większości związane z układem detektora przesterowań lub z niskoprądowymi szynami zasilającymi stopnie wejściowe układu.

Styki +57 V i -57 V złącza wejściowego zasilania CON3 są przylutowane do ścieżek

warstwy górnej, które biegną odpowiednio do podstawek bezpiecznikowych SMD F1 i F2 na górnej stronie płytki. Następnie łączą się one z dwoma kolejnymi ścieżkami warstwy wierzchniej, które idą do końcówek kolektorów tranzystorów wyjściowych.

Prąd wyjściowy na końcówkach emiterów płynie następnie ścieżkami na spodzie płytki drukowanej do rezystorów emiterowych SMD 0,1 Ω , które są zamontowane bezpośrednio pod oprawkami bezpieczników. Prąd ten jest następnie kierowany do kolejnej ścieżki na dolnej warstwie, która łączy prąd ze wszystkich czterech tranzystorów wyjściowych z powierzoną cewką indukcyjną L2. Następnie



Rysunek 6. Należy postępować zgodnie z tym schematem, aby zainstalować elementy na górnej części płytki drukowanej. Najpierw zamontujemy elementy SMD w kolejności podanej w tekście, a następnie odwróćmy płytkę drukowaną i zamontujemy elementy SMD od spodu, jak pokazano na rysunku 7. Następnie można zamontować pozostałe elementy przewlekane. Należy pamiętać, że Q7-Q13 są przylutowane do płytki drukowanej dopiero po przymocowaniu ich do radiatora

ścieżka dolnej warstwy łączy się z wyjściowym zaciskiem głośnika CON2.

Montaż

Dwustronna płytkę drukowaną, na której zbudowany jest Ultra-LD Mk.4, nosi kod 01107151 i ma wymiary 135×93 mm. Tranzystory wyjściowe są zamontowane na radiatorze z odlewane go ciśnieniowo aluminium z wykorzystaniem tego samego rozkładu, co we wzmacniaczu Ultra-LD Mk.3. Najłatwiejszym sposobem zbudowania modułu jest umieszczenie większości elementów SMD na górnej warstwie, następnie ośmiu elementów SMD na spodzie, potem pozostałych elementów SMD i elementów przelotowych na płycie, a na końcu tranzystorów na radiatorze.

Wszystkie elementy SMD można przylutować za pomocą zwykłej lutownicy (np. precyzyjnej) i drutu lutowniczego, jeśli ma się

trochę knota lutowniczego i pasty topnikowej. W zależności od naszego wzroku może być również potrzebna lampa powiększająca lub przyłbica. Można także użyć stacji do lutowania na gorące powietrze lub pieca rozplwowego, ale w takim przypadku należy ręcznie wlutować oprawki bezpieczników i LEDy, ponieważ mogą one ulec uszkodzeniu pod wpływem zbyt wysokiej temperatury.

Jeśli składamy urządzenie z zestawu, w którym elementy SMD są wstępnie przylutowane, możemy pominąć następną część artykułu.

Lutowanie elementów SMD

Na rysunku 6 pokazano rozmieszczenie elementów na górnej powierzchni płytki drukowanej. Rozpocznijmy od zainstalowania tranzystora Q2. Ma on najmniejszy rozstaw spośród elementów – 0,95 mm, ale nie jest to specjalnie trudne do lutowania. Najpierw wyjmijmy HN2C51F z opakowania (nie upuśćmy go!) i obejrzymy pod powiększeniem, aby zlokalizować pin nr 1 na górze.

Umieścimy go w pobliżu miejsca montażu, zachowując prawidłową orientację. Upewnijmy się, że jest ustawiony we właściwy sposób; wyprowadzenia powinny stykać się z płytką drukowaną. Następnie nanieśmy niewielką ilość lutu na jeden z narożnych padów na płycie, nie nakładając go na inne pady. Wyczyścimy lutownicę, drugą ręką delikatnie chwycimy część skośną pęsety,

podgrzejmy lut na tym padzie i wsuńmy element na miejsce.

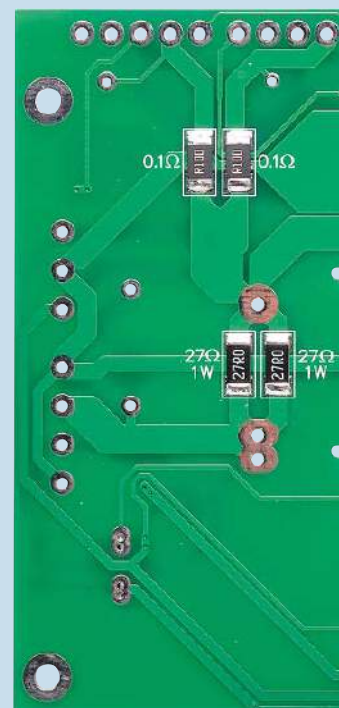
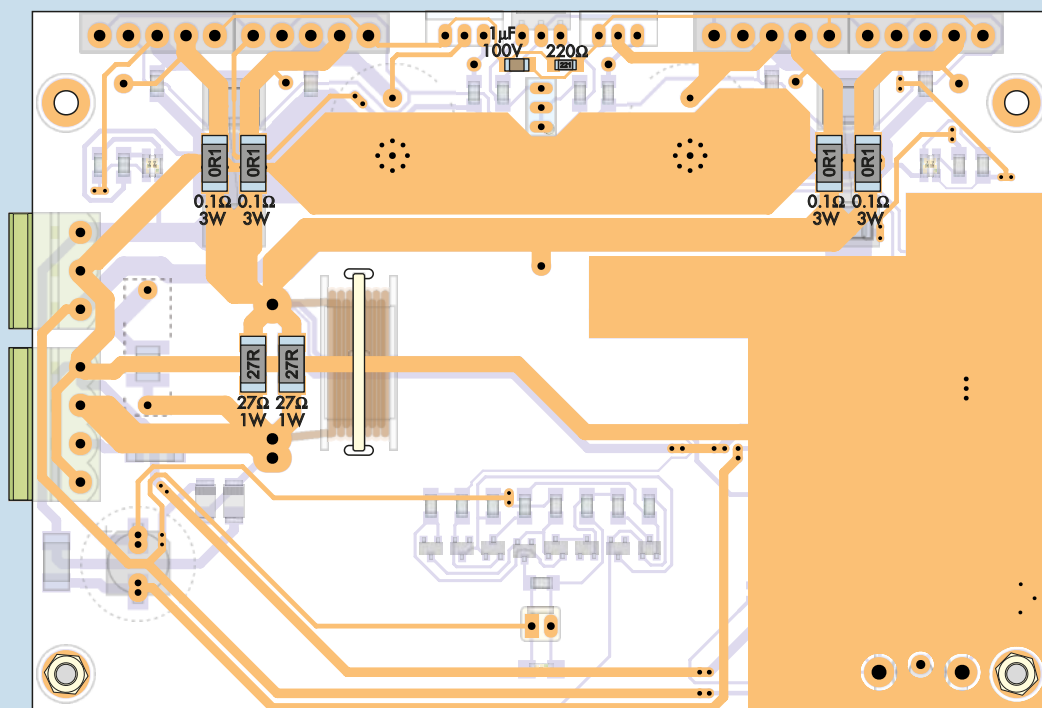
Odlóżmy lutownicę i sprawdźmy pod lupą lub podobnym narzędziem, czy wszystkie sześć pinów jest prawidłowo umieszczonych na swoich miejscach, czy kropka przy pinie nr 1 jest we właściwym miejscu i czy element leży płasko na płycie. Jeśli tak nie jest, podgrzejmy ponownie luty i usuńmy problem, delikatnie szturchając element. Na przykład, jeśli element nie leży płasko na płycie, naciśnijmy na niego (nie za mocno) pęsetą podczas podgrzewania lutów, a powinien wpaść na swoje miejsce. Ewentualnie, jeśli jest źle ustawiony, ostrożnie obróćmy lub przesunijmy go podczas podgrzewania lutu.

Po umieszczeniu go na miejscu przylutujemy styki po drugiej stronie układu. Nie przejmujmy się zbyttno ich zwarcie. Upewnijmy się tylko, że lut spłynie na wszystkie trzy piny i odpowiadające im styki. Następnie wykonajmy te same czynności dla trzech styków po drugiej stronie, w tym dla tego, który został przylutowany na początku.

Teraz wystarczy nanieść niewielką ilość pasty topnikowej po obu stronach układu scalonego, a następnie usunąć nadmiar lutu za pomocą knota lutowniczego. Wyczyścimy je za pomocą zmywacza do topników (wystarczy spirytus metylowy lub alkohol do nacierania) i obejrzymy pod powiększeniem, aby upewnić się, że wszystkie sześć lutów uformowało się prawidłowo. Należy pamiętać, że lut może przylgnąć do styku, nie spływając na płytkę drukowaną poniżej.

UWAGA!

Po włączeniu zasilania, w module wzmacniacza występują wysokie napięcia prądu stałego (np. ±57 V). W szczególności należy zwrócić uwagę, że pomiędzy dwoma szynami zasilającymi jest napięcie 114 V prądu stałego. Podczas pracy wzmacniacza nie należy dotykać przewodów zasilających (w tym bezpieczników), gdyż grozi to śmiertelnym porażeniem prądem.



Rysunek 7. Po zamontowaniu wszystkich elementów SMD na górnej stronie, należy odwrócić płytkę drukowaną i zgodnie z tym schematem rozmieścić osiem elementów SMD na spodniej stronie. Należy pamiętać, że cztery rezystory 0,1 Ω muszą mieć moc znamionową 3 W, a dwa rezystory 27 Ω moc znamionową 1 W (nie należy pomylić tych elementów). Tabela 1 na sąsiedniej stronie przedstawia kod wartości nadrukowany na wierzchu każdego rezystora SMD

Po uzyskaniu zadowalającego efektu zamontujemy Q1 i Q3, które znajdują się w identycznych obudowach.

Następnie należy dopasować 11 elementów pakietu SOT-23: Q4, Q14–Q16, D1–D2, D5–D7 i ZD1–ZD2. Są one podobne do Q1–Q3, ale z trzema szeroko rozstawionymi nóżkami. Zastosujemy tę samą podstawową procedurę; prawidłowa orientacja powinna być oczywista, ponieważ po jednej stronie obudowy znajduje się tylko jedno wyprowadzenie, a po drugiej dwa.

Należy jednak uważać, aby nie umieścić elementów w niewłaściwym miejscu; w razie wątpliwości należy zapoznać się z rysunkiem 6.

Następnie można zamontować tranzystory Q5 i Q6. Są one w większych obudowach

z trzema wyprowadzeniami i wypustką do odprowadzania ciepła i są przylutowane do dużych miedzianych płaszczyzn, więc będzie to wymagało dość gorącego grotu. Rozsmarujemy niewielką ilość pasty lutowniczej na dużym placku lutowniczym, a następnie umieścimy element na płytce drukowanej i przylutujemy jedną z mniejszych nóżek na obu końcach. Następnie można przylutować wypustkę i zakończyć pracę z dwoma pozostałymi wyprowadzeniami. Upewnijmy się, że FZT796A znajduje się po lewej stronie, a FZT696B po prawej.

Teraz możemy wlutować diody D3 i D4, tak aby ich paski katodowe były skierowane do góry płytki. Paski te są zwykle dość słabo widoczne i aby je dostrzec, może być potrzebne szkło powiększające.

Następnie można zamontować cztery LEDy. Jeśli używamy dokładnie tych samych typów, które podaliśmy na wykazie elementów w poprzedniej części, każdy z nich będzie miał zielone oznaczenie katody. Jednak niektóre inne diody SMD mają podobne oznaczenia na anodach. Dlatego w przypadku stosowania różnych typów diod należy sprawdzić ich dane techniczne lub użyć multimetru cyfrowego w trybie testowania diod, aby ustalić, który koniec jest anodą, a który katodą.

LEDy SMD można przylutować stosując podobną procedurę jak poprzednio, tzn. przymocować jedną stronę, a następnie przylutować drugą. Nie należy mieszać różnych typów.

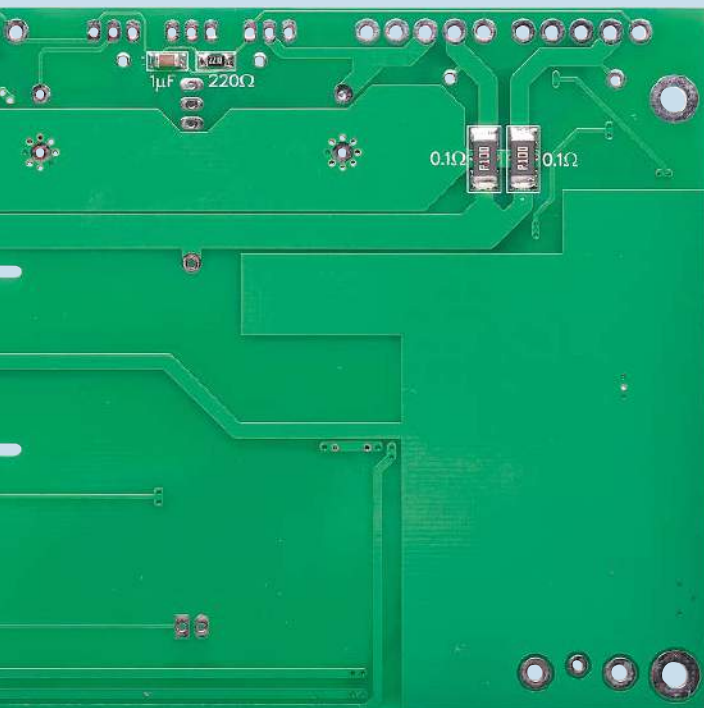
Należy pamiętać, że diody LED2 i LED3 mają po cztery nóżki, dlatego należy unikać zwierania dwóch z nich na każdym końcu. Jeśli tak się stanie, użyjmy topnika i knota lutowniczego, aby je oczyścić. Należy również szybko wykonać złącza, aby nie podgrzewać ich zbyt długo; LEDy są dość małe i mogą ulec uszkodzeniu pod wpływem ciepła.

Plastikowe soczewki LEDów SMD w szczególności mogą ulec uszkodzeniu, jeśli grot lutownicy będzie trzymany na nich zbyt długo lub jeśli temperatura powietrza będzie zbyt wysoka. Zachowajmy również ostrożność, używając gorącego powietrza lub podcierwieni. Trymer VR2 może zostać zamontowany w następnej kolejności. Należy unikać kontaktu lutu z metalową płytką regulacyjną. Następnie można wlutować rezystory SMD po górnej stronie. Na każdym z nich wydrukowana jest jego wartość w postaci 3- lub 4-cyfrowego kodu. Jedyne rezystory, które nie zostały zamontowane na tym etapie, to jeden 220 Ω 0,5 W, dwa 27 Ω 1 W i cztery 0,1 Ω 3 W. Następnie zamontujemy wszystkie ceramiczne elementy SMD, z wyjątkiem jednego kondensatora 1 μF, który zostanie umieszczony na spodzie płytki. Obecnie można również zamontować ferrytowe koraliki L1. Na spodniej stronie płytki znajduje się osiem pasywnych elementów SMD (patrz rysunek 7). Należy je zamontować teraz, stosując tę samą metodę co poprzednio. Wracając do górnej strony, można teraz zamontować kondensatory

Errata do listy elementów

Na liście części zamieszczonej w poprzedniej części dwie diody VS-3EJH02 były błędnie wymienione jako D2 i D4. Są to D3 i D4.

Ponadto w specyfikacji szpulki dla cewki indukcyjnej z rdzeniem powietrznym o pojemności 2,2 μH (L2) błędnie podano, że ma ona 10 mm średnicy wewnętrznej. Powinna ona mieć 13 mm średnicy wewnętrznej. Na koniec, w zależności od sposobu zamocowania tranzystorów na radiatorze, mogą być potrzebne dodatkowe elementy, których nie wymieniono w poprzedniej części, w tym trzy śruby maszynowe M3×10 mm i cztery M3×15 mm.



Powyżej: ten widok przedstawia dolną część płytki drukowanej z zamontowanymi wszystkimi elementami SMD, przed zamontowaniem jakichkolwiek elementów przewlekanych. Łatwiej jest zamontować te elementy przed zamontowaniem większych na górnej stronie, tak aby płytka drukowana nadal leżała płasko na stole

Tabela 1. Kody oznaczeń rezystorów

Nr	Wartość	Kod trzycyfrowy (E24)	Kod czterocyfrowy (E96)
1	1 MΩ	105	1004
6	100 kΩ	104	1003
2	68 kΩ or 68,1 kΩ	683	6812
4	47 kΩ	473	4702
1	33 kΩ	333	3302
3	12 kΩ or 12,1 kΩ	123	1212
2	6,2 kΩ or 6,49 kΩ	622	6491
4	2,2 kΩ or 2,21 kΩ	222	2211
2	1 kΩ	102	1001
1	510 Ω or 511 Ω	511	5110
1	390 Ω 1 W	391	nie dotyczy
2	330 Ω or 332 Ω	331	3320
1	220 Ω or 221 Ω	221	2210
1	120 Ω or 121 Ω	121	1210
1	100 Ω 1 W	101	nie dotyczy
3	100 Ω	101	1000
3	68 Ω or 68,1 Ω	680	68R1
2	47 Ω or 47,5 Ω	470	47R5
4	27 Ω 1 W	270	nie dotyczy
1	10 Ω	100	10R0
4	0,1 Ω 3 W	0R1	nie dotyczy

elektrolityczne SMD. Składają się one z metalowej obudowy na plastikowej podstawie z dwiema płaskimi końcówkami. Wszystkie obudowy, z wyjątkiem jednej, mają czarny pasek na górze, oznaczający końcówkę ujemną i ściętą podstawę z boku końcówki dodatniej. Ustawmy każdy kondensator tak, jak pokazano na rysunku 6, a następnie przylutujemy je na miejscu, stosując podobną procedurę, jak w przypadku kondensatorów ceramicznych.

Ostatnimi elementami SMD, które należy zamontować na górze, są dwie oprawki bezpieczników. Są to dość duże elementy o dużej bezwładności cieplnej, ponieważ są przylutowane do dużych miedzianych ścieżek. Konieczne będzie zastosowanie sporej ilości ciepła, ale procedura jest podobna do tej, którą stosuje się w przypadku innych elementów. Należy pamiętać, że plastikowa część może ulec uszkodzeniu pod wpływem zbyt wysokiej temperatury.

Elementy przewlekane

Teraz można zamontować elementy przewlekane, poza dużymi tranzystorami, w zwykły sposób. Najlepiej zacząć od trymera VR1, a następnie wybrać kolejno: CON4 (jeśli jest montowany), CON2, CON3 i CON1. W przypadku złącza CON2 i CON3 zaleca się, aby po ich podłączeniu wejście przewodów znajdowało się z prawej strony płytki. Najłatwiejszym sposobem jest tymczasowe podłączenie nóżek tuż przed lutowaniem, aby sprawdzić ich orientację.

Należy pamiętać, że w zależności od układu obudowy wzmacniacza można je zamontować odwrotnie, tak aby przewody wchodziły na płytkę drukowaną. Nie próbowaliśmy tego jednak. Teraz możemy zainstalować opcjonalne kondensatory przewlekane, jeśli ich używamy, z wyjątkiem kondensatorów 1000 µF, które zostawimy na później. Na pewno trzeba będzie zamontować kondensator elektrolityczny 47 µF 63 V, jeśli nie zamontowano jeszcze jego odpowiednika SMD w prawym dolnym rogu płytki.

Podobnie, jeśli w filtrze wyjściowym używany jest kondensator polipropylenowy, a nie ceramiczny SMD NPO, należy go teraz zamontować.

W razie potrzeby do punktów testowych można przymocować kołki PC. Ułatwia to nieco regulację, ponieważ można do nich przypiąć przewody z krokodylkami. W takim przypadku należy jednak uważać, aby przypadkowo nie doprowadzić do zwarcia sąsiednich elementów.

Następnie należy włożyć cewkę indukcyjną, ale najpierw trzeba ją nawinąć.

Nawijanie cewki indukcyjnej

Najłatwiej jest ją nawinąć, jeśli przygotowuje się przyrząd do nawijania, jak pokazano w załączonej ramce. Potrzebujemy tylko kilku tanich i łatwych do zdobycia elementów, które przysłużą się za każdym razem, gdy będzie trzeba nawinąć mały dławik z rdzeniem powietrznym.

Do nawinięcia induktora użyto emaliowanego drutu miedzianego o średnicy 1,25 mm o długości ~1 m, umieszczonego na plastikowej szpulce o szerokości 10 mm i średnicy wewnętrznej 13 mm. Dopasujemy szpulkę do przyrządu lub, jeśli nie mamy przyrządu, owińmy taśmę elektryczną wokół śruby lub kołka, tak aby mocno przylegała do środka szpulki, co zapobiegnie pękaniu plastiku podczas nawijania drutu miedzianego.

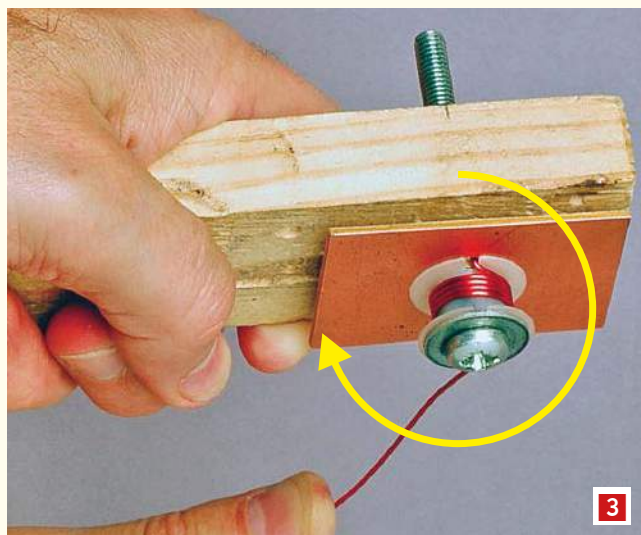
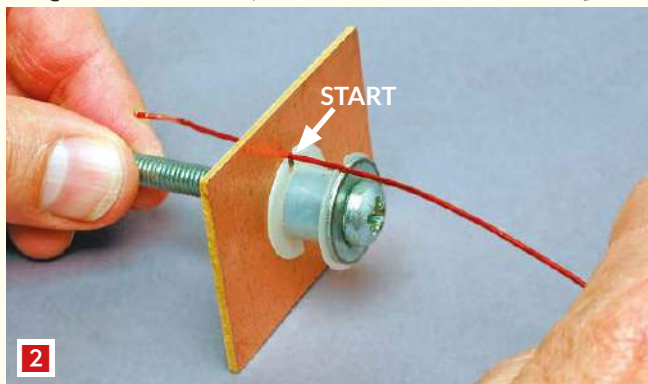
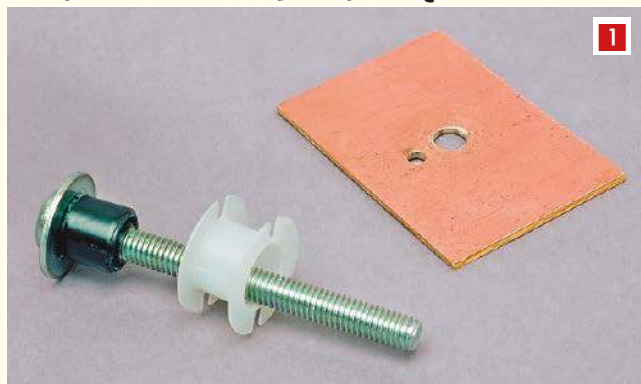
Aby uzyskać staranny efekt, można najpierw wyprostować drut, mocując jeden koniec w imadle i mocno ciągnąc za drugi koniec za pomocą dużych kombinerek. Wymaga to użycia sporej siły, więc należy zachować ostrożność na wypadek, gdyby kombinierki lub imadło puściły.

W odległości 25 mm od jednego końca drutu wykonajmy zagięcie pod kątem prostym, a następnie przełożymy ten koniec przez jedną ze szczelin szpulki i nawińmy siedem ściśle przylegających zwojów, które powinny wypełnić szerokość szpulki.

Ponieważ kierunek nawijania wpływa na parametry cewki, zalecamy nawijanie w tym samym kierunku co my, jak pokazano na zdjęciach.

Po wykonaniu tej warstwy, nawińmy na nią kolejne 6,5 zwojów, ponownie blisko siebie i w tym samym kierunku, a następnie zagniemy drut przez przeciwległą szczelinę i odetniemy go w odległości 25 mm od szpulki.

Wykonanie przyrządu do nawinięcia uzwojenia cewki 2,2 μH



Te zdjęcia pokazują, jak przyrząd do nawijania jest używany do wykonania induktora 2,2 μH . Najpierw należy wsunąć szpulkę na kołnierzyk na śrubie (1), następnie założyć końcówkę i przewlec drut przez szczelinę wyjściową (2). Następnie należy zamocować uchwyt i nawinąć cewkę na szpulkę, używając 13,5 zwojów emaliowanego drutu miedzianego o średnicy 1,25 mm (3). Gotowe uzwojenie (4) zabezpiecza się jedną lub dwiema opaskami termokurczliwej na zewnątrz.



Przyrząd do nawijania składa się ze śruby M5x70 mm, dwóch nakrętek M5, podkładki płaskiej M5, kawałka złomu z płytki drukowanej (ok. 40x50 mm) i kawałka drewna (ok. 140x45x20 mm) na uchwyt. Podczas użytkowania płaska podkładka przylega do łba śruby, po czym na śrubę nakłada się kołnierzyk, który służy do mocowania szpulki. Kołnierzyk ten powinien mieć szerokość nieco mniejszą niż szerokość (wysokość) szpulki i można go nawinąć za pomocą taśmy izolacyjnej. Należy nawinąć odpowiednią

ilość taśmy, tak aby szpulka ściśle przylegała do kołnierza, ale nie była zbyt ciasna. Następnie wywieramy otwór o średnicy 5 mm w środku płytki drukowanej, a następnie otwór wyjściowy o średnicy 1,5 mm w odległości około 8 mm, który będzie pokrywał się z jedną ze szczelin w szpulce. Szpulkę można nasunąć na kołnierzyk, a następnie na śrubę nasunąć „policzek końcowy” płytki drukowanej (tzn. szpulka jest umieszczona pomiędzy podkładką a płytką drukowaną). Należy ustawić szpulkę w taki sposób, aby jeden z jej rowków pokrywał się z otworem wyjściowym w policzku końcowym, a następnie zamontować pierwszą nakrętkę i porządnie ją zabezpieczyć. Następnie można zamontować uchwyt, wierząc otwór o średnicy 5 mm na jednym końcu, nasuwając go na śrubę i montując drugą nakrętkę.

Aby utrzymać uzwojenia na miejscu, należy odciąć 10 mm odcinek opaski termokurczliwej o średnicy 20 mm i nasunąć ją na szpulkę, a następnie delikatnie obkurczyć za pomocą pistoletu na gorące powietrze na niskim ustawieniu. Przytnijmy dwa wystające druty na długość dokładnie 20 mm od podstawy szpulki, a następnie usuńmy po 5 mm emalii z każdego końca, używając papieru ściernego lub noża hobby-stycznego/skalpela, i pocynujemy przewody.

Aby uzyskać określoną jakość, należy zamontować induktor w sposób pokazany na rysunkach 6, 9 i na zdjęciach. Są dwa otwory, w których można umieścić opaskę zaciskową, która przytrzyma go na miejscu. Wygnijmy jego wyprowadzenia w dół o 90°, aby pasowały do punktów lutowniczych na płycie drukowanej, a następnie założymy i zaciśniemy opaskę zaciskową przed przylutowaniem i obcięciem

wyprowadzeń. Zwróćmy uwagę na sposób rozmieszczenia przewodów: każdy przewód z płytki drukowanej biegnie do cewki, a następnie pod nią.

Wiercenie i przytwierdzenie radiatora

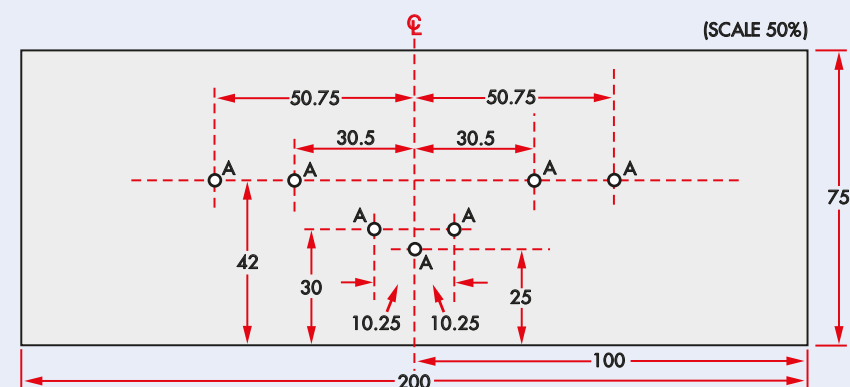
Jeśli modernizujemy wcześniejszą wersję modułu lub budujemy go z zestawu, być może mamy już wywiercony lub nagwintowany radiator. W przeciwnym razie należy zapoznać się z załączoną ramką i schematem wiercenia (rysunek 8). Najlepiej byłoby, gdyby siedem otworów montażowych tranzystorów było gwintowanych na gwint M3. Nie należy się spieszyć, ponieważ dość łatwo jest „przekalibrować” otwór w aluminium. W takim przypadku konieczne może być rozpoczęcie od nowa z nowym radiatorem (lub wywiercenie otworu na wylot).

Jeśli nie chcemy gwintować otworów, możemy przewiercić się przez cały radiator i użyć dłuższych wkrętów maszynowych (wprowadzonych między żebra) oraz nakrętek do zamocowania tranzystorów. Otwory muszą być jednak wywiercone z dużą dokładnością, w przeciwnym razie śruby mogą nie zmieścić się między żebrowaniem.

Po wywierceniu i nagwintowaniu otworów montażowych tranzystora należy zająć się jego zamocowaniem w obudowie. Preferowaną przez nas metodą jest wywiercenie i nagwintowanie trzech dodatkowych otworów wzdłuż dolnej części radiatora, jak pokazano na zdjęciu. Można również zamontować wsporniki kątowe do żebrowania po obu stronach radiatora. Można to zrobić, przewiercając się przez żebra i używając śrub oraz nakrętek do zamocowania wsporników.

Po wywierceniu wszystkich otworów należy usunąć zadziory za pomocą wiertła o większym

Wiercenie i gwintowanie aluminium radiatora



HOLES A: DRILL 3mm DIAMETER OR DRILL 2.5mm DIAMETER & TAP FOR M3 SCREW. DEBURR ALL HOLES.

Rysunek 8. ten rysunek przedstawia szczegóły wiercenia w radiatorze. Otwory można wywiercić i nagwintować (za pomocą gwintownika M3) lub nawiercić do 3 mm i zamontować tranzystory za pomocą śrub maszynowych, nakrętek i podkładek

Na rysunku 8 powyżej pokazano szczegóły wiercenia radiatora. W przypadku gwintowania otworów, należy je wywiercić do średnicy 2,5 mm na wylot przez płytę radiatora, a następnie gwintować do średnicy 3 mm. Alternatywnie można wywiercić otwory wiertłem o średnicy 3 mm i zamontować tranzystory za pomocą śrub, nakrętek i podkładek.

Nagwintowanie otworów wymaga nieco więcej pracy, ale znacznie ułatwia montaż tranzystorów (nie trzeba nakręcać nakrętek) i nadaje schludny wygląd.

Przed przystąpieniem do wiercenia radiatora należy dokładnie zaznaczyć miejsca otworów za pomocą bardzo ostrych ołówka. Po wykonaniu tych czynności należy użyć małej wiertarki ręcznej z wiertłem o średnicy 1 mm, aby rozpocząć rozmieszczanie każdego otworu. Jest to ważne, ponieważ umożliwia precyzyjne ustawienie otworów (ich rozmieszczenie jest bardzo

ważne) przed przejściem na większe wiertła w wiertarce. Do wywiercenia otworów należy użyć wiertarki (bez niej nie uda się uzyskać otworów idealnie prostokątnych do powierzchni mocowania). Na początek należy użyć małego wiertła prowadzącego (np. 1,5 mm), a następnie ostrożnie zwiększać rozmiar wiertła do 2,5 mm lub 3 mm. Otwory muszą przechodzić między żebrami, dlatego bardzo ważne jest ich dokładne rozmieszczenie. Dodatkowo można wywiercić (i nagwintować) trzy otwory w podstawie radiatora, aby później można go było przykręcić do obudowy. Podczas wiercenia otworów należy użyć odpowiedniego środka smarnego. Zalecany środkiem smarnym do aluminium jest nafta, ale stwierdziliśmy, że lekki olej maszynowy (np. Singer lub 3 w 1) również dobrze sprawdza się w tego typu zadaniach. Nie próbujmy wiercić otworów za jednym razem.

Podczas wiercenia w aluminium ważne jest, aby regularnie wyjmować wiertło z otworu i usuwać opiłki metalu. Jeśli tego nie zrobimy, opiłki aluminium mają nieprzyjemny zwyczaj zakleszczania wiertła i łamania go. Za każdym razem przed wznowieniem wiercenia należy ponownie nasmarować otwór i wiertło olejem.

Gwintowanie

Do gwintowania otworów potrzebny będzie gwintownik pośredni (początkowy) M3 (nie końcowy). Sztuczka polega na tym, aby robić to powoli i ostrożnie. Należy utrzymywać odpowiednią ilość smaru i regularnie wyciągać gwintownik, aby usunąć opiłki metalu z otworu. Za każdym razem przed wznowieniem pracy należy ponownie nasmarować gwintownik.

Na żadnym etapie nie należy wywiercać na gwintownik nadmiernej siły. Łatwo jest złamać gwintownik na pół, jeśli używa się dużej siły, a jeśli pęknięcie nastąpi na powierzchni radiatora lub poniżej, można porysować zarówno gwintownik, jak i radiator (oraz około 25 USD). Podobnie, jeśli podczas wyciągania gwintownika z radiatora napotkamy opór, delikatnie obróćmy go w przód i w tył i pozwólmy mu się wykleszczyć. Krótko mówiąc, nie wymuszajmy tego.

Po zakończeniu gwintowania należy usunąć z powierzchni montażowej opiłki metalu z wszystkich otworów za pomocą wiertła o zbyt dużym rozmiarze. Powierzchnia montażowa musi być idealnie gładka, aby zapobiec przebiciu podkładek izolacyjnych tranzystora.

Na koniec radiator należy dokładnie wyszorować wodą z detergentem i pozostawić do wyschnięcia.

rozmiarze i oczyścić je z cząstek aluminium lub opiłków. Sprawdźmy, czy powierzchnie wokół otworów są idealnie gładkie, aby uniknąć możliwości przebicia podkładek izolacyjnych.

Mocowanie radiatora

Teraz nadszedł czas na połączenie płytki drukowanej z radiatorem, ale najpierw należy ponownie sprawdzić front radiatora. Wszystkie otwory muszą być wygładzone i muszą być idealnie czyste i pozbawione zadziórów lub opiłków metalu.

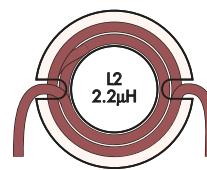
Rozpocznijmy montaż radiatora, montując elementy Q7, Q8 i Q9. Pomiędzy każdym z tych tranzystorów a radiatorem znajduje się podkładka z gumy silikonowej. Q7 i Q8 wymagają również zastosowania tulei izolacyjnej pod każdym łbem śruby. Na rysunku 10 (A i B) pokazano ułożenie montażowe.

Dla Q9 wybraliśmy podkładkę izolacyjną TO-126/TO-225, ponieważ jest

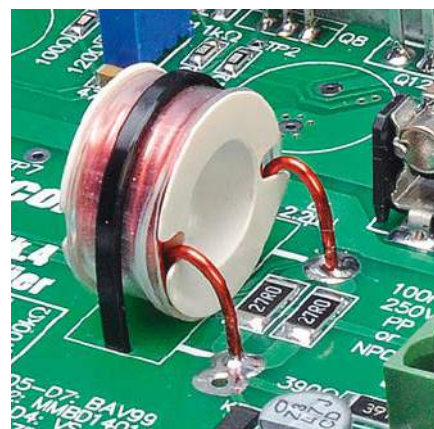
ona mniejsza niż TO-220 dla Q7 i Q8, ale jeśli nie mamy takiej podkładki, zawsze możemy przyciąć podkładkę TO-220 do odpowiedniego rozmiaru. Upewnijmy się, że jest ona wystarczająco duża, aby całkowicie zakryć metalową podkładkę Q9, uwzględniając ewentualne luzy w otworze na śruby.

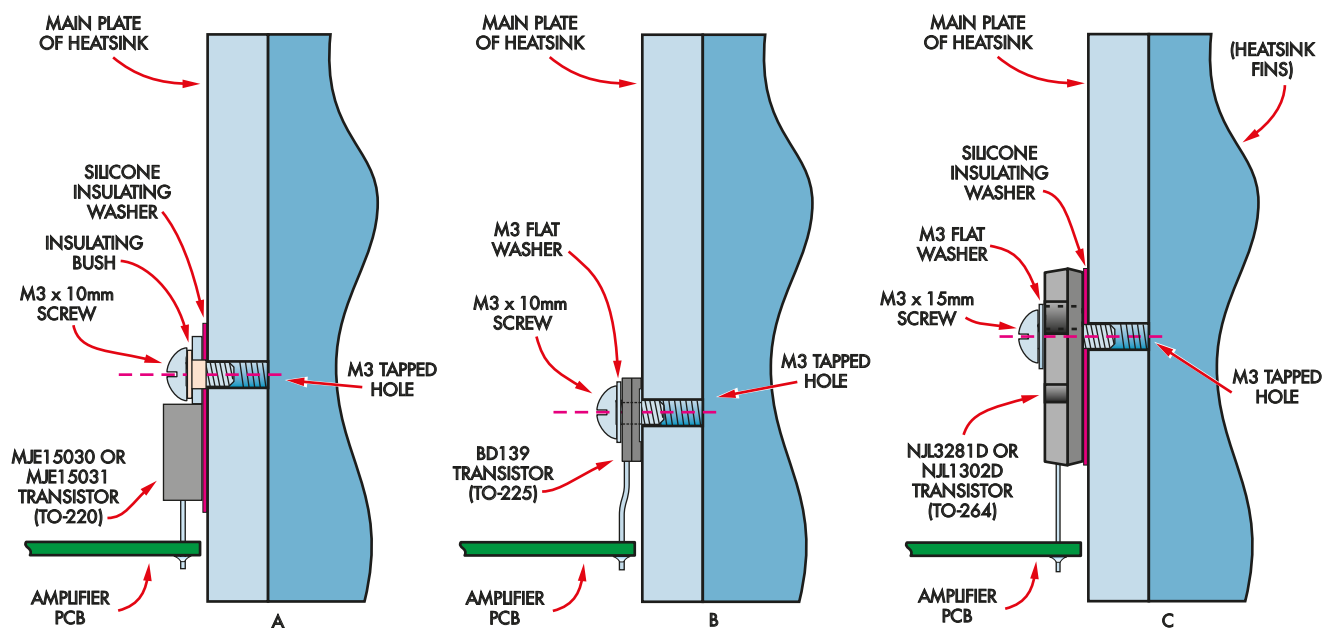
Należy uważać, aby nie pomylić Q7 i Q8, ponieważ numery ich typów są podobne. Jeśli otwory są nagwintowane, tranzystory można zamocować za pomocą wkrętów maszynowych M3×10 mm. Alternatywnie, jeśli wywiercono otwory bez gwintowania, należy użyć wkrętów maszynowych M3×15 mm, przy czym wkręty powinny przechodzić od strony radiatora (tzn. łby wkrętów powinny znajdować się między żebrami radiatora).

Upewnijmy się, że trzy tranzystory i ich izolatory są ustawione prawidłowo w pionie, a następnie dokręćmy śruby do końca. Nie



Rysunek 9. Należy odgąć wyprowadzenia cewki indukcyjnej L2 i zamontować ją na płytce drukowanej w sposób pokazany na rysunku, aby uzyskać jak najlepsze parametry pracy





Rysunek 10. Ten schemat pokazuje szczegóły montażu tranzystorów sterujących TO-220 (A), mnożnika VBE BD139 (B) i czterech tranzystorów wyjściowych (C). Po zamontowaniu tych tranzystorów należy użyć multimetru (przełączonego na zakres niskich Ohmów), aby sprawdzić, czy są one prawidłowo odizolowane od radiatora – szczegóły w tekście

dokręcajmy ich jednak jeszcze teraz, tzn. powinna istnieć możliwość obracania tranzystorów w każdym kierunku.

Następnym krokiem jest zamocowanie gwintowanego odstępnika M3×9 mm (lub 10 mm) w każdym z narożnych otworów montażowych płytki drukowanej. Przymocujemy je za pomocą wkrętów maszynowych M3×6 mm. Po ich założeniu usadzimy płytkę na odstępnikach i opuścimy radiator tak, aby końcówki tranzystorów przechodziły przez odpowiadające im punkty lutownicze na płycie drukowanej. Zauważmy, że prawdopodobnie będzie trzeba odgiąć wyprowadzenia Q9 od radiatora, jak pokazano na rysunku 10.

Montaż tranzystorów wyjściowych

Teraz można zamontować cztery tranzystory wyjściowe (Q10–Q13). Zastosowane są tu dwa różne typy, więc należy uważać, aby ich nie pomylić (należy sprawdzić schemat układu). Jak pokazano na rysunku 10(C), tranzystory te muszą być również odizolowane od radiatora za pomocą silikonowych podkładek izolacyjnych.

Zamontujmy najpierw Q10. Procedura polega na włożeniu końcówek do otworów montażowych płytki drukowanej, a następnie odchyleniu elementu do tyłu i częściowym przełożeniu przez śrubę

montażową M3×15 mm z podkładką płaską (lub M3×20 mm w przypadku niegwintowanych otworów). Po wykonaniu tych czynności należy zawiesić podkładkę izolacyjną na końcu śruby, a następnie luźno przykręcić zestaw do radiatora.

Pozostałe trzy tranzystory instaluje się w dokładnie taki sam sposób, ale należy zwrócić uwagę, aby w każdym miejscu zamontować odpowiedni typ elementu. Po ich założeniu dociśnijmy lekko płytkę, tak aby wszystkie cztery odstępniki (i radiator) stykały się z blatem stołu. Powoduje to automatyczne dopasowanie długości wyprowadzeń tranzystorów i zapewnia, że dolna część płytki drukowanej znajduje się 9–10 mm nad dolną krawędzią radiatora.

Teraz ustawmy montowaną płytkę drukowaną w poziomie tak, aby każda strona była oddalona o 32,5 mm od sąsiedniego końca radiatora. Po upewnieniu się, że tranzystor jest prawidłowo umieszczony, dokręćmy wszystkie śruby tranzystora na tyle, aby przytrzymały go na miejscu, a podkładki izolacyjne były prawidłowo ułożone.

Następnym krokiem jest lekkie przylutowanie zewnętrznych wyprowadzeń Q10 i Q13 do ich padów w górnej części płytki. Następnie obracamy zestaw do góry nogami, aby można było przylutować wyprowadzenia tranzystorów.

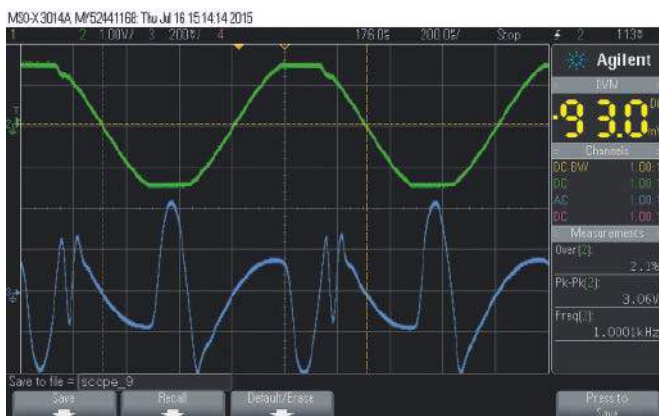
Jako podpór można użyć kilku tekturowych cylindrów przyciętych na długość 63 mm (np. po jednym w każdym rogu obok CON1 i CON3). Po ich założeniu sprawdzimy, czy płytkę jest prawidłowo wyśrodkowana na radiatorze, a następnie przylutujemy



W podstawie radiatora można wywiercić i nagwintować trzy otwory M3 lub M4, tak aby można go było później przymocować do obudowy. Wykonajmy je na głębokość ok. 10 mm.



Zakres 1: sygnał wyjściowy wzmacniacza dla fali prostokątnej o częstotliwości 1 kHz przy obciążeniu 4-ohmowym. Jak widać, występuje niewielkie przesterowanie (około 5%), ale szybko powraca do normalnego stanu, z bardzo niewielkim dzwonieniem



Zakres 3: tutaj wzmacniacz dostarcza falę sinusoidalną o częstotliwości 1 kHz do 8-ohmowego obciążenia z mocą około 150 W, a więc w punkcie granicznym. Jak widać, powrót z przeciążenia na minusie jest dość czysty. Powrót z przeciążenia na plusie ma niewielki skok ze względu na wysokie wzmocnienie pętli otwartej, ale po około 25 μ s powraca do normalnego nachylenia z niewielkim dzwonieniem

wszystkie 29 wyprowadzeń. Upewnijmy się, że luty są dobre, ponieważ przez niektóre z nich może przechodzić wiele amperów przy pełnej mocy.

Po zakończeniu lutowania należy przyciąć nóżki, używając stalowego liniału jako prostej krawędzi, aby zapewnić stałą długość nóżek. Po wykonaniu tych czynności należy odwrócić płytkę w prawą stronę i dokręcić śruby mocujące tranzystory, aby zapewnić dobre połączenie termiczne między elementami a radiatorem.

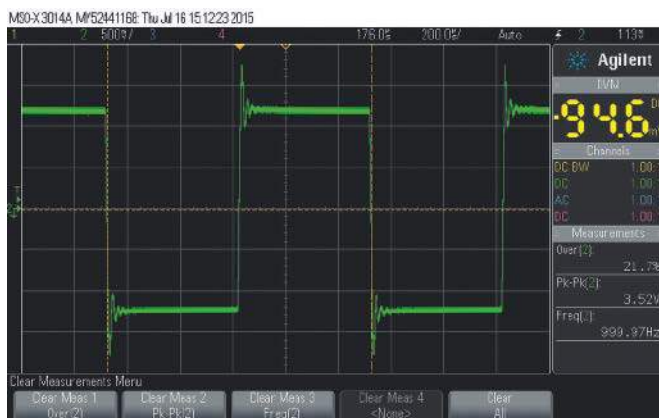
Nie należy jednak zbyt mocno dokręcać śrub mocujących. Pamiętajmy, że radiator jest wykonany z aluminium, więc jeśli będziemy zbyt nieuważni, możemy przekręcić gwintowanie.

Kontrola izolacji elementów

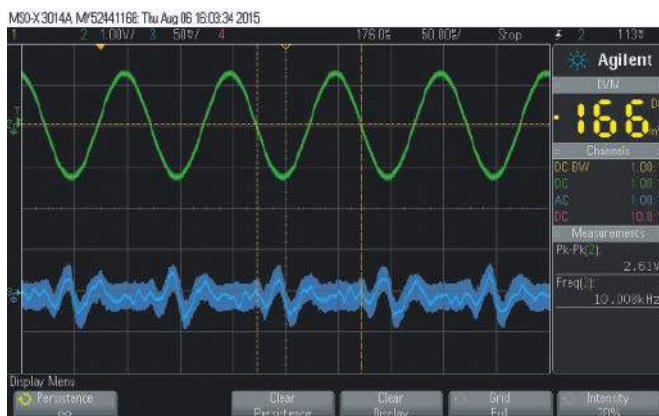
Należy sprawdzić, czy wszystkie tranzystory są elektrycznie odizolowane od radiatora. Można to zrobić, przełączając multimetr na zakres wysokich Ohmów i sprawdzając, czy nie ma zwarcia między powierzchnią montażową radiatora a kolektorami

tranzystorów (uwaga: kolektor każdego tranzystora jest połączony z jego metalową powierzchnią lub wypustką).

W przypadku tranzystorów Q7–Q8 i Q10–Q13 należy po prostu sprawdzić, czy nie ma połączeń pomiędzy każdą z wkładek bezpiecznika znajdujących się najbliżej radiatora a samym radiatorem (tzn. po obu stronach wzmacniacza). Dzieje się tak, ponieważ kolektory tranzystorów w każdej połowie stopnia wyjściowego są połączone razem i biegną do odpowiednich bezpieczników.



Zakres 2: ten sam test co w Zakresie 1, ale z kondensatorem 2 μ F na obciążeniu. Powoduje to większe przekroczenie zakresu (~20%) i pewne dzwonienie, ale wszystko jest pod kontrolą. Jest to standardowy test stabilności wzmacniacza



Zakres 4: zniekształcenia resztkowe przy 100 W przy 8 Ω i 10 kHz. Poziom zniekształceń jest tak niski, że znaczna ich część jest szumem nawet przy tym poziomie mocy i częstotliwości, o czym świadczy widoczny ślad na wyświetlaczu. Zniekształcenia występują głównie wokół najbardziej ujemnej części przebiegu, dlatego są jeszcze mniej znaczące przy niższych poziomach mocy.

Wartości rezystorów cienkowarstwowych SMD

Być może dało się zauważyć, że w wykazie elementów opublikowanym w poprzedniej części podaliśmy kilka rezystorów o dziwnej wartości. Na przykład 6,49 k Ω , 332 Ω , 47,5 Ω itd. Jak już wyjaśnialiśmy, wiele rezystorów w układzie musi być cienkowarstwowych, aby zapewnić dobre działanie (wiele rezystorów SMD ma konstrukcję grubowarstwową, która nie jest odpowiednia).

Najlepsze cienkowarstwowe rezystory SMD, jakie udało nam się znaleźć, są produkowane przez firmę Stackpole Electronics. Oprócz tego, że są cienkowarstwowe, mają też stosunkowo wysoką moc znamionową, jak na swoje

rozmiary (3,2×1,6 mm), wynoszącą 0,5 W. Jednak ta seria rezystorów (RNC1206FTD) występuje w serii wartości E96, a nie w serii E24, do której jesteśmy przyzwyczajeni.

Seria E24 jest następująca: 10, 11, 12, 13, 15, 16, 18, 20, 22, 24, 27, 30, 33, 36, 39, 43, 47, 51, 56, 62, 68, 75, 82, 91. Powtarza się to o rząd wielkości wyżej lub niżej. Innymi słowy, w każdym stosunku rezystancji 10:1 są 24 wartości do wyboru, a każda z nich jest o około 10% wyższa od następnej.

Jak można się domyślić, seria E96 zawiera 96 różnych wartości dla każdej dekady. O ile

jednak seria E24 zawiera wszystkie wartości serii E12 i po prostu dodaje nowe wartości pomiędzy nimi, o tyle seria E96 nie zawiera wszystkich wartości serii E24. Tak więc seria rezystorów RNC1206FTD nie oferuje rezystorów 6,2 k Ω , 330 Ω ani 47 Ω . W praktyce nie ma to znaczenia, ponieważ wybraliśmy po prostu zbliżone wartości; ten układ toleruje wartości o kilka procent wyższe lub niższe, o ile wszystkie rezystory o tej samej wartości nominalnej są ściśle dopasowane.

Poprawa zniekształceń i stabilności

W naszym pierwszym prototypie Ultra-LD Mk.4 wprowadziliśmy szereg zmian, które – jak się spodziewaliśmy – pozwolą zmniejszyć zniekształcenia w porównaniu z poprzednią wersją. Na przykład lepsze tłumienie pola magnetycznego dzięki nowemu układowi płytki drukowanej, nieindukcyjne rezystory emiterowe do montażu powierzchniowego oraz większe wzmocnienie w pętli otwartej zapewnione przez nowe tranzystory powinny przynieść korzyści. Z rozczarowaniem stwierdziliśmy więc, że poziom zniekształceń był początkowo bardzo zbliżony do wersji Mk.3.

Przekonani, że parametry powinny być znacznie lepsze, sprawdziliśmy, co może je ograniczać. W trakcie tego procesu dokonaliśmy wielu interesujących i ważnych odkryć. Jednym z nich był fakt, że zastosowanie różnych rezystorów obciążenia miało znaczący wpływ na pomiary zniekształceń, zwłaszcza przy wyższych częstotliwościach.

Impedancja cewki wyjściowej wzrasta wraz z częstotliwością i tworzy ona dzielnik napięcia z obciążeniem. W przypadku obciążenia czysto rezystancyjnego spowoduje to jedynie pogorszenie charakterystyki częstotliwościowej. Jeśli jednak obciążenie ma jakieś nieliniowości, to nawet jeśli sygnał ze wzmacniacza jest idealnie czysty, spowoduje to powstanie zniekształceń na obciążeniu.

Do testowania wzmacniaczy używamy Dummy Load Box (sztucznego obciążenia) opisanego w numerze z sierpnia 1992 r. Zakładaliśmy, że jest to urządzenie liniowe, ponieważ do tej pory dawało dobre wyniki. Kiedy jednak doprowadziliśmy sygnał o wartości 14 V RMS z generatora o ultra niskich zniekształceniach Audio Precision System Two do jednego końca sztucznego obciążenia i podłączyliśmy kondensator polipropylenowy z drugiego końca do masy sygnału, tworząc filtr dolnoprzepustowy, okazało się, że tak nie jest.

Przeprowadzenie tego testu z rezystorem, o którym myśleliśmy, że będzie bardzo liniowy (typ 5 W wirewound), dało wynik 0,00025% THD+N przy 10 kHz i paśmie pomiarowym 80 kHz. Jednak zastosowanie naszego sztucznego obciążenia jako rezystora dało wyższy odczyt, około 0,0008%, czyli trzykrotnie wyższy. Jest więc prawdopodobne, że to samo sztuczne obciążenie przyczyniło się do zwiększenia zniekształceń odczytywanych przez wzmacniacz.

Aby ustalić przyczynę, przylutowaliśmy kilka przewodów bezpośrednio do szeregu rezystorów w układzie sztucznego obciążenia i powtórzyliśmy test. Odczyt spadł do 0,00025%. Dlatego uważamy, że problem tkwi albo w złączach, albo w układzie przetaczania przekładników w układzie sztucznego obciążenia. Musieliśmy więc kontynuować

testy z wykorzystaniem połączeń lutowanych, ponieważ był to jedyny sposób na uzyskanie prawdziwego odczytu parametrów wzmacniacza (w późniejszym terminie będziemy musieli dokładniej zbadać źródło zniekształceń w układzie sztucznego obciążenia).

Dostosowywanie filtra wyjściowego

Następnie zmierzaliśmy, że wzmacniacz osiąga poziom około 0,0015% THD+N przy 10 kHz, co stanowi niewielką poprawę w stosunku do modułu Ultra-LD Mk.3 w tych samych warunkach (około 0,002%). Uznaliśmy jednak, że nowy moduł powinien przynieść większą poprawę, a następnie odkryliśmy, że jeśli zmierzmy zniekształcenia przed filtrem wyjściowym, będą one znacznie niższe i wyniosą około 0,0008% przy 10 kHz.

Ponieważ filtr był nadal w obwodzie, a prąd obciążenia nadal płynął przez cewkę indukcyjną L2, oznaczało to, że nie było to spowodowane żadną interakcją między filtrem wyjściowym a częścią wejściową urządzenia. Musiało to więc dotyczyć samego filtra wyjściowego; albo kondensator lub rezystory SMD nie były wystarczająco liniowe, albo z sygnałem w cewce indukcyjnej działo się coś dziwnego.

Następnie osobno przetestowaliśmy kilka różnych rezystorów i kondensatorów, stosując podobną metodę jak poprzednio, tzn. podłączając je jako filtry RC i wykorzystując sprzęt Audio Precision do sprawdzenia działania. W ten sposób rezystory SMD otrzymały nasz certyfikat jakości, ponieważ cztery połączone równolegle działały tak samo dobrze jak rezystor drutowy 6,8 Ω.

Jednak kondensator polipropylenowy X2, którego używaliśmy w tym prototypie, dawał w tym teście zniekształcenia na poziomie około 0,0006%. Przetestowaliśmy trzy inne kondensatory polipropylenowe, dwa inne typy X2 i jeden MKP. MKP i jeden z kondensatorów X2 uzyskały pozytywny wynik (odczyt około 0,00025%), natomiast jeden z pozostałych kondensatorów X2 dał zniekształcenia wyższe niż oczekiwano. Dlatego umieściliśmy na płytce lepszy kondensator, ale to tylko nieznacznie poprawiło jej działanie.

Po wykluczeniu kondensatora i rezystora jako przyczyny problemu, podejrzenie padło na cewkę indukcyjną. Ale czy to możliwe, że sposób poprowadzenia połączeń na naszej płytce drukowanej nie był w 100% poprawny, zwłaszcza w przypadku ścieżek masy? Aby to wykluczyć, usunęliśmy filtr RLC z płytki drukowanej i zamontowaliśmy go całkowicie poza płytką, między złączem wyjściowym a obciążeniem testowym.

Tranzystor Q9 (mnożnik VBE) jest inny. W takim przypadku należy sprawdzić, czy nie ma zwarcia między środkowym przewodem (kolektorem) a radiatorem.

W obu przypadkach należy uzyskać odczyt obwodu otwartego. Jeśli znajdziemy zwarcie, odkręmy po kolei każdą śrubę mocującą tranzystor, aż zwarcie zniknie. Wówczas wystarczy

zlokalizować przyczynę problemu i ponownie zamontować problematyczny tranzystor.

Jeśli podkładka izolacyjna została w jakikolwiek sposób uszkodzona (np. przebita), należy ją wymienić.

Zakończenie montażu

Montaż płytki drukowanej można teraz zakończyć, instalując dwa kondensatory 1000 μF 63 V – zakładając, że zdecydowaliśmy się je zamontować. Jak wspomniano w poprzedniej części, można je pominąć pod warunkiem, że przewody zasilające będą krótkie i wykonane z grubego drutu. W przeciwnym razie maksymalna moc wyjściowa nieco spadnie ze względu na straty w tych przewodach, ale nie powinno to mieć wpływu na parametry.

Jedną ze zmian, jakie wprowadziliśmy przy projektowaniu tej płytki drukowanej, było umieszczenie tych kondensatorów w taki sposób, aby nie utrudniały dostępu do śrub mocujących radiator w takim stopniu, jak miało to miejsce w wersjach Mk.2 i Mk.3. Jednak praca na płytce drukowanej jest łatwiejsza, jeśli duże kondensatory nie są zamontowane,

a ze względu na ich bliskość do radiatora prawdopodobnie w końcu wyschną (choć prawdopodobnie po ponad 10 000 godzin pracy, zakładając, że są to kondensatory dobrej jakości).

Teraz usuńmy dwa odstępniki z krawędzi płytki przylegającej do radiatora. W rzeczywistości bardzo ważne jest, aby tylna krawędź płytki trzymała się tylko na wyprowadzeniach tranzystorów na radiatorze. W ten sposób unika się ryzyka pęknięcia ścieżek płytki drukowanej i podkładek wokół tranzystorów na radiatorze z powodu rozszerzania i kurczenia się ich wyprowadzeń w miarę nagrzewania i stygnięcia.

Krótko mówiąc, tylne odstępniki są montowane tylko podczas montażu wyprowadzeń tranzystorów na radiatorze, a następnie muszą zostać usunięte. Nie odgrywają one żadnej roli w zabezpieczaniu modułu. Zamiast tego ta krawędź modułu jest mocowana przez przykręcenie samego radiatora do obudowy.

Jak już wcześniej wspomniano, można to zrobić, nagwintowując otwory M3 (lub M4) w głównej płycie od spodu radiatora lub używając wsporników kątowych. Przód płytki jest mocowany za pomocą dwóch



Ten widok pokazuje szczegóły montażu tranzystora mnożnika VBE (Q9) oraz dwóch tranzystorów sterujących (Q7 i Q8). Sprawdźmy, czy wszystkie te tranzystory oraz cztery tranzystory wyjściowe (Q10-Q13) są odizolowane od radiatora

Całkowicie rozwiązało to problem, zapewniając doskonałe osiągi, które wykazano w poprzedniej części na wykresach Audio Precision. Ale dlaczego? Przenieśliśmy cewkę i rezystory z powrotem na płytce drukowanej, ale połączenia pozostały takie same, a zmierzone zniekształcenia ponownie się podwoiły. To w zasadzie wykluczyło problem ze ścieżkami. Zamontowaliśmy więc cewkę indukcyjną na krótkich odcinkach elastycznego drutu i eksperymentowaliśmy ze zmianą jej położenia i orientacji.

Zarówno położenie, jak i orientacja cewki indukcyjnej miały wpływ na wynik, jednak miejsce montażu miało znacznie mniejsze znaczenie, gdy cewka była obrócona tak, aby spoczywała na boku. Prawdopodobnie jest to spowodowane oddziaływaniem pola magnetycznego na płaszczyznę ortogonalną do ścieżek na płytce drukowanej. W ten sposób uzyskaliśmy ostateczny układ montażowy. Jedynym powodem, dla którego możemy się domyślać, że ma to znaczenie, jest to, że impulsy wysokoprądowe w ścieżkach zasilających płytki drukowanej były odbierane przez cewkę indukcyjną i wprowadzały sygnał zniekształcający do obciążenia. Efekt ten jest większy przy wyższych częstotliwościach, ponieważ wyższa impedancja cewki indukcyjnej w przypadku tych sygnałów skuteczniej izoluje wyjście głośnika od niskoimpedancyjnego złącza rezystorów emiterowych tranzystorów wyjściowych.

Przy okazji, jesteśmy pewni, że ten wzmacniacz ma mniejsze zniekształcenia niż wzmacniacz 20 W pracujący w klasie A, który został opublikowany w numerze maj-wrzesień 2007. Główną zaletą wzmacniacza klasy A w porównaniu ze wzmacniaczem klasy B lub klasy AB jest to, że nie występują w nim żadne zniekształcenia skrośne, ponieważ wszystkie tranzystory wyjściowe przewodzą przez cały czas. Cóż, jeśli nowa konstrukcja Ultra-LD ma jakiegokolwiek zniekształcenia skrośne, to z pewnością nie możemy ich wykryć!

W rzeczywistości, jeśli dokona się bezpośredniego porównania charakterystyk zniekształceń przedstawionych w wydaniu z lipca 2011 r. z charakterystykami opublikowanymi w poprzedniej części, można zauważyć, że Ultra-LD Mk4 jest zdecydowanie lepszym urządzeniem niż poprzednia konstrukcja. Czy istnieje prawdopodobieństwo, że różnica będzie słyszalna? Uważamy, że jest to bardzo mało prawdopodobne!

Podejrzewamy, że niewielka ilość zniekształceń wynika głównie z nieliniowości wejściowej części wzmacniacza, z którą wzmacniacz pracujący w klasie A borykałby się w równym stopniu. Wniosek jest taki, że nie ma już sensu budować wzmacniacza pracującego w klasie A. Równie dobrze można zbudować ten model i uzyskać znacznie większą moc, wyższą sprawność i mniejsze rozpraszanie ciepła.

odstępników M3×9 mm (lub 10 mm), które zostały zamontowane wcześniej.

Moduły zasilania i ochrony głośników

Na tym kończy się montaż modułu wzmacniacza mocy. Następnym krokiem jest zbudowanie modułu zasilacza (widocznego na zdjęciu obok), a sposób wykonania tego zadania opiszemy w następnej części. Wyjaśnimy także, jak włączyć zasilanie i przetestować wzmacniacz oraz podamy kilka podstawowych informacji na temat umieszczania go w metalowej obudowie.

Na koniec przedstawimy poprawiony moduł ochrony głośników, który może również monitorować temperaturę radiatora. Będzie on potrzebny, aby zapobiec uszkodzeniu wzmacniacza, które mogłoby zniszczyć głośniki i spowodować pożar. ■

Nicholas Vinen

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

Poprawki stabilności

Podczas udoskonalania wzmacniacza zmieniliśmy niektóre elementy, co pogorszyło stabilność i czasami powodowało oscylacje, ale w rezultacie nie doszło do żadnych uszkodzeń. Pozwoliło nam to jednak odkryć kilka sposobów na poprawę ogólnej stabilności.

Stało się to niemal przez przypadek. Okazało się, że gdy wzmacniacz był w niestabilnym stanie i zaczynał oscylować, dotknięcie pewnych elementów na płytce drukowanej powodowało chwilowe ustanie oscylacji.

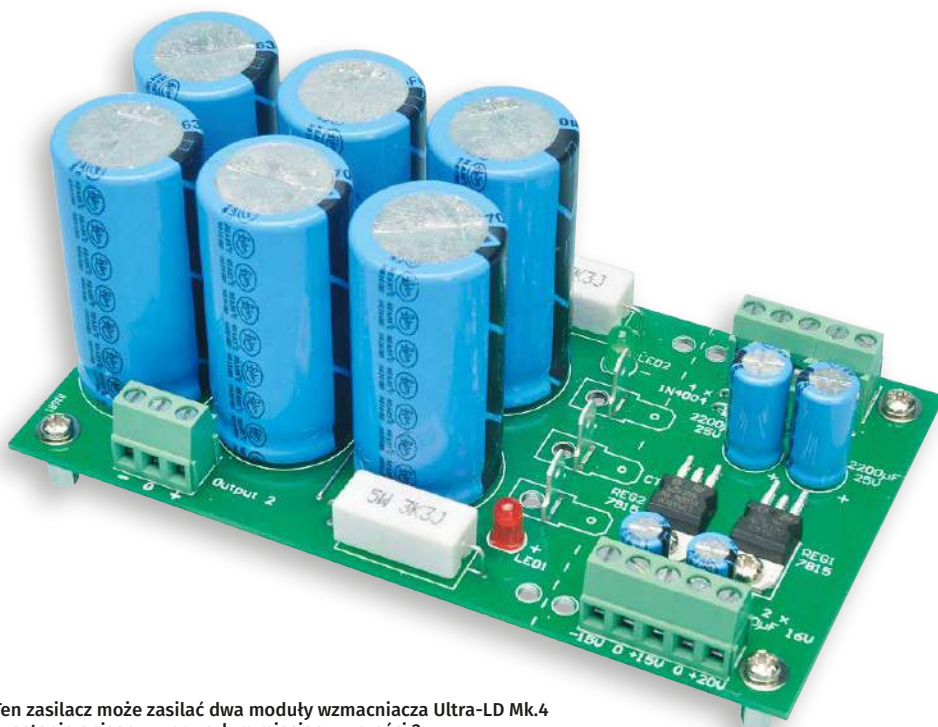
Wyodrębniliśmy ten efekt do dwóch konkretnych składników: rezystora kolektora Q4 i rezystora 2,2 kΩ od złącza dwóch kondensatorów 150 pF do ujemnej szyny (część sieci kompensacyjnej wokół Q4/Q6). Uznaliśmy, że podłączenie kondensatorów do tych rezystorów będzie miało podobny wpływ na stabilność jak dotknięcie ich palcem i udowodniliśmy to, modyfikując prototyp w ten sposób i całkowicie eliminując oscylacje.

Powód wprowadzenia pierwszego z tych dwóch ulepszeń wyjaśniliśmy w punkcie 1: eliminuje ono zjawisko Early'ego dla Q4, które powoduje rodzaj lokalnego sprzężenia zwrotnego. Już sama ta zmiana wydaje się czynić wzmacniacz znacznie bardziej tolerancyjnym i umożliwia zmniejszenie kompensacji bez uszczerbku dla stabilności w zakresie wysokich częstotliwości.

Korzyść z zastosowania kondensatora 15 pF w obwodzie kompensacyjnym jest mniej oczywista. Symulacje sugerują, że nieznacznie zmniejsza on przesunięcie fazowe w pobliżu VAS przy bardzo wysokich częstotliwościach, a jednocześnie ma pomijalny wpływ na wzmocnienie. Wydaje się jednak, że łączny efekt tych dwóch zmian jest taki, że jeśli wzmacniacz „źle się zachowuje”, jest o wiele mniej prawdopodobne, że wpadnie w szkodliwe oscylacje wysokiej częstotliwości.

Przy okazji, przetestowaliśmy wszystkie zmiany w układzie w symulacjach SPICE, aby sprawdzić, czy są one sensowne, ale ostatecznie musieliśmy wypróbować je wszystkie na prototypie, aby sprawdzić ich wpływ na zniekształcenia i stabilność. Symulacja jest dobrym sposobem na szybkie sprawdzenie, czy zmiana jest złym pomysłem, bez wysadzania wzmacniacza w powietrze, ale jeśli symulacja pokazuje, że coś powinno działać, to wcale nie jest pewne, czy rzeczywiście tak będzie.

Jednym z obszarów, w którym symulacja jest najlepsza, jest możliwość zobaczenia, co dzieje się w układzie. Na przykład można łatwo wyświetlić natężenie prądu przepływającego przez dowolny element układu, podczas gdy zrobienie tego na prawdziwym prototypie wymagałoby odlutowania elementu i wstawienia bocznika, co mogłoby zakłócić działanie układu.



Ten zasilacz może zasilać dwa moduły wzmacniacza Ultra-LD Mk.4 i zostanie opisany w przyszłym miesiącu w części 3

Szkoła Konstruktorów



W Szkole Konstruktorów może wziąć udział każdy Czytelnik EdW, także i Ty!

Możesz zostać stałym uczestnikiem Szkoły, ale możesz tylko jednorazowo nadesłać pojedyncze rozwiązanie jednego zadania, które Cię najbardziej zainteresowało. Nie trzeba się zapisywać, nie ma żadnych zobowiązań – można tylko zyskać. Co miesiąc przydzielane są punkty, upominki, nagrody i kupony do Sklepu AVT, a raz na rok najaktywniejsi uczestnicy Szkoły Konstruktorów są nagradzani dodatkowo. W każdym numerze zamieszczane są zadania trzech klas (*Zadanie główne, Co tu nie gra?* oraz *Policz*).

W terminie dwóch miesięcy możesz więc nadesłać e-mailem na adres: szkola@elportal.pl (szkoła, a nie szkoła), rozwiązanie jednego, dwóch albo wszystkich trzech zadań Szkoły z danego numeru.

Potwierdzam otrzymanie rozwiązań, nadsyłanych e-mailem. Jeśli w terminie dwóch tygodni nie otrzymasz mojego potwierdzenia, prześlij rozwiązanie jeszcze raz.

Bardzo proszę: dla ułatwienia segregacji niech tytuł Twojego e-maila (i nazwa każdego ewentualnego załącznika), oprócz **nazwy konkursu** oraz **numeru zadania**, zawiera też **Twoje nazwisko** (najlepiej bez typowo polskich liter), na przykład: *Szko300Kowalski, Policz300Zielinski, NieGra300Malinowski, Jak02Krzyzanowski*. Chodzi o to, żeby w tytule e-maila i w nazwach wszystkich załączników była zarówno informacja o zadaniu, jak i o Autorze. Bardzo też proszę, żeby jeden Twój

e-mail zawierał rozwiązanie tylko jednego konkursu, a nie kilku, co znacznie mi ułatwi segregowanie poczty.

Do wysyłki nagród i upominków potrzebny jest Twój adres pocztowy. Oszczędzisz mi sporo niepotrzebnej pracy, jeśli podasz go w jednej linii: **imię i nazwisko ulica i numer domu kod pocztowy miejscowość e-mail**.

Jeśli na łamach czasopisma nie chcesz ujawniać imienia i nazwiska – napisz, a zachowam dyskrecję, podając albo pseudonim, albo imię i pierwszą literę nazwiska, ewentualnie miejscowość zamieszkania. Jeśli nadesłesz rozwiązanie zadania głównego, możesz dołączyć swoją fotografię (portret), która będzie zamieszczona przy rozwiązaniu zadania. Zachęcam też do podawania **roku urodzenia, a w przypadku uczniów i studentów także informacji o szkole/klasie lub uczelni**. Jest to pomocne przy opracowywaniu i ocenie rozwiązań (Twoje dane nie są nigdzie przekazywane, tylko wykorzystywane w redakcji EdW wyłącznie w związku z oceną prac i przydzielanymi nagrodami).

Najbardziej cieszę się z krótkich i zwięzłych rozwiązań, bo to ułatwia ich opracowanie. Ale jeżeli Twoje rozwiązanie będzie obszerniejsze, mam prośbę dotyczącą kwestii technicznych: Nie umieszczaj ilustracji w tekście! Wszystkie ilustracje (fotografie i rysunki) prześlij w e-mailu jako oddzielne pliki – załączniki. Bardzo proszę też o przysyłanie schematów, projektów płytek

i wszelkich innych rysunków w popularnych formatach, na przykład PDF, SVG, JPG, GIF czy PNG, i to także wtedy, gdy przysyłasz oryginalny, źródłowy plik z danego programu projektowego (.sch, .pcb, .brd, .ddb, itp.).

Jeżeli w ramach zadania głównego zrealizujesz rozwiązanie praktyczne, czyli zbudujesz konkretny układ-model, mam następujące wskazówki i prośby:

Nie przysyłaj modelu do redakcji! Nie ma też potrzeby nadsyłania papierowych wydruków, płyty CD/DVD, ani modelu – całkowicie wystarczy załączone do e-maila pliki i fotografie zrobione przez Ciebie.

Przygotowując opis **skorzystaj z szablonu** dostępnego pod adresem:

<http://edw.elportal.pl/szablon>

Więcej wskazówek na temat przygotowania materiałów i prawidłowego fotografowania modeli znajdziesz w Elportalu na stronie: <https://edw.elportal.pl/zostan-wspolautorem-elektroniki-dla-wszystkich/>.

Twoje praktyczne rozwiązanie głównego zadania Szkoły może być później opublikowane jako artykuł w EdW, za który otrzymasz honorarium. Dlatego w treści e-maila umieść wtedy tekst: *Oświadczam, że materiał, który przesyłam w tym e-mailu do redakcji „Elektroniki dla Wszystkich”, jest moim osobistym opracowaniem i nie był wcześniej nigdzie publikowany.*

REKLAMA

młody m.technik

Ciekawi świata są zawsze młodzi

przejrzyj i kupisz na
www.ulubionykiosk.pl

eprasa.pl 0d581f01a1



Zadanie główne 315

Pomysłodawcą zadania głównego numer 315 jest **Ryszard Nowak**, którego e-maile były cytowane w rubryce Poczta, ostatni raz bodaj w numerze 3/2022. Oto fragment jednego z listów:

(...) moduł, o którym napisałem zainteresowałem mnie głównie ze względu na wyświetlacz LCD 1602 i wymiary, co bardzo dobrze by się komponowało w zasilaczu, który przerabiam. Może strzałem w dziesiątkę byłoby opracowanie przez EdW modułu na wzór tego, o którym wspomniałem, opartego na wspomnianym przez Pana MCP 3421 i wyświetlaczu znakowym 1602. [Wskazania] modułu z TM7705 „startują” od 5–8 mA, bo tak wychodzi mi z pomiarów (...) Pozdrawiam

Korespondencja dotyczyła modułowych mierników prądu stałego, przede wszystkim woltomierzy. Temat ten był też omawiany w serii artykułów cieszących się bardzo dużym zainteresowaniem, a dotyczącej budowy, możliwości zastosowań oraz dokładności tego rodzaju modułów. Ryszard proponuje opracowanie w EdW modułu pomiarowego: woltomierza i amperomierza do akumulatora. Modułu przeznaczonego jako dodatkowe wyposażenie zasilacza, który nie ma mierników napięcia i prądu wyjściowego.

Projektowanie modułu to poważne i kosztowne zadanie, dlatego ja rozszerzam pierwotną propozycję. Nie trzeba projektować modułu, bowiem można wykorzystać gotowe moduły woltomierzy i amperomierzy, dostępne bez problemu w atrakcyjnych cenach.

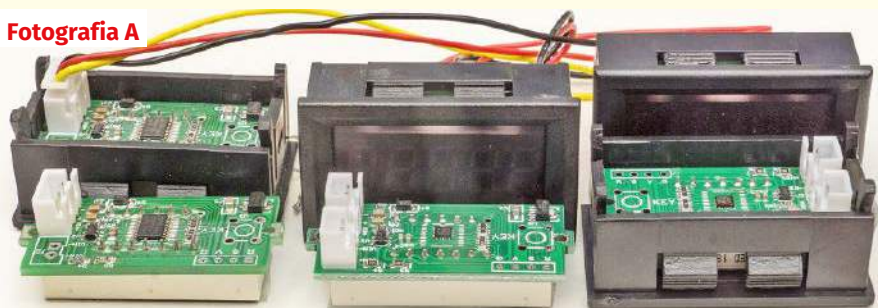
Oto temat zadania 315:

Zaproponuj wykorzystanie i modyfikacje gotowych modułów woltomierzy i amperomierzy, przede wszystkim tych zawierających precyzyjny 18-bitowy przetwornik ADC typu MCP3421.

Czyli nie projektujemy od zera, tylko wykorzystujemy, ewentualnie modyfikujemy moduł fabryczny. Tym sposobem zadanie 315 staje się dostępne dla każdego. Propozycje mogą być praktyczne oraz teoretyczne.

W razie potrzeby wróćcie do niedawnej serii artykułów na temat takich modułów. Podkreślam tylko, że przytłaczająca większość tanich modułów wykorzystuje marnej

Fotografia A



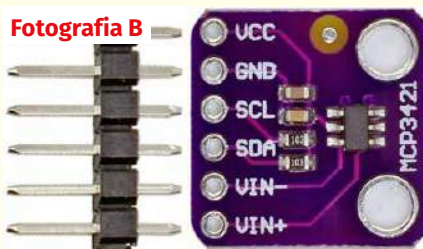
jakości przetwornik ADC wbudowany w tani mikroprocesor – dokładność takich mierników, w szczególności amperomierzy, jest słaba.

Opisywałem też moduły znacznie dokładniejsze, i zewnętrznymi przetwornikami ADC. Autor propozycji testował tego rodzaju moduły z przetwornikiem ADC typu TM7705 i LM385 (LM285) w roli źródła napięcia odniesienia. Nie zapewni on wysokich parametrów. Najlepsze okazują się moduły z 18-bitowym przetwornikiem $\Sigma\Delta$ typu MCP3421, który ma wbudowane przyzwoite źródło napięcia odniesienia o stabilności 15 ppm/K.

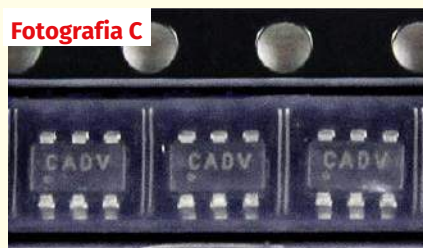
Proponuję skupić się na modułach z tym właśnie przetwornikiem, bo wtedy sensowne jest wykorzystanie nawet wyświetlacza 5-cyfrowego. Zwracam uwagę na trzy główne kierunki, z których polecam przede wszystkim dwa pierwsze:

1. Dostępne są kompletne, gotowe moduły woltomierzy i amperomierzy z kostką MCP3421 (**fotografia A**). Tego rodzaju moduły można wykorzystać bez żadnych przeróbek albo po małych przeróbkach wewnątrz.
2. Dostępne są też moduły zawierające tylko przetwornik MCP3421 (**fotografia B**). Taki przetwornik można bez trudu podłączyć do mikroprocesora, choćby do Arduino – dostępne są stosowne biblioteki. Tu ogromną zaletą jest możliwość pełnego wykorzystania możliwości przetwornika ADC. W jednym z ostatnich artykułów o modułowych miernikach przeanalizowałem niemożliwość pełnego wykorzystania potencjału MCP3421 przy wykorzystaniu gotowego woltomierza z fabrycznym programem. Dołączając

Fotografia B



Fotografia C



do Arduino jeden lub kilka modułów (dwa bez problemu), można tak napisać program i tak ustawić zakresy pomiarowe, by „wycisnąć” z bardzo porządnego MCP3421 wszystko, co się da. Polecam tę drogę!

W ten sposób najprościej można zrealizować naprawdę precyzyjny woltomierz czy amperomierz napięcia i prądu stałego.

Jest i trzeci kierunek – najtrudniejszy, realizujący pomysł Ryszarda: dostępne są same małe układy scalone MCP3421 (**fotografia C** z Aliexpress). Można dodać mikroprocesor, zaprojektować płytkę i stworzyć własny moduł, skrojony na miarę konkretnych potrzeb.

Zachęcam do udziału w tym bardzo praktycznym zadaniu 315. Gwarantuję, że satysfakcja z realizacji we własnym zakresie naprawdę dokładnego miernika jest ogromna! ■

Piotr Górecki

Nadsyłajcie propozycje zadań!

Autorzy propozycji zadań, które zostaną wykorzystane w Szkole, otrzymują jako nagrodę kupon 100 zł na zakupy w sklepie AVT: <http://sklep.avt.pl>
Koszty przesyłki pokrywa AVT. Dobra propozycja nie powinna być ani zbyt trudna, ani zbyt ogólna, ani zbyt wąsko ukierunkowana. Dobre zadanie Szkoły powinno mieć na tyle szeroki zakres, żeby mogli w nim wziąć udział zarówno doświadczeni elektronicy, jak i początkujący, w tym najmłodszy.

Zachęcam do nadsyłania propozycji następnych zadań Szkoły!



Rozwiązanie zadania głównego 310

Temat styczniowego zadania 310 brzmiał: **Zbadaj zagadnienie kontroli wilgotności gleby i zaproponuj sposób sygnalizacji wysychania trawnika wykorzystujący elektronikę.**

Pomysłodawca tego zadania, **Dariusz** (prosił o anonimowość) napisał: (...) Mam trawnik od południa posiany na piachu, który szybko wysycha na pył, jak dłużej nie ma deszczu. (...) Czas leci i już parę razy było tak latem, że zaczął podsychać. Zaczęły się robić żółte plamy na środku (...) [chciałbym] na razie tylko zamontować na środku trawnika w najgorszym miejscu czujnik wilgoci, żeby mi dał znać w tygodniu, że koniecznie trzeba podlać. Już trochę myślałem i szukałem w Internecie (...) jak pociętałem, to wychodzi na to, że to nie jest aż tak proste, jak się wydaje z początku. I pomyślałem, że napiszę do EdW (...) bo jak było o kompozycji, to może też o trawniku. (...) na razie mi chodzi tylko o alarm, jak trawnik wysycha podczas upałów (...)

Zadanie było wbrew pozorom trudne, specyficzne, a temat jest ważny dla wielu elektroników.

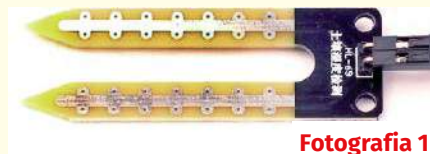
20-letni **Łukasz** (...) z Warszawy przedstawił następujący pomysł: (...) też myślę o czymś takim (...) czujnik jest prosty (...) mogą być dwa gołe miedziane druty elektryczne wbite w ziemię (...) są też w sklepach dla elektroników gotowe czujniki wilgotności (...) z czujnikiem nie ma kłopotu (...) tylko żeby nie zerwać kabli (...) a w domu jakiś układ

mierzący określoną rezystancję. Jak ziemia wysychając zwiększa rezystancję, spowoduje zadziałanie alarmu. Moim zdaniem to trzeba jakoś sprawdzić eksperymentalnie, przy jakiej rezystancji ma być alarm. (...) [Mam] kilka propozycji (...) wszystko analogowe (...) [W układzie według rysunku 1] trzeba tylko nastawić potencjometr (...) tu nie ma potrzeby komplikacji i wystarczy (...) [układ] najprostszy z jednym wzmacniaczem operacyjnym, ale można też na mikroprocesorze (...) jak wszyscy na Arduino (...) jak na rysunku [2] (...) Czujnik taki sam, ale przez przetwornik analogowo-cyfrowy można dokładnie zmierzyć wilgotność, żeby (...) pokazywał na wyświetlaczu i do tego też alarm wysychania (...) lepiej z czujnikiem światła, albo z zegarem RTC, żeby nie obudził całego domu o północy (...)

Autorowi przydzielam upominek i trzy punkty. Jednak zaproponowane układy najprawdopodobniej nie spełnią przewidzianej roli. Można byłoby omawiać różne błędy i usterki występujące na rysunkach 1, 2, ale wspomnijmy tylko o jednym. Otóż wielu hobbystów przekonało się, że w najtańszych i najpopularniejszych czujnikach rezystancyjnych (fotografia 1), wykonanych z płytki drukowanej (laminatu pokrytego warstwą miedzi), po dłuższym czasie ta miedź... znika. Jest to najprościej biorąc, wynik elektrolizy. Aby takie „zjadanie miedzi” nie wystąpiło, tego rodzaju czujniki powinny pracować przy zasilaniu

napęciem przemiennym, a nie stałym. To jednak jest drobny szczegół. Skupmy się na najważniejszym,

Otóż poszczególne rodzaje gleb mają różne właściwości jeżeli chodzi o przewodzenie prądu elektrycznego. W omawianym przypadku sprawa jest prostsza o tyle, że chodzi



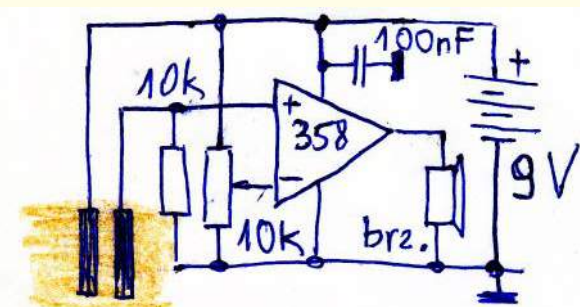
Fotografia 1

o pomiar jednego rodzaju gleby, o zmiany wilgotności w jednym jedynym punkcie (prawdopodobnie na środku trawnika). Jednak z drugiej strony chodzi o glebę piaszczystą, której rezystancja szczególnie silnie zależy będzie nie tylko od wilgotności, ale też od zawartości soli mineralnych, a konkretnie nawozów (sztucznych), którymi okresowo są zasilane trawniki. I właśnie **silny wpływ nawozów mineralnych na rezystancję wręcz przekreśla możliwość sensownego określenia wilgotności gleby na podstawie jej rezystancji.**

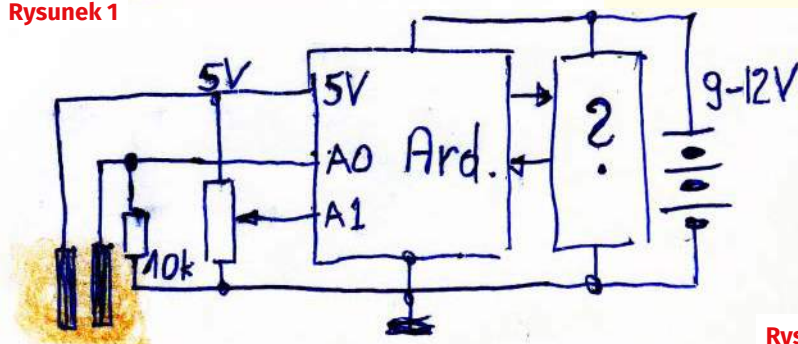
Wilgotność gleby trzeba mierzyć inaczej, nie prymitywną metodą rezystancyjną, o czym szerzej za chwilę. Jej rezystancję (przewodność elektryczną), owszem, warto mierzyć, ale dodatkowo, co przy uwzględnieniu wilgotności pozwoli ocenić właśnie stopień fertylizacji i konieczność zasilenia trawnika nawozem. Te tematy są mało znane, a warto je poznać, ponieważ dziś dostępne są niedrogie czujniki, pozwalające w sensowny sposób mierzyć aż trzy parametry: wilgotność (*humidity*), przewodność elektryczną (*conductivity* – EC) oraz temperaturę gleby. Przykład na fotografii 2.

O potrzebie szerszego omówienia tematu kontroli wilgotności gleby świadczą też liczne projekty w Internecie, które okazują się zupełnie nieprzydatne w praktyce, ponieważ ich Autorzy koncentrują się tylko na najprostszych sposobach pomiaru rezystancji gleby i na tej podstawie chcą określać jej wilgotność.

A co do innych sposobów pomiaru wilgotności, to **Radosław Nowak** z miejscowości Miętne napisał: *Dzień dobry! Chciałbym wziąć udział w rozwiązaniu zadania (...) tylko teoretycznie. Mam taki pomysł (...) w roli dwóch sond dwa takie same kawałki nierdzewki albo blachy miedzianej czy podobnej z przewodami, wbite w izolującą rączkę (...) jak na rysunku [3] (...) wcisnąć w ziemię (...) uniwersalny czujnik, co może mierzyć i rezystancję między sondami, ale w pierwszej kolejności pojemność (...) i wtedy można je polakierować. (...) po informacjach [z poprzednich numerów EdW wiem] (...), że trzeba mierzyć zmiany pojemności. Dlatego mój czujnik ma sondy w formie jak największych*



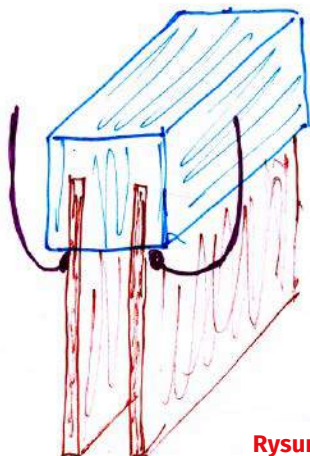
Rysunek 1



Rysunek 2



Fotografia 2



Rysunek 3

blach. Żeby pojemność była względnie duża, blachy powinny być najbliżej siebie (...) obawiam się, czy się uda (...) nie można za blisko (...) między nimi ma być ziemia (...) nie wiem czy wystarczy blachy 5x5 cm, czy może trzeba 50x50 cm, ale wydaje mi się, że zmiany pojemności będą duże, bo przenikalność wody jest 80 razy większa od powietrza (...) jakby pojemność wyszła duża, to można z takiego kondensatora zmiennego zbudować generator RC na 555 albo LC na 1 tranzystorze i mierzyć mikroprocesorem częstotliwość. (...) tylko pomysł, nie zrobiłem (...)

Pomysł jest interesujący, warto byłoby poeksperymentować z różnej wielkości czujnikami pojemnościowymi, choć z góry wiadomo, że w praktyce wykorzystanie dużego czujnika tego rodzaju byłoby kłopotliwe.

Krzysztof z Krakowa przy okazji korespondencji na temat Radiowej Oślej Łączki (ROŁ) napisał do mnie między innymi: (...) Miałem nawet w planie wziąć udział w Szkole Konstruktorów, odpowiadając na zadanie o wilgotności gleby i opisać zalety metody opartej o mikrofalę. To mógłby być przyszły projekt wykorzystujący wiedzę z ROŁ. (...)

Zapewne byłoby to bardzo interesujący układ i przy okazji można byłoby się wiele nauczyć z dwóch tak odległych dziedzin. Mam nadzieję, że pomysł zostanie zrealizowany właśnie w związku z tak szeroko dyskutowanym kursem Radiowej Oślej Łączki.

Rafał Orodziński z Białegostoku w rozwiązaniu tego zadania napisał: Dzień Dobry (...) Z metod pomiaru wilgotności gleby najczęściej używana jest metoda wykorzystująca zmianę pojemności kondensatora, w którym funkcję dielektryka pełni gleba. Taki pomiar jest możliwy ze względu na dużą różnicę pomiędzy

wartością stałej dielektrycznej wody, a reszty składników gleby. Metoda ta mimo pewnych wad jest najczęściej używaną metodą pomiaru wilgotności gleby. Ja pierwsze eksperymenty przeprowadziłem z wykorzystaniem pojemnościowych czujników wilgotności zakupionych na portalu aukcyjnym – **fotografia 3**. Czujniki te wykorzystują układ 555 generujący częstotliwość nieco powyżej 1 MHz. Elektrody sondy tworzą dzielnik rezystancyjno-pojemnościowy pomiędzy generatorem a masą. Wilgotność gleby określana jest na podstawie zmierzonego napięcia w.c.z. Wykorzystując te układy w warunkach polowych, należy osłonić część elektroniczną od warunków zewnętrznych, umieszczając ją w obudowie, ewentualnie zalewając żywicą.

Zastosowane czujniki wilgotności gleby wykazują dość duży rozrzut parametrów oraz wpływ temperatury na wynik pomiaru. Wpływ temperatury na wynik pomiaru można zminimalizować, realizując pomiar i ewentualnie następnie podlewanie w nocy – czujniki nie są wtedy narażone na nagrzewanie przez promieniowanie słoneczne.

Głównym problem przy użyciu tego typu czujników okazał się problem wrażliwości czujnika na przyleganie gleby do czujnika. Poprawę kontaktu sond z glebą można uzyskać stosując sondę z elektrodami prętowymi. (...) wykonany czujnik pokazany jest na **fotografii 4**. Układ wytwarza kilka częstotliwości z zakresu od 30 do 75 MHz. Częstotliwość stabilizowana jest rezonatorem kwarcowym. Elektrody wykonane są z stali kwasoodpornej, można je zabezpieczyć dodatkową odpowiednią farbą antykorozyjną. Elektrody są odizolowane od napięć stałych. Przeprowadzone testy na różnych rodzajach gleb pokazały wpływ struktury gleby na wynik pomiaru, jednak przy pomiarach porównawczych w jednym miejscu nie ma to w praktyce większego znaczenia. Możemy ocenić kiedy



Fotografia 3

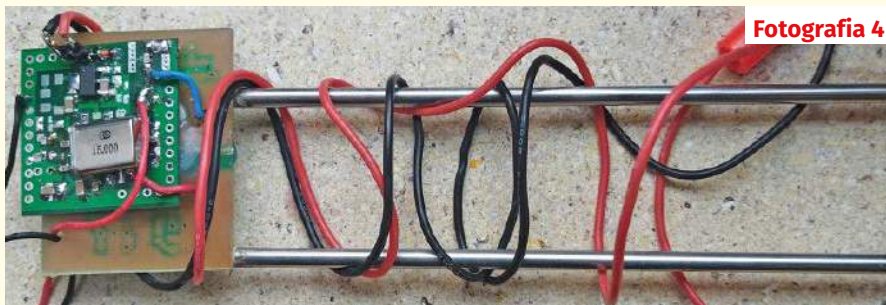
gleba jest odpowiednio wilgotna wizualnie i taką wartość uznać za referencyjną. Warto zauważyć, że woda podczas podlewania potrzebuje czasu by „rozejść” się w glebie, widać to doskonale na wynikach pomiarów, gdy od chwili końca podlewania wilgotność wzrasta nawet przez 24 h (czas ten zależy silnie od rodzaju podłoża).

Docelowo układ będzie pełnym sterownikiem podlewania w ogrodzie, sterującym sześcioma sekcjami. Panel regulatora będzie umieszczony w domu tak, by łatwo móc kontrolować wilgotność. Komunikacja z modułem wykonawczym odbywa się bezprzewodowo, docelowy układ już w trakcie obudowywania pokazany jest na **fotografii 5**.

Za opisane eksperymenty, a także za prace nad wspomnianym systemem przydzielam duży kupon. I to nie tyle Rafałowi – ojcu, co jego synowi, który jak wiem realizuje bardzo ważną część tego przedsięwzięcia.

A co można byłoby finalnie doradzić pomysłodawcy zadania – Dariuszowi, posiadaczowi wysychającego trawnika?

Na pewno nie ma rozwiązań bardzo prostych i bardzo tanich. Czujniki rezystancyjne to ślepy zaułek. Trzeba wykorzystać czujnik pojemnościowy. Jeżeli instalacja miałaby być realizowana samodzielnie i w grę wchodzi późniejsza rozbudowa systemu, to raczej należałoby polecić mikroprocesor (np. Arduino) oraz fabryczne czujniki pojemnościowe (FDR) mierzące też przewodność i temperaturę



Fotografia 4



(fotografia 2), dołączone przewodowo do modułu Arduino. Program trzeba by napisać samodzielnie. Dedykowana biblioteka do obsługi czujnika nie jest potrzebna – warto wykorzystać wersję czujnika z interfejsem RS485, który wykorzystuje komendy MODBUS. A wtedy potrzebna jest tylko informacja, spod których adresów MODBUS można odczytać zmierzone przez czujnik wilgotności, przewodność i temperaturę. Te informacje zawarte są w ulotce, która dostarcza producent czujnika. Wtedy oczywiście oprócz modułu Arduino trzeba dodać konwerter RS485. Taka wersja powoli na dalszą rozbudowę i to w dowolnym zakresie.

Przewody ułożone w ziemi powinny wytrzymać tam przez wiele lat. Dlatego w zasadzie należałoby zastosować czarny, żelowany kabel ze skrętką. Z powodzeniem może to być czarna, żelowana skrętka komputerowa UTP. A jeśli grunt jest piaszczysty, to może wystarczyć zwykła skrętka UTP?

Jeżeli ktoś obawia się Arduino albo nie chciałby stosować przewodów, to powinien wiedzieć, że dostępne są też zasilane bateryjnie czujniki wilgotności gleby bezprzewodowe, wykorzystujące Wi-Fi, współpracujące z elektrozaaworem, realizujące automatyczne nawadnianie, obsługiwane też przez aplikację na smartfonie z opcją podłączenia do systemu inteligentnego domu – rysunek 4 (z Aliexpress).

Zachęcam do samodzielnego zgłębiania bardzo interesującego i praktycznego tematu pomiaru parametrów gleby i nawadniania!

Aktualne informacje o punktacji oraz rozdziale nagród, upominków i kuponów podane są w tabelkach. Znak zapytania oznacza, że ewentualna publikacja nastąpi dopiero po nadesłaniu ostatecznych materiałów. Osoby nagrodzone kuponami otrzymują z AVT stosowny e-mail z informacją i wskazówkami, a dopiero potem zamawiają w sklepie AVT (wrzucając do koszyka pod adresem www.sklep.avt.pl) towary za przydzieloną sumę, a w uwagach piszą, że jest to kupon ze Szkoły Konstruktorów. Kupon za zadania z kolejnych miesięcy można sumować, by kupić sprzęt o większej wartości. Istnieje też możliwość dopłaty różnicy cen w przypadku zamówienia na sumę większą niż przydzielony kupon. Ale **uwaga: kupon ważny jest tylko 12 miesięcy – po tym terminie traci ważność i przepada.**

Serdecznie zapraszam do udziału w zadaniu głównym 315, a także w drugiej i trzeciej klasie naszej Szkoły Konstruktorów! Zachęcam uczestników, żeby praktyczne rozwiązania zadań Szkoły przygotowywali według Szablonu ze strony:

<http://edw.elportal.pl/zostan-wspolautorem-elektroniki-dla-wszystkich/> ■

Piotr Górecki



Fotografia 5

Imię	Nazwisko	Miejscowość	Punkty	Publikacja	Nagroda	Talon AVT PLN
Ryszard	Nowak		—	—	—	—
Łukasz	—	Warszawa	3	—	U	—
Radostaw	Nowak	Miętne	2	—	U	—
Krzysztof	—	Kraków	—	—	—	—
Rafał	Orodziński	Białystok	6	?	—	300

Punktacja Szkoły Konstruktorów



Sławomir Węgrzyn Dziekanowice.....	92	Szymon Wójtowicz Warszawa.....	13
Michał Stach Kamionka.....	88	Piotr Wyderski.....	13
Daniel Turbasa Kraków.....	88	Michał Zięba Poznań.....	13
Łukasz Dachowski Cymbark.....	72	Andrzej Adamczyk Ostrowiec Św.....	11
Artur Bereit Barcin Wieś.....	69	Jakub Jakubczyk Kluczbork.....	11
Aleksander Bernaczek Magnuszowice.....	69	Piotr Świerczek Bielsko-Biała.....	11
Krzysztof Smoliński Poznań.....	68	Zygmunt Flisak Opole.....	10
Szymon Trygar Szczecin.....	66	Michał Lis Gdynia.....	9
Radostaw Smalec Zabrze.....	64	Maciej Skrodzewicz Szczecin.....	9
Paweł Hoffmann Wrocław.....	62	Paweł Błaszczak.....	8
Rafał Orodziński Białystok.....	59	Adam Sosnowski Kolutki.....	8
Robert Szolc Bytom.....	58	Andrzej Kubiak Rumia.....	7
Andrzej Herbut Siekierczyn.....	52	Michał Stomkowski Poznań.....	7
Łukasz Olszok Tarn. Góry.....	51	Marcin Bambynek Kalety.....	6
Adam Ples Jaworzno.....	51	Piotr Chrobok Piekary Śląskie.....	6
Sebastian Jarmosiewicz Motwica.....	50	Wojciech Goliszewski Szczecin.....	6
Adam Sobczyk Warszawa.....	50	Piotr Graffstein Warszawa.....	5
Circuit Chaos Warszawa.....	48	Michał Grzemski Grudziądz.....	5
Michał Pędzimąż Stara Słupia.....	48	Mariusz Hejto Łowczówek.....	5
Rafał Równiak Maciejów.....	46	Janusz Pańczyk Poręba.....	5
Krzysztof Kawa Lubcza.....	44	Tomasz Zygmunt Szczecin.....	5
Dawid Placha Rdzawa.....	44	Dominik Badura Ustroń.....	4
Szymon Czepiel Piszarowice.....	43	Marek Czechyra Dzierżonów.....	4
Piotr Gajdosz Grybów.....	41	Łukasz Warszawa.....	3
Maciej Zieliński Kraków.....	41	Wojciech Bagiński Biała Podlaska.....	3
Teodor Woźniak Łódź.....	35	Przemysław Kamiński Warszawa.....	3
Jarostaw Węgliński Warszawa.....	34	Krzysztof Łos Hubenice.....	3
Tomasz Zaorski Kalinówka.....	34	Zbigniew Ryba Orzechowo.....	3
Łukasz Nowak Gdańsk.....	33	Krzysztof S. Adamówek.....	3
Andrzej Nowicki Warszawa.....	29	Marcin Sosnowski Pólko.....	3
Jacek Konieczny Poznań.....	26	Michał Szymaniak Poznań.....	3
Piotr Grzegorzczak Siedlce.....	25	Marcin Wierzbicki Nowa Wieś.....	3
Jacek Rączka Połomia.....	25	Paweł Marchewka Pionki.....	2
Marian Gabrowski Polkowice.....	23	Radostaw Nowak Miętne.....	2
Roman Braumberger Bytom.....	21	Marek Woś Łazy.....	2
Jakub Gajda Kraków.....	20	Leszek Rzesutek Warszawa.....	1
Marian Caruk Luban.....	17		
Bogdan Kosiński Szczecin.....	16		
Łukasz Kojro Gdańsk.....	15		
Marcin Malich Wodzisław Śl.....	13		
Paweł Sablik Piszarowice.....	13		



Co tu nie gra? Zadanie 315

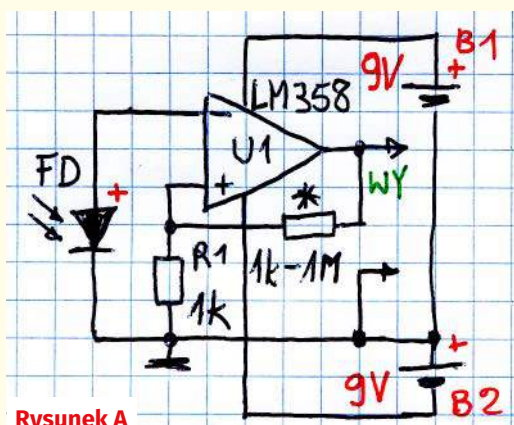
Na **rysunku A** przedstawiony jest schemat prostego eksperymentalnego odbiornika sygnałów świetlnych i podczerwieni z fotodiodą, który może wzmacniać i sygnały stałe, i zmienne.

Jak zwykle pytanie brzmi:

Co tu nie gra?

Nawet gdy w układzie jest kilka usterek, możesz zgłosić tylko jedną. Bardzo proszę o możliwie krótkie odpowiedzi.

Odpowiedź oznacz **NieGra315** i nadeślij w terminie 60 dni od ukazania się tego numeru EdW. Od razu podaj też swój adres pocztowy, żebym nie musiał pytać, gdy przydzielę upominek. Możesz jeszcze przysłać rozwiązania zadania **NieGra** z poprzedniego miesiąca. Uczestnicy konkursu otrzymują upominki, a najaktywniejsi uczestnicy są co rok nagradzani bezpłatnymi prenumeratami EdW lub innego wybranego czasopisma AVT.



Rysunek A

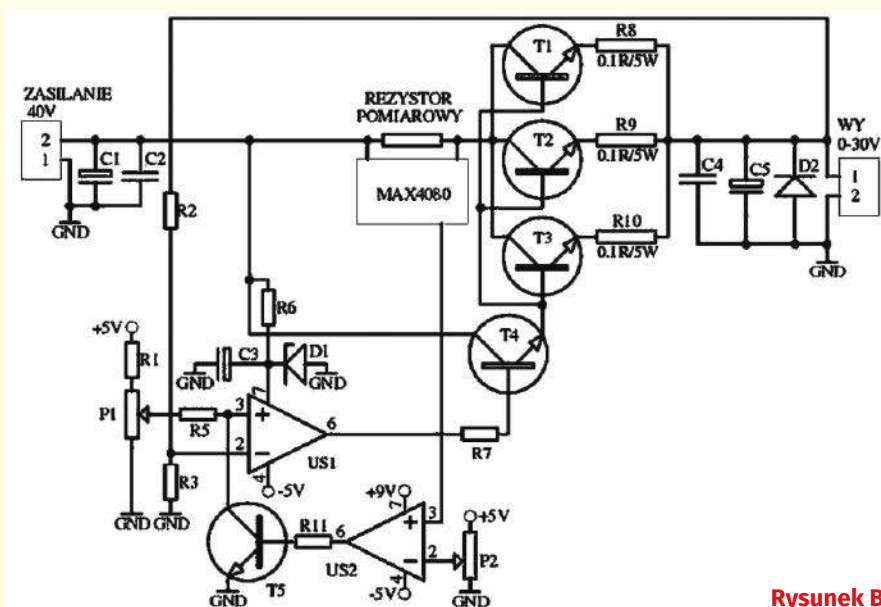
Co tu nie gra? Rozwiązanie zadania 310

Na **rysunku B** pokazany jest zamieszczony w EdW 1/2022 „nieco uproszczony schemat zasilacza ze stabilizacją napięcia i prądu”.

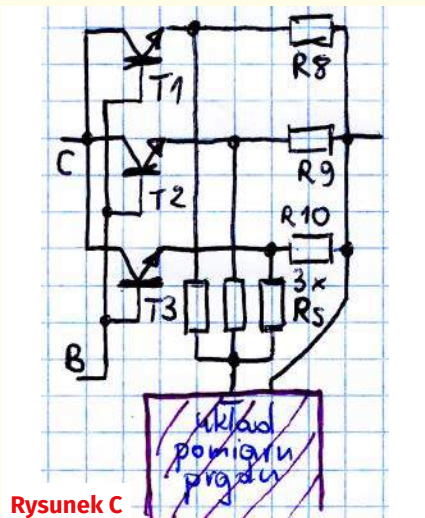
Zadanie było bardzo trudne, jeżeli ktoś postanowił znaleźć wszystkie błędy i usterki. Było łatwe, jeżeli ktoś chciał znaleźć tylko jeden błąd, pozwalający na wzięcie udziału w zadaniu **NieGra**.

O błędach za chwilę. Ale najpierw o czymś innym. Schemat został nadesłany przez jednego z Czytelników, który prosił o konsultacje i opinię. I bardzo słusznie, ponieważ projektowanie zasilaczy tylko na pozór wydaje się łatwe i proste. Duży szacunek dla Autora schematu, który zaproponował interesujące, godne uwagi rozwiązanie! Wymagające jednak dopracowania.

Warto dokładniej przeanalizować tę propozycję, ponieważ z podobnymi problemami ma też do czynienia wielu innych elektroników, którzy próbują samodzielnie zaprojektować układ, a nie



Rysunek B



Rysunek C

korzystać z gotowych schematów z Internetu, z których większość ma zresztą znikomą wartość praktyczną, ponieważ też zawiera mnóstwo błędów i niedoróbek. O to omówienie nadesłanych rozwiązań. Zaczę od drobiazgów, wypatrzonych przez najbardziej uważnych.

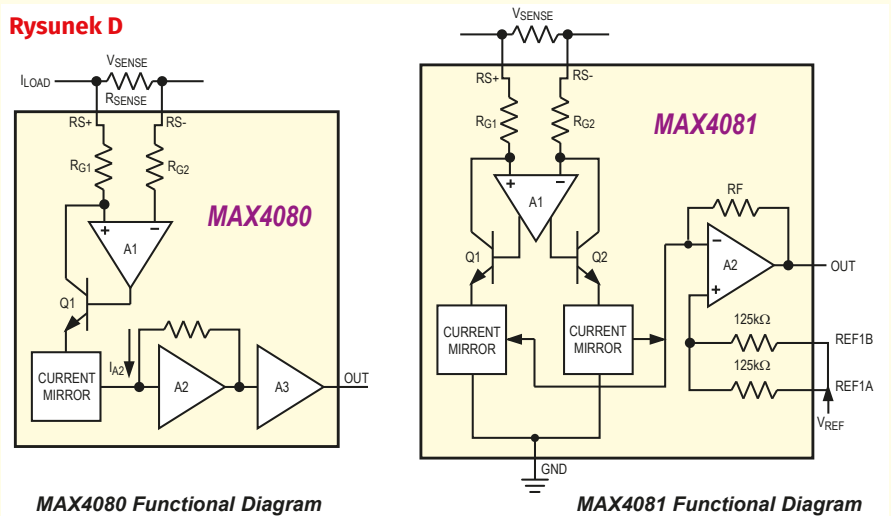
Ktoś napisał: D2 jest niepotrzebna, co zresztą może być podstawą jałowej dyskusji, bo można byłoby szukać egzotycznych przypadków, gdy byłaby potrzebna (np. przy wykorzystaniu zasilacza w roli ładowarki przy błędnym, odwrotnym dołączeniu akumulatora). Generalnie dioda nie wygląda na potrzebną, ale jej obecność nie przeszkadza, nie jest błędem.

Jeżeli chodzi o drobiazgi, jeden z uczestników napisał: (...) Złącze zasilania 40 V jest oznaczone na odwrót, czyli 1 jest masą, a 2 jest +, natomiast wyjście jest oznaczone na odwrót. Czyli masa jest tam na 2. (...) Użycie MAX4080 jest dla mnie dziwne, bo to bardziej układ do zastosowań w motoryzacji stosowane do ładowania ogniów (...) A nawet nie unikając w działanie Maxa nie jest on podłączony do masy, co jest już dziwne, poza tym brak (...) [sposobu] uzyskania -5 V w układzie. Poza tym dziwnie zrobiony rezystor pomiarowy, te rezystory za emiterami tranzystorów to są rezystory pomiarowe, i od nich powinien być uzależniony obwód ograniczenia prądu oparty na opampach. Pozdrawiam

Brak źródła zasilania -5 V i brak symboli masy przy MAX4080 można uznać za wspomniane uproszczenie schematu. Jednak co do rezystora pomiarowego, to święta racja! Można i warto wykorzystać rezystory emiterowe R8...R10, na przykład według też uproszczonej koncepcji z rysunku C. Szczegóły będą zależeć, czym jest ten układ pomiaru prądu.

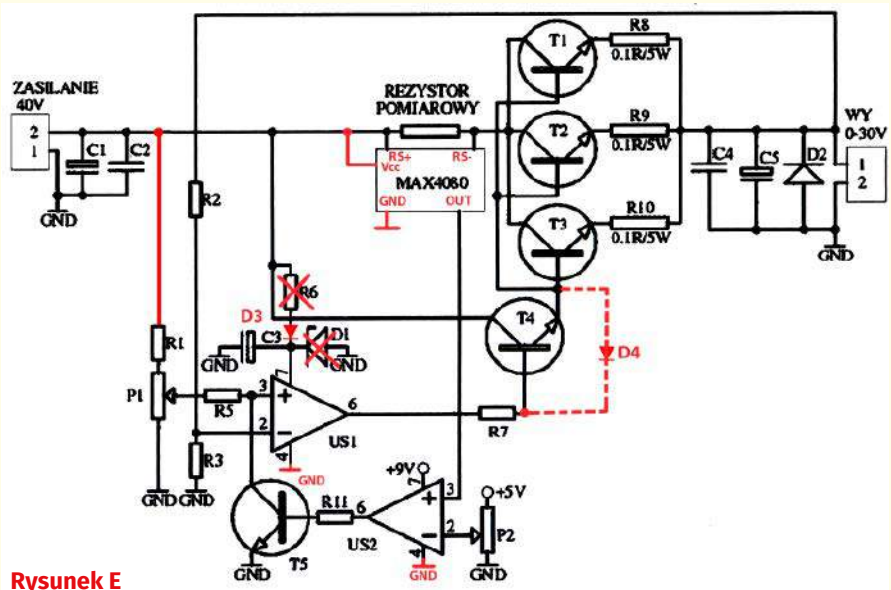
A jeżeli chodzi o problem braku symbolu masy, to także w karcie katalogowej przy

Rysunek D



MAX4080 Functional Diagram

MAX4081 Functional Diagram



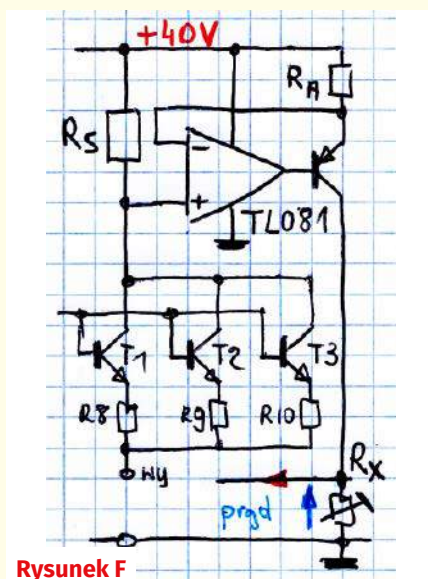
Rysunek E

schemacie blokowym MAX4080 tego symbolu nie ma, jak pokazuje rysunek D.

Przy okazji widać tu zasadę działania kostki MAX4080, która jest miernikiem jednokierunkowym, natomiast kostka MAX4081 może mierzyć przepływ prądu w obu kierunkach.

Jeden ze stałych uczestników przysłał propozycję poprawek – rysunek E. Napisał: (...) 1. Układ MAX4080 należy podłączyć jak pokazano na rysunku. Układ ten występuje w trzech odmianach różniących się wzmacnieniem. Dlatego jego dobór należy powiązać z maksymalnym prądem obciążenia i rezystorem pomiarowym. Tak, aby napięcie na pinie IN+ wzmacniacza US2 było nie większe niż 5 V (przynajmniej w tym przypadku). A najlepiej aby było nieco mniejsze, np. 4,9 V. Wtedy uzyskamy szeroki zakres regulacji. Od stanu pełnego otwarcia do stanu całkowitego ograniczenia prądu wyjściowego. 2. W tej roli powinien się sprawdzić LT1013. Ponieważ ten WO występuje

w obudowie typu „dual” to w tym przypadku możemy wykorzystać oba WO. 3. Kolejną kwestią pozostaje zasilanie WO a szczególnie US1. Jeżeli miałoby tak być jak na oryginalnym rysunku zadania (bezpośrednio) to należy uwzględnić diodę D3, która uniemożliwi rozładowywanie kondensatora C3 przez główny tor prądowy. Wyrzucam rezystor R6 i diodę Zenera D1. Nie jestem za takim rozwiązaniem zasilania WO bezpośrednio z głównej szyny zasilającej (prądowej). W moim wypadku wymaga to doboru WO na napięcie zasilania powyżej 40 V (44 V) co zawęży nam krąg doboru WO. Zdecydowałbym się na niższe napięcie zasilania 33 V. Nie mniej niż 33 V, gdyż aby układ regulacji działał prawidłowo, na wyjściu US1 potrzebujemy 32 V. Wspomniany LT1013 spełnia te warunki. Aby zasilac US1 napięciem 33 V potrzebujemy oddzielnego zasilania. Możemy zrobić prosty stabilizator z układem LM317LZ. Drugi WO operacyjny również zasilimy ze stabilizatora. Jeżeli byłby to LT1013,



to sprawa zasilania się upraszcza, gdyż zasilanie jest wspólne dla obu WO. 4. Na rysunku linią czerwoną przerywaną zaznaczono diodę D4. Może ona być konieczna aby napięcie na złączu B-E tranzystora T4, przy zmianie polaryzacji (w trakcie regulacji napięcia lub prądu), nie spowodowało uszkodzenia tranzystora (napięcie U_{BE} większe niż -5 V). Ale to należałoby sprawdzić w docelowym rozwiązaniu po dobraniu wszystkich elementów. Ostatnia kwestia to podłączenie dzielnika R1, P1. Jak na rysunku. Pozdrawiam

Cóż, tu otwiera się kolejne pytanie o błędy w tak zmienionym układzie...

W każdym razie rzeczywiście, pierwotny układ z rysunku B trzeba zmodyfikować. Jeden z Kolegów zaproponował układ, który po przerysowaniu przeze mnie pokazany jest na rysunku F. Napięcie na rezystorze R_x jest wprost proporcjonalne do prądu płynącego przez R_s . Autor słusznie stwierdził też, że prawdopodobnie dałoby się go zmodyfikować, żeby w roli rezystorów pomiarowych wykorzystać rezystory emiterowe tranzystorów mocy.

Wracamy do rysunku B. Kilka osób zaproponowało zasilanie wzmacniaczy operacyjnych US1, US2 napięciem pojedynczym, bez napięcia ujemnego -5 V . Ogólnie biorąc, to dobry pomysł, ale w tym przypadku należałoby być może rozważyć kwestię odwrotną: czy aby nie dołączyć emitera T5 do źródła napięcia -5 V ?

Problem w tym, że rzeczywiste wzmacniacze operacyjne mają jakieś napięcie niezrównoważenia, a tu oba pracują z otwartą pętlą, z maksymalnym wzmocnieniem.

W przypadku wzmacniacza US2 będzie to tylko oznaczało ewentualne kłopoty z minimalną wartością prądu ograniczania.

Gorzej z głównym wzmacniaczem US1. Drobną kwestią to analogiczne kłopoty z ustawieniem najniższego napięcia wyjściowego.

Główny problem z tym, czy w ogóle zadziała ogranicznik prądu?

Dotyczy to nie tylko wersji z zasilaniem napięciem pojedynczym, ale obu wersji! Idea realizacji ogranicznika (stabilizatora) prądu jest taka, że przy nadmiernym wroście prądu zaczyna przewodzić tranzystor T5 i zmniejsza napięcie wyjściowe, by przez to zmniejszyć też prąd wyjściowy. Pomysł zasadniczo prawidłowy, ale w tym przypadku ryzykowny, bo skuteczność działania zależy od biegunowości i wartości napięcia niezrównoważenia US1!

Mianowicie tranzystor T5 w najlepszym przypadku zewrze wejście „dodatnie” US1 do masy, zakładając, że napięcie nasycenia U_{CEsat} bipolarnego T5 jest równe zero (co nie jest prawdą). Gdy trafimy na egzemplarz wzmacniacza US1 o „niekorzystnej” w tym przypadku biegunowości, to nawet „żywe” zwarcie tego wejścia US1 nie spowoduje zmniejszenia napięcia wyjściowego zasilacza do zera, tylko do jakiejś niewielkiej wartości. A jeżeli wystąpi tam jakieś niewielkie napięcie, nawet rzędu miliwoltów, to jeżeli wtedy na wyjściu zasilacza występowałoby czyste zwarcie, to popłynie maksymalny prąd. Ogranicznik teoretycznie zadziała, ale w tym przypadku nie ograniczy prądu.

Podkreślił – tylko przy czystym zwarcu na wyjściu, ale niestety takie sytuacje zdarzają się często.

Podsumowując, koncepcja ogranicznika prądu ze zwieranie do masy przez tranzystor T5 jest ryzykowna: statystycznie w połowie tak zrealizowanych zasilaczy ogranicznik prądu nie będzie działał przy czystym zwarcu wyjścia zasilacza, a to z powodu „niekorzystnej” biegunowości napięcia niezrównoważenia US1. Ten problem trzeba jakoś rozwiązać!

Następna podniesiona przez niektórych uczestników kwestia: co jest źródłem napięcia odniesienia?

W wersji z rysunku B jest to jakieś źródło napięcia oznaczone $+5\text{ V}$. I można się domyślać, że jest ono odpowiednio stabilne. Gorzej w układzie z rysunku E, gdzie w torze napięciowym napięciem odniesienia jest... napięcie wyjściowe z jego tętnieniami, szumami i inną niestabilnością.

Idziemy dalej. Jeden z uczestników napisał: (...) Na przedstawionym schemacie znalazłem dwa błędy. Pierwszym problemem jest miejsce podłączenia kolektora T4. W przedstawionym układzie prąd płynący przez ten tranzystor trafia do obciążenia, ale nie jest mierzony – prąd przekazywany do obciążenia jest więc większy niż ustawiony limit.

Drugi problem – wciąż związany z limitem prądu – to miejsce podłączenia kolektora T5. W momencie zadziałania zabezpieczenia nadprądowego tranzystor ten

zewrze wejście komparatora US1 do masy, napięcie na wyjściu szybko obniży się i ze względu na duże wzmocnienie układu US1 mogą rozpocząć się oscylacje na wyjściu zasilacza. Lepiej podłączyć kolektor T5 pomiędzy R_7 a bazę T4 aby uniknąć takiego zachowania.

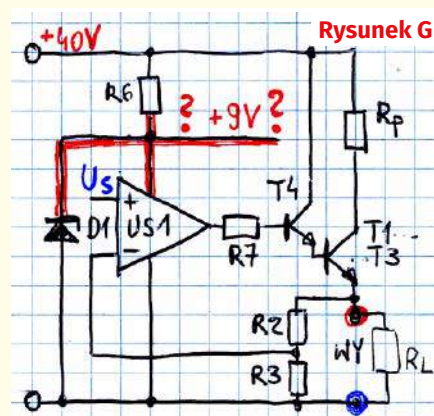
Uwaga dotycząca kolektora T4 jest słuszna. Inny dociekliwy uczestnik stwierdził, że działanie ogranicznika i w założeniu stabilizatora prądu opiera się na kontroli prądu płynącego przez „rezystor pomiarowy”, a nie prądu wyjściowego zasilacza. Zastanawiał się też, czy może wystąpić przypadek, na przykład przy zwarcu, że proponowany układ ogranicznika prawidłowo ograniczy prąd rezystora pomiarowego do jakiejś małej wartości, a popłynie niekontrolowany prąd przez T4?

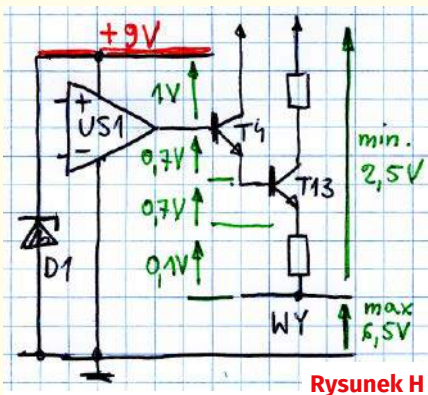
To kwestia do głębszej analizy, ponieważ oznaczałoby to niekontrolowany przepływ prądu przez obwody baza-emiter tranzystorów mocy T1, T2, T3. Jednak działanie ogranicznika opiera się na zatykaniu tych głównych tranzystorów mocy, więc przy prawidłowym działaniu proponowanego ogranicznika wydaje się to nieprawdopodobne. Jednak generalnie wszelkie tego rodzaju przypuszczenia i wątpliwości warto rozważyć dokładnie.

Idziemy dalej: nikt z uczestników nie napisał wprost o podstawowym będzie, dyskwalifikującym pomysł z rysunku B. Tylko dwóch uczestników wspomniało o tym przy okazji, a próba rozwiązania jest na rysunku E.

Otóż kluczowym, najważniejszym błędem koncepcji z rysunku B jest całkowity brak możliwości uzyskania na wyjściu wyższych napięć, zbliżonych do napięcia zasilającego 40 V !

Rysunek G pokazuje mocno uproszczony schemat tego zasilacza bez obwodów kontroli prądu. Nie wiadomo, jakim napięciem jest zasilany wzmacniacz US1. Jest to napięcie wyznaczone przez diodę Zenera D1. Ponieważ drugi wzmacniacz US2 zasilany jest napięciem dodatnim $+9\text{ V}$, więc można sądzić, iż właśnie takie jest napięcie na diodzie Zenera D1. Nawet gdyby było wyższe, problem





Rysunek H

pozostaje: napięcie na wyjściu zasilacza zawsze będzie znacząco niższe, niż napięcie zasilania wzmacniacza operacyjnego US1!

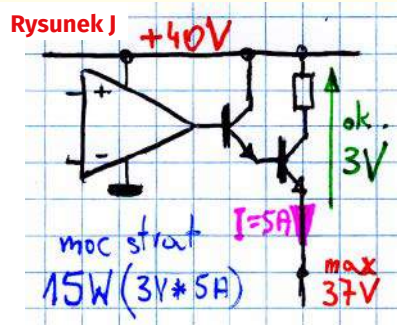
W grę wchodzi: wyjściowe dodatnie napięcie nasycenia wyjścia US1, dwa spadki napięcia U_{BE} na tranzystorach oraz spadek napięcia na rezystancjach emiterowych. Te niepożądane spadki napięć to minimum 2,5 V, w praktyce znacznie więcej, zapewne ponad 3 V. Problem pokazuje rysunek H.

To jest absolutnie dyskwalifikujący błąd. Jednak problemu nie rozwiązuje usunięcie D1, zwarcie R6 i zastosowanie wzmacniacza US1 o napięciu zasilania do 44 V, na przykład bardzo popularnego NE5532.

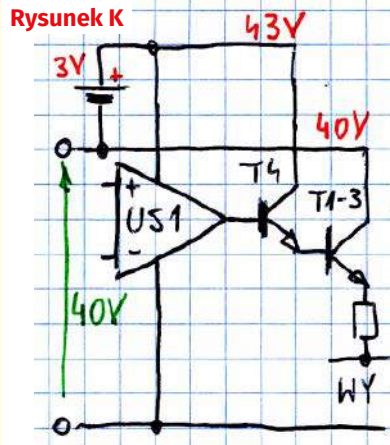
Obecność na rysunku B trzech tranzystorów mocy i rezystorów 0,1 oma wskazuje, że prąd wyjściowy zasilacza ma być duży, zapewne 5 A lub więcej. Na marginesie – wartości R8...R10 mogłyby być mniejsze, co zasygnalizował tylko jeden z uczestników.

Problem sygnalizuje **rysunek J**. Jakie maksymalne napięcie wyjściowe uzyskamy przy napięciu wejściowym 40 V? Na pewno niższe od 37 V, na dodatek jeżeli maksymalny prąd wyjściowy zasilacza ma być rzędu 5 A, to przy niepotrzebnej stracie mocy na tranzystorach około 15 W!

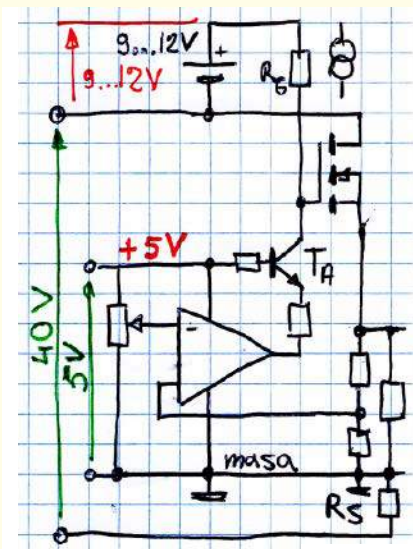
Każdy liniowy zasilacz niemiłosiernie grzeje się przy większych prądach, co akceptujemy, ale tu chodzi głównie o stratę napięcia co najmniej 3 V zgodnie z rysunkiem J. Co najmniej 3 V. Wprawdzie według rysunku H teoretycznie mogłoby to być 2,5 V, ale w praktyce będzie to dużo więcej, choćby z uwagi na tętnienia



Rysunek J



Rysunek K



Rysunek L

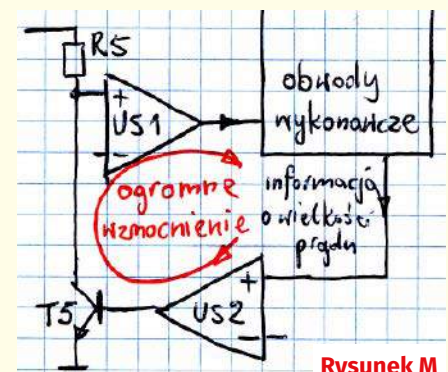
i wahania napięcia wejściowego 40 V. Warto byłoby to poprawić.

Częściową poprawą byłoby dodanie pomocniczego źródła zasilania według **rysunku K**. Problem jest poważny i trudny. W moim cyklu o „Takim zwyczajnym zasilaczu...” zaproponowałem rozwiązanie według **rysunku L**. Zastosowałem wzmacniacze operacyjne zasilane niskim napięciem około 5 V. Jednak omawianego teraz problemu tam nie ma, bo po pierwsze dodałem tranzystor pośredniczący TA pracujący w układzie wspólnej bazy. Po drugie zamiast połączenia bipolarnych tranzystorów mocy, zastosowałem jednego odpowiednio dużego MOSFET-a. W przypadku MOSFET-a mocy nie ma problemu napięcia nasycenia, tylko pokrewny, w tym przypadku nieporównanie mniejszy problem spadku napięcia na rezystancji wewnętrznej R_{DSon} . Aby w pełni otworzyć MOSFET-a i nie tracić napięcia, dodałem pomocnicze „górne” źródło zasilania 9...12 V, które gwarantuje pełne jego otwarcie.

To jeszcze nie koniec błędów na rysunku B. Wracamy do zasygnalizowanego w sposób obrazowy problemu oscylacji

i samowzbudzenia. Na ten problem zwrócił uwagę jeszcze tylko jeden (stały) uczestnik: (...) Układ ograniczania prądu z MAX4080 może wpaść w samowzbudzenie: gdy przez rezystor pomiarowy popłynie za duży prąd, układ MAX4080 wytworzy odpowiednie napięcie, które wysteruje US2, zaś ten włączy T5, który zewrze wejście + US1. Na jego wyjściu spadnie napięcie i T4 wyłączy się, a ten wyłączy trójkę T1-T2-T3. Prąd pomiarowy spadnie do zera. Napięcie na MAX spadnie i ostatecznie zniknie zwarcie we+ na S1 i układ wróci do stanu wejściowego, za duży prąd pomiarowy i tak w kółko...

Słusznie! Problem samowzbudzenia bardzo często dotyczy właśnie ogranicznika prądu, ale może dotyczyć też zasilacza pracującego jako stabilizator napięcia. To bardzo słabo rozumiany temat: kiedy wzmacniacz staje się generatorem? Najprościej biorąc: gdy wzmacnienie jest bardzo duże i gdy występują nieuniknione opóźnienia wskutek obecności pasożytniczych obwodów RC. Problem dotyczy każdego układu wzmacniającego z pętlą sprzężenia zwrotnego, a konkretnie sytuacji, gdy sprzężenie z ujemnego, staje się dodatnie. A dodatnie sprzężenie (plus odpowiednio duże wzmacnienie) oznacza powstanie generatora. Tego problemu nie da się omówić w ramach zadania NieGra, więc potwierdzę tylko, że zasygnalizowane obawy dotyczące oscylacji są jak najbardziej słuszne. **Zasilacz według koncepcji z rysunku B na pewno wzbudzi się przy pracy w trybie ograniczania prądu.** Wskazuje na to ogromne wzmacnienie w pętli sprzężenia, która w trybie ograniczania prądu (**rysunek M**) obejmuje dwa wzmacniacze operacyjne pracujące z maksymalnym wzmacnieniem i tranzystor T5 w układzie wspólnego emitera, czyli też o dużym wzmacnieniu. Całkowite wzmacnienie napięciowe w tej pętli będzie rzędu miliardów lub milionów, zależnie od częstotliwości, co zapewne zaowocuje samowzbudzeniem. Poprawa tego problemu i w ogóle uzyskanie sensownej charakterystyki dynamicznej to najtrudniejszy problem przy projektowaniu i budowie zasilaczy. Problem bardzo słabo rozumiany,



Rysunek M

wymagający starannej kompensacji częstotliwościowej układu.

Wracamy do rozwiązań zadania NieGra310. Pojawiły się w nich też opinie nietrafne i błędne. Oto przykłady.

(...) Zasilanie wzmacniacza operacyjnego nie jest symetryczne, a jest $+9\text{ V}$ i -5 V (...) To akurat nie jest błąd.

(...) Dodatkowo błąd że jest powtórzone na schemacie nóżka 4 OPampa nie ma takiej potrzeby (...) To nie jest błąd, bowiem ewidentnie mamy tu dwa pojedyncze wzmacniacze operacyjne, a nie jeden podwójny. Świadczy o tym numeracja nóżek, a najbardziej numeracja wyjścia.

(...) Użycie MAX4080 jest dla mnie dziwne, bo to bardziej układ do zastosowań

w motoryzacji stosowane do ładowania ogniw, więc tutaj moim zdaniem nie ma sensu (...) Układy „motoryzacyjne” nie są gorsze, mogą nawet pracować w szerszym zakresie temperatur niż układy „zwykłe”.

(...) Piny 4 US1 i US2 są niewłaściwie zasilane z ujemnego napięcia -5 V . Takie zasilanie może spowodować pojawienie się napięcia ujemnego na wyjściach US1 i US2. Bazy tranzystorów T4 i T5 wymagają sterowania dodatnimi napięciami. (...) Tranzystory krzemowe bez problemu wytrzymują obecność napięcia 5 V , ujemnego względem emitera – po prostu pozostaną zatkane.

(...) Wzmacniacze operacyjne nie muszą mieć zasilania napięciem ujemnym. Zasilanie może być na poziomie masy układu. Jedyne

muszą to być wzmacniacze, które mogą pracować przy potencjale masy. (...) Owszem, ale w grę wchodzi omówiony wcześniej problem z T5 i napięciem niezrównoważenia US1.

Jednak ogólnie biorąc, mogłem uznać wszystkie nadesłane rozwiązania, ponieważ w każdym prawidłowo został wskazany co najmniej jeden błąd. Nagrody-upominki za zadanie NieGra310 otrzymują:

Zdzisław Landowski – Gdańsk,

Łukasz Fortuna – Wołowa,

Andrzej Kalinowski – Poznań.

Wszystkich uczestników dopisuję do listy kandydatów na bezpłatne prenumeraty. ■

Piotr Górecki

Policz – zadanie 315

Zadanie Policz315 jest kontynuacją rozwiązane zadania 310. Otóż do wstępnego sprawdzenia, czy w ogóle wzmacniacz operacyjny działa, czy jest uszkodzony chcemy zrobić przystawkę z podstawką według rysunku A.

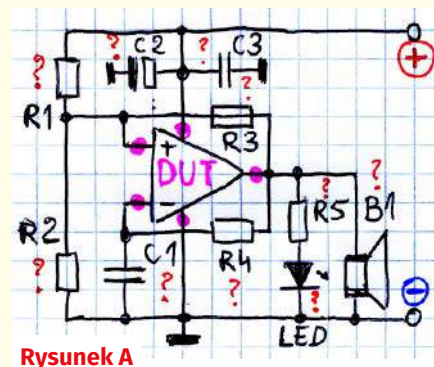
W ramach zadania Policz315 należy:

– zaproponować wartości elementów takiego prostego testera.

Zapraszam do udziału zarówno elektroników doświadczonych, jak i początkujących, którzy jeszcze nie potrafili przeanalizować

wszystkich subtelności układu. Z uwagi na specyfikę zadania proszę o podawanie swojego wieku oraz miejsca nauki czy pracy.

Odpowiedź nadesłaj w terminie 60 dni od ukazania się tego numeru EdW. Tytuł e-maila powinien zawierać nazwę konkursu i numer zadania oraz Twoje nazwisko (**Policz315_Nazwisko**). Jeżeli chcesz uczestniczyć w podziale upominków, w e-mailu podaj od razu swój adres pocztowy. Możesz też jeszcze przysłać rozwiązanie zadania Policz z poprzedniego miesiąca.



Rysunek A

Policz – rozwiązanie zadania 310

W EdW 1/2022 przedstawione było zadanie Policz310, które brzmiało: Wykonaliśmy Miernik wzmacniaczy operacyjnych z EdW 11/2021. Zmierzyliśmy napięcia niezrównoważenia posiadanych wzmacniaczy i według tych pomiarów kilka z nich ma niewygodnie duże napięcie niezrównoważenia

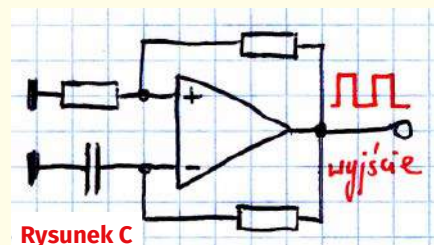
ponad 20 mV . Aż się wierzyć nie chce, więc chcemy to potwierdzić w prostym układzie z rysunku B.

W ramach zadania Policz310 należy:

– zaproponować wartości elementów.

Zadanie to wprost zostało „z życia wzięte”: jeden z Czytelników zgłosił problem niewygodnie dużego napięcia niezrównoważenia niektórych posiadanych wzmacniaczy operacyjnych. Zastanawiał się, czy nie jest to aby wina „Miernika wzmacniaczy operacyjnych” zrealizowanego według opisu w EdW.

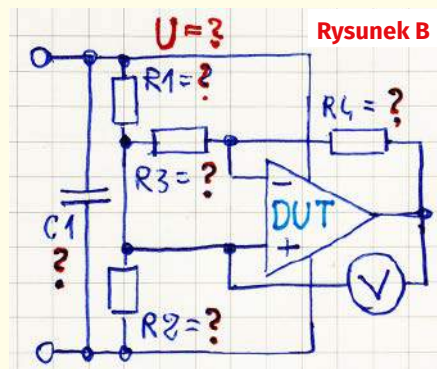
Aby wyjaśnić wątpliwości, zaproponowałem budowę dwóch przyrządów – testerów. Otóż w przypadku wzmacniaczy operacyjnych, które już były gdzieś używane, to po pierwsze należałoby sprawdzić, czy one w ogóle są sprawne i czy prawidłowo pełnią swą podstawową rolę. Można to sprawdzić w różny sposób, ale popularną metodą jest sprawdzenie, czy taki



Rysunek C

wzmacniacz w ogóle działa jako generator w układzie z rysunku C. To zresztą jest związane z zadaniem Policz 315.

Dopiero gdy taka próba wykaże, że wzmacniacz wytwarza przebieg prostokątny i jest sprawny, wtedy można i warto zmierzyć jego napięcie niezrównoważenia. Na przykład za pomocą oddzielnego narzędzia: prostego testera właśnie o schemacie takim, jak pokazany na rysunku B. Taki tester byłby przeznaczony tylko do zgrubnego sprawdzenia,



Rysunek B



czy rzeczywiście napięcie niezrównoważenia danego wzmacniacza wykracza poza wartości podane w katalogu. A może tak być! I to niekoniecznie z winy producenta.

Zdarza się, że podczas rozmaitych testów wzmacniacz operacyjny nie zostanie całkowicie uszkodzony, tylko właśnie pogorszą mu się niektóre parametry z powodu przeciążenia jego obwodów wejściowych. Jest to prawdopodobne i zdarza się dość często. Wprawdzie w katalogach są stosowne ostrzeżenia i zalecenia dotyczące montażu i serwisu, ale hobbyści powszechnie lekceważą problem ładunków statycznych. A właśnie mikrowyładowania spowodowane ładunkami statycznymi mogą być przyczyną pogorszenia parametrów, a niekoniecznie całkowitego uszkodzenia.

Kiedyś dawno, elementami najdelikatniejszymi i najczęściej ulegającymi uszkodzeniom wskutek ładunków statycznych, były tranzystory JFET (krajowe BF245) oraz układy cyfrowe MOS i CMOS. Z czasem polepszone ich odporność oraz dodano obwody zabezpieczające w wielu elementach. Jednak w przypadku wzmacniaczy operacyjnych, zwłaszcza tych o najwyższych parametrach, dodawanie obwodów ochronnych na wejściach może pogorszyć niektóre cenne właściwości, więc producenci tego unikają. W efekcie obwody wejściowe wzmacniaczy operacyjnych zwykle są słabo zabezpieczone przed uszkodzeniem przez ładunki statyczne. A praktyka pokazuje, że ładunki statyczne o niedużej energii często zamiast całkowicie uszkodzić układ scalony, tylko pogarszają parametry wejściowe. A niefrasobliwy

hobbysta tego nie widzi i nie stosuje żadnych środków ochrony przed ładunkami statycznymi, jak choćby uziemianie i rozładowywanie ciała przez kontaktem z elementami półprzewodnikowymi.

Pomiar napięcia niezrównoważenia bodaj najprościej można wykonać w układzie zbudowanym według rysunku D. I to pomiar nie tylko wzmacniaczy podejrzewanych o uszkodzenie, ale dowolnych, całkowicie sprawnych, także tych o małym napięciu niezrównoważenia.

U góry rysunku mamy dwie podstawowe konfiguracje wzmacniacza: odwracającego i nieodwracającego. A gdy napięcie wejściowe wynosi zero, daje to układ pokazany na dole rysunku D.

Tester z rysunku D jest zwyczajnym wzmacniaczem, którego robocze wzmocnienie wyznaczone jest przez stosunek rezystorów R_A , R_B . Teoretycznie nie ma on czego wzmacniać, ale w praktyce wzmacnia własne napięcie niezrównoważenia. Nie wchodząc w szczegóły i definicje można przyjąć, że wzmocnienie jest równe stosunkowi $G=R_A/R_B$, a w układzie z rysunku B wzmocnienie $G=R_4/R_3$, czyli że woltomierz pokaże napięcie niezrównoważenia wzmocnione G razy.

Wersja z rysunku B działa na dokładniej tej samej zasadzie, jak pokazuje rysunek D, tylko ma dodatkowo obwód sztucznej masy z rezystorami R_1 , R_2 , co zwiększa wygodę użytkowania, bo pozwala zasilać tester napięciem pojedynczym.

Tłumacząc to tylko dlatego, że w rozwiązaniach pojawiły się wątpliwości co do sposobu włączenia woltomierza.

Jeżeli tester miałby być uniwersalny, przewidywany zakres napięć zasilania wyniesie od 5 V do nawet 40 V i więcej. Na pewno warto więc zastosować kondensator C_1 o odpowiednio dużym napięciu dopuszczalnym.

W zadaniu była mowa o napięciu niezrównoważenia ponad 20 mV. A to prowadzi do wniosku, że wzmocnienie G nie powinno być zbyt duże, żeby przy zasilaniu napięciem 5 V nie nasyciło wzmacniacza, tylko pozwoliło na pomiar przy wartościach nie większych niż te 20 mV. Natomiast napięcie

niezrównoważenia powyżej 20 mV ewidentnie wskazuje na jakieś uszkodzenie.

Większość uczestników słusznie stwierdziła, że wartość G nie powinna przekraczać 100x.

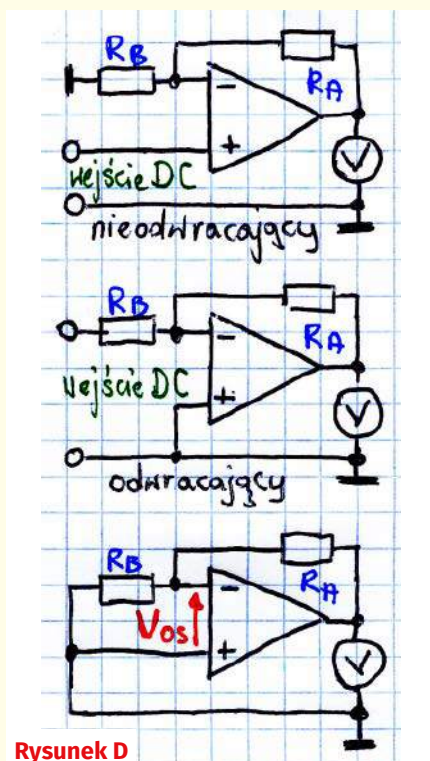
Wspomniany Czytelnik, z którym prowadziłem korespondencję, też w związku z zadaniem Policzn310 przysłał e-mail, zaznaczając na wstępie: *Dzień dobry. Na początku chcę zaznaczyć, że nie biorę udziału w tym zadaniu. (...) Na wysłanie swojej analizy zdecydowałem się dlatego, że obiecałem sobie zająć się elektroniką poważniej (...) Tester (...) pracuje w układzie wzmacniacza odwracającego. Znając wartości rezystorów R_3 i R_4 możemy wyliczyć jego wzmocnienie ze wzoru $G=-R_4/R_3$. (...) z powyższej zależności na wzmocnienie możemy jedynie określić stosunek tych wartości, a nie możemy obliczyć ich bezwzględnych wartości. Według informacji zawartych w EdW dotyczących wzmacniaczy operacyjnych oraz materiałów na YouTube (...) wywnioskowałem, że nie ma jednej słusznej metody określającej wartości bezwzględne rezystorów. Takie układy działają poprawnie dla szerokiego zakresu wartości rezystorów. Nie można stosować bardzo małych rezystancji, ponieważ (...) prąd z wyjścia wzmacniacza będzie powodował dodatkowe straty mocy, które są niepotrzebne. Nie mogą to być również rezystancje bardzo duże, ponieważ wtedy pojawi się problem z prądem polaryzacji i prądem niezrównoważenia.*

Dla potrzeb określenia napięcia niezrównoważenia należy tak dobrać rezystory R_4 i R_3 , aby ich stosunek wynosił 100 (1000, a nawet więcej przy pomiarze najlepszych wzmacniaczy). Mogą to być pary 10 k/100 Ω ; 100 k/1 k itd.

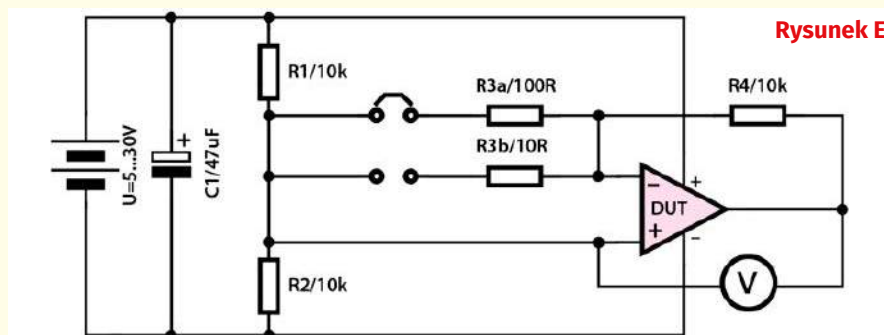
Jednakowe rezystory R_1 , R_2 ustalają potencjał sztucznej masy w połowie napięcia zasilania i można je również przyjąć w szerokich granicach (k Ω).

Kondensator C_1 zapobiega samowzbudzeniu WO oraz jako filtrowanie zasilania. Jego wartość przyjmuję równą 100 nF (zgodnie z zastosowanym w mierniku wzmacniaczy operacyjnych).

Co do napięcia zasilania, to najlepiej zastosować takie, przy jakim WO pracuje w realnym układzie. (...).



Rysunek D



Rysunek E



Tabela 1.

Wartości napięcia niezrównoważenia dla różnych WO i wybranych wartości napięcia zasilania uzyskanych w doświadczalnym układzie pomiarowym [mV]

	6V	10V	20V	30V
1	2	3	4	5
TL071	0,95	1,5	0,85	0,63
LF356	1	1,6	1,33	1,15
NE5534	0,370	0,39	0,430	0,470
LM201	0,870	1,3	2,02	2,02
LM2904	1,04	1,03	1,008	1,005 (26V)
LM358P	0,111	0,146	0,193	0,160
LM358N	0,526	0,538	0,552	0,552
TL082	1,44	0,106	0,163	0,193
LM833	0,916	0,790	0,670	0,603
NE5532	0,120	0,096	0,088	0,100
RC4558	1,50	1,45	1,43	1,40
LF355	1,42	1,25	1,22	1,19
uA741	0,900	1,80	1,90	2,00
LT1013	0,020	0,016	0,290	0,120
LM258N	0,40	0,43	0,48	0,52

Jeden ze stałych uczestników zaproponował: (...) R1 i R2 mają tę samą wartość w celu utworzenia sztucznej „ziemi” w połowie

napięcia U. Proponuję wartość 1 kOhm. R3 i R4 są w obwodzie ujemnego sprzężenia mierzonego DUT (...) proponuję wzmocnienie 100, gdyż napięcie niezrównoważenia 20 mV wytworzy na wyjściu napięcie 2 V. Dla wzmocnienia 100 proponuję więc albo R3=100 Ohm i R4=10 k, albo R3=1 k i R4=100 k. Napięcie zasilające zależy od parametrów wzmacniacza DUT (...) C1 (...) tradycyjnie 100 nF.

Inny stały uczestnik napisał: (...) Moja propozycja prostego układu pomiarowego przedstawiona jest na rysunku [E]. Wprowadziłem drobną modyfikację dotyczącą rezystora R3 rozdzielając go na dwa zakresy. (...) do 20 mV korzystamy z rezystora R3a=100 R. Na wyjściu wtedy możemy się spodziewać napięcia 2 V, a więc nie przekroczymy wartości napięcia zasilania np. dla 5 V. (...) Pierwszy pomiar można dokonać z rezystorem R3a, a jeżeli wynik będzie na tyle niski, że nie przekroczy napięcia zasilania, powtórzyć z rezystorem R3b=10 R.

(...) napięcie zasilania (...) należy dobrać zgodnie z zakresem pracy WO (...) w docelowym układzie. Mamy WO pracujące z napięciem nie większym niż 6 V, a również

takie na 44 V (to ten najczęściej spotykany przez Czytelników EdW zakres, ale może być i 900 V np. PA94) (...) przyjąłem ten zakres w granicach 5...30 V. (...) dokonałem pomiaru kilkunastu WO. Niestety (ze względów poznawczych) wśród tych WO nie trafiłem na egzemplarz o napięciu niezrównoważenia tak dużym jak 20 mV. Z tabeli 1 widać, że wartość napięcia zasilania ma wpływ na napięcie niezrównoważenia. Widać również różnice między poszczególnymi WO. Wartości podane w tabeli dotyczą konkretnego egzemplarza WO. Nie jest to wartość średnia obliczona na podstawie pomiarów kilku WO danego typu. Pozdrawiam

Myślę, że podane informacje wyczerpują temat. Wszystkie nadesłane odpowiedzi były prawidłowe. Nagrody-upominki za zadanie **Policz310** otrzymują:

Stanisław Turek – Chełm,
Andrzej Szulda – Olsztyn,
Damian Ż. – Lipie.

Wszystkich uczestników dopisuję do listy kandydatów na bezpłatne prenumeraty. ■

Piotr Górecki

Jak to działa

W numerze 2/2022 przedstawiony był, pokazany na **rysunku A**, nieskomplikowany układ elektroniczny.

Jest to... **realizacja innowacyjnej metody zmniejszania rozmiarów głównego kondensatora filtrującego w zasilaczu sieciowym**. Układ realizuje też inne pożyteczne funkcje.

Analizę schematu warto poprzedzić informacjami wprowadzającymi w temat.

Otóż wielu z nas dziwi się, jak błyskawicznie, na przykład przez krócej niż pół godziny można naładować najnowsze smartfony

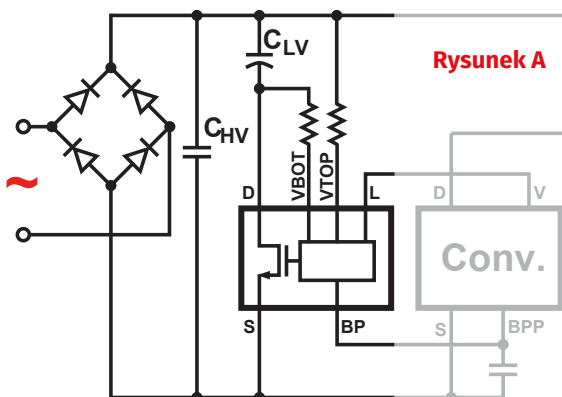
i tablety. Nie ulega wątpliwości, że zawarty w nich akumulator musi być ładowany bardzo dużym prądem. Na pewno zastosowany tam akumulator litowy musi być przystosowany do tak szybkiego ładowania, ponieważ jak wiadomo, większość akumulatorów litowo-jonowych wymaga ładowania niezbyt dużym prądem w czasie ponad dwie godziny.

To jednak odrębny temat. Druga sprawa to fakt, że takie błyskawiczne

ładowanie odbywa się przy wykorzystaniu łącza USB. Przez delikatny kabelek USB musi być przekazywana duża moc. I tu nasuwa się pytanie o moc i prąd.

Jak wiadomo, włączu USB od początku oprócz żył sygnałowych, przewidziano też dwie żyły zasilania napięciem 5 V. Pierwotnie przewidywano wydajność prądową takiego obwodu zasilania USB równą maksymalnie 500 mA, a w niektórych urządzeniach tylko 100 mA. Kolejne

REKLAMA



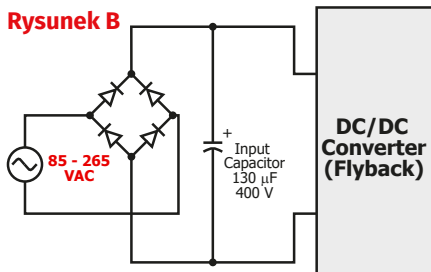
Certyfikat Underwriters Laboratories
94V-0 E480148 TYPE 1

Zakład produkcyjny:
05-660 Warka
ul. M. Rzepiełkowskiej 17
tel. 22 781 63 95
22 761 95 60
fax. 22 781 63 95 w 23
www.elmax.waw.pl
elmax@elmax.waw.pl

OBWODY DRUKOWANE

Produkcja, Projektowanie, Montaż

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

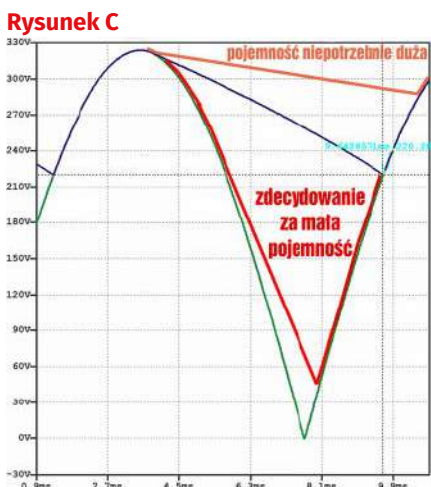


zmiany i nowelizacje standardu przewidywały zwiększenie wydajności prądowej łącza USB przy zachowaniu napięcia 5 V. Nie wchodząc w dalsze szczegóły trzeba stwierdzić, że przy napięciu 5 V nie można przekazać dużych mocy, ponieważ konieczny do tego duży prąd powodowałby niedopuszczalnie duże grzanie się przewodów. Dlatego w najnowszych wersjach standardu USB przewiduje się możliwość inteligentnej komunikacji i negocjacji między urządzeniem zasilanym a zasilającym, pozwalające na zwiększenie napięcia zasilania dużo ponad 5 V. Pozwala to bez obawy o przegrzanie kabli i styków przesłać moc kilkudziesięciu watów. I to jest konieczny warunek błyskawicznego ładowania za pomocą łącza USB. Dla dociekliwych – chodzi o PPS (Programmable Power Supply) – takie hasło można wpisać w wyszukiwarkę, żeby zapoznać się ze szczegółami).

I właśnie najnowsze smartfony mają dedykowane ładowarki o dużej mocy kilkudziesięciu watów. Na przykład tanie i coraz bardziej popularne chińskie smartfony pochodzące od Xiaomi mają ładowarkę o sumarycznej mocy do 33 W. A dostępne są podobne zasilacze/ładowarki o jeszcze większej mocy

I tu pomalutkę przechodzimy do zadania Jak2: dziwimy się, jak małe rozmiary mają takie ładowarki, będące zasilaczami sieciowymi o mocy kilkudziesięciu watów (i jak mało się grzeją przy takiej mocy).

Zadanie Jak2 ma ścisły związek właśnie ze zmniejszaniem rozmiarów współczesnych impulsowych zasilaczy sieciowych.

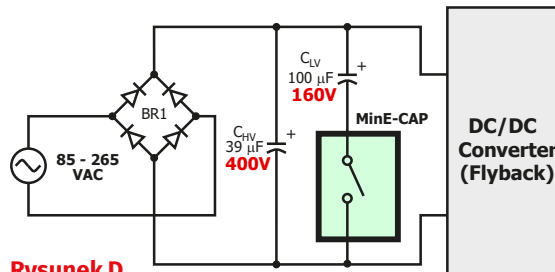


Okazuje się bowiem, że można zminiaturyzować transformator oraz kondensatory wyjściowe przetwornicy stosując wysoką częstotliwość pracy.

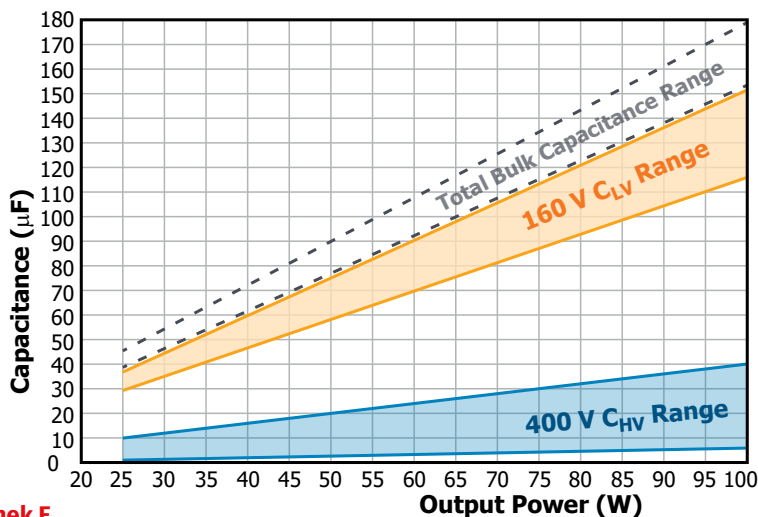
Wąskim gardłem okazuje się jednak główny kondensator filtrujący wyprostowane napięcie sieci. Każdy zasilacz sieciowy musi mieć prostownik oraz

główny kondensator filtrujący (nazywany zwykle *bulk capacitor*), zamieniające sinusoidalne napięcie sieci na napięcie stałe, które zostanie podane na wejście przetwornicy impulsowej DC/DC. Ilustruje to **rysunek B** (z materiałów Power Integrations).

Ten główny kondensator filtrujący musi mieć odpowiednio wysokie napięcie, najczęściej 400 V oraz odpowiednio dużą pojemność,



zmniejsza się. Jeżeli pojemność jest niepotrzebnie duża, to zmiany napięcia na wejściu przetwornicy są małe. Z jednej strony to korzystne, ale niepotrzebnie zwiększa pojemność filtrującą i rozmiary kondensatora. Przy zbyt małej pojemności, napięcie na wejściu przetwornicy DC/DC zanadto zmniejsza się i w tych chwilach przetwornica nie może poprawnie pracować.



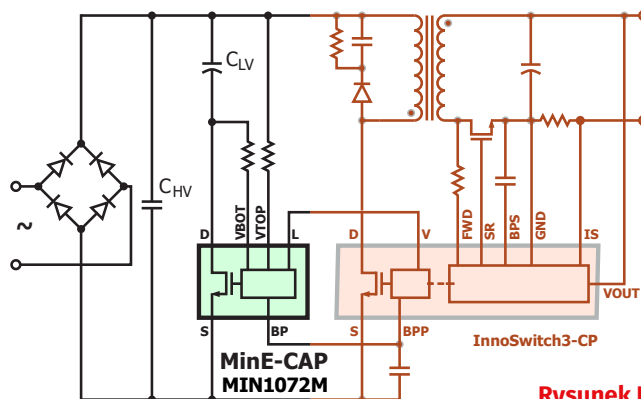
Rysunek E

zależnie od maksymalnej mocy zasilacza. Chodzi o to, że wewnętrzna przetwornica DC/DC (zwykle typu flyback, ale coraz częściej jednor tranzystorowy forward) pracuje prawidłowo w ograniczonym zakresie napięć wejściowych. W szczególności jeżeli napięcie wejściowe przetwornicy (i głównego kondensatora filtrującego) będzie zbyt niskie, to przetwornica nie dostarczy w tych chwilach potrzebnego napięcia i prądu.

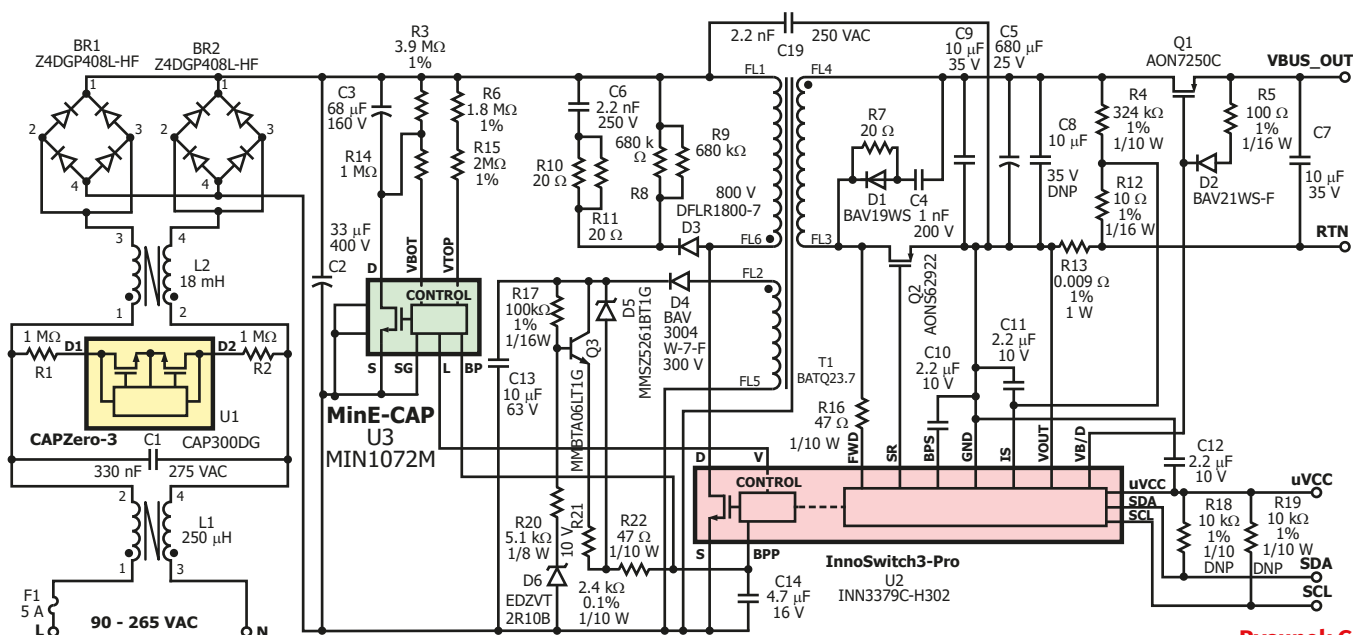
Problem ilustruje **rysunek C**: w szczytach sinusoidy sieci, główny kondensator ładuje się do ok. 325 V (przy napięciu 230 VAC). Zależnie od pojemności kondensatora i od poboru prądu przez przetwornicę, między połówkami sinusoidy napięcie zasilające przetwornicę DC/DC

Należy zastosować możliwie małą pojemność głównego kondensatora filtrującego, która jednak zagwarantuje, że przetwornica DC/DC będzie pracować prawidłowo także w dolinach między szczytami napięcia sieci.

To tylko połowa zagadnienia. Problem w tym, że **rysunek C** pokazuje sytuację przy zmiennym napięciu sieci 230 V. A współczesne uniwersalne zasilacze i ładowarki



Rysunek F



Rysunek G

muszą prawidłowo pracować także w sieciach 100...120-woltowych (w Japonii i Ameryce), i to też przy obniżonym napięciu 85 VAC, co sygnalizuje rysunek B. Przy napięciu w sieci 85 VAC napięcie szczytowe na kondensatorze wyniesie tylko około 120 VDC. I jeszcze dodatkowo będzie zmniejszać się pomiędzy szczytami sinusoidy, stosownie do pojemności kondensatora. Co bardzo ważne, przy niskim napięciu sieci dla zapewnienia prawidłowej pracy przetwornicy DC/DC niezbędna jest dużo większa pojemność filtrująca, niż przy wyższym napięciu sieci.

W zasilaczu uniwersalnym według rysunku B mamy więc mocno niekorzystną sytuację. Po pierwsze kondensator filtrujący musi mieć wysokie napięcie nominalne. Przy 265 VAC szczytowe napięcie na kondensatorze wyniesie nieco ponad 370 V, co wymaga kondensatora o napięciu nominalnym co najmniej 400 V. Po drugie, do pracy przy minimalnym napięciu sieci potrzebna jest duża pojemność.

W sumie oznacza to konieczność zastosowania kondensatora o dużej pojemności i wysokim napięciu nominalnym. Czyli kondensatora o dużych rozmiarach.

I właśnie schemat z rysunku A pokazuje sposób znaczącego zmniejszenia tego problemu. Idea, pokazana na **rysunku D** jest w sumie bardzo prosta. Mianowicie przy wysokim napięciu sieci klucz – wyłącznik pozostaje rozarty, a do prawidłowej pracy przetwornicy DC/DC wystarcza kondensator C_{HV} o niedużej pojemności i wysokim napięciu nominalnym 400 V. Jeżeli natomiast napięcie sieci jest niskie, niższe niż powiedzmy 100 V, to zostaje równolegle dołączony drugi kondensator filtrujący, który ma znacznie większą pojemność, ale zdecydowanie niższe napięcie nominalne.

Pochodzący z materiałów Power Integrations **rysunek E** pokazuje (ba gorzej) jaką pojemność powinien mieć pojedynczy kondensator wysokonapięciowy w koncepcji z rysunku B, a jakie powinny być pojemności przy wykorzystaniu

omawianego rozwiązania z układem scalonym MinE-Cap.

Co ważne, sumaryczna objętość dwóch tak dobranych kondensatorów jest tu znacznie mniejsza, niż w rozwiązaniu klasycznym

REKLAMA

KEY PRODUCENT AUTOMATYKI GRZEWCZEJ
11-200 Bartoszyce ul. Bohaterów Warszawy 67 pwkey@onet.pl
tel. (89)7635050 fax (89)7635051

TANIE REGULATORY

DO KOTŁÓW WĘGLOWYCH I NA DREWNO

z wbudowanym termostatem pokojowym
zapewniającym komfort i oszczędność



REGULATORY DO KOTŁÓW Z PODAJNIKIEM

REGULATORY POGODOWE

- Prosta obsługa, bogate możliwości programowania
- Możliwość dopasowania do każdego kotła i rodzaju paliwa
- Wysoka jakość
- Gwarancja 24 miesiące

www.pwkey.pl

według rysunku B. To jest jeden ze sposobów znacznego zmniejszenia objętości i przyczyna naszego zdziwienia rozmiarem zasilaczy o dużej mocy.

Rysunek B to zmodyfikowana wersja schematu z materiałów Power Integrations. Jeden z oryginalnych schematów z przetwornicą flyback pokazany jest na **rysunku F**.

Z kolei realizacja z przetwornicą forward pokazana jest na **rysunku G** (z noty AN-92). Tu koniecznie trzeba zwrócić uwagę na obecność interesującego elementu CAPZero-3, który też omawialiśmy w EdW oraz na obecność scalonej przetwornicy InnoSwitch3-Pro. Intrygująco wyglądające dwa mostki prostownicze zastosowano tylko dla rozłożenia mocy strat.

Na **fotografii H** (z firmowego raportu DER-822) widać trzy ujęcia zadziwiająco małego 60-watowego zasilacza o rozmiarach 52×29×26 mm, co daje gęstość mocy ponad 33 W na cal sześcienny!

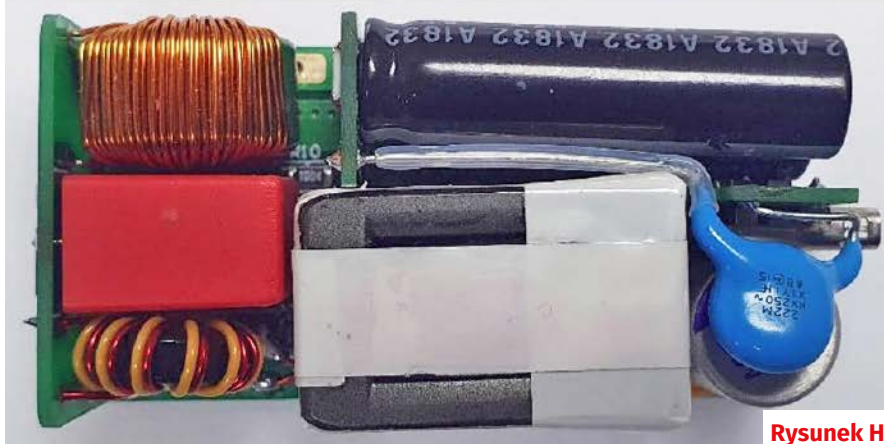
Omawiany teraz układ scalony MinE-CAP pełni także dodatkową funkcję. Pozwala na rezygnację z typowego dla zasilaczy impulsowych termistora NTC, ponieważ podczas startu wbudowany klucz (MOSFET) włącza się tylko na króciutkie chwile, przez co impuls prądowy przy włączaniu jest zdecydowanie mniejszy – spowodowany jest tylko przez kondensator C_{HV} a nie C_{LV} . Rezystory dołączane do nóżek VBOT, VTOP pozwalają na bieżąco mierzyć napięcie na C_{LV} i tak otwiera na krótko klucz, by na C_{LV} występowało napięcie bliskie nominalnemu także wtedy, gdy nie jest on wykorzystywany – jest to potrzebne w przypadku „alumiowych elektrolitów”.

Najbardziej zainteresowani dalsze informacje znajdują w materiałach firmowych.

Jeżeli chodzi o nadesłane rozwiązania, większość była prawidłowa. Oto fragmenty dwóch najbardziej trafnych rozwiązań nadesłanych przed stałych uczestników tego konkursu:

(...) Jest to innowacyjna metoda zwiększenia efektywności kondensatorów wygładzających w zasilaczach sieciowych. Zamiast jednego kondensatora o dużej pojemności, stosuje dwa; CHV posiadający 20% całkowitej pojemności i pełnym napięciem, oraz CLV mający 80% całkowitej pojemności, ale o zredukowanym max napięciu, np. 3 krotnie. Dla przykładu; pierwszy ma napięcie 400 V, a drugi 160 V. (...)

(...) Na rysunku widzimy fragment przetwornicy AC/DC z układem U1 MinE-CAP.



Rysunek H

Rolą tego układu jest zmniejszenie kondensatorów wejściowych, które są na napięcie 400 V i więcej. Na wejściu mamy kondensator CHV na większe napięcie ale mniejszej pojemności. Ten kondensator jest połączony równolegle z kondensatorem na niższe napięcie (CLV) o większej pojemności, który jest połączony szeregowo z układem scalonym MinE-CAP.

Nagrody-upominki za zadanie **JakDziała2** otrzymują:

Janusz Wrona – Kielpin,
Mikołaj Czarnecki – Tarnów,
Henryk Biały – Laskowizna.

Wszyscy uczestnicy konkursu zostają dopisani do listy kandydatów na bezpłatne prenumeraty. ■

Piotr Górecki

REKLAMA

www.ulubionykiosk.pl

Pomiary pH wody i gleby, część 1

Artykuł powstał na prośbę Czytelników, którzy zainteresowani są elektronicznymi sposobami pomiaru pH różnych substancji, w tym wody w akwariach i gleby. Pierwotnie miały to być dwa krótkie odcinki, jednak zagadnienie okazało się obszerne, trudne, ale też bardzo interesujące z punktu widzenia elektroniki i elektronika.

...pehametry...

Teoretycznie wszystko jest beznadziejnie proste: kupujemy przyrząd elektroniczny ze wskaźnikiem cyfrowym albo analogowym, tak zwany pH-metr (pehametr), zanurzamy jego głowicę pomiarową w badanej cieczy, odczytujemy wynik, wartość pH i... koniec zabawy – wszystko jasne.

Pehametry mają ceny rzędu kilkuset, nawet kilku tysięcy złotych i więcej (przykład na **fotografii 1**).

Jednak ktoś może zaprotestować: najtańsze pH-metry, a przynajmniej przyrządy, na których widnieje napis **pH meter**, można kupić już za kilkanaście złotych. Niestety, większość z nich słusznie budzi wątpliwości, co w ogóle takie przyrządy mierzą (przykład na **fotografii 2**). Problem w tym, że typowe pehametry mają delikatną szklaną elektrodę. Nie każdy wie, jak działają, ale sonda pH nieodłącznie kojarzy się z delikatnym szkłem. A na fotografii 2 widzimy jakieś metalowe szpikulce, co ewentualnie pasowałoby do pomiaru wilgotności, ale nie odczynu pH.

Takich wątpliwości nie budzą inne zaskakująco tanie przyrządy – przykład na **fotografii 3**. Tu zaufanie wzbudzają szaszetki, pozwalające na dokładną kalibrację.

Niezależnie od kwestii cen wielu elektroników zainteresowanych jest samodzielną budową pehametru, nie tylko ze względów ekonomicznych, ale także z ciekawości. Jednak „od zawsze” wiadomo, że **czujnika**,



2

czyli sondy pH, nie sposób zrobić we własnym zakresie, tylko trzeba ją kupić. Ale układ elektroniczny można wykonać samodzielnie. Trzeba tylko zadbać o to, żeby miał ogromną rezystancję wejściową. Zasadniczo budowa pH-metru jest prosta: kupujemy sondę z przewodem i wtyczką BNC (wersja z **fotografii 4** dostępna jest od około 40 zł) i dołączamy ją do woltomierza o ogromnej oporności wejściowej. Przecież pomiar pH w istocie polega na pomiarze wytwarzanego przez sondę niedużego napięcia stałego (około 60 mV/pH). Problem tylko w tym, że rezystancja sondy jest ogromna, rzędu stu megaomów, więc woltomierz, będący obciążeniem czujnika, musi mieć rezystancję znacznie większą! Dlatego absolutnie nie nadają się do tego klasyczne multimetry cyfrowe, których rezystancja wejściowa zwykle wynosi 10 MΩ (w tańszych 1 MΩ), czyli co najmniej 100...1000 razy za mało.

<https://www.mera-sp.com.pl/katalog-produktow/przyrzady-pomiarowe/ph-m>

Lekki, kompaktowy miernik pH/mV/°C HI 8314

Numer katalogowy: **HI8314** HANNA

Cena netto: **1325.00 PLN**

[Zamów](#)

[Zadaj pytanie](#)

Funkcje/cechy

- Pomiar pH, mV i temperatury
- Kompaktowa, ergonomiczna obudowa
- Ręczna kalibracja w jednym lub dwu punktach
- Automatyczna kompensacja temperatury
- Łatwy do czyszczenia i utrzymania w czystości
- Wskaźnik niskiego poziomu baterii
- Prosta obsługa

Cena obejmuje

- miernik HI8314
- elektrodę pH HI 1217D
- roztwór buforowy w saszetce HI70004 pH 4.01
- roztwór buforowy w saszetce HI70007 pH 7.01
- roztwór do czyszczenia elektrod HI700601 - 2 sz
- śrubokręt do kalibracji
- bateria

1



3

4



Dawniej realizacja odpowiedniego miliwoltomierza o oporności wejściowej rzędu 1 gigaoma była dużym wyzwaniem, ale dziś wśród mnóstwa wzmacniaczy operacyjnych z powodzeniem znajdziemy wiele, mających odpowiednie parametry wejścia. Dostępne są też tanie gotowe moduły, przeznaczone do współpracy z sondami pH – płytkę pokazaną na **fotografii 5** można kupić za równowartość kilkunastu złotych.

Sytuacja wydaje się komfortowa, bo dostępna jest szeroka gama gotowych mierników i sond pH. Owszem, ale w przypadku gleby, która płynem nie jest, bezpośredni pomiar pH jest niemożliwy. Najpierw trzeba do próbki gleby dodać wody (destylowanej lub demineralizowanej), w proporcji od 1:1 do 1:5 (gleba:woda), dobrze rozmieszać, by uzyskać konsystencję płynu (a co najmniej rzadkiego błota) i dopiero wtedy zanurzyć sondę w takiej próbce i zmierzyć pH. Też bardzo prosto! Jeśli chcesz, możesz już kupić sprzęt i zacząć mierzyć.

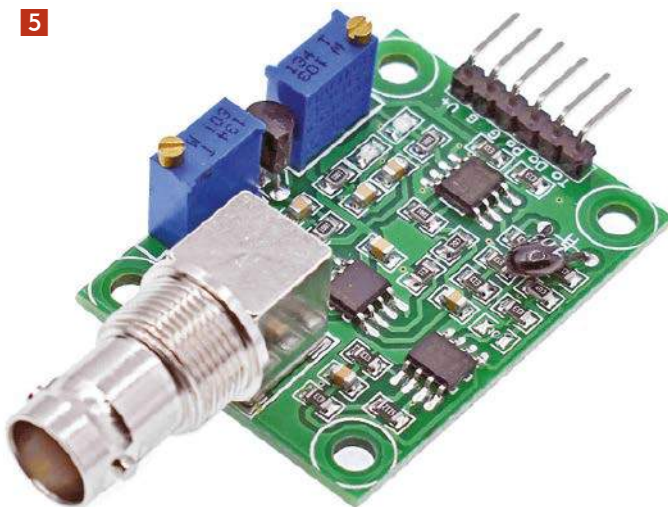
Tak, ale wskazania przyrządów bywają niestabilne, wyniki pomiarów (nawet kolejnych pomiarów tej samej próbki) niepewne, dziwne i mało wiarygodne, przyrządy nie dają się wykalibrować, a w przypadku pomiaru pH gleby nasuwają się liczne dodatkowe poważne wątpliwości...

Na pozór prosty pomiar pH okazuje się skomplikowany i dziwny. W grę wchodzi szereg zaskakujących aspektów. Omówimy je po kolei, zaczynając od podstaw.

Skala pH

Powszechnie wiadomo, że skala pH obejmuje wartości 0...14 i że wartość $\text{pH}=7$ oznacza odczyn obojętny. Wartości 0...7 to odczyn kwaśny, wartości 7...14 – odczyn zasadowy. Wartości bliskie 0 oraz 14 to silne kwasy i zasady. Czysta woda ma neutralny odczyn $\text{pH}=7$. **Rysunek 6**

5



pokazuje przybliżone wartości pH niektórych płynów (wodnych rozтворów rozmaitych substancji).

Niestety, podręcznikowe definicje współczynnika pH rodzą u elektronika więcej pytań, niż wyjaśniają. Według Wikipedii: *Skala pH – ilościowa skala kwasowości i zasadowości rozтворów wodnych związków chemicznych. Skala ta jest oparta na aktywności jonów hydroniowych $[\text{H}_3\text{O}^+]$ w rozтворach wodnych. Tradycyjnie pH definiuje się jako: $\text{pH} = -\log_{10}[\text{H}_3\text{O}^+]$ czyli ujemny logarytm dziesiętny aktywności jonów hydroniowych wyrażonych w molach na decymetr sześcienny. Współcześnie jednak nie jest to ścisła definicja tej wielkości.*

Dla nas ważne jest, że skala pH ma sens w odniesieniu do **wodnych rozтворów** substancji chemicznych i określa ich kwasowość/zasadowość, cokolwiek to znaczy. Dla wielu Czytelników cała reszta to czarna magia. W (nadmiernie) uproszczonych wyjaśnieniach sensu współczynnika pH pojawia się informacja, że stężenie jonów wodorowych w idealnie czystej wodzie to 10^{-7} . Wykładnik potęgi to minus 7, dlatego ma ona współczynnik $\text{pH}=7$. Jeżeli stężenie dodatnich jonów wodorowych jest znikome, wynosi 10^{-14} , czyli 1/100000000000000, wtedy analogicznie, współczynnik pH tego rozтворu wynosi 14 i mamy do czynienia z silną zasadą. A jeżeli stężenie wynosi 10^0 , wtedy pH wynosi 0 i mamy do czynienia z bardzo silnym kwasem. Co prawda w przypadku liczb mniejszych od jedności ich logarytm dziesiętny jest ujemny, ale można się umówić i bez żadnej szkody znak minus po prostu pominąć. **Współczynnik pH to wykładnik potęgi określającej zawartość dodatnich jonów wodorowych (tylko bez znaku minus).** Wydawałoby się, że zrozumieliśmy sens współczynnika pH. Wprowadził go na początku XX wieku duński biochemik Søren Sørensen, który zapewne nie chciał posługiwać się maleńkimi stężeniami z wieloma zerami po przecinku i sprytnie uprościł sprawę, wykorzystując znacznie wygodniejszą w tym przypadku miarę logarytmiczną. Niestety, takie podejście jest zbyt uproszczone, o czym później.

https://pl.wikipedia.org/wiki/Skala_pH **6**

Przykładowe wartości pH	
Substancja	pH
1 M kwas solny	0
0,1 M kwas solny	1
Sok żołądkowy	1,5 – 2
Sok cytrynowy	2,4
Coca-Cola	2,5
Oceł	2,9
Sok pomarańczowy	3,5
Piwo	4,5
Kawa	5,0
Herbata	5,5
Kwaśny deszcz	< 5,6
Mleko	6,5
Chemicznie czysta woda	7
Ślina człowieka	6,5 – 7,4
Krew	7,35 – 7,45
Woda morska	8,0
Mydło	9,0 – 10,0
Woda amoniakalna	11,5
Wodorotlenek wapnia	12,5
1 M rozтвор NaOH	14

Pomiary pH

Zanim omówimy szczegóły, powinniśmy rozszerzyć temat i wspomnieć o metodach kolorymetrycznych (nie kalorymetrycznych). Otóż prostym i popularnym sposobem pomiaru pH było (i nadal jest) użycie różnych substancji, które zmieniają kolor zależnie od wartości pH rozтворu. Zapewne na myśl przychodzi przysłowiowe **papierki lakmusowe**. Wskaźnikiem może być nie tylko lakmus



(pozyskiwany dawniej z pewnych gatunków porostów), ale też szereg innych substancji. Dostępnym w każdym domu wskaźnikiem pH jest napar czarnej herbaty, który jaśnieje pod wpływem kwasów i ciemnieje pod wpływem zasad. Znacznie szerszy zakres zmian barwy mają substancje barwiące z czerwonej kapusty. Bez problemu można kupić papierki wskaźnikowe, a także zestawy do pomiaru pH gleby zawierające naczynie do przygotowania roztworu. **Fotografia 7** pokazuje przyrząd zwany kwasomierzem Helliga. Do zagłębienia wkłada się grudkę badanej gleby, zalewa się tzw. płynem Helliga (woda, alkohol, czerwien metylowa, błękit tymolowy, błękit bromotymolowy). Po rozmieszaniu i odczekaniu 1...2 minut płyn zmienia kolor. Wtedy wystarczy przechylić naczynie, by płyn wypełnił rowek i wzrokowo określić odczyn w zakresie 4...8 pH z dokładnością około 0,5 pH.

Nie warto byłoby przypominać takich archaicznych (choć wciąż popularnych metod), gdyby nie fakt, że opracowano ich nowocześniejsze, „elektroniczne odmiany”. Wprawdzie z uwagi na stopień skomplikowania i koszty nie upowszechniły się przyrządy do elektrooptycznego pomiaru pH. Ale są wykorzystywane w pewnych specyficznych

<https://www.deltatrak.com/isfet-ph-meters>

New
Products

Products

Solutions

Support

Pocket ISFET pH Meter, Model 24006

Key Features:

- Low cost, non-glass ISFET pH system
- Instant response time
- Able to read a single drop of

zastosowaniach, m.in. przy produkcji żywności i napojów. A zasada pomiaru jest bardzo interesująca: wykorzystuje się fakt, że zależnie od odczynu pH zmieniają się właściwości optyczne niektórych substancji: niektóre zmieniają kolor, inne... świecą (fluorescencja), w jeszcze innych zmienia się współczynnik pochłaniania światła. Te parametry światła przechodzącego, odbitego albo emitowanego są mierzone za pomocą nowoczesnej elektroniki, częstokroć „bezdotykowo i w biegu”. Na przykład w jednej z takich metod wykorzystuje się fakt, że odczyn pH dość mocno zmienia pochłanianie światła o długości fali 620 nm, natomiast zupełnie nie zmienia pochłaniania fal o długości 510 nm i dłuższych niż 750 nm. W tym przypadku pomiar pH polega więc na porównaniu pochłaniania światła o różnych barwach.

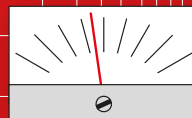
Dla elektronika interesujące jest, że odczyn pH można mierzyć za pomocą... tranzystorów. Konkretnie za pomocą specyficznej odmiany tranzystorów MOSFET, zwanych ISFET (Ion-Sensitive FET – czuły na jony tranzystor polowy). Nie wnikając w szczegóły – prąd ISFET-a jest zależny od współczynnika pH roztworu. Czujniki z tranzystorami ISFET jak na razie nie upowszechniły się z uwagi na wysokie ceny, niemniej jak najbardziej dostępne są pH-metry z takimi czujnikami (kosztujące kilkaset dolarów/euro) – przykład na **fotografii 8**. Zaletą jest duża dokładność, stabilność i trwałość, a podstawową wadą – wysoka cena. ■

Piotr Górecki

REKLAMA

KOMPUTERY RASPBERRY PI I MODUŁY ARDUINO

<http://sklep.avt.pl>

AUDIO
OUT

By Jake Rothman

Projektowanie mini monitorów Hi-Fi, część 1

Brytyjska branża audio ma długie tradycje w tworzeniu małych, lecz dobrej jakości głośników Hi-Fi, zwanych mini monitorami. Ich koncepcja wywodzi się z BBC, która w latach sześćdziesiątych i siedemdziesiątych XX wieku prowadziła pionierskie badania i prace nad systemami głośnikowymi. W ich rezultacie powstał dwudrożny, małych wymiarów zestaw głośnikowy przeznaczony do użytku w wozach nadawczych i studiach radiowych, znany jako „LS3/5A”. Ten właśnie wypróbowany i sprawdzony model jest inspiracją niniejszej serii artykułów.

W tym miesiącu zaczniemy od krótkiego tła historycznego opisującego kulisy powstania zestawu LS3/5A i przeprowadzimy techniczne rozważania nad procesem projektowania głośników (co zaskakujące, dokumenty powstałe w wyniku tej finansowanej przez płatników pracy są dostępne do pobrania ze strony internetowej BBC. Jest to obowiązkowa lektura dla każdego aspirującego projektanta głośników. Obecnie BBC znacznie ograniczyło prace badawcze i zadowala się wykorzystywaniem technologii konsumenckich, takich jak Skype czy głośniki komercyjne). W przeciwieństwie do większości elektroniki użytkowej, głośniki są nadal dość drogie, mimo że w niektórych obszarach parametry dźwięku nawet się pogorszyły. Nie wiem, co jest przyczyną, ale wygląda

na to, iż dobre głośniki kupują teraz tylko elitarne kręgi entuzjastów Hi-Fi i profesjonalistów zajmujących się dźwiękiem. Solidne, drewniane głośniki zdobiące salony w latach osiemdziesiątych poprzedniego wieku zostały zastąpione przez plastikową tandetę. W ciągu najbliższych kilku miesięcy pokażę, jak wykonać własne głośniki i nadać im nieco niższy poziom koloryzacji, jaki wykazywały starsze konstrukcje BBC i KEF. Do zbudowania będzie sześć wariantów, z których większość będzie hołdem dla BBC LS3/5A – patrz ramka obok. Nietypowa nazwa LS3/5A wywodzi się z konwencji BBC. Oznaczenie „LS3” oznacza, że był to model przeznaczony (głównie) do wozów transmisyjnych, natomiast „LS5” był używany jako monitory studyjne. Cyfra po ukośniku określa numer modelu (5), a ostatnia litera oznacza wersję specyfikacji, która była tylko jedna – stąd też „A”. Niecodzienną specyfikacją modelu LS3/5A jest fakt, iż od 1975 roku BBC licencjonowało produkcję niewielkiej liczbie firm brytyjskich, ponieważ samo nie było zainteresowane produkcją głośników. Współczesne wcielenie oryginalnych LS3/5A firmy Rogers pokazano na rysunku 1, natomiast rysunki 2a i 2b przedstawiają reklamy starszych mini monitorów z lat 60., Maxim firmy Goodmans i Ditton 10 firmy Celestion. Wyjątkowe, okrągłe głośniki JR 149s (rysunek 3), jak również inspirowane modelem LS3/5A głośniki Proacs, Spendors i Harbeth HL-P3 osiągają obecnie wysokie ceny

na eBayu. Produkowane przez firmę Rogers modele LS3/5A były w powszechnej opinii uznawane za najlepsze a ostatnio pojawiły się ponownie wyposażone w chińskie wersje kultowych głośników KEF. Ich cena wynosi około 2800 funtów, co sprawia, że z punktu widzenia inżynierskiego stają się czymś w rodzaju zegarków Rolex.

Synergia systemu

Wersje projektowe mini monitorów inspirowanych modelem są opisane w spisie Mini monitory PE – opcje budowy. Pasywny system głośnikowy to zbiór części i podsystemów, które wzajemnie na siebie oddziałują i łączą się, tworząc spójną całość. Oprócz głośników jest jeszcze obudowa i specjalny układ zwany zwrotnicą, który kieruje odpowiednie zakresy częstotliwości do każdego z głośników. Podobnie jak w dobrze zestrojonym samochodzie, wszystkie te części muszą być tak dobrane, aby ze sobą współpracowały. Nie ma przecież sensu montować głośnika niskotonowego za 100 funtów w pudełku po butach. Jeśli chcesz stworzyć własną konstrukcję z wykorzystaniem alternatywnych głośników, będziesz musiał zainwestować w oprogramowanie do pomiaru charakterystyki częstotliwościowej lub użyć starego plotera analogowego, takiego jak Neutrix Audiograph (rysunek 4) w pomieszczeniu bezechowym (w moim przypadku było to poddasze wyłożone materacami) dla przykładu pokazanym



Rysunek 1. LS3/5A produkowany obecnie przez firmę Rogers



Rysunek 2. Reklamy z lat 60. przedstawiające pierwsze udane brytyjskie mini monitory: po lewej Goodmans Maxim, po prawej zaś Ditton 10 firmy Celestion



Rysunek 3. Okrągły Jim Rogers JR149 wywodzący się z modelu LS3/5A. Kolumna o bardzo niskim podkoloryzowaniu

na rysunku 5. Na początku musimy omówić jednak pewne podstawy i terminologię związaną z głośnikami. Projektowanie i budowa głośników ma charakter multidyscyplinarny – obejmuje bowiem akustykę, mechanikę, materiałoznawstwo, technologię klejenia, magnetyzm, a nawet stolarkę. W przypadku konstrukcji pasywnych elektronika stanowi niewielką część tego, co wistocie jest systemem elektroakustycznym. Nadal jest to w dużej mierze rzemiosło, które może dać prawdziwą satysfakcję z budowy (przekonałem się, że budowa głośników to doskonały sposób na uczenie techniki zniechęconych, młodych ludzi!).

Zakres częstotliwości, głośniki niskotonowe i wysokotonowe

Dziedzina audio jest niezwykle spośród inżynierskich technologii, ponieważ obejmuje trzy dekady zakresu częstotliwości w stosunku 1000 do 1 – od 20 Hz do 20 kHz, z odpowiadającymi im długościami fal od 17 m do 17 mm. Dla porównania – stosunek całości widma światła widzialnego wynosi zaledwie

1,8 do 1. Komplikuje to fizyczne wymagania dotyczące optymalnego rozchodzenia się dźwięku, ponieważ na przeciwnych krańcach spektrum dźwiękowego należy brać pod uwagę zupełnie różne czynniki. Aby zapewnić odpowiednią reprodukcję niskich częstotliwości, głośnik musi mieć dużą powierzchnię – co najmniej 70 cm² (10 cali kwadratowych) – i być w stanie poruszać część ruchomą do przodu i do tyłu o około 12,5 mm (pół cala). Jeśli głośnik ma większą powierzchnię, nie musi się tak bardzo wychylać – liczy się bowiem całkowite przemieszczenie powietrza. W przypadku wysokich częstotliwości ważna jest z kolei masa ruchoma, ponieważ membrana musi poruszać się szybko i bez bezwładności. Niewielka powierzchnia około 2 cm² jest wystarczająca, ponieważ przy wysokich częstotliwościach drgania są skutecznie sprężane z powietrzem, co można uznać za pewną formę dopasowania impedancji akustycznej. Te dwa bardzo różne wymagania konstrukcyjne są powodem, dla którego wszystkie systemy głośników Hi-Fi składają się z co najmniej



Rysunek 4. Do wykreślenia charakterystyki częstotliwościowej głośnika potrzebny jest co najmniej generator sygnałowy i multimetr z odczytem dB. Po lewej – osobiście używam Neutrika Audiograph 3300 z 1981 roku; po prawej – Neutrik i generator stałoprądowy używany do pomiaru krzywej impedancji



dwóch głośników – jednego zoptymalizowanego pod kątem niskich tonów i drugiego zoptymalizowanego pod kątem wysokich tonów. Te głośniki są nazywane odpowiednio „głośnikiem niskotonowym” i „głośnikiem wysokotonowym”, a cały system jest określany jako „dwudrożny”. Zazwyczaj głośnik niskotonowy jest zasilany częstotliwościami z zakresu od 20 Hz do 3 kHz, a głośnik wysokotonowy od 3 kHz do 20 kHz.

Zwrotnice

Aby system dwudrożny działał efektywnie, potrzebne są filtry, które skierują sygnały o określonych częstotliwościach do odpowiedniego głośnika. Najprostszym tego typu układem jest kondensator, zwykle o wartości 2,2 μ F, połączony szeregowo z głośnikiem wysokotonowym, tworzący filtr górnoprzepustowy o częstotliwości 5 kHz w przypadku głośnika wysokotonowego 8 Ω . Zapobiega to niszczeniu cewki i membrany głośnika wysokotonowego o małej masie przez prądy o niskich częstotliwościach (należy pamiętać, że w muzyce jest znacznie więcej energii o niskiej częstotliwości niż o wysokiej. Z tego powodu cewka głośnika wysokotonowego może mieć obciążalność cieplną zaledwie kilku watów, nawet jeśli jest używana w systemie zasilanym przez wzmacniacz o mocy 50 W). W przypadku basu, filtrem dolnoprzepustowym może być po prostu mechaniczne ograniczenie głośnika niskotonowego. Innymi słowy, w przypadku bardzo prostej zwrotnicy głośnik niskotonowy nie ma żadnych elementów elektrycznych. Taki system może działać zaskakująco dobrze, jak na przykład w modelu Acoustic Research AR7. Jednak w przypadku dobrej jakości zwrotnic kondensator głośnika wysokotonowego nie wystarcza, a dla jednostki basowej korzystne jest zastosowanie filtra. Dzieje się tak dlatego, że mechaniczna granica

Mini monitory PE („Practical Electronics”) – opcje budowy

1. Prawdziwy BBC LS3/5A – wszystkie komponenty z Falcon Acoustics (FA)

Obudowa – obudowa z zestawu, specyfikacja BBC ze sklepu internetowego FA
Głośnik/driver – specyfikacja BBC ze sklepu internetowego FA (KEF B110A)
Głośnik wysokotonowy – specyfikacja BBC ze sklepu internetowego FA (KEF T27)
Zwrotnica – specyfikacja BBC ze sklepu internetowego FA
Całkowity koszt elementów to około 850 funtów – łączny koszt głośników to około 2300 funtów.

2. Konstrukcja głośnika Vifa (duńska) – dobry, mały mini monitor – nie „LS3/5A”

Obudowa – wstępnie zmontowana, niepomalowana obudowa z MDF-u od Earworm Records
Głośnik/driver – driver Vifa (BC14SG49-08), od Earworm Records
Głośnik wysokotonowy – głośnik wysokotonowy Vifa (TC26SF05-06), od Earworm Records
Zwrotnica – zwrotnica zabudowana na drewnie od Earworm Records
Ta konstrukcja jest zasadniczo zmontowana, włączając w to głośnik, głośnik wysokotonowy i zwrotnicę, ale mam kilka wskazówek odnośnie ulepszeń. Koszt to około 200 funtów za parę zmontowanych urządzeń + koszty wysyłki: łączny koszt to około 600 funtów za głośniki.

3. Projekt DIY Vifa – opcja 2, ale z własną konstrukcją PE dla obudowy i zwrotnicy

Obudowa – konstrukcja ze sklejki PE – o głębokości 9 cali z przegrodą dla modułów Vifa
Głośnik/driver – jednostki napędowe Vifa (BC14SG49-08)
Głośnik wysokotonowy – głośnik wysokotonowy Vifa (TC26SF05-06)
Zwrotnica – płytki zwrotnicy projektu PE

4. Mini monitor PE, w stylu LS3/5A (wersja FA)

Obudowa – konstrukcja PE ze sklejki (głębokość 9 cali) z przegrodą dla driver'a FA
Głośnik/driver – ten sam przetwornik, co w zestawie LS3/5A (nr 1) (KEF B110A)
Głośnik wysokotonowy – inny niż w zestawie nr 1 – wykorzystany jest przetwornik SEAS
Zwrotnica – płytki zwrotnicy projektu PE

5. Mini monitor PE, w stylu LS3/5A (wersja Wavecor)

Obudowa – konstrukcja PE ze sklejki (głębokość 9 cali) z przegrodą dla modułów Wavecor
Głośnik/driver – Wavecor
Głośnik wysokotonowy – Wavecor
Zwrotnica – płytki zwrotnicy projektu PE

6. Aktywny mini monitor PE, w stylu LS3/5A (układ działa ze wszystkimi powyższymi)

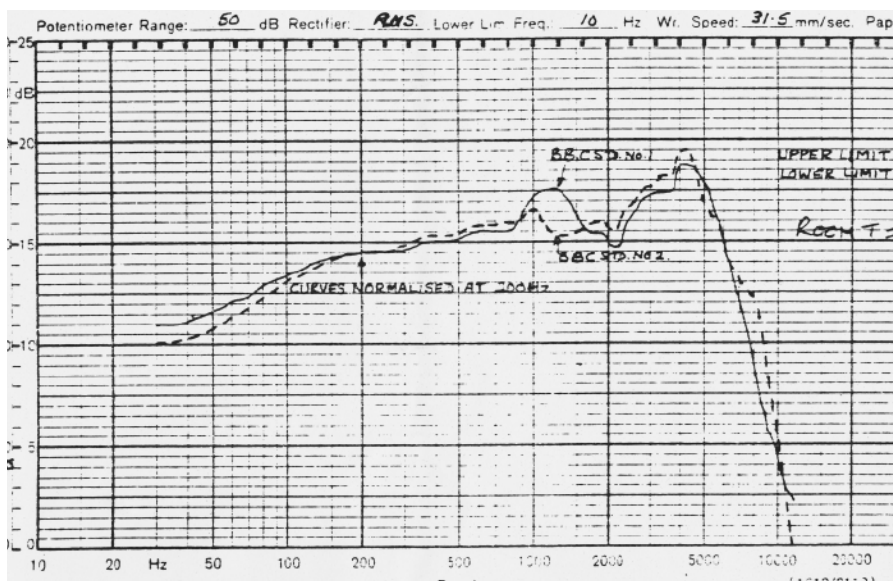
Opcja aktywna będzie zbudowana na pojedynczej płytce drukowanej i będzie zawierała wiele szczegółów dotyczących rozmieszczenia przewodów, wzmacniaczy mocy i zasilaczy, które opisałem wcześniej w Audio Out. Układ ten będzie w stanieysterować nie tylko mini monitor, ale także praktycznie każdy mały dwudrożny zestaw głośnikowy. Konstrukcja obudowy PE ze sklejki zapewni, iż dodanie opcji aktywnej wymagać będzie jedynie wymiany tylnej ścianki, na której znajdzie się wycięcie na radiator, złącze sieciowe IEC i gniazdo XLR.

wysokich częstotliwości głośnika niskotonowego jest słabo zdefiniowana. Może się ona znacznie różnić w zależności od głośnika

i często wzrasta wraz z poziomem głośności. Do skutecznego filtrowania potrzebna jest więc co najmniej cewka w szeregu. Ponadto konieczne są dodatkowe elementy reaktywne do uzyskania bardziej stromego zbocza filtrowania oraz do dostrojenia obwodów w celu usunięcia pików. Cewki muszą być stosowane, ponieważ filtr dolnoprzepustowy RC miałby niedopuszczalne straty ze względu na rezystancję szeregową. Cewka zaś ma niską impedancję przy niskich częstotliwościach przepuszczając je, podczas gdy jej rosnąca reaktancja przy wysokich częstotliwościach blokuje sygnały o wysokiej częstotliwości bez żadnych strat mocy/ciepła. Niestety, cewka jest najmniej „doskonałym” elementem pasywnym; ma znaczną rezystancję przy prądzie stałym, zniekształcenia (z powodu histerezy i wibracji), koszt i duże rozmiary fizyczne. Te wady są zwykle o rząd wielkości gorsze niż w przypadku kondensatorów. Jednak w przypadku zwrotnicy pasywnej nie mamy wyboru i musimy użyć



Rysunek 5. Aby zminimalizować odbicia dźwięku w pomieszczeniu podczas wykonywania pomiarów charakterystyki częstotliwościowej, konieczne jest wyłożenie ścian materiałem pochłaniającym dźwięk. Przedstawione tutaj materace na moim strychu pomagają zapewnić niemal bezchłowe warunki dla wartości powyżej kilkuset herców. Na zdjęciu widać testowany system aktywny Meridian M2



Rysunek 6. Charakterystyka częstotliwościowa głośnika niskotonowego B110A zastosowanego w LS3/5A w obudowie. Wymaga on znacznej korekcy przez filtr zwrotnicy. Ten stary wykres z BBC opracowany przez Malcolma Jonesa pokazuje dopuszczalne różnice charakterystyki częstotliwościowej pomiędzy urządzeniami. Należy pamiętać, że krzywe charakterystyki częstotliwościowej głośników zawsze wykresła się na osiach logarytmicznych, przy czym na osi Y znajduje się odpowiedź w decybelach (dB), a na osi X częstotliwość sygnału. Niestety BBC oparło swoją konstrukcję na wyjątkowo dobrej próbce, co spowodowało problemy z masową produkcją w momencie pojawienia się zamówień w firmie KEF

ale oczywiście potrzebny jest zasilacz, który nie ma zastosowania w systemie pasywnym. W przypadku głośników aktywnych jest to jednak preferowana droga projektowania i dzięki temu unika się cewek).

Częstotliwość zwrotnicy

W zależności od zakresu częstotliwości głośników należy wybrać rozsądny punkt podziału pasma pomiędzy. Ten punkt podziału określa się jako „częstotliwość zwrotnicy”. W idealnej sytuacji, odpowiedź głośników powinna być szersza o oktawę po obu stronach częstotliwości zwrotnicy. Jednak rzadko się to zdarza, szczególnie w przypadku głośników basowych, gdzie zazwyczaj po osiągnięciu szczytu charakterystyka gwałtownie się obniża dla wyższych częstotliwości. Głośnik niskotonowy KEF B110A zastosowany w LS3/5A nie jest wyjątkiem, jak pokazano na rysunku 6 (odpowiednia krzywa dla głośnika wysokotonowego KEF T27 jest przedstawiona później na rysunku 31). Częstotliwość zwrotnicy jest zwykle ustawiana między 2 kHz a 4 kHz dla głośników o standardowych rozmiarach. Takie głośniki niskotonowe mają zwykle średnicę 80–200 mm (4–8 cali), a głośniki wysokotonowe 18–30 mm (0,75–1,25 cala).



cewek (w filtrach aktywnych nie ma potrzeby stosowania cewek i można wykorzystać tanie, dokładne i małe rezystory i kondensatory,

Rysunek 7. Przerosty i uszkodzenia ujawnione na lekkim i prostym głośniku 8-calowym Audax z papierową membraną. Membrana została posypana talkiem, a następnie wystawiona na fale sinusoidalne o mocy 5 W_{rms} i stopniowo zwiększającej się częstotliwości, aż do momentu wystąpienia rezonansu. Proszek pozostał na nieruchomych częściach membrany, ujawniając mapę drgań

Rysunek 8. Pierwsze promieniste przełamanie przy 820 Hz

Rysunek 9. 1,3 kHz

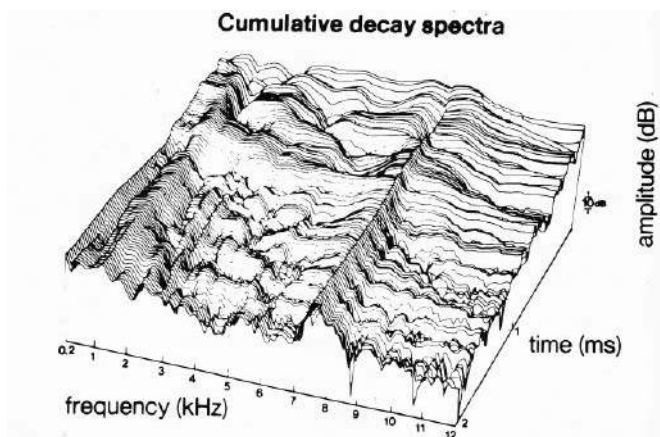
Rysunek 10. 2,1 kHz

Rysunek 11. 2,75 kHz

Rysunek 12. Kołowe przełamanie przy 4,25 kHz

Rysunek 13. 6,6 kHz





110mm moving coil bass/mid-range unit with lightweight coil in 7 litre closed box.

Rysunek 14. Wykres wodospadowy odpowiedzi częstotliwościowej w czasie. Firma KEF była pionierem w stosowaniu trójwymiarowych wykresów wodospadowych z trzecią osią Z dla czasu w celu pokazania opóźnionego rezonansu. Ten eksperymentalny głośnik średniotonowy wykazuje rezonans przy 7,4 kHz, który nie pojawia się w początkowej charakterystyce częstotliwościowej, ale staje się dość oczywisty po pewnej chwili



Rysunek 16. Montaż głośników niejako „chowający je” w płaszczyźnie obudowy pomaga zminimalizować efekty filtrowania grzebieniowego w paśmie górnych częstotliwości, spowodowane narastającym efektem dyfrakcji. Obudowy dla przetworników Wavecor mają fronty wykonane z dwóch sklejonych ze sobą kawałków sklejk

Jeśli w celu uzyskania wyższych poziomów dźwięku stosowane są większe głośniki niskotonowe (>200 mm), konieczne staje się dalsze podzielenie pasma częstotliwości poprzez dodanie jednostki średniotonowej (zwanej „squawker” w amerykańskich kregach audio). Punkt zwrotnicy średniej częstotliwości znajduje się wtedy zwykle w zakresie od 200 do 700 Hz. Należy pamiętać, że w systemach dwudrożnych bardziej poprawne technicznie jest opisywanie głośnika niskotonowego jako jednostki „nisko-średniotonowej”. Ważna jest jeszcze jedna kwestia dotycząca konstrukcji zwrotnicy – ponieważ zwrotnica to w zasadzie dwa filtry: dolnoprzepustowy i górnoprzepustowy, sensowne jest modyfikowanie tych filtrów w celu spłaszczenia lub „wyrównania” wszelkich złych tendencji lub szczytów w paśmie przenoszenia głośników, dla których zostały zaprojektowane.

Koloryzacja

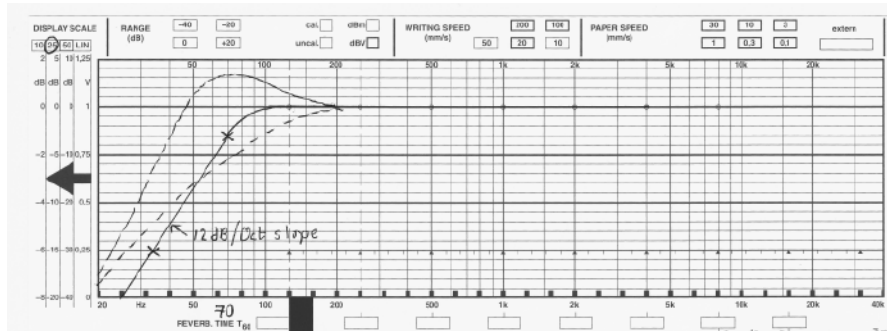
Głośnik, który wykazuje szczególną aberrację tonalną, zwykle w wyniku niepożądanych rezonansów, określa się jako „podkolorizowany”.

Do opisanie tego efektu często używa się subiektywnych określeń, w zależności od obszaru częstotliwości, którego on dotyczy. Wszystkie materiały, z których wykonane są membrany i obudowy, mają charakterystyczną koloryzację. Typową koloryzację głośnika można odtworzyć, mówiąc do plastikowego kubka. Stukanie w polipropylenowy pojemnik wywoła dźwięk charakterystyczny dla membrany polipropylenowej. Foliowe opakowanie na ciastka będzie brzmiało jak membrana aluminiowa. Legenda głosi, że Dudley Harwood został zainspirowany do wynalezienia głośnika z plastikową membraną przez plastikowy kubek w stołówce BBC. Zauważyłem też, że konstruktorzy głośników ciągle stukają i słuchają różnych rzeczy. Budzi to powszechne zdziwienie, gdy wraz z żoną wybieramy meble!

Przełamywanie

W teorii, membrany głośników powinny drgać niczym tłok, aby uzyskać przewidywalne rezultaty. Powyżej pewnej częstotliwości jest to niemożliwe, ponieważ żaden materiał nie jest nieskończenie sztywny. Uformowanie

materiału w kształt stożka/trójkąta znacznie zwiększa sztywność przeciw odkształceniom. Typowe nieprzewidywalne przełamywania zginające występują przy częstotliwości około kilkuset herców dla membran 200 mm i około 700 kHz dla membran 110 mm. Rysunki od 7 do 13 przedstawiają serię przełamań pojawiających się wraz ze wzrostem częstotliwości. Każde z nich jest zwykle związane z załamaniem krzywej odpowiedzi głośnika. Te przełamywania (break up), jak się je powszechnie nazywa, dają początek charakterystycznemu brzmieniu głośników. Można je wykorzystać w sposób artystyczny, na przykład w głośnikach gitarowych. W przypadku głośników Hi-Fi jest to poważny problem. W Stanach Zjednoczonych inżynierowie audio stosują bardzo opisowy termin; określają ten dźwięk jako „płacz membrany”. Podobnie jak w przypadku turbulencji, matematyka tego zjawiska jest złożona i chaotyczna. Istnieją modele, ale nie jestem jeszcze przekonany o ich skuteczności. Mam program komputerowy, który symuluje odpowiedź głośnika gitarowego Celestion G12, ale jest on zbyt uproszczony. Przypomina tylko podstawowe wzorce drgań korpusów skrzypiec i wyjaśnia, dlaczego nie ma aplikacji Stradivarius na iPhone’a. Kształt membrany jest także istotny. W większości zestawów nisko-średniotonowych Hi-Fi stosuje się membranę zakrzywioną, ponieważ mimo że membrana rezonuje przy niższej częstotliwości, przełamywania są mniej dotkliwe. Ponadto obszar emisji stożka zmniejsza się wraz ze wzrostem częstotliwości, co rozszerza pasmo przenoszenia wysokich częstotliwości i zwiększa dyspersję. Prostoboczne stożki są lepsze w przypadku zestawów niskotonowych, ponieważ są sztywniejsze i mogą



Rysunek 15. Idealne obciążenie pasma uzyskane przy użyciu obciążenia zamkniętego. Dla Q równego 0,707 punkt -3 dB znajduje się przy 70 Hz. Linie przerywane pokazują efekt nadmiernego tłumienia (utrata basów) i niedotłumienia, powodującego ich podbicie

przenosić większe siły na wierzchołku. Kompensując te efekty, korekcja dźwięku (EQ) może jedynie zmniejszyć wzbudzenie mód rezonansowych. Po ich wzbudzeniu mogą one brzmieć jeszcze przez pewien czas po ustaniu sygnału. Nazywa się to „rezonansem opóźnionym” i można go przedstawić na wykresie wodospadowym, jak pokazano na rysunku 14. Czasami można je usłyszeć 30 dB niżej. Nie ma większego sensu wyrównywanie charakterystyki głośnika, aby była płaska. Aby głośnik brzmiał płasko, korektor musiałby uwzględniać energię zmagazynowaną. W wielu konstrukcjach głośników po prostu wprowadza się większe niż teoretycznie konieczne obniżenie pasma przenoszenia w okolicach załamania średnich tonów, aż do uzyskania „realistycznego” brzmienia. Możliwość szybkiego nagrywania i odtwarzania próbek dźwięków na komputerze jest bardzo przydatna w przypadku tej techniki testowania na żywo kontra odwzorowanie”.

Obudowa

Obudowa głośnika jest istotną częścią konstrukcji (a prawdopodobnie także częścią, która jest najmniej rozumiana i której entuzjaści elektroniki obawiają się najbardziej). Nie jest to tylko skrzynia, w której umieszcza się głośniki, lecz pełni ona również ważną funkcję akustyczną, polegającą na zapewnieniu odpowiednich basów. Głośnik nie jest w stanie emitować niskich tonów ani długich fal, jeśli nie zapewni się środków zapobiegających znoszeniu dźwięków w przeciwnej fazie z tyłu membrany przez dźwięki z przodu. W większości małych głośników osiąga się to przez zamontowanie jednostki basowej w szczelnej obudowie. Zapobiega to wygaszaniu fazy, ale ma pewną wadę: poduszka powietrzna z uwięzionego powietrza usztywnia zawieszenie membrany. Ma to wpływ na podwyższenie częstotliwości rezonansowej. Poniżej częstotliwości rezonansowej reprodukcja basów gwałtownie spada. Dobrą ilustracją tego efektu jest umieszczenie plastikowego pojemnika, takiego jak miska do ciasta drożdżowego, odwróconego do góry dnem w wodzie i zapukanie. Będzie on rezonował trochę jak basowy bęben. Częstotliwość będzie tym wyższa, im mniej powietrza będzie pod nim uwięzione. Spróbuj tego następnym razem, gdy będziesz zmywał naczynia.

Szczelna skrzynia

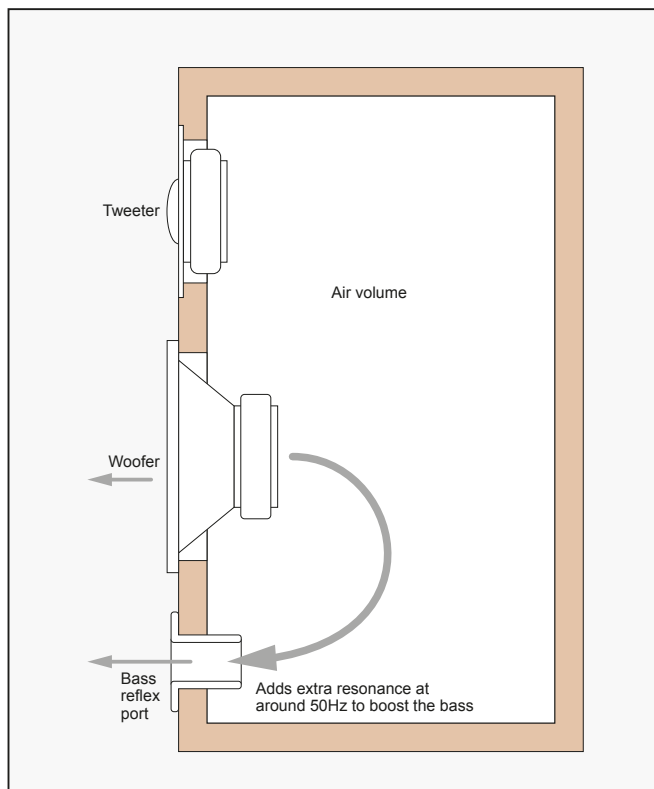
Głośnik w zamkniętej obudowie charakteryzuje się charakterystyką filtra górnoprzepustowego o tłumieniu 12 dB/oktawę, jak pokazano na rysunku 15. W idealnej sytuacji w punkcie odcięcia nie powinno być podbicia, które powodowałoby dudnienie, ani też

spadek nie powinien zaczynać się zbyt wcześnie, co powodowałoby utratę basów. Efekty te są określane przez całkowite tłumienie lub tzw. „Q” systemu. Q, czyli od angielskiego słowa „quality” – „jakość”, jest miarą tego, jak mocno system rezonuje. Wysokie Q oznacza, że rezonans spada bardzo gwałtownie po obu stronach częstotliwości rezonansowej. Niska Q oznacza, że jest odwrotnie, a efekty rezonansowe są widoczne w znacznie szerszym zakresie częstotliwości. Ogólna wartość współczynnika Q głośnika zależy od kilku czynników: współczynnika

Q elektromagnetycznego przetwornika, współczynnika Q mechanicznego zawieszenia przetwornika oraz stopnia, w jakim „sprężyna powietrzna” wypycha rezonans. Jeśli system jest nieco nadmiernie wytłumiony, na przykład przy Q równym 0,5, podbicie basów za pomocą zwykłego regulatora basów Baxandalla naprawi to. Takie podbicie basów nie może być zastosowane, jeśli Q jest wysokie (znacznie powyżej 0,5). Optymalna wartość dająca najbardziej płaską charakterystykę jest taka sama jak w przypadku filtra Butterwortha, z Q równym 0,7.

Odbicia

Kolejnym aspektem projektowania obudowy są odbicia i fale stojące, które występują w obudowie. Podobnie jak w przypadku pomieszczeń, najgorszym kształtem jest sześcian, ponieważ szczyty rezonansowe z każdego kierunku pokrywają się. Bardziej równomierny rozkład szczytów uzyskuje się, gdy każdy wymiar obudowy (wysokość, szerokość, głębokość) jest inny. Jeszcze lepiej jest, gdy wymiary odpowiadają złotej proporcji (1:1,618...). Ponadto obudowa powinna być głębsza niż szerokość, aby zminimalizować odbicia od tyłu. Jednak odbicia szcztkowe nadal muszą być tłumione, dlatego większość obudów głośnikowych jest wyłożona pianką. Pianka ta musi być miękka, typu otwartokomórkowego, która pochłania



Rysunek 17. Obudowy typu reflex oferują więcej basu, lecz o gorszej jakości. Pozwalają one na użycie jednostek basowych o wyższym rezonansie (>45 Hz) i niższym Q (<0,3), które nie nadawałyby się do skrzyni zamkniętej

wodę. Na rynku dostępne są pianki specjalnie zoptymalizowane do zastosowań akustycznych. Mieszkam w Walii, więc wspieram lokalną gospodarkę, wypełniając moje kolumny wełną owczą, która ma dobrze udokumentowane, doskonałe właściwości akustyczne. Kolejną korzyścią jest to, że wełna sprawia, iż objętość obudowy wydaje się nieco większa, obniżając częstotliwość rezonansową o kilka herców. Dzieje się tak dzięki spowolnieniu prędkości dźwięku ze względu na jej dużą masę w stosunku do powietrza. Włókna również poruszają się, zwiększając masę ruchomą systemu. Firma KEF odkryła ostatnio, że wypełnienie obudowy woreczkami z węglem aktywnym wzmacnia te efekty, sprawiając, że obudowa wydaje się o jedną trzecią większa. Z tymi technikami można przesadzić – możliwe jest nadmierne zatłumienie obudowy, co prowadzi do utraty wydajności i basów.

Schodek dyfrakcyjny

Wymiary przedniej ścianki skrzynki są ważne ze względu na ich związek z długością fali odtwarzanego dźwięku. Gdy szerokość skrzynki zbliża się do długości fali dźwięku, wyjście systemu zmienia charakterystykę z dookólnej na kierunkową, podbijając wyjście mierzone przed głośnikiem, na jego osi promieniowania. To podbicie jest często nazywane schodkiem dyfrakcyjnym („baffle step”)



Rysunek 18. Dopięcie głośnika nisko-średniotonowego wykonane za pomocą warstwy tłumiącej wiskoelastycznej o wysokiej histerezie zmniejsza ilość przesterowań. Powłoka ta jest widoczna na tylnej ścianie membrany B110A

i występuje zwykle przy częstotliwości około 300 Hz. Jest ono niepożądane w dobrych głośnikach Hi-Fi i można je skorygować za pomocą filtra dolnoprzepustowego, co jest techniką zwaną „korekcją charakterystyki skokowej”.

Dyfrakcja głośnika wysokotonowego

Gdy długość fali sygnału dla wysokich częstotliwości wzrasta powyżej (powiedzmy) 1,5 kHz, dźwięk zaczyna rozpraszać się na wystających krawędziach przedniej części obudowy. Powoduje to powstawanie interferencji,



Rysunek 20. Przegroda LS3/5A. Filcowe paski wokół głośnika wysokotonowego redukują zakłócenia dyfrakcyjne od krawędzi ramki maskownicy. W przeciwieństwie do większości głośników, LS3/5A pracuje lepiej z założoną maskownicą. Po jej zdjęciu pojawia się wcięcie w charakterystyce na częstotliwości 6,5 kHz



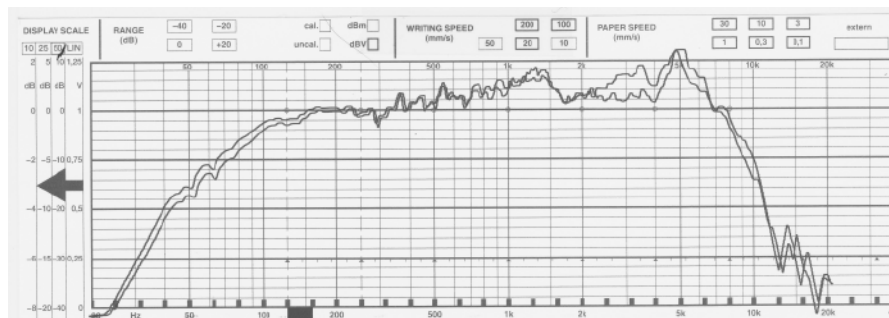
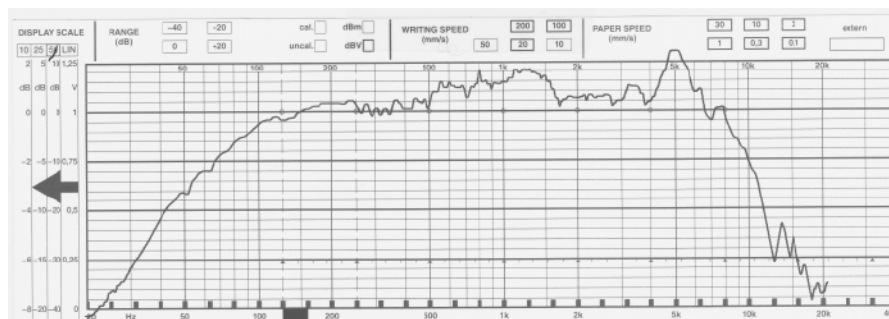
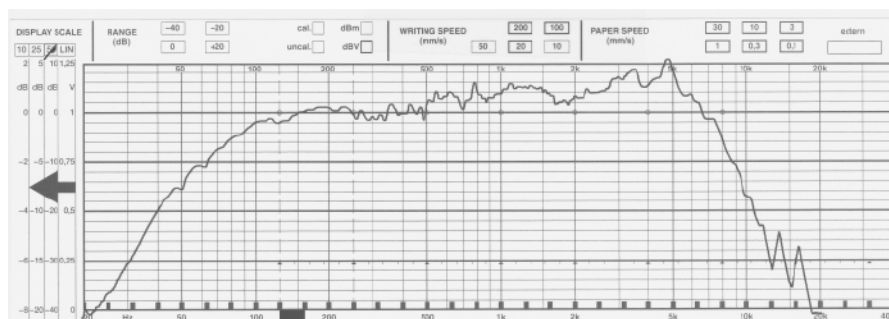
Rysunek 21. Tylne części przegrody LS3/5A z widoczną zwrotnicą. Została ona zamontowana nad magnesem głośnika wysokotonowego przy użyciu filcowej podkładki o grubości 3/8 cala, aby zapobiec grzechotaniu. Obawiałem się, że magnes będzie wpływał na cewkę ale nie udało mi się zmierzyć żadnej istotnej różnicy



Rysunek 19. Tłumienie paneli głośnikowych za pomocą podkładek tłumiących karoserię samochodową zmniejsza vibrację. Na zdjęciu prezentowane jest wnętrze obudowy KEF 104ab. (W moich demonstracjach w sali wykładowej, skradzionych z wykładów Lauri Fincham, znajdują się dwa kawałki sklejek – jeden z tłumieniem, drugi bez. Po uderzeniu patką jeden z nich wibruje, drugi wydaje tępy stukot)

co skutkuje nieregularnościami w odpowiedzi częstotliwościowej w przypadku ostrych krawędzi (jak je zmniejszyć, patrz rysunek 20). Jeśli krawędzie są zaokrąglone, dyfrakcja jest łagodna. Z tego samego powodu konieczne

jest montowanie głośników wysokotonowych w przednim panelu głośnika – przegrodzie. Jeśli nie zostanie to zrobione, ostra okrągła krawędź przedniej płyty głośnika wysokotonowego będzie często powodować spadki –6 dB,



Rysunek 22. Porównanie przedniego (górna krzywa) i tylnego (środkowa krzywa) montażu jednostki basowej w LS3/5A. Krzywa jest bardziej pofalowana w przypadku montażu tylnego, ale wyrównanie fazowe na zwrotnicy koryguje ten efekt. Najniższa krzywa prezentuje obie charakterystyki nałożone na siebie

i działać jak filtr grzebieniowy od 4 kHz w górę. Montaż głośników wysokotonowych we wgłębieniach w przegrodzie jest zazwyczaj wykonywany za pomocą frezarki i szablonu. W przypadku pojedynczych konstrukcji używam zwykle dwóch kawałków sklejk, jak pokazano na rysunku 16.

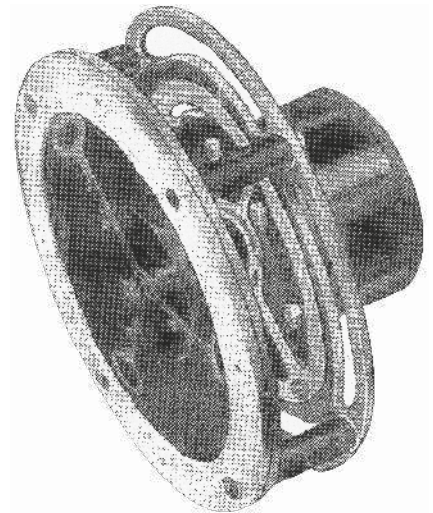
Skrzynia basowa

Głośnik w szczelnej obudowie był kiedyś rozwiązaniem stosowanym powszechnie, ale obecnie większość nowoczesnych konstrukcji to systemy typu reflex, co oznacza, że w obudowie znajduje się otwór, zwykle z przymocowaną rurką – patrz na rysunku 17. Dzięki temu uzyskuje się więcej basów i większą wydajność, ponieważ sumuje się rezonans Helmholtza. Zjawisko to można łatwo zademonstrować, stukając w plastikową butelkę po napełnieniu i przykładając otwór do ucha. Zmniejszając rozmiar otworu wentylacyjnego i wydłużając jego długość ręką, można zmniejszyć częstotliwość rezonansu. Obudowy typu reflex dobrze się sprzedają, ponieważ są głośniejsze i można w nich stosować jednostki o wyższym rezonansie. Uważam jednak, że są one mniej dokładne niż skrzynki zamknięte, ponieważ promieniują bas po zakończeniu nuty. Ponadto mają one znacznie ostrzejsze odcienie czwartego rzędu i dwukrotnie większą szybkość

zmiany fazy (zwaną czasem opóźnieniem grupowym). Ostatnio pojawiły się dowody na to, że ucho może słyszeć duże opóźnienie grupowe przy niskich częstotliwościach. Mikserzy dubbingu filmów i animacji skarżą się, że przy użyciu głośników typu reflex trudno jest precyzyjnie przyciąć dźwięk do huków i eksplozji. Innym problemem jest podbarwiony dźwięk w średnim zakresie, który może przedostawać się przez tubę typu reflex. Z tego powodu zwykle umieszczam otwory z tyłu obudowy. Kolejną zaletą głośnika w obudowie zamkniętej jest to, że łatwiej jest go zintegrować z subwooferami. W przypadku konstrukcji typu reflex potrzebna byłaby zwrotnica szóstego rzędu, która ma ogromne opóźnienie grupowe dokładnie tam, gdzie nie jest ono pożądane.

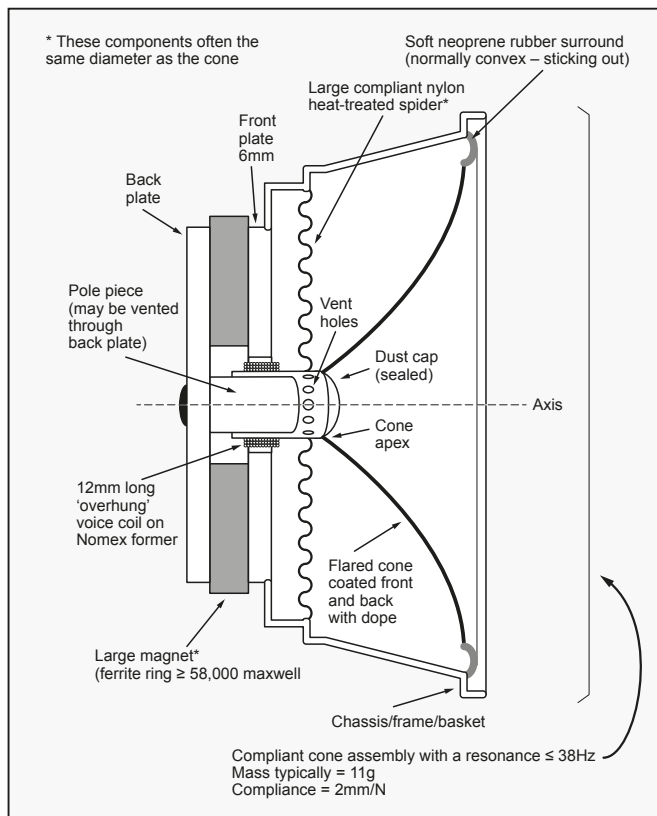
Dlaczego dzisiejsze głośniki nie są tak dobre jak stare

Skuteczną techniką tłumienia rezonansów membrany jest pomalowanie membrany miękką żywicą PVA, taką jak Plastiflex P1200 (klej introligatorski). Nazywa się to „dopingiem” i nadaje membranie błyszczący, «mokry» wygląd. Trudno jest jednak sprawić, aby żywica ta przylegała do membrany z tworzywa sztucznego. Firma KEF nakładała najpierw pośrednią warstwę kleju Bostik 4141, a dopiero na nią nakładała drugą warstwę Plastiflexu.



Rysunek 23. Driver głośnika Goodmans typu „Infinite baffle” – pierwszego komercyjnego głośnika w obudowie zamkniętej, wprowadzonego na rynek w 1940 roku

Na przedniej części membrany nakładano tylko jedną warstwę Plastiflexu, która czasami odpada w okolicy osłony. Doping nie jest obecnie zbyt często stosowany, ponieważ głośnik staje się dużo cichszy, co jest niedopuszczalne w dzisiejszym świecie komercyjnym zorientowanym na decybele. Doping jest także praktyczny, ponieważ zazwyczaj nakłada się go ręcznie, co z kolei powoduje różnice między



Rysunek 24. Ogólna konstrukcja modułu nisko-średniotonowego odnosząca się do mini monitora. Należy zwrócić uwagę na stosunkowo duży magnes i zawieszenie („magnet”, „spider”), a także długą cewkę drgającą („over_hung voice coil”)



Rysunek 25. Pierwszy prawdziwy głośnik nisko-średniotonowy mini monitor zastosowany w modelu Goodmans Maxim z 1965 roku. Posiadał on zawieszenie membrany wykonane z tkaniny impregnowanej lateksem, w celu uzyskania odpowiedniej elastyczności. W późniejszej produkcji zmieniono je na zawieszenie gumowe. Sam kiedyś wykonałem podobne dla głośnika nisko-tonowego z płótna i kleju lateksowego Copydex



Rysunek 26. Po lewej – widok z przodu głośnika Goodmans Maxim (dzięki uprzejmości Bernarda Pattersona, wiernego czytelnika Practical Electronics); po prawej – wzór tkaniny maskownicy z lat 60.



poszczególnymi przetwornikami. Kiedy pracowałem w Castle Acoustics w 1990 roku zautomatyzowaliśmy ten proces, nakładając odmierzoną ilość środka i obracając membranę. Wygląda na to, że tylko Falcon, Spendor i ATC wciąż stosują doping. A szkoda, bo w ten sposób pozbywamy się większości tej „kubkowej, krzykliwej” jakości głośnika. Ten sam trend utraty rzemiosła dotyczy zresztą obudów. W dobrej jakości konstrukcjach z lat 70. i 80.



Rysunek 27. Głośnik niskotonowy B110A firmy KEF



Rysunek 28. Zespół membrany B110A. Należy zwrócić uwagę na otwory w cewce z Nomexu, aby zapobiec wzrostowi ciśnienia pod osłoną przeciwpływową (co jest częstą przyczyną zniszczeń)

normą było nakładanie bitumicznych podkładek tłumiących o grubości pół cala na wewnętrzne strony paneli, co było okropną pracą z użyciem lepkiego kleju dekarckiego Aquaseal No.5. Zmniejszyło to podkoloryzowanie dźwięku wydobywającego się z obudowy (ścianki obudowy również mają tryby rezonansowe). W konstrukcji LS3/5A przewidziano podkładki tłumiące na skrzynce wykonanej z 12 mm sklejk brzoazowej z listwami z twardego drewna bukowego wzdłuż każdej wewnętrznej krawędzi. Rozmiary płyt zostały dobrane tak, aby rezonanse nie pokrywały się ze sobą. Nowoczesne obudowy głośników są wykonywane z płyty MDF, ponieważ jest ona łatwa w obróbce i dobrze wygląda w sprzedaży. Nie brzmi ona jednak tak dobrze, ponieważ jest sztywna, ciężka i trudna do wytłumienia. Domowy konstruktor nie jest ograniczony takimi nowoczesnymi naciskami komercyjnymi, więc śmiało można wytłumiać obudowy i zalewać membrany. Jak to zrobić, pokażę później, gdy będziemy składać własną obudowę w stylu LS3/5A.

Originalne i najlepsze

Po ponad dziesięcioletniej przerwie w produkcji, w 2015 roku firma Falcon Acoustics wznowiła wytwarzanie czcigodnych, 40-letnich LS3/5A, a nawet przeniosła produkcję oryginalnych jednostek napędowych KEF do własnego zakładu. Założyciel firmy Falcon Acoustics, Malcolm Jones, wciąż jest w nią zaangażowany. Spotkałem go kilka razy na przestrzeni lat i mogę ręczyć za jego inżynierską uczciwość. Niestety, nowa para LS3/5A kosztuje dziś około 2300 funtów, co oznacza wzrost z około 150



Rysunek 29. Widok z przodu/z tyłu na tweeter KEF T27



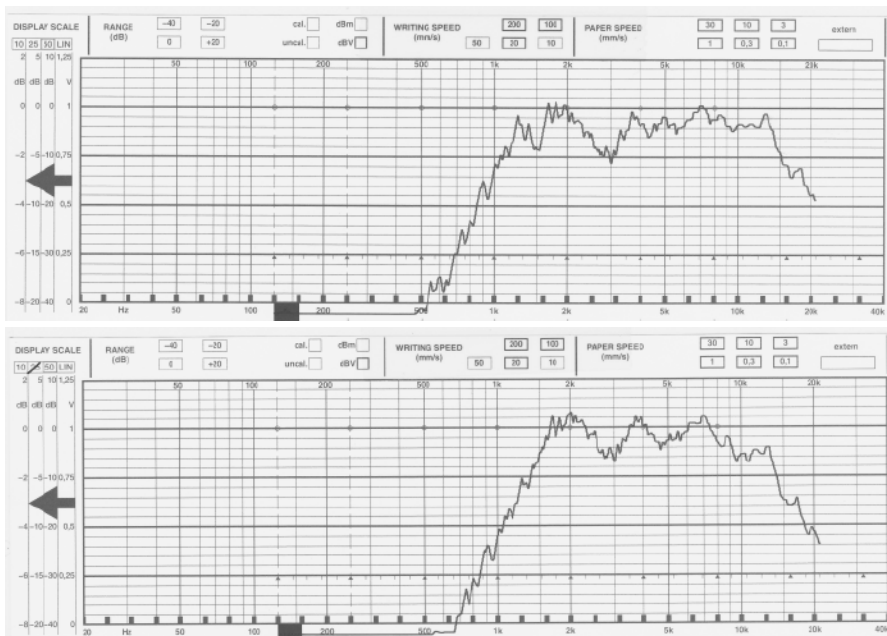
Rysunek 30. Maskownica membrany T27

funtów od czasu gdy pojawiły się po raz pierwszy w 1975 roku – inflacja w branży głośnikowej jest więc niemal dwukrotnie wyższa niż w całej gospodarce! Wygląda to jeszcze drożej, ponieważ względna cena elementów elektronicznych na bazie krzemu znacznie spadła w ujęciu realnym. Koszt części elektromechanicznych, które stanowią podstawę ręcznie budowanych głośników podąża za ogólnym wzrostem cen towarów, więc prawo Moore'a nie ma tu zastosowania. Istnieje także „efekt inflacji mody retro z eBay'a”, który spowodował wzrost cen analogowych automatów perkusyjnych i syntezatorów. Z drugiej strony, to właśnie ten nowy entuzjazm dla starego, dobrego sprzętu analogowego sprawił, że ponowne wprowadzenie tych klasycznych konstrukcji na rynek stało się opłacalne. Falcon sprzedaje jednostki napędowe, zwrotnice i większość części, więc możliwe jest zbudowanie własnych LS3/5A za około

850 funtów. Aby było jasne, warunki licencji od BBC uniemożliwiają im oferowanie „kompletnych zestawów”, ale sprzedają wszystkie potrzebne części oddzielnie. Swego czasu technicy BBC mogli kupować takie zestawy a mój wujek zbudował ich parę. Po jego śmierci zostały szybko wyciągnięte ze sklepu charytatywnego. W ciągu ostatnich 30 lat sam też zbudowałem kilka takich zestawów. Można zauważyć, że jak na standardy głośnikowe, konstrukcja wewnętrzna typowego głośnika LS3/5A jest dość skomplikowana, jak pokazano na rysunku 21. Nietypowo, jednostka niskotonowa jest zamontowana za przegrodą. Zwykle zaleca się montowanie głośników od przodu, aby uniknąć efektów rezonansu wnęki, chociaż nadal możliwe jest uzyskanie efektów wnęki tylnej przy grubej przegrodzie przedniej. Rysunek 22 pokazuje, że występuje mieralne pogorszenie pasma przenoszenia, ale zostało ono uwzględnione w zwrotnicy. Podejrzewam, że zrobiono to po to, aby głośnik wysokotonowy i niskotonowy znalazły się jak najbliżej siebie. W przeciwnym razie ich ramki nachodziłyby na siebie.

Trochę historii

W latach 60. ubiegłego wieku istniało tylko kilka udanych komercyjnych mini monitorów, Goodmans Maxim i Celestion Ditton 10 (rysunek 2). Maxim został zaprojektowany przez doświadczonych konstruktorów głośników Teda Jordana i Lauri Finchama. Firma Goodmans była pionierem w dziedzinie głośników niskotonowych do zamkniętych skrzynek basowych i reklamowała głośniki typu „Infinite baffle” w numerze Wireless World z kwietnia 1940 roku (rysunek 23). W czterocalowym głośniku niskotonowym Maxim zastosowano magnes Alnico, który Ted Jordan zaprojektował w 1962 roku dla 8-calowego, pełnozakresowego głośnika Axiette. Urządzenie Maxim definiowało wymagania konstrukcyjne stawiane pierwszemu prawdziwemu mini-woofrowi, który z powodzeniem pracował w małej obudowie zamkniętej. Jest on przedstawiony na rysunku 24. Oryginalną jednostkę basową Maxim pokazano na rysunku 25. BBC początkowo doceniło Maxim i uznało go za najlepszy,



Rysunek 31. Wpływ maskownicy na charakterystykę częstotliwościową głośnika T27 (podłączzonego bez filtra zwrotnicy). U góry – bez maskownicy; u dołu – z maskownicą, podnoszącą pasmo przenoszenia w zakresie od 9 do 13 kHz i zapewniającą użyteczny efekt wygładzenia

ale nie dość dobry. Prawdopodobnie zawiódł ich stożkowy głośnik wysokotonowy (patrz rysunek 26). Raport z 1965 roku (<https://bbc.in/2Ghzn5s>) jest obiektywną odskocznią od new age’owego subiektywizmu dzisiejszej prasy Hi-Fi. Na tej podstawie prace działu badawczo-rozwojowego BBC doprowadziły do stworzenia w 1975 roku słynnego mini monitora BBC LS3/5A. Jest on opisany w raporcie RD 1976/29 (<https://bbc.in/2y33SYJ>). BBC wybrało głośnik niskotonowy B110 i głośnik wysokotonowy T27 zaprojektowane przez Malcoma Jonesa z firmy KEF. W 1966 roku były to najnowocześniejsze urządzenia z plastikowymi membranami. W modelu B110 zastosowano Bextrene, mieszanke polistyrenu i neoprenu opracowaną na potrzeby wnętrz samochodów, materiał po raz pierwszy wykorzystany przez BBC. Membranę B110A pokazano na rysunku 27, a zespół membran B110A jest pokazany na rysunku 28. W głośniku wysokotonowym zastosowano folię poliestrową firmy ICI o nazwie Melinex. Zdjęcie głośnika T27 pokazano na rysunku 29.

W LS3/5A była ona przykryta osłoną kopułki z głośnika wysokotonowego Celestion HF2000, pokazaną na rysunku 30, a jej korzystny wpływ pokazano na rysunku 31. W modelu T27 we wgłębieniu w nabiegunkniku umieszczano filtr z papierosów, który pochłaniał promieniowanie tylne. Firma KEF wyprodukowała własny system głośnikowy o nazwie Coda Mk1 wykorzystujący te głośniki, który miał konkurować z systemem Ditton 10. Nie wykorzystywał on w pełni potencjału głośników, ponieważ posiadał prostą czteroelementową zwrotnicę drugiego rzędu i obudowę z płyty wiórowej.

W przyszłym miesiącu

W następnym odcinku przyjrzymy się dokładnie zwrotnicom LS3/5A i rozpoczniemy rozważania o opcji aktywnej i o wyborze driver’ów. ■

Jake Rothman

Ten artykuł był wcześniej opublikowany na łamach „Everyday Practical Electronics”, wrzesień 2019 (www.epemag3.com)

ERRATA

1. Artykuł „Dwukierunkowy interkom” (EdW05/2022 str. 84, 85)

Na schemacie elektrycznym (rysunek 2) oraz na schemacie montażowym (rysunek 4) odwrotnie zaznaczono polaryzacje kondensatorów elektrolitycznych C1, C2 oraz C7 i C8.

2. Artykuł „Generator sygnałów i inwerter wykorzystujący timery NE555” (EdW05/2022 str. 80)

Na schemacie elektrycznym (rysunek 2) oraz na schemacie montażowym (rysunek 6) błędnie podano wartości rezystorów R14, R20. Powinno być 2 Ω.

W tym samym artykule został przedstawiony nawias we wzorze na stronie 79. Prawidłowy zapis tego wzoru:

$$F = \frac{1}{(0,7 \cdot (R7 + 2 \cdot (R8 + VR2)) \cdot Cx)}$$

W tym samym artykule w zdaniu „Oznacza to, że częstotliwość wyjściowa IC2 jest dzielona...”, tak samo jak w zdaniu poprzednim i następnym chodzi o podział amplitudy sygnału.

Elektrozawory nie tylko do nawadniania, część 2

Po omówieniu w pierwszej części elektrozaworów kulowych omówimy teraz...

Elektrozawory membranowe

W domowych systemach nawadniania powszechnie stosuje się stosunkowo duże elektrozawory 1-calowe (DN25). Kosztują one poniżej 100 zł i mają specyficzną, interesującą budowę. **Rysunek 1**, pochodzący z materiałów Bermad, pokazuje wewnętrzną budowę najpopularniejszej odmiany elektrozaworów membranowych.

Z kolei **fotografia 2** pokazuje popularny elektrozawór o takiej budowie. Solenoid – cewka nie kontroluje tu bezpośrednio głównej drogi przepływu wody.

Miedzy górną i dolną częścią membrany jest możliwość niewielkiego przepływu wody, co pozwala na powolne wyrównanie ciśnienia z dwóch stron membrany. Dlatego w stanie spoczynku z obu stron membrany panuje to samo ciśnienie i ta elastyczna membrana pozostaje zamknięta – jest dociskana do gniazda sprężyną.

Elektryczne pobudzenie cewki – elektromagnesu otwiera tylko niewielki kanał przepływu wody z górnej części nad membraną do wylotu. To powoduje zmniejszenie ciśnienia w przestrzeni nad membraną, bo wypływ wody z obszaru nad membraną jest szybszy niż wspomniane wcześniej wyrównywanie ciśnienia. A jeżeli ciśnienie nad membraną się zmniejsza, to siła ciśnienia wody pod membraną, na wejściu elektrozaworu, powoduje przezwyciężenie siły sprężyny dociskającej i membrana zostanie podniesiona, umożliwiając przepływ wody w głównym kanale.

Trwa to przez czas, gdy solenoid przepuszcza wodę przez mały kanałek pomocniczy. Gdy elektromagnes zostaje wyłączony i ten przepływ ustaje, ciśnienie wody nad membraną zaczyna wzrastać i w ciągu około sekundy wzrośnie na tyle, że sprężyna znów docisnie membranę do gniazda i zamknie główny kanał.

To tylko nieco uproszczony opis, a w rzeczywistości stosowane są nieco inne rozwiązania, w szczególności zapewniające też regulację wielkości przepływu.

Z punktu widzenia elektronika najważniejsze jest to, że elektromagnes steruje przepływem wody tylko w niewielkim kanale pomocniczym, co pozwala zmniejszyć ilość energii potrzebnej do sterowania. A jeszcze bardziej zmniejsza ilość tej energii zastosowanie elektromagnesu bistabilnego, który jest sterowany krótkimi impulsami załącz/wyłącz.

W mniejszych sterownikach nawadniania zasilanych bateryjnie stosowane są właśnie elektromagnesy bistabilne. W większych, obsługujących kilka linii – sekcji, powszechnie wykorzystywane są cewki – elektromagnesy na prąd przemienny 24 V 50 Hz.

Wielu elektronikom może się wydawać, że lepsze byłyby elektromagnesy na prąd stały, a zastosowanie elektromagnesów na prąd zmienny i to na napięcie 24 V jest anachronizmem i skutkiem trzymania się archaicznych rozwiązań. To błędne wyobrażenie!



2

Praca elektromagnesów przy prądzie zmiennym jest celowa i uzasadniona. Elektromagnesy takie mogą pracować przy prądzie stałym, ale jest to zdecydowanie bardziej kłopotliwe. Zasilanie napięciem stałym jest możliwe i nietrudne do realizacji, jednak należy dobrze zrozumieć wchodzące w grę ograniczenia.

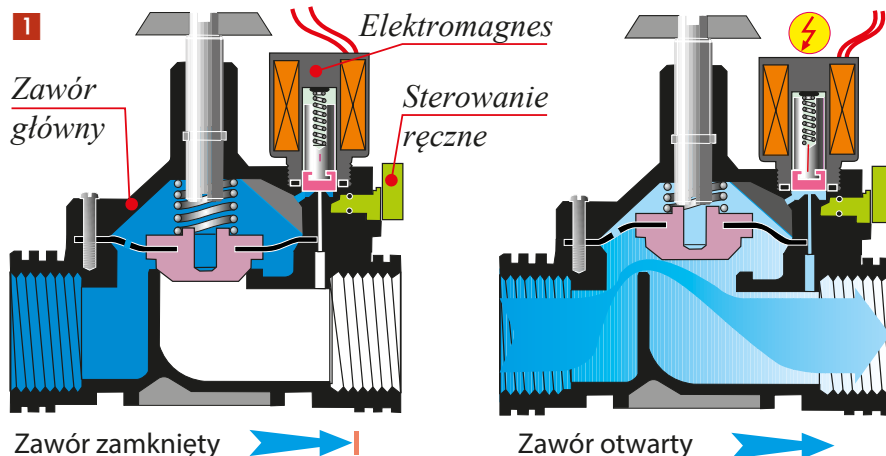
AC czy DC?

Popularne cewki – elektromagnesy 24 VAC amerykańskiej firmy Hunter (**fotografia 3**) mają rezystancję około 24 Ω i znaczną indukcyjność.

Współczesny elektronik bardzo rzadko ma do czynienia z elektromagnesami prądu zmiennego, dlatego trzeba przypomnieć, że ich prąd pracy jest ograniczony nie tylko przez rezystancję R uzwojenia, ale też przez reaktancję indukcyjną X_L uzwojenia cewki. My jako elektronicy mamy do czynienia głównie z cewkami o indukcyjności rzędu mikrohenrów, a w omawianych elektrozaworach pracują elektromagnesy zawierające wiele zwojów, przez co ich indukcyjność jest niewiele mniejsza od 1 henra. Nawet przy częstotliwości sieci 50 Hz daje to znaczną reaktancję indukcyjną.

Rysunek 4 pochodzi ze strony producenta i wskazuje, że po dołączeniu napięcia (24 VAC) w pierwszej chwili prąd rozruchu wynosi 370 mA, a po chwili zmniejszy się on zdecydowanie, do spoczynkowej wartości tylko 210 mA.

Na początku przez chwilę prąd jest większy, co jest typowe dla elektromagnesów.





3



POLSKI

Hunter®

Built on Innovation®

Dane techniczne cewki elektromagnetycznej

cewka 24 V

prąd rozruchowy 350 mA, prąd podtrzymania 190 mA, 60 Hz
prąd rozruchowy 370 mA, prąd podtrzymania 210 mA, 50 Hz

4

Gdy elektromagnes zadziała i przyciągnie (wciągnie) ferromagnetyczną kotwicę, tym samym zwiększy indukcyjność, a tym samym jego reaktancję, co zmniejszy prąd.

Prąd o finalnej wartości 0,21 A płynie przy napięciu zmiennym 24 V wtedy, gdy obciążenie ma impedancję 114 omów. W skład tej impedancji wchodzi szeregowo połączone rezystancja drutu 23,5 oma oraz reaktancja indukcyjna uzwojenia. Z uwagi na wektorowe dodawanie oporności można od razu pominąć dużo mniejszą rezystancję, ale kto chciałby przeprowadzić dokładniejsze obliczenia, dowie się, że reaktancja wynosi 112 Ω. Taką reaktancję ma przy częstotliwości 50 Hz indukcyjność 0,356 henra.

Zasilanie prądem zmiennym ma istotne zalety. Siła elektromagnesu jest wyznaczona przez wartość prądu, a ten zaraz po włączeniu ma prawie 0,4 A. Daje to znaczną siłę i pewność włączania. Z kolei po zadziałaniu, podczas długotrwałej pracy, duża siła nie jest już potrzebna i wtedy samoczynnie prąd się zmniejsza do wartości 0,21 A. Ta wartość prądu z dużym zapasem wystarcza do utrzymania przyciągniętej kotwicy solenoidu.

W stanie ustalonym przy napięciu 24 VAC i prądzie 0,21 A moc wynosi 5 W (24 V·0,21 A), ale jest to moc pozorna, z której większość to moc bierna, która krąży między źródłem zasilania a cewką. Korzystne jest to, że w uzwojeniu występują bardzo małe straty ciepłne. Mianowicie straty ciepłne (moc czynna) występują tylko na rezystancji drutu ($P=I^2 \cdot R$). Dla prądu 0,21 A i rezystancji około 24 Ω daje to niecałe pół wata: $P=(0,21)^2 \cdot 24=0,48$ W. Przy tak małej mocy strat elektromagnes praktycznie się nie grzeje.

Nieporównanie gorzej jest przy prądzie stałym. Aby uzyskać odpowiednio dużą siłę w pierwszej chwili po włączeniu, powinniśmy zapewnić prąd około 370 mA. W przypadku cewki 24-omowej uzyskamy go przy napięciu stałym bliskim 9 V. Czyli podczas włączania napięcie stałe nie może być znacząco mniejsze niż 9 V, a dla pewności działania powinno być trochę większe.

Jeżeli nie chcemy komplikować obwodów sterowania, to samo napięcie stałe, rzędu 9...12 V, będziemy podawać na cewkę elektromagnesu przez cały czas jego pracy. Czyli w przeciwieństwie do sytuacji przy prądzie zmiennym, teraz cały czas przez cewkę będzie płynął jednakowy prąd, wyznaczony przez rezystancję drutu. Zapewni on po zadziałaniu siłę przyciągania wielokrotnie większą, niż potrzeba, ale co najważniejsze, spowoduje wydzielanie się dużej mocy strat. Przy napięciu zasilania 12 VDC prąd wyniesie 0,5 A, czyli moc strat sięgnie 6 W! Dwanaście razy więcej niż przy zasilaniu 24 VAC!

Nawet przy zasilaniu minimalnym napięciem 9 V moc strat w uzwojeniu wyniesie 3,4 W.

To są dane dotyczące cewki firmy Hunter. Inni wytwórcy oferują elektromagnesy – cewki o innych parametrach, co znacząco zmienia tego rodzaju wyliczenia.

Przykładowo jeżeli cewka nawinięta jest grubszym drutem i ma mniejszą rezystancję, to przy prądzie zmiennym straty ciepłne są jeszcze mniejsze. Z kolei nietypowe zasilanie takiej cewki prądem stałym stwarza jeszcze większe problemy, bo łatwo ją przegrzać. To co jest dobre przy prądzie zmiennym, jest problemem przy prądzie stałym.

Z kolei, jeżeli rezystancja uzwojenia jest większa (cienki drut), to przy prądzie zmiennym straty ciepłne są trochę większe, ale może się okazać, że łatwiejsze jest jej wykorzystanie przy prądzie stałym.

Dostępne w Internecie informacje o zasilaniu takich elektromagnesów napięciem stałym w większości pisane są przez osoby nierozumiejące problemu. Dlatego informacje takie oraz podawane tam wskazówki są niepełne, fragmentaryczne i nie zawsze wiarygodne.

W każdym razie kluczowym problemem jest właśnie moc strat ciepłnych ($P=I^2 \cdot R$), powodująca grzanie cewki. Nie chodzi o oszczędność mocy, tylko o ryzyko przegrzania cewki i trwałego uszkodzenia izolacji drutu.

Najprościej biorąc, **cewka elektrozaworu może pracować przy napięciu stałym, ale podczas pracy drut nie może osiągnąć temperatury, która mogłaby uszkodzić izolację.** Przy prostych testach trzeba pamiętać, że temperatura drutu uzwojenia zawsze jest znacząco wyższa od temperatury obudowy tego elektromagnesu.

Nie ma jednak prostej recepty, jakim napięciem stałym można zasilac cewki elektrozaworów 24 VAC. Na pewno napięcie stałe musi być niższe niż 24 VDC. Nie może być za duże, żeby nie przegrzać uzwojenia podczas długotrwałej pracy.

Ale nie może być za małe! Nawet jeżeli cewka podczas testów na biurku zadziała przy napięciu znacznie niższym niż 10 V, to w realnym systemie z wodą pod ciśnieniem wymagane minimalne napięcie i prąd załączenia będą znacząco wyższe niż podczas „suchych” testów na biurku.

Może się też okazać, że w elektrozaworze opory zadziałania będą rosły wraz z upływem czasu. Gdyby dobrany został prąd pracy niewiele większy niż prąd minimalny, to z czasem elektrozawór może przestać działać. Na pewno dobierając napięcie stałe i prąd pracy, trzeba uwzględnić znaczący zapas, żeby elektrozawór zadziałał także po zwiększeniu oporów ruchu. Ale zwiększenie napięcia stałego i prądu pracy zwiększy moc strat i temperaturę.

Informacje, dość łatwe do znalezienia w Internecie, wskazują, że przynajmniej w przypadku niektórych elektrozaworów 24 VAC możliwe jest dobranie sensownego stałego napięcia zasilania, które zapewni zarówno pewność włączania, jak i nie doprowadzi do przegrzania.

Ale w przypadku niektórych odmian o małej rezystancji może się to okazać trudne lub niemożliwe. Wtedy trzeba wykorzystać inne możliwości. W grę wchodzi też inne niebezpieczeństwo – trwałe namagnesowanie rdzenia.

Tymi kwestiami zajmiemy się dokładniej w następnej części artykułu. ■

Piotr Górecki

Współczesne akumulatory

5. Akumulatory litowe – historia

W kilkuodcinkowym artykule przedstawione są informacje o właściwościach stosowanych dziś akumulatorów, w tym wskazówki dotyczące ich prawidłowej obsługi, zapewniających długą żywotność oraz bezpieczeństwo.

Akumulatory litowe znajdziemy w praktycznie wszystkich współczesnych urządzeniach przenośnych, dlatego warto poznać także ich historię.

Trochę historii

Eksperymenty z bateriami litowymi podjęto już ponad 100 lat temu(!), w roku 1912 (G.N. Lewis), jednak dopiero na początku lat 70. pojawiły się na rynku pierwsze *jednorazowe baterie litowe*. Natomiast *akumulatory litowe* weszły na rynek dopiero w roku 1991, jako źródło zasilania kamery Sony CCD TR1. Eksperymentalne *akumulatory litowe* realizowano wprawdzie już od początku lat 70., jednak kluczowym praktycznym problemem była i nadal jest duża aktywność chemiczna metalicznego litu, który w razie przeładowania czy rozszczelnienia może spowodować pożar. W prototypowych akumulatorach, gdzie jedną z elektrod był metaliczny lit, tworzyły się tzw. dendryty – cienkie włókna metalowe, które potrafiły zewrzeć baterię, a nawet doprowadzić do jej pożaru. Dopiero zastosowanie zamiast metalicznego litu jego związków, które mogły łatwo oddawać lub przyjmować jony litu, otworzyło drogę do realizacji bezpiecznych dla użytkownika akumulatorów.

Dziś mamy do dyspozycji szereg rodzajów, odmian i wykonanń akumulatorów litowych (jednorazowe baterie pomijamy). Łatwo się pogubić w szczegółach. Koniecznie trzeba jednak wiedzieć, że **podstawowa zasada działania wszystkich dostępnych akumulatorów litowych jest taka sama.**

Dziś dostępnych jest wiele odmian, które choć noszą różne nazwy, **wszystkie są akumulatorami litowo-jonowymi.**

Podstawą magazynowania energii są dodatnie jony litu i ich ruch między anodą i katodą w przewodzącym elektrolicie, co związane jest z przemianami chemicznymi. Napięcie ogniwa zależnie jest od użytych związków chemicznych.

Dodatnia elektroda (katoda) wykonana może być z różnych związków litu. Nie jest zbudowana z metalicznego litu, tylko

z takich związków litu, które łatwo mogą oddawać i przyjmować jony litu. Natomiast elektroda ujemna (anoda) zbudowana jest najczęściej z jakiejś odmiany węgla. Elektrody muszą być od siebie odseparowane, żeby nie nastąpiło zwarcie, natomiast między nimi musi być umieszczony taki elektrolit, który pozwoli na przemieszczanie się dodatnich jonów litu między elektrodami.

Ponieważ podstawą działania jest wymiana jonów litu między elektrodami, badano najróżniejsze związki litu, by wykorzystać takie, które łatwo oddają i przyjmują te jony. Po latach badań okazało się, że do budowy katody nadają się związki litu, mające specyficzną budowę siatki krystalicznej. **Rysunek 1** pokazuje zasadę budowy akumulatorów litowo-jonowych.

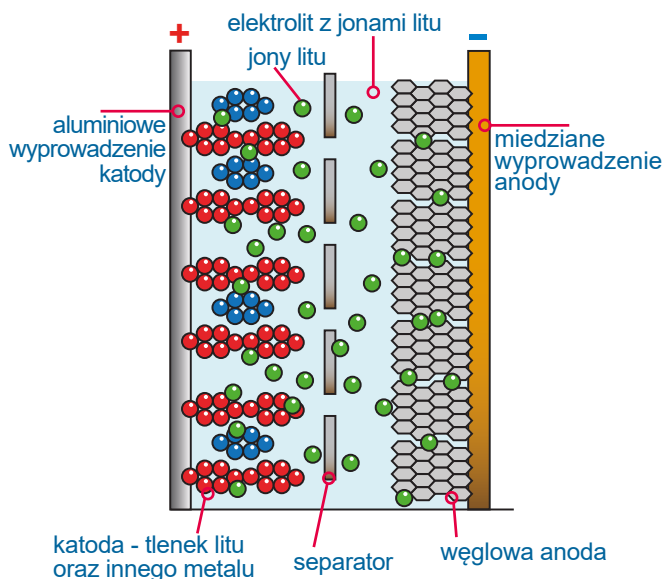
Naukowcy zbadali dziesiątki i setki różnych związków litu (a także licznych innych pierwiastków). Najlepsze do budowy katod okazały się tlenki metali, ale nie tlenki jednego metalu (litu), tylko dwóch, a nawet trzech. Odpowiednią przestrzenną strukturę krystaliczną tworzą wiązania atomów tlenu z mieszkanką atomów metali: na pewno litu, a do tego innego metalu lub częściej metali.

Co ważne, ubytek jonów litu podczas ładowania nie niszczy struktury krystalicznej katody, a jedynie zmienia jej skład chemiczny. Po pierwsze zazwyczaj nie wszystkie atomy litu są uwalniane z katody (co akurat nie jest zaletą), po drugie atomy pozostałych pierwiastków nie biorą udziału w reakcji i nadal zachowują wcześniejszą strukturę krystaliczną materiału.

W pierwszych praktycznie użytecznych akumulatorach litowych (Sony 1991) katoda zbudowana była z tlenku litu i kobaltu – LiCoO_2 . Do dziś akumulatory litowe zawierające kobalt są bardzo popularne. Węglowa elektroda ujemna na początku budowana była z... koksu, który ma porowatą strukturę, a obecnie wykonywana jest z odpowiednio preparowanego grafitu.

Podczas ładowania dodatnie jony litu przechodzą z katody do węglowej anody, gdzie łączą się z atomami węgla w specyficzną warstwową strukturę opisywaną wzorem chemicznym LiC_6 . Podczas rozładowania jony litu wracają z anody do katody. Ilustruje to **rysunek 2**.

Akumulatory zawierające kobalt mają dobrą zdolność magazynowania energii i inne cenne właściwości, jednak problemem są wysokie ceny kobaltu i pewne trudności technologiczne. Dlatego od lat eksperymentowano z substancjami, głównie tlenkami o podobnej budowie krystalicznej, ale zawierającymi inny metal zamiast kobaltu. Dość szybko



1

okazało się, iż porównywalne właściwości ma zdecydowanie łatwiej dostępny mangan w związku LiMn_2O_4 .

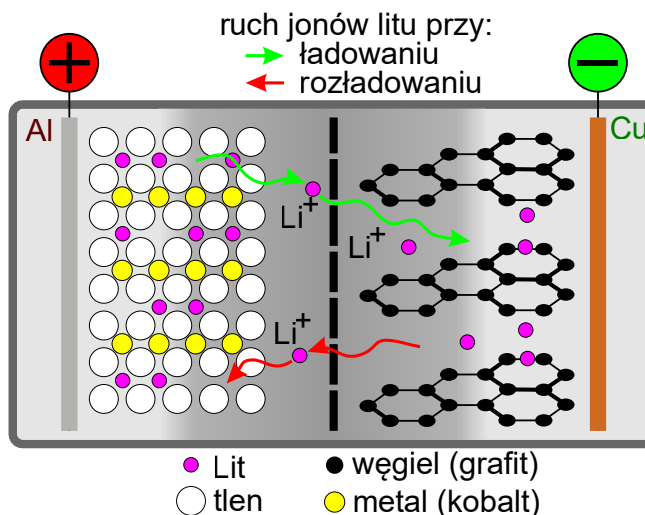
W wersji akumulatorów z LiCoO_2 oznaczane bywają LCO. Opracowane trochę później akumulatory litowo-jonowe z katodą z LiMn_2O_4 , oznaczane są LMO i nazywane manganowymi lub spinelowymi. Zasadniczą ich wadą jest mała gęstość energii, o połowę mniejsza od kobaltowych, a istotną zaletą – możliwość oddawania dużych prądów. Z czasem okazało się, że korzystniejsze niektóre cechy mają kryształy tlenku, zawierające atomy litu oraz manganu i kobaltu LiMnCoO_2 (oznaczenie LMC), a jeszcze korzystniejsze: niklu, manganu i kobaltu LiNiMnCoO_2 (oznaczenie NMC).

Prowadzono też badania nad różnymi innymi związkami litu, w tym zawierającymi tanie i popularne żelazo. Niestety, bodaj najprostszy pokrewny związek LiFeO_2 słabo oddaje i przyjmuje jony litu. Wnikliwe badania wykazały, iż dużo lepsze parametry oferują związki o budowie krystalicznej typu oliwiny, a konkretnie LiFePO_4 , czyli tlenek żelaza zawierający też fosfor.

Zasada działania akumulatora „żelazowego” jest taka sama jak we wcześniejszych rozwiązaniach: kluczowe znaczenie ma krystaliczna budowa LiFePO_4 , który podczas ładowania chętnie oddaje jony litu i pozostaje jako fosforan żelaza FePO_4 o takiej samej budowie krystalicznej. **Akumulatory żelazowo-fosforanowe, zwane też fosfatowymi** (LFP – LiFePO_4) pojawiły się na rynku dość późno, po przezwycięzeniu szeregu problemów technologicznych. Są to też jak najbardziej akumulatory litowo-jonowe. Mają napięcie nominalne 3,2 V, natomiast wcześniej omawiane Li-Ion i Li-Po (LCO, LMO, LMC, NMC) mają napięcia nominalne w zakresie 3,6...3,8 V.

We wszystkich katoda jest zbudowana z materiału łatwo oddającego (ładowanie) i przyjmującego (rozładowanie) jony litu, a anoda zawiera węgiel. Trwają też intensywne badania nad innymi związkami litu. Dostępne są informacje o akumulatorach litowo-jonowych zawierających glin (NCA – LiNiCoAlO_2), tytan (LTO – $\text{Li}_4\text{Ti}_5\text{O}_{12}$), gdzie zaletą jest możliwość bardzo szybkiego ładowania, czy siarkę (LIS – Li_2S_8), gdzie zaletą jest teoretyczna gęstość magazynowanej energii do 500 Wh/kg. Jednak oprócz zalet mają też istotne wady, dlatego na razie, w wersji dostępnych na rynku największą gęstość energii, do 200 Wh/kg, oferują wersje kobaltowe (LiCoO_2).

I jeszcze jeden ważny szczegół. W akumulatorach litowych typu LCO, LMO, LMC stosowano i nadal stosuje się ciekłe elektrolity o różnym składzie, w tym wodne roztwory różnych bardziej i mniej bezpiecznych substancji zawierających lit, jak choćby LiPF_6 ,



2

czyli sześćfluorofosforan litu, czy LiBOB , którego polska nazwa to bis(szczawiano) boran litu.

Istotnym wynalazkiem okazało się zastąpienie cieczy stałym elektrolitem w postaci przewodzących polimerów, zawierających sole litu. Tak powstały **akumulatory litowo-polimerowe**, oznaczane Li-Po, LiPo lub LIP. Są to też akumulatory litowo-jonowe, zawierające kobalt lub mangan, gdzie ciekły elektrolit zastąpiono stałym. Zasadniczo polimerowy elektrolit niekorzystnie zwiększył rezystancję wewnętrzną, ale polepszył bezpieczeństwo przez eliminację możliwości pożaru.

Tu warto wspomnieć o pewnych niejasnościach. Otóż skrót **Li-Ion** obejmuje wszystkie akumulatory litowe, ponieważ wszystkie są litowo-jonowe. Jednak w praktyce często określenie Li-Ion odnosi się do akumulatorów z ciekłym elektrolitem, natomiast skrót **Li-Po** lub **LiPo** oznacza akumulatory z elektrolitem stałym, polimerowym.

Ponadto niedostateczna świadomość użytkowników i względy marketingowe spowodowały, że polimerowymi nazywa się też akumulatory z ciekłym elektrolitem, mające kształt prostokątnych poduszek z tworzywa. Dla niektórych wystarczającym powodem, by mówić o akumulatorach polimerowych, była też obecność w środku separatora – porowatej folii z tworzywa sztucznego, uznanego za polimer. Ponieważ wcześniejsze wersje miały metalową obudowę, dla niektórych takim powodem użycia określenia Li-Po była obecność miękkiej, plastikowej obudowy. W rezultacie określenia Li-Po (LiPo) spotykane w materiałach reklamowych, a nawet artykułach technicznych, często są nieprecyzyjne.

Pomimo tego rodzaju niejasności podstawowe zasady są proste: akumulatory polimerowe (Li-Po) to też akumulatory litowo-jonowe (Li-Ion), zwykle zawierające kobalt lub mangan i ich podstawowe parametry, w tym parametry

ładowania, są podobne, jak klasycznych z ciekłym elektrolitem.

Obecnie na rynku dominują akumulatory litowo-jonowe LCO, LMO, LMC, NMC, czyli zawierające kobalt, mangan i nikiel, w wykonaniach klasycznych z ciekłym elektrolitem i ze stałym elektrolitem polimerowym. Coraz większą popularność zyskują akumulatory LFP, czyli żelazowo-fosforanowe, oznaczane też LiFePO_4 (LiFePO_4). W ofercie rynkowej można też znaleźć ulepszone akumulatory fosforanowe, domieszkowane itrem, oznaczane LiFeYPO_4 .

Podsumowanie

W krótkim podsumowaniu podanych informacji można powiedzieć: wszystkie akumulatory litowe działają na tej samej zasadzie.

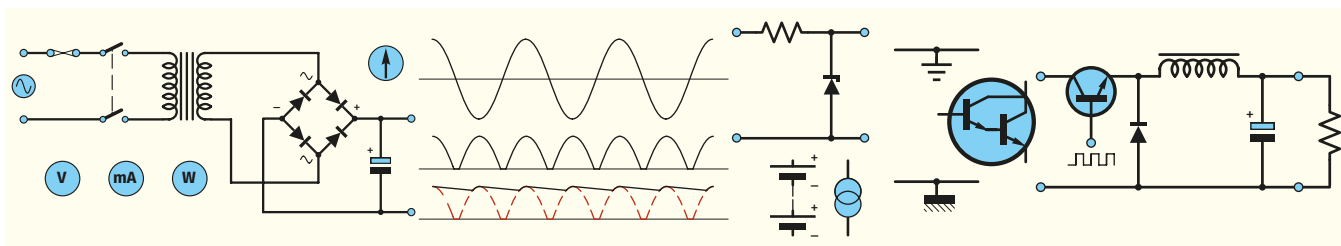
Różnią się po pierwsze elektrolitem, który pełni funkcję pośredniczącą, bowiem umożliwia przemieszczanie jonów litu między elektrodami akumulatora podczas ładowania i rozładowania. Zwykle elektrolit jest ciekły. Część akumulatorów ma elektrolit polimerowy w postaci pasty.

Po drugie różne odmiany akumulatorów litowych różnią się materiałem czynnym anody, która oprócz litu zawiera inne pierwiastki, które nie biorą udziału w reakcjach chemicznych, a jedynie są swego rodzaju ramą, szkieletem anody. Zależnie od składu chemicznego anody, akumulatory mogą mieć nieco inne napięcie i zakres dopuszczalnych prądów.

Akumulatory litowe pozwalają uzyskać dużą pojemność, trwałość oraz bezpieczeństwo, jednak zawsze jest to wynikiem po pierwsze pewnych kompromisów, po drugie wymaga zaawansowanej technologii produkcji. Po trzecie: w kwestiach ładowania i rozładowania konieczne jest ściśle trzymanie się zaleceń producenta.

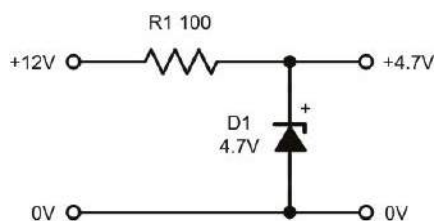
W następnym odcinku omówimy, w jakich obudowach umieszczane są akumulatory litowe. ■

Piotr Górecki

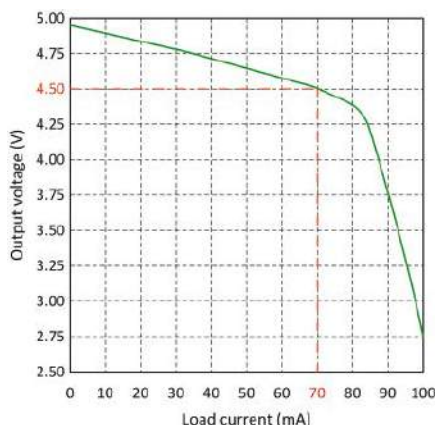


Zasilanie do twojego projektu, część 3. Stabilizatory liniowe

Twój projekt wydaje się być skończony i gotowy do uruchomienia, lecz zadanie nie jest jeszcze wykonane dopóki nie zastosujesz odpowiedniego źródła zasilania. Rozwiązanie tego problemu może być tak proste, jak zastosowanie gotowego zasilacza o odpowiednich parametrach, bądź też tak złożone, jak zbudowanie własnego systemu zasilania z zasilaniem impulsowym, kilkoma wyjściami oraz podtrzymaniem baterijnym. Nasz kurs ma za zadanie pomóc w rozwiązaniu tej kwestii i dostarczyć wiedzy na temat ważnych aspektów zasilania układów elektronicznych. W trzeciej części niniejszego kursu zajmiemy się stabilizatorami liniowymi. Te relatywnie tanie, łatwe w użyciu układy stanowią podstawę dla każdego prostego, lecz efektywnego źródła zasilania. W tym miesiącu zaprezentujemy również dwa projekty praktyczne. Pierwszym z nich jest moduł ze stabilizatorem napięcia o stałym napięciu wyjściowym, podczas gdy drugi umożliwi regulację napięcia wyjściowego w zakresie od 1,5 do 13,5 V i uzyskania wydajności około 0,6 A. W połączeniu z prezentowanym w zeszłym miesiącu praktycznym projektem prostego zasilacza prądu stałego, mogą spełniać funkcję zasilacza laboratoryjnego wspomagającego proces budowy i testowania układów elektronicznych.



Rysunek 3.1. Prosty stabilizator napięcia z diodą Zenera

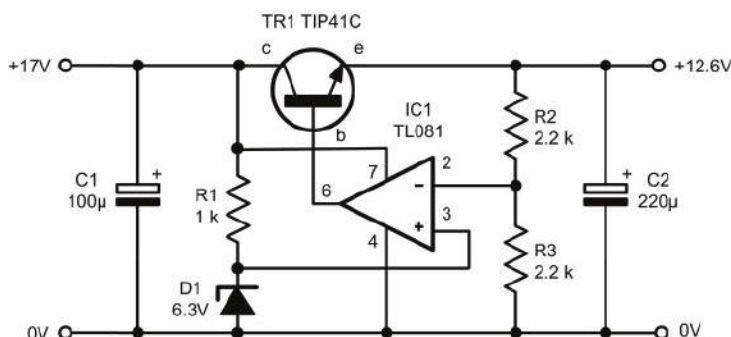


Rysunek 3.2. Wykres stabilizacji obciążenia dla stabilizatora napięcia z diodą Zenera z rysunku 3.1

Stabilizacja napięcia

Prawdopodobnie najbardziej oczywistym rozwiązaniem problemu zapewnienia zasilania o stałym napięciu jest użycie stabilizatora liniowego. Może on bazować na elementach dyskretnych, być dedykowanym układem scalonym, bądź też stanowić kombinację obu technik. Stabilizacja napięcia wyjściowego źródła zasilania jest istotna niemal w każdych zastosowaniach. W zależności od potrzeb, można uzyskać ją na kilka sposobów, lecz najprostszą z nich jest zastosowanie diody Zenera bocznikującej obciążenie, jak pokazano na rysunku 3.1.

Niestety ten prosty sposób posiada istotne ograniczenia i jest przydatny wyłącznie dla zastosowań o niskim poborze prądu (poniżej 100 mA). W praktyce taki układ stabilizuje napięcie (patrz nasze zeszłomiesięczne rozważania) z dokładnością 10% lub lepszą, z rezystancją wyjściową mniejszą niż 20 Ω . Aby zrozumieć lepiej te wartości, warto przyrzeć się prostemu przykładowi. Na rysunku 3.1 dioda Zenera 4,7 V połączona jest równolegle z wyjściem, więc suma prądów diody i obciążenia będzie płynąć przez podłączony szeregowo rezystor R1 o wartości 100 Ω . Wykres



Rysunek 3.3. Prosty układ stabilizatora napięcia

stabilizacji obciążenia przedstawiony jest na rysunku 3.2. Wskazuje on, że maksymalny prąd obciążenia takiego układu (w punkcie, w którym stabilizacja „kończy się”) to około 70 mA. W tymże punkcie napięcie wyjściowe spada do około 4,5 V (zauważ, że napięcie wyjściowe bez obciążenia wynosi niewiele poniżej 5 V). Stosując zależności omówione w pierwszej części kursu oraz informacje wynikające z rysunku 3.2 możemy określić wartości stabilizacji obciążenia oraz rezystancji wyjściowej, wynoszących odpowiednio 9,1% oraz 6,4 Ω . Dla niektórych zastosowań takie wartości mogą okazać się wystarczające, lecz, podsumowując, zastosowanie takiego prostego układu z diodą Zenera posiada następujące ograniczenia:

- Układ jest przeznaczony dla małych prądów wyjściowych (mniejszych niż 100 mA),
- Niska sprawność (istotna moc tracona jest na diodzie Zenera oraz na połączonym szeregowo rezystorze – nawet przy braku obciążenia),
- Rezystancja wyjściowa jest relatywnie wysoka (dla wielu aplikacji pożądane jest, aby rezystancja wyjściowa była bardzo niska, nie większa niż 1 Ω),
- Współczynniki stabilizacji, zarówno napięciowej jak i prądowej są dość słabe (lecz można je poprawić używając kilku dodatkowych komponentów).

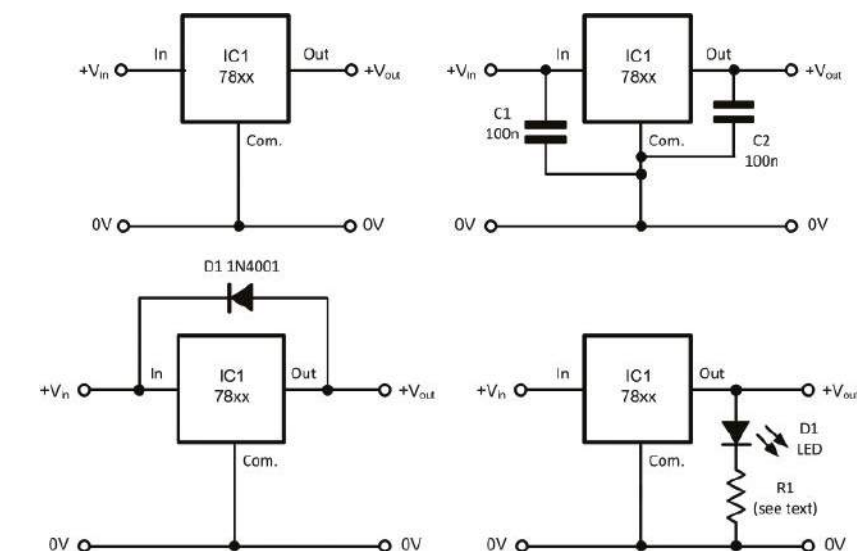
Bardziej skomplikowany układ stabilizatora napięcia pokazano na rysunku 3.3. Bazuje on na wzmacniaczu operacyjnym pracującym jako wzmacniacz błęd IC1 oraz na połączonym szeregowo tranzystorze TR1, który działa jako wzmacniacz prądu i jest niekiedy nazywany tranzystorem szeregowym. Zauważ, że wzmacniacz IC1 musi być w stanie dostarczyć odpowiedni prąd doysterowania bazy

tranzystora TR1, który z kolei ma za zadanie dostarczyć maksymalny prąd do obciążenia. Dla wysokich wartości prądu kolektora można oczekiwać raczej niskiego wzmocnienia prądowego tranzystora TR1, co może ograniczać wydajność układu. Warto zilustrować ten ważny punkt przykładem. Załóżmy, iż musimy zapewnić prąd obciążenia o wartości 1 A, a wzmacniacz operacyjny jest w stanie dostarczyć maksymalnie 20 mA. Wartość wzmocnienia prądowego tranzystora TR1 może typowo spaść z 80 dla prądu kolektora o wartości dziesiątych części mA do mniej niż 50 dla prądu kolektora wynoszącego 1 A. Nasz przykład wymaga, aby wzmocnienie prądowe wyniosło co najmniej 50 (1000/20). Z tego też powodu tranzystor TR1 musi być starannie dobrany. Alternatywnie można użyć tranzystor Darlingtona lub MOSFET aby zapewnić dużo wyższe wzmocnienie prądowe. Zagadnienie to omówimy szerzej w dalszej części kursu, gdy zajmiemy się zagadnieniami zasilaczy wysokoprądowych i wysokonapięciowych. Na rysunku 3.3 napięcie referencyjne jest uzyskane z prostego układu stabilizatora napięcia z diodą Zenera (R1 oraz D1), który jest nam już znany z rysunku 3.1. Napięcie referencyjne (w miarę stabilne 6,3 V) odcłada się na diodzie D1 i jest podawane na nieodwracające wejście komparatora. Napięcie wyjściowe jest z kolei podane na dzielnik potencjału tworzony przez rezystory R2 i R3. Warto zauważyć, iż w tym przykładzie wartości R2 oraz R3 zostały dobrane w taki sposób, aby 50% napięcia wyjściowego pojawiło się na wejściu odwracającym komparatora. Wzmacniacz operacyjny IC1 porównuje podzielone napięcie wyjściowe z napięciem referencyjnym z diody Zenera. Jeśli napięcie wyjściowe przekracza napięcie referencyjne, wyjście wzmacniacza IC1 opadnie i mniej prądu zostanie podane na bazę tranzystora

TR1. W konsekwencji napięcie wyjściowe emitera spadnie. W odwrotnym przypadku, jeśli napięcie wyjściowe z dzielnika spadnie poniżej napięcia referencyjnego, wyjście wzmacniacza IC1 wzrośnie i więcej prądu będzie podawane na bazę tranzystora TR1. W konsekwencji napięcie wyjściowe emitera wzrośnie. W ten sposób układ automatycznie kompensuje zmiany w zapotrzebowaniu obciążenia, utrzymując napięcie wyjściowe na poziomie bardzo zbliżonym do dwukrotności napięcia referencyjnego. Napięcie wyjściowe można bardzo łatwo zmienić, odpowiednio zmieniając stosunek R2 do R3. Alternatywnie można użyć potencjometru, aby zapewnić sobie możliwość ciągłej regulacji napięcia wyjściowego.

Trzykońcówkowe stabilizatory napięcia

Na szczęście nie ma potrzeby budować układów takich, jak zaprezentowany na rysunku 3.3, gdyż problem ten doskonale rozwiązuje szeroka gama powszechnie dostępnych trzykońcówkowych stabilizatorów napięcia. Te zintegrowane układy są łatwe w użyciu i dostępne są w wersjach o ustalonym lub dobieranym napięciu. Najpopularniejsza seria stabilizatorów o ustalonym napięciu dostępna jest w obudowie TO220 i posiada prefix 78 (dla dodatniego napięcia wejściowego i wyjściowego) lub 79 (dla ujemnego napięcia wejściowego i wyjściowego). Elementy te są dostępne dla szerokiego zestawu napięć (5 V, 9 V, 12 V, 15 V, 18 V i 24 V) i zwykle przeznaczone są dla maksymalnego prądu obciążenia o wartości 1 A. Wewnętrzna budowa tych elementów bazuje na analizowanym przez nas schemacie z rysunku 3.3, dodatkowo uzupełniona o zabezpieczenie termiczne wraz z określonym tzw. SOA – obszarem bezpiecznej pracy elementu. Zauważ, że napięcie referencyjne jest stałe i wbudowane w układ, a w konsekwencji nie może być zmieniane. Kilka podstawowych aplikacji stabilizatorów z serii 78xx zostało pokazanych na rysunku 3.4. Na rysunku 3.4a nóżka masy stabilizatora podłączona jest do linii 0 V wejścia oraz wyjścia. Należy mieć na uwadze, że w najgorszym przypadku niestabilizowane napięcie wejściowe prądu stałego powinno być co najmniej o 2,5 V wyższe niż nominalne, stabilizowane napięcie wyjściowe. Innymi słowy, różnica pomiędzy napięciem wejściowym a wyjściowym powinna być większa niż 2,5 V dla uzyskania efektywnej stabilizacji, jednakże dostępne są na rynku również stabilizatory LDO, które mogą pracować przy mniejszej różnicy napięć. Negatywną konsekwencją niezachowania minimalnej różnicy pomiędzy napięciem wejściowym a wyjściowym jest słaba stabilizacja wraz z nieakceptowalnym poziomem zakłóceń przenoszonych na wyjście. Innym ważnym czynnikiem brany pod uwagę podczas projektowania jest



Rysunek 3.4. Trzykońcówkowe stabilizatory napięcia

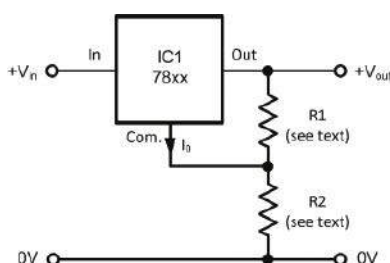
fakt, iż niestabilizowane napięcie wejściowe nie może być zbyt duże, gdyż moc, którą stabilizator będzie musiał rozproszyć przekroczy jego możliwości wytrzymałościowe. Z tego też powodu w tabeli 3.1 zawarto zalecane zakresy napięć wejściowych dla powszechnych typów stabilizatorów o ustalonym napięciu.

Warianty układu

Na rysunku 3.4b dodaliśmy dwa kondensatory o niskiej wartości (C1 oraz C2) aby zapobiec niestabilnościom oraz oscylacjom. Pamiętaj, że elementy te muszą znajdować się możliwie blisko nóżek stabilizatora, gdyż w przeciwnym razie będą nieefektywne. Ponadto powinny być one zawsze stosowane, gdy używany jest stabilizator w postaci układu scalonego. Na rysunku 3.4c dodano diodę pomiędzy wyjściem a wejściem stabilizatora. Dioda ta zabezpiecza układ w przypadku, gdy napięcie wyjściowe przekroczy napięcie wejściowe. Sytuacja taka może powstać, jeśli wyjście będzie podłączone do źródła zasilania (np. baterii) podczas gdy wejście zostanie odłączone. Na rysunku 3.4d widać, jak prostym jest dodanie wskazania obecności napięcia wyjściowego za pomocą diody LED. Wartość rezystora R1 można dobrać na podstawie tabeli 3.1. Jeśli to konieczne, napięcie wyjściowe ze stabilizatora napięcia może być ustalone na wyższą wartość niż nominalne. Dzięki temu można zbudować stabilizator na właściwie każde napięcie, jak na przykład 7,5 V, 11 V czy 13,8 V. Na rysunku 3.5 przedstawiono schemat takiej konstrukcji. Dwa rezystory (R1 oraz R2) tworzą dzielnik. Jeśli stabilizator posiada stałe napięcie nominalne V_{xx} , napięcie wyjściowe dla układu z rysunku 3.5 będzie miało wartość zgodną z następującym równaniem:

$$V_{out} = V_{xx} + \left(\left(\frac{V_{xx}}{R1} + I_o \right) \cdot R2 \right)$$

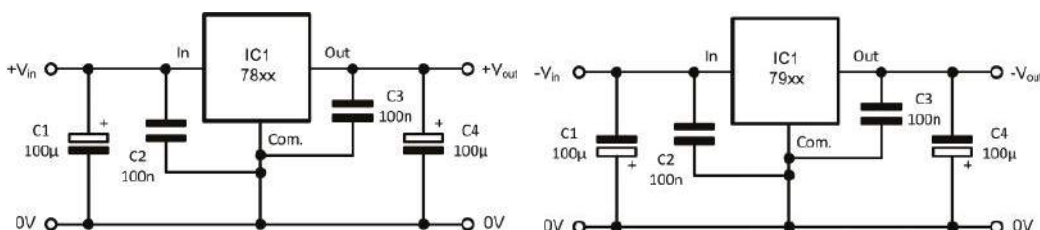
Prąd masy różni się w zależności od elementu, lecz wynosi zwykle około 4,3 mA. Załóżmy więc, że potrzebujemy stabilizowanego źródła zasilania o napięciu wyjściowym 7,5 V oraz



Rysunek 3.5. Stabilizator napięcia z regulacją napięcia wyjściowego

Tabela 3.1. Dobór komponentów dla powszechnie używanych stabilizatorów napięcia (patrz rysunek 3.6)

Output voltage	IC1	Input voltage			Voltage rating for C1 and C4 in Fig.3.6	R1 in Fig.3.4(d)
		Max.	Min.	Typ.		
5 V	7805	20 V	7,5 V	9 V	25 V	330 Ω
9 V	7809	25 V	12 V	13 V	25 V	680 Ω
12 V	7812	30 V	15 V	17 V	35 V	1,2 kΩ
15 V	7815	30 V	18 V	20 V	35 V	1,5 kΩ
-5 V	7905	-20 V	-7,5 V	-9 V	25 V	330 Ω
-9 V	7909	-25 V	-12 V	-13 V	25 V	680 Ω
-12 V	7912	-30 V	-15 V	-17 V	35 V	1,2 kΩ
-15 V	7915	-30 V	-18 V	-20 V	35 V	1,5 kΩ



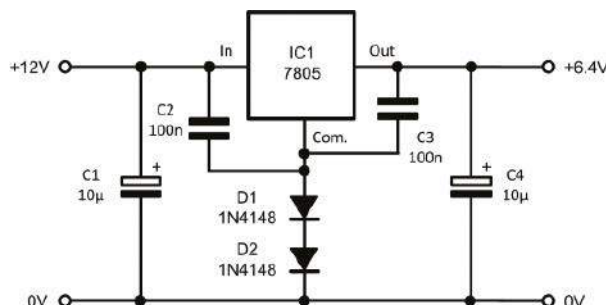
Rysunek 3.6. Porównanie obwodów z trzykońcówkowym stabilizatorem napięcia dla dodatniego i ujemnego napięcia

wydajności 1 A. Jeśli wybierzemy rezystor R1 o popularnej wartości 1 kΩ możemy przeliczyć wartość rezystora R2 z następującego równania:

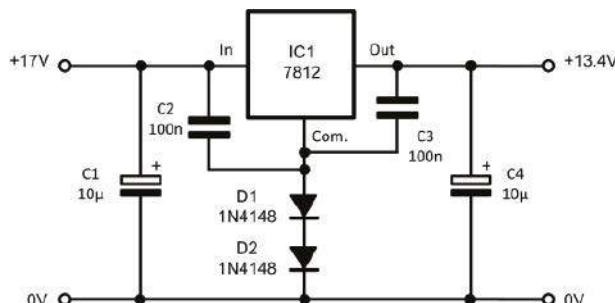
$$R2 = \frac{(V_{out} - V_{xx})}{\left(\frac{V_{xx}}{R1} + I_o \right)} = \frac{(7,5 - 5)}{\left(\frac{5}{1000} + 0,0045 \right)} = \frac{2,5}{0,0093} = 269 \Omega$$

Należy wybrać więc najbliższą dostępną z szeregu wartość czyli 270 Ω. Jeżeli wymagana jest regulacja (lub prąd masy jest inny od oczekiwanego), R2 może być zastąpiony potencjometrem. W takim przypadku wystarczający będzie rezystor o wartości 500 Ω. Stabilizatory napięcia są dostępne w wersjach do pracy z dodatnimi lub ujemnymi napięciami wejściowymi i wyjściowymi. Na rysunku 3.6 zilustrowano porównanie obwodów z trzykońcówkowym stabilizatorem napięcia w wersji dla dodatniego i dla ujemnego napięcia. Dodano dwa kondensatory C1 oraz C4 w celu redukcji

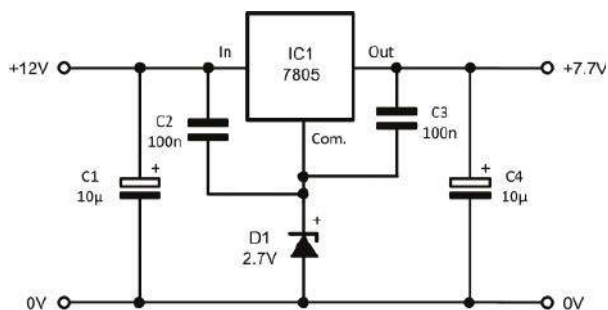
szumów. Powinny być one odpowiednio dobrane pod kątem napięcia pracy oraz, z uwagi na to, iż są to kondensatory elektrolityczne, odpowiednio spolaryzowane (patrz tabela 3.1 z rekomendowanymi wartościami). Inny, prosty sposób zwiększenia napięcia wyjściowego stabilizatora o krok o wartości 0,7 V jest pokazany na rysunku 3.7. W tej konfiguracji dwie krzemowe diody spolaryzowane w kierunku przewodzenia połączone są szeregowo



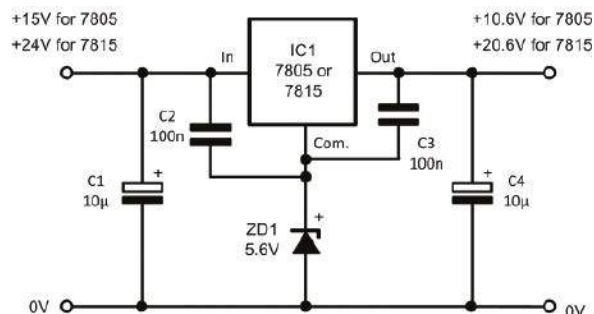
Rysunek 3.7. Użycie diod do zwiększenia napięcia wyjściowego stabilizatora napięcia 5 V



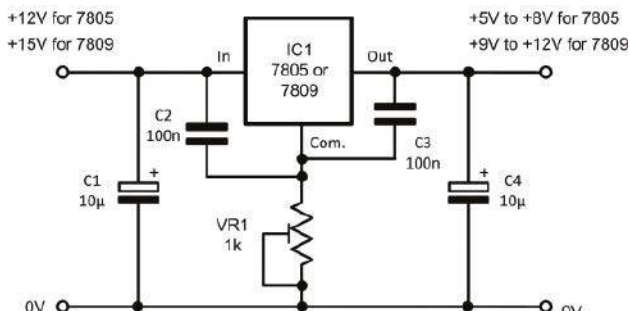
Rysunek 3.8. Schemat zasilania o napięciu 13,4 V z użyciem stabilizatora 12 V



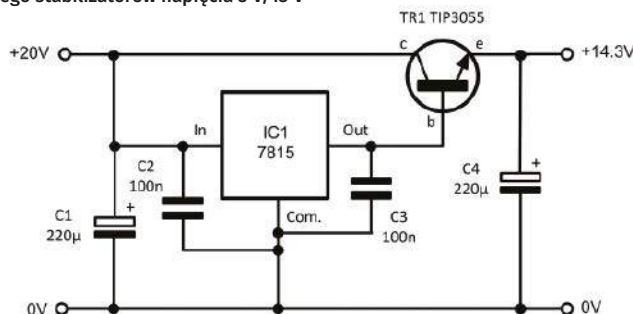
Rysunek 3.9. Użycie diody Zenera 2,7 V do zwiększenia napięcia wyjściowego stabilizatora napięcia 5 V



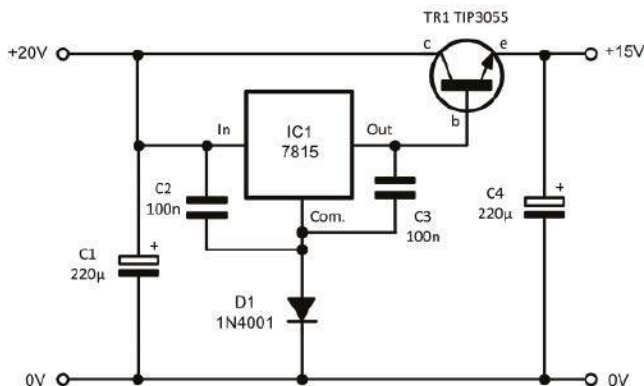
Rysunek 3.10. Użycie diody Zenera 5,6 V do zwiększenia napięcia wyjściowego stabilizatorów napięcia 5 V/15 V



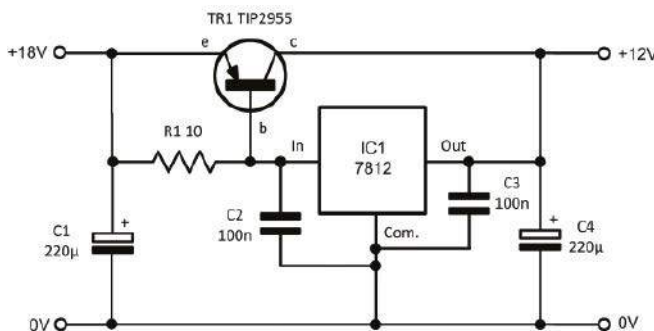
Rysunek 3.11. Regulowane napięcie wyjściowe stabilizatora napięcia



Rysunek 3.12. Układ trzykońcówkowego stabilizatora ze zwiększoną wydajnością prądową do około 5 A



Rysunek 3.13. Dodanie diody w celu kompensacji spadku napięcia wyjściowego w układzie z rysunku 3.12



Rysunek 3.14. Ulepszony układ stabilizatora ze zwiększoną wydajnością prądową

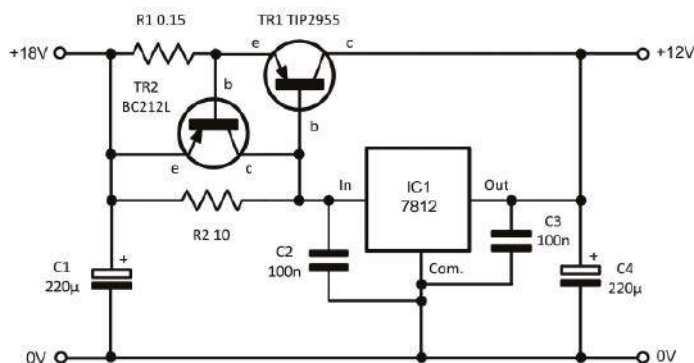
z nóżką masy stabilizatora. Napięcie wyjściowe stabilizatora zwiększa się o $(2 \times 0,7 \text{ V})$ dając w rezultacie $(5 + 1,4) = 6,4 \text{ V}$. Ten drobny trik może być użyty w każdym z dotychczas prezentowanych układów. Przykładowo, rysunek 3.8 prezentuje w jaki sposób standardowy stabilizator 12 V może być zmodyfikowany do uzyskania napięcia wyjściowego 13,4 V przy użyciu jedynie dwóch diod. Dla wyższych napięć wyjściowych może być wygodniejsze zastąpienie kilku diod jedną diodą Zenera, jak pokazano na rysunku 3.9. Schemat ten pokazuje również, w jaki sposób użyć stabilizatora 5 V aby przy użyciu diody Zenera 2,7 V uzyskać napięcie wyjściowe 7,7 V. Schemat na rysunku 3.10 pokazuje zaś w jaki sposób można otrzymać napięcia wyjściowe 10,6 V i 20,6 V przy użyciu diody Zenera 5,6 V wraz ze stabilizatorami odpowiednio 7805 oraz 7815. Podobna technika

„podbijająca” może być zastosowana do uzyskania regulowanego napięcia wyjściowego poprzez dodatnie potencjometru do linii masy stabilizatora, jak pokazano na rysunku 3.11. Obwód ten dostarcza regulowane napięcie w zakresie od 5 V do 8 V z elementu 7805 lub 9 V do 12 V z elementu 7809.

Zwiększenie prądu wyjściowego

Gdy potrzebny jest nam prąd o wartości większej niż 1 A, standardowe możliwości stabilizatora można rozszerzyć poprzez dodanie tranzystora szeregowego, jak pokazano na rysunku 3.12. Napięcie wyjściowe w tej konfiguracji będzie o około 0,7 V mniejsze niż nominalne napięcie stabilizatora. W prezentowanym przypadku napięcie wyjściowe wyniesie więc około 14,3 V. Zauważ, że taki układ nie będzie miał takich zabezpieczeń

termicznych i nadprądowych jak stabilizator użyty „samodzielnie”. W niektórych zastosowaniach ten fakt może nie mieć znaczenia, lecz dla innych może okazać się kluczowy. Zauważ również, że spadek napięcia wyjściowego o 0,7 V może być łatwo skompensowany poprzez dodanie krzemowej diody do nóżki masy stabilizatora, jak opisano już wcześniej (patrz rysunek 3.13). Kolejna modyfikacja układu z trzykońcówkowym stabilizatorem pokazana jest na rysunku 3.14. Układ ten używa tranzystora PNP zamiast NPN, jak to było w poprzednich układach. Nie ma potrzeby użycia dodatkowej diody podłączonej do nóżki masy stabilizatora. Aby zabezpieczyć tranzystor TR1 i ograniczyć prąd wyjściowy do bezpiecznej wartości, należy dodać jeszcze jeden tranzystor, jak pokazano na rysunku 3.15. W tej konfiguracji prąd obciążenia jest mierzony poprzez R1



Rysunek 3.15. Zastosowanie dodatkowego tranzystora w celu ograniczenia prądu

i w konsekwencji steruje tranzystorem TR2, który przewodzi, gdy prąd obciążenia przekracza około 4,5 A. Jeśli założymy, że I_{lim} jest wymaganym limitem prądu, to wymagana wartość R1 może być wyliczona z formuły: $I_{lim} = 0,7/R1$.

Projekt praktyczny: Moduł stabilizatora napięcia 1 A

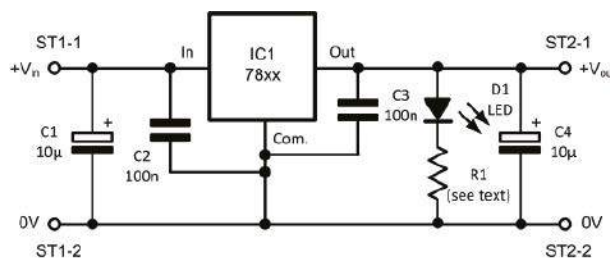
Pierwszy z dwóch projektów praktycznych stanowi kolejny moduł, który może znaleźć zastosowanie w wielu projektach. Jest nim niskokosztowy moduł stabilizatora napięcia o wydajności prądowej do 1 A. W zależności od użytego modelu stabilizatora, może być zbudowany w odmianach dostarczających 5 V, 9 V, 12 V oraz 15 V. Schemat modułu jest pokazany na rysunku 3.16. Zbudowany jest on zgodnie z koncepcjami prezentowanymi w niniejszej części kursu i zawiera wskaźnik działania w postaci diody LED. Właściwa wartość rezystora R1 będącego w szeregu z diodą LED powinna być dobrana odpowiednio z tabeli 3.1.

Będziesz potrzebował...

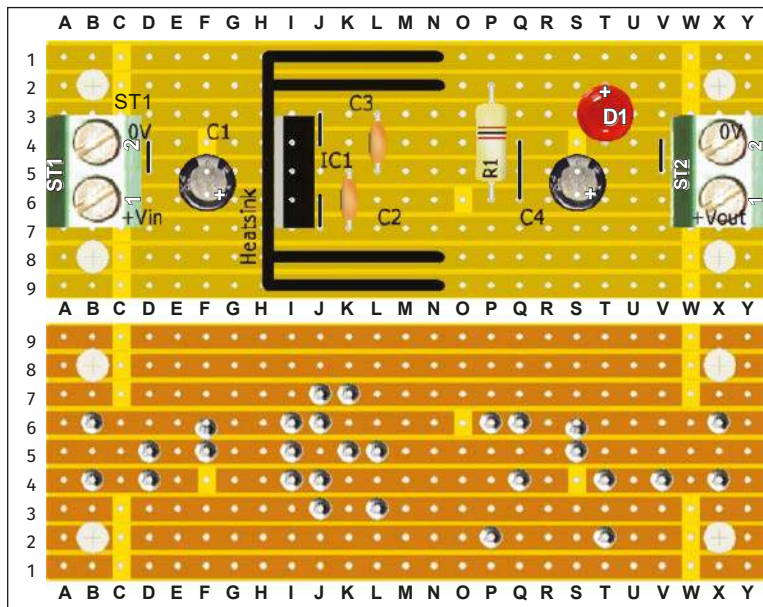
- 1 płytkę uniwersalną z miedzianymi paskami (9 paszków po 25 otworów)
- 2 listwy zaciskowe 2-torowe do druku (ST1 i ST2)
- 1 rezystor 1kΩ (R1) (patrz tekst)
- 2 kondensatory 10μF 35V (C1 i C4)
- 2 kondensatory 100nF (C2 i C3)
- 1 78xx trójnóżkowy stabilizator napięcia (patrz tekst)
- 1 czerwony LED
- 1 mały radiator TO220 (patrz tekst)
- 4 tulejki dystansowe i śruby montażowe.

Budowa

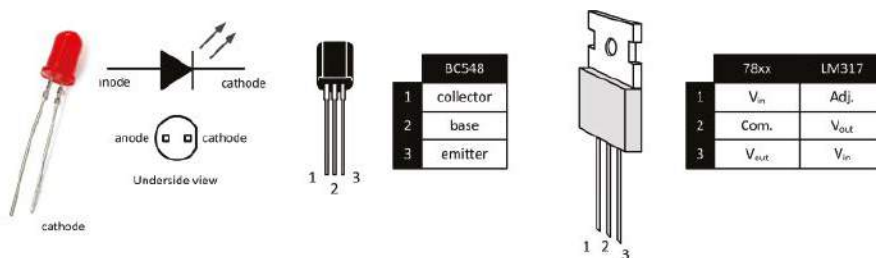
Mozaika ścieżek układu jest pokazana na rysunku 3.17. Zwróć uwagę, iż posiada on 15 przerw ścieżek i 5 zworek, a opis wyprowadzeń elementów półprzewodnikowych znajduje się na rysunku 3.18. Kondensatory C2 oraz C3 zostały zamontowane bardzo blisko pinów stabilizatora IC1. Jak już wspomniano,



Rysunek 3.16. Schemat modułu stabilizatora napięcia 1 A

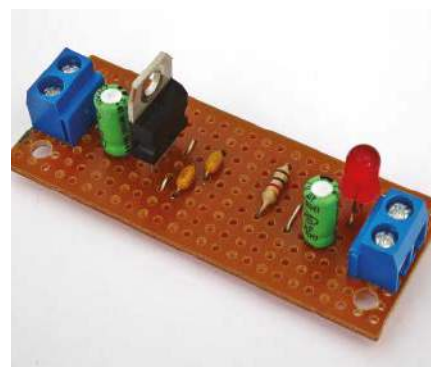


Rysunek 3.17. Ułożenie elementów na płytce prototypowej układu modułu stabilizatora napięcia 1 A: (wyżej) widok z góry z komponentami, (poniżej) widok ścieżek

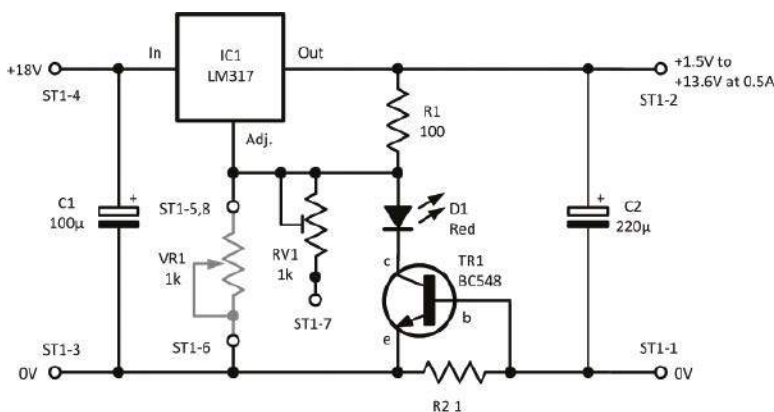


Rysunek 3.18. Wyprowadzenia pinów półprzewodników użytych w projektach praktycznych. Od lewej: LED, BC548, 78xx oraz LM317

te dwa kondensatory są niezbędne aby uniknąć możliwych oscylacji spowodowanych przez pasożytniczą reaktancję ścieżek i przewodów. Trzykońcówkowe stabilizatory powinny być montowane do małych radiatorów do obudów TO220 (dla czytelności został usunięty na rysunku 3.19). Do większości zastosowań odpowiedni będzie radiator o rezystancji termicznej lepszej niż 13,5°C/W, lecz jeśli moduł ma pracować w trybie ciągłym na pełnym obciążeniu, rekomendowany jest radiator o rezystancji 7,8°C/W. Radiator taki może być zamontowany przy użyciu śruby i nakrętki lub



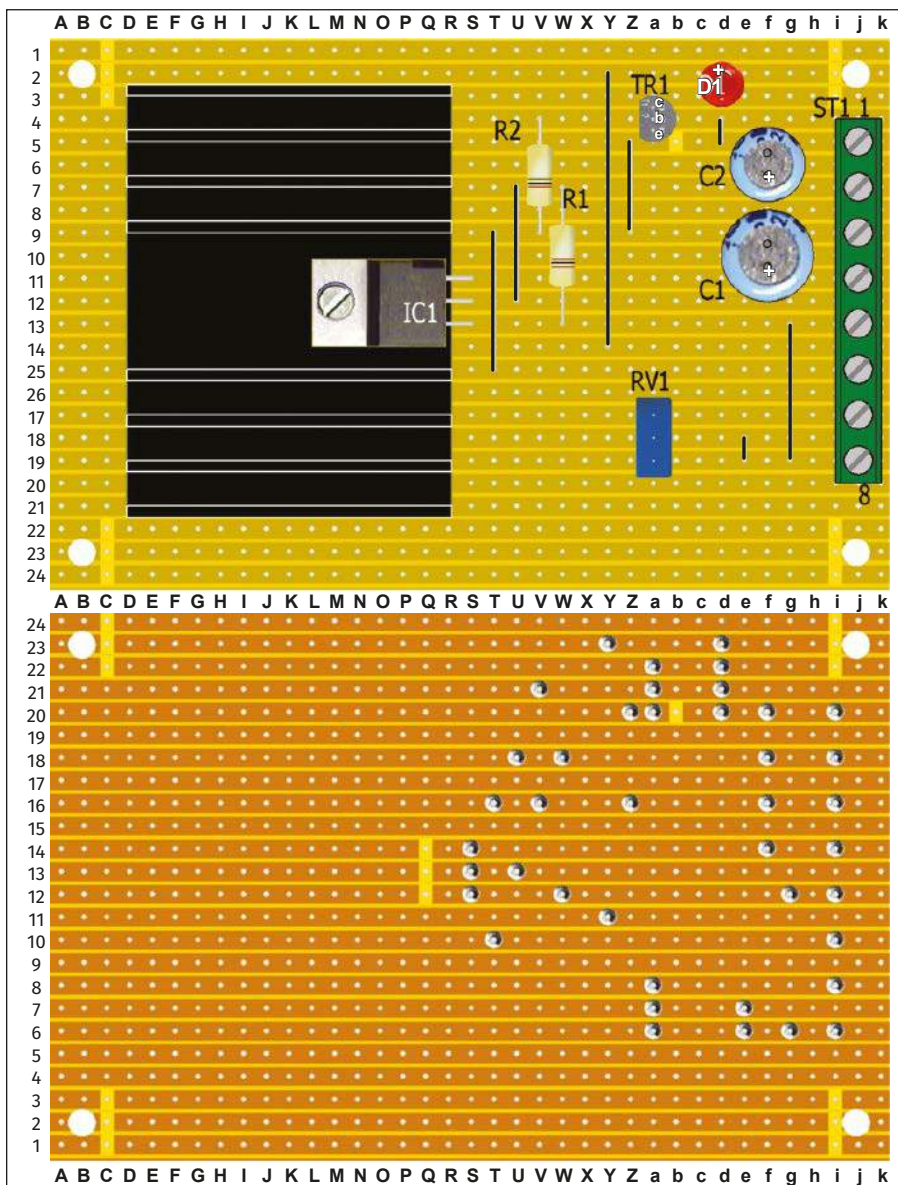
Rysunek 3.19. Gotowy moduł stabilizatora napięcia 1 A



Rysunek 3.20. Schemat regulowanego zasilacza stabilizowanego

odpowiedniego mocowania. W dalszej części kursu będziemy rozważać ten temat bardziej szczegółowo. Aby polepszyć przewodność cieplną zrezygnowaliśmy z podkładek izolacyjnych

i przytwierdziliśmy metalową obudowę stabilizatora bezpośrednio do radiatora. Obudowa ta jest wewnętrznie połączona z masą więc potencjał na radiatorze będzie zerowy.



Rysunek 3.21. Ułożenie elementów na płytce prototypowej regulowanego zasilacza stabilizowanego: (wyżej) widok od strony komponentów, (poniżej) widok ścieżek (przerwane ścieżki zaznaczono na żółto)

Projekt praktyczny: Prosty, regulowany zasilacz stabilizowany

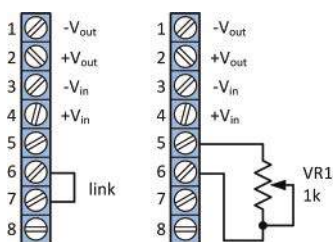
Naszym drugim projektem praktycznym w tym miesiącu jest prosty zasilacz stabilizowany prądu stałego. Moduł ten nadaje się do użycia w połączeniu z opisanym w zeszłym miesiącu zasilaczem i jest w stanie dostarczyć napięcie wyjściowe regulowane w ciągłym zakresie od 1,5 do 13,5 V z prądem do 0,6 A. Może on być podstawą do stworzenia użytecznego stanowiska zasilania zawierającego zabezpieczenie nadprądowe ograniczające prąd wyjściowy. Schemat stabilizowanego zasilacza pokazano na rysunku 3.20. Trójnóżkowy stabilizator napięcia (IC1) został zaprojektowany do dostarczenia regulowanego napięcia a wyjście masy układu LM317 staje się pinem regulacyjnym. Jest to konfiguracja bardzo podobna do prezentowanej na rysunku 3.5. Obwód ten może posiadać stałą wartość napięcia wyjściowego lub być regulowany za pomocą potencjometru wpiętego w złącze ST1, zgodnie z rysunkiem 3.22. Zawiera on też proste zabezpieczenie prądowe tworzone przez elementy R1, TR1 i D1, w którym to prąd wyjściowy jest mierzony przez rezystor 1 Ω podpięty do linii 0 V. Gdy spadek napięcia na rezystorze przekroczy 0,6 V, tranzystor TR1 zaczyna przewodzić ustawiając stan niski na pinie regulacyjnym IC1. Gdy tranzystor TR1 przewodzi, D1 zapala się dostarczając nam informacji o przeciążeniu.

Będziesz potrzebował...

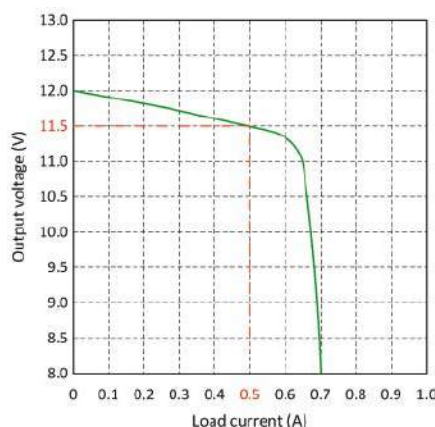
- 1 płytkę uniwersalną z miedzianymi paskami (24 paski po 37 otworów)
- 4 listwy zaciskowe 2-torowe do druku (ST1)
- 1 rezystor 100 Ω 0,5 W (R1)
- 1 rezystor 1 Ω 0,5 W (R2)
- 1 potencjometr 1 kΩ (RV1) (patrz tekst)
- 1 potencjometr 1 kΩ (VR1) (patrz tekst)
- 1 kondensator 100 µF 35 V (C1)
- 1 kondensator 220 µF 35 V (C2)
- 1 BC548 tranzystor NPN (TR1)
- 1 czerwona dioda LED (D1)
- 1 radiator TO220 (patrz tekst)
- 4 tulejki dystansowe i śruby montażowe

Budowa

Konfiguracja płytki prototypowej stabilizowanego zasilacza regulowanego jest pokazana na rysunku 3.21. Zwróć uwagę, iż posiada ona 16 przerw ścieżek oraz 7 zwerek, a opis wyprowadzeń elementów półprzewodnikowych znajduje się na rysunku 3.18. Trzykońcówkowy stabilizator musi być zamontowany na żeberkowym radiatorze. Ponownie, aby polepszyć przewodność cieplną, unikamy użycia izolacyjnej podkładki i zamiast tego przykręcamy metalową obudowę IC1 bezpośrednio do radiatora, który musi być

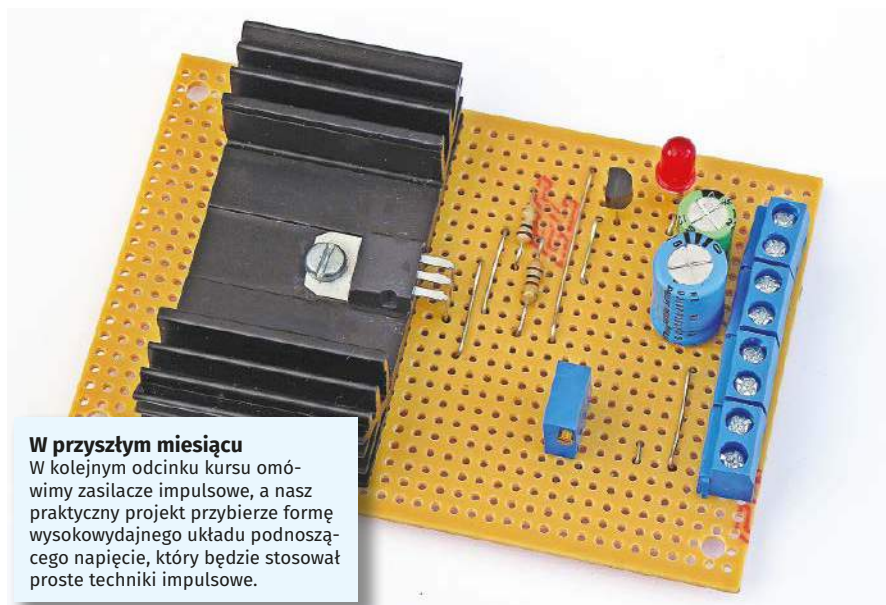


Rysunek 3.22. Konfiguracja połączeń na listwie zaciskowej ST1 dla uzyskania stałego napięcia wyjściowego (po lewej) oraz dla uzyskania regulowanego napięcia wyjściowego (po prawej)



Rysunek 3.24. Wykres stabilizacji obciążenia regulowanego zasilacza stabilizowanego

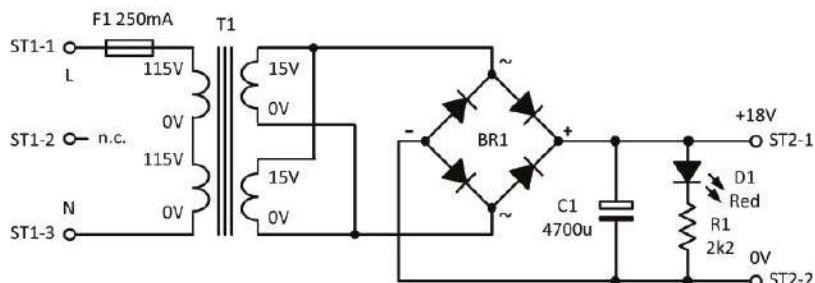
odizolowany od metalowych części obudowy. Zauważ, że w tym przypadku obudowa TO-220 jest bezpośrednio połączona z dodatnim wyjściem. Jak zawsze, po montażu sugerowane jest wykonanie uważnego sprawdzenia zmontowanego układu, zwłaszcza prawidłowego podłączenia zworek oraz połączeń w listwie ST1 (patrz rysunek 3.22). Typowy wykres stabilizacji obciążenia zasilacza jest pokazany na rysunku 3.24. Zauważ, że ograniczenie prądowe zaczyna działać przy wartości 0,6 A i wraz ze wzrostem prądu powyżej tej



W przyszłym miesiącu

W kolejnym odcinku kursu omówimy zasilacze impulsowe, a nasz praktyczny projekt przybierze formę wysokowydajnego układu podnoszącego napięcie, który będzie stosował proste techniki impulsowe.

Rysunek 3.23. Widok gotowego regulowanego zasilacza stabilizowanego



Errata

Schemat pokazany na rysunku 2.18 w poprzednim numerze jest nieprawidłowy. Dwa zaczepty wtórne powinny być połączone równolegle, a nie szeregowo (patrz powyżej – poprawiony schemat). Układ na płytce uniwersalnej (rysunek 2.19) został narysowany prawidłowo.

wartości, napięcie wyjściowe jest gwałtownie zredukowane. ■

Mike Tooley

Ten artykuł był wcześniej opublikowany na łamach „Everyday Practical Electronics”, luty 2019 (www.epemag3.com)

REKLAMA

Świat projektantów i programistów dla elektroniki w nowej odsłonie. Odwiedź nowy

ELPORTAL.pl

Obserwuj nas również na Facebooku:
www.facebook.com/Elportalpl

Zrób to dobrze! – Podłączanie zasilania sieciowego



Rysunek 3.25. Zawsze używaj wtyczki typu IEC z zalewaną obudową i prawidłowo dobranym bezpiecznikiem

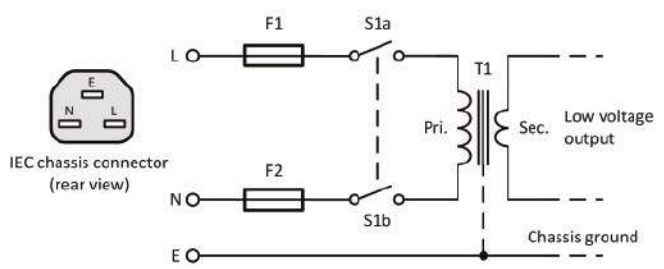


Rysunek 3.26. Męskie gniazdo typu IEC z wbudowanym bezpiecznikiem



Rysunek 3.27. Trzy rodzaje męskich gniazd typu IEC – widok od tyłu. Należy unikać wersji z połączeniami śrubowymi (po prawej), gdyż trudno takie połączenia odpowiednio zaizolować

Rysunek 3.28. Schemat zalecanego podłączenia napięcia sieciowego poprzez dwupolowy przełącznik i dwa bezpieczniki



Jak wspomniano w części 2, podłączanie prądu zmiennego z sieci elektrycznej do naszych układów wymaga szczególnej uwagi. Większość projektów prezentowanych w tym kursie używa bezpiecznego, niskiego napięcia stałego, lecz praca z zasilaniem sieciowym wymaga specjalnych środków ostrożności, aby zminimalizować ryzyko porażenia. W tym odcinku serii Zrób to dobrze! opiszemy podstawowe zasady, które są niezbędne aby zachować bezpieczeństwo przy podłączaniu budowanych urządzeń do sieci.

1. Zawsze używaj wtyczki typu IEC z zalewaną obudową (rysunek 3.25) oraz z prawidłowo dobranym bezpiecznikiem (domyślnie są one dostarczane z bezpiecznikiem 5 A i jest to wystarczające dla większości zastosowań). Zauważ, że standard kolorów używanych w Wielkiej Brytanii dla prądu zmiennego 230 V jest następujący: brązowy dla przewodu fazowego (L), niebieski dla przewodu neutralnego (N) oraz żółto-zielony dla przewodu ochronnego (E). W Polsce, zgodnie z normą Międzynarodowej Komisji Elektrotechnicznej IEC 60446, stosowane są identyczne kolory. W przeszłości dla przewodu fazowego (L) używano zarówno koloru brązowego jak i czarnego.
2. Po stronie urządzenia używaj męskich

gniazd typu IEC do montażu na obudowach. (patrz rysunki 3.26 oraz 3.27). Dostępne są one w różnych formach, włącznie z takimi o wbudowanych bezpiecznikach i włącznikach.

3. Jeśli używasz włącznika na obudowie czy panelu urządzenia i oddzielne bezpieczników, powinny one być podłączone zgodnie ze schematem na rysunku 3.28. Zauważ, że włącznik zasilania sieciowego powinien być odpowiednio dobrany do pracy z napięciem sieciowym zmiennym. Włącznik powinien być dwupolowy typu włącz/wyłącz przełączając obie linie (L i N). Może to być istotne w przypadku, gdy przewody fazowy i neutralny zostaną zamienione. Aby uniknąć kontaktu z potencjałem sieci, przełączniki i bezpieczniki powinny być odpowiednio zabezpieczone koszulkami termokurczliwymi, jak pokazano na rysunku 3.29.
4. Upewnij się, że wszystkie metalowe części obudowy są prawidłowo uziemione (patrz rysunek 3.30). Połączenie uziemienia musi być połączone z bolcem (E) wtyczki/złącza typu IEC.

W następnym miesiącu przyjrzymy się tematyce doboru prawidłowych bezpieczników do projektowanych modułów i urządzeń

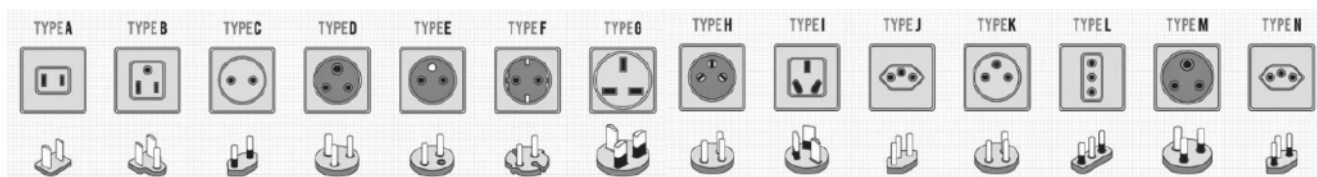


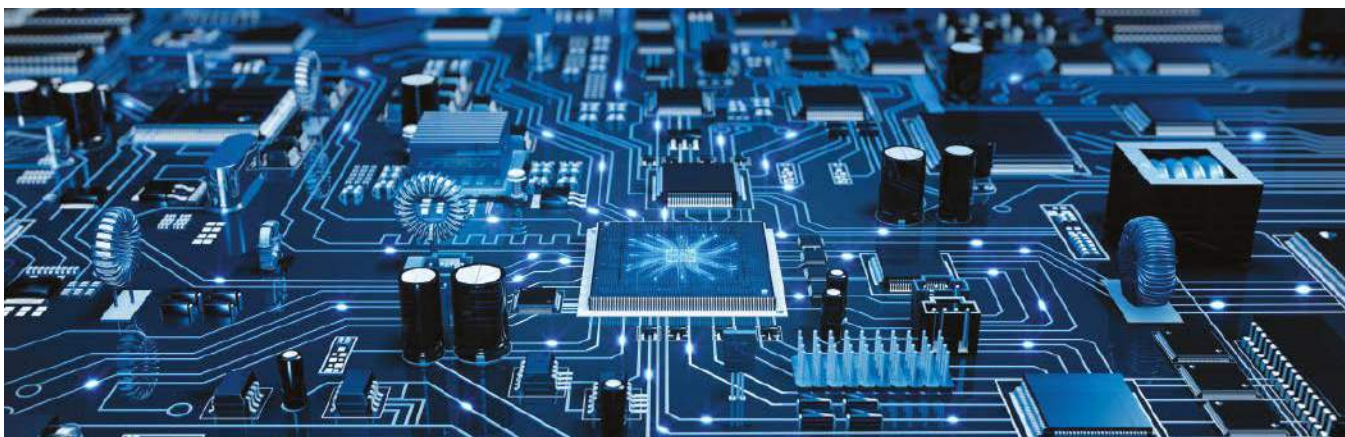
Rysunek 3.29. Używanie koszulek termokurczliwych do zabezpieczenia przełączników, bezpieczników i konektorów



Rysunek 3.30. Pewne połączenie z uziemieniem powinno być wykonane dla wszystkich metalowych części obudowy

Od Redakcji EdW. Na świecie, a nawet w samej Europie, jest wiele różnorodnych rozwiązań gniazdek i wtyczek sieciowych. Omówiony wyżej model brytyjski jest zgodnie uważany za najlepszy, najbardziej bezpieczny. Mnogość rozwiązań obrazuje rysunek poniżej, pochodzący z bloga <https://bit.ly/3yACfgf>, gdzie można przeczytać świetny artykuł o różnych standardach w świecie wtyczek i gniazdek.





Projektowanie układów porównujących

Jakiś czas temu na forum EEWeb pojawiło się kilka postów odnoszących się do budowy układów porównujących. Qasim Ahtesham zapytał: „Czy ktokolwiek z was miał do czynienia z układem LM741 i używał go do porównywania napięć? Próbuję użyć go do takiego właśnie porównania dwóch napięć i zastosowania w roli termostatu. Mam jednak problem, gdyż w przypadku, gdy $V_+ > V_-$, napięcie wyjściowe wynosi między 1,3 a 1,6 V – nigdy zaś 0 V. Czy potrzebny mi w takim razie offset? Póki co, używam tylko 5 pinów. Jakakolwiek pomoc w tym temacie lub dokładny schemat do porównywania napięć będzie mile widziany”. Graham Rounce napisał zaś: „Szukam prostego sposobu do określenia, czy dane napięcie znajduje się poza zdefiniowanym zakresem. Założmy, że ten zakres zawiera się pomiędzy +2 V a +3 V. Chcę zbudować układ, który poda stan wysoki na wyjściu (5 V) jeśli napięcie wejściowe spełni warunek ≤ 2 V lub ≥ 3 V. Ponadto potrzebne mi będą dwa oddzielne wyjścia: jedno dla ≤ 2 V, drugie zaś dla ≥ 3 V. Wzmacniacz (lub wzmacniacze) operacyjne bez wątpienia mogą mi w tym pomóc, lecz niestety jestem nieco zagubiony, gdy używam ich z pojedynczym zasilaniem 5 V”.

Poruszymy więc dziś temat komparatorów. Przyjrzymy się podstawowym zasadom działania tych układów, a także przerzutnikowi Schmitta oraz histerezie – często stosowanej w układach komparatorów i również szeroko dyskutowanej na wyżej wspomnianym forum. Oba wspomniane w wstępie posty sugerują użycie wzmacniaczy operacyjnych (LM741 jest również wzmacniaczem), nie zaś dedykowanych układów komparatorów. Mimo to w dalszej dyskusji pod pytaniem Grahama został wspomnian właśnie taki dedykowany układ, jakim jest LM393. Kolejną zatem kwestią, której się przyjrzymy będą różnice

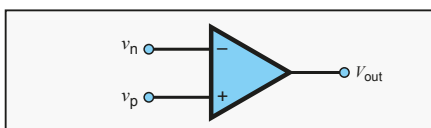
między wzmacniaczami operacyjnymi, które mogą pracować jako komparatory, a dedykowanymi do tego celu układami. Użyjemy symulacji z programu LTSpice do zilustrowania pracy obu z nich. W poprzednich artykułach serii znalazło się kilka informacji na temat podstaw programu LTSpice, lecz w niniejszym odcinku będziemy symulować działanie wzmacniaczy, które nie są elementem podstawowej biblioteki. Nie będziemy jednak zagłębiać się w techniczne szczegóły symulacji, gdyż temat ten poruszymy szerzej w kolejnym artykule.

Definicje

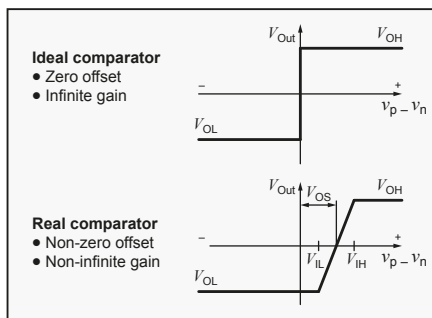
Komparator jest układem, który porównuje dwa sygnały analogowe i wystawia na wyjściu sygnał binarny na podstawie rezultatu tego porównania. Działa więc jak jednobitowy konwerter analogowo-cyfrowy. Wzmacniacz operacyjny bez ujemnego sprzężenia zwrotnego (z otwartą pętlą) ma bardzo wysokie wzmocnienie,

więc dla każdej, nawet bardzo małej różnicy sygnałów wejściowych otrzymamy na wyjściu najniższe lub najwyższe możliwe napięcie (nasylenie wyjścia), zwykle w pobliżu skrajnych wartości napięć linii zasilania. Te dwa stany nasycenia napięcia wyjściowego mogą reprezentować logiczne 0 lub 1. W takiej konfiguracji wzmacniacz operacyjny działa jak komparator, lecz jego działanie nieco różni się od dedykowanego układu komparatora. Zdefiniujemy więc charakterystyki komparatora i przeanalizujemy, jak wygląda ich porównanie w stosunku do charakterystyk wzmacniaczy operacyjnych. Symbol komparatora pokazano na rysunku 1 – jest on właściwie taki sam jak symbol wzmacniacza. Matematycznie, napięcie wyjściowe można zapisać w następujący sposób:

$$V_{\text{out}} = \begin{cases} V_{\text{OH}} & \text{jeśli } v_p > v_n \rightarrow \text{logiczne 1} \\ V_{\text{OL}} & \text{jeśli } v_p < v_n \rightarrow \text{logiczne 0} \end{cases}$$



Rysunek 1. Symbol komparatora – taki sam jak wzmacniacza operacyjnego

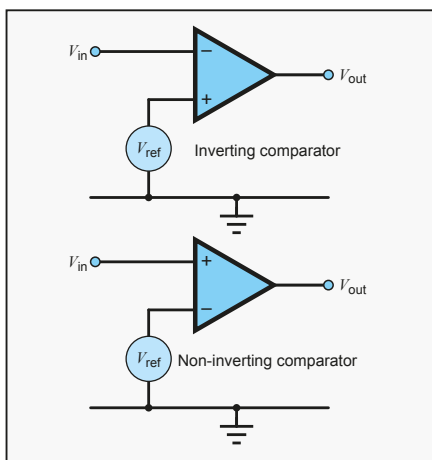


Rysunek 2. Charakterystyki przejściowe komparatora

gdzie v_p oraz v_n są napięciami wejściowymi (jak na rysunku 1), V_{OH} jest logiczną jedynką lub wysokim napięciem wyjściowym, V_{OL} jest logicznym zerem lub niskim napięciem wyjściowym. Powyższe równanie zakłada nieskończone wzmocnienie oraz offset V_{OS} wynoszący zero. Oznacza to, że nieskończenie mała zmiana napięcia około $v_p = v_n$ spowoduje przełączenie wyjścia (nieskończone wzmocnienie), a przełączenie to nastąpi dokładnie w punkcie $v_p = v_n$ (zerowy offset). Rysunek 2 ilustruje nam jednak wpływ skończonego wzmocnienia oraz niezerowego offsetu na charakterystykę przejściową (zależność napięcia wyjściowego od różnicy napięć wejściowych). Efektem działania tych dwóch parametrów jest redukcja rozdzielczości komparatora sprawiająca, iż różnica pomiędzy sygnałami wejściowymi musi być o pewne minimum większa aby umożliwić niezawodną detekcję momentu przełączenia. Rysunek 2 ponadto przedstawia sytuację dla statycznych wejść, gdzie rozdzielczość wynosi około $V_{OS} + V_{IH} = V_{OS} + V_{OH} / \text{wzmocnienie}$. Zwykle jednak rozdzielczość pogarsza się dla zmiennych sygnałów wejściowych.

Podstawowe konfiguracje

W wielu zastosowaniach komparatory są używane do porównania sygnału wejściowego z napięciem referencyjnym generowanym

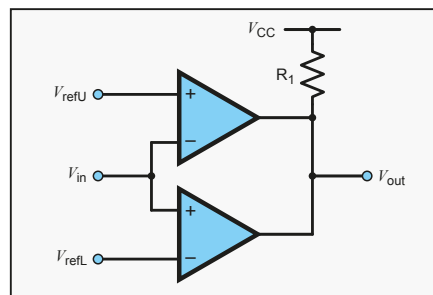


Rysunek 3. Konfiguracje komparatora – odwracająca i nieodwracająca

przez układ. W takich przypadkach możemy zastosować układ odwracający lub nieodwracający, w zależności od tego, które wejście komparatora podłączymy do napięcia referencyjnego (patrz rysunek 3). Komparator nieodwracający ustawia na wyjściu stan wysoki (logiczne 1) gdy napięcie wejściowe jest większe od napięcia referencyjnego. Dla komparatora odwracającego, na wyjściu ustawiany jest stan wysoki, gdy napięcie wejściowe jest poniżej napięcia referencyjnego. Wspomniane we wstępie wymagania Grahama, aby komparator posiadał dwa wyjścia wskazujące, czy napięcie jest ≤ 2 V lub ≥ 3 V wymusza użycie dwóch komparatorów, jednego odwracającego (z napięciem referencyjnym 2 V, oraz jednego nieodwracającego, z napięciem referencyjnym 3 V. Zdarza się również, iż zachodzi potrzeba zastosowania pojedynczego wyjścia sygnalizującego, iż napięcie wejściowe znajduje się w pewnym określonym zakresie (na przykład powyżej dolnego limitu V_{refL} , lecz poniżej górnego limitu V_{refH}). Taki układ nazywany jest komparatorem okienkowym. On również wymaga użycia dwóch komparatorów (ponownie jednego odwracającego i jednego nieodwracającego), lecz wyjścia są połączone funkcją logiczną AND („powyżej dolnego” AND „poniżej górnego”). Funkcja AND może zostać zaimplementowana w postaci bramki logicznej, lecz powszechnie używa się także komparatorów z wyjściami o otwartym kolektorze, które mogą być podłączone do wspólnego rezystora podciągającego (patrz rysunek 4), tworząc galwaniczną funkcję AND (tzw. iloczyn montażowy).

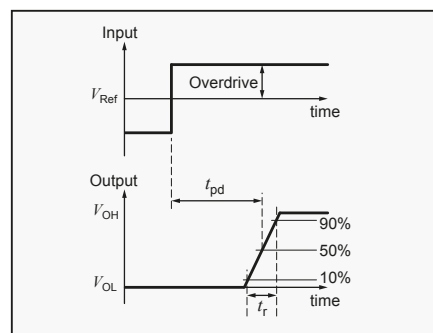
Wzmacniacze operacyjne i komparatory

Wzmocniacze operacyjne zostały zaprojektowane do użycia z ujemnym sprzężeniem zwrotnym (realizowanym na przykład za pomocą rezystorów ustalających wzmocnienie układu). Wszystkie wzmocniacze posiadają pewne opóźnienie od wejścia do wyjścia, co skutkuje wzrostem przesunięcia fazowego wraz ze wzrostem częstotliwości sygnału. W pewnym momencie przesunięcie fazowe osiąga 180° , co jest równoznaczne z odwróceniem sygnału, a tym samym układ ujemnego sprzężenia zwrotnego zaczyna tak naprawdę dostarczać dodatnie sprzężenie. Jeśli wzmocnienie wzmocniacza i układu sprzężenia zwrotnego jest większe niż 1 dla takiej częstotliwości, pojawią się oscylacje. Wzmocnienie wielu wzmocniaczy operacyjnych jest celowo zmniejszane wraz ze wzrostem częstotliwości, aby uniknąć niestabilności – zjawisko to nazywamy kompensacją. Komparatory używane są zaś z otwartą pętlą lub z dodatnim sprzężeniem zwrotnym, więc kompensacja nie jest wymagana, co jest



Rysunek 4. Komparator okienkowy używający wyjść z otwartym kolektorem

znaczącą różnicą pomiędzy tymi dwoma typami układów. Wzmocniacze operacyjne są liniowymi wzmacniaczami o dużym wzmocnieniu – podczas normalnej pracy różnica napięć pomiędzy jego wejściami jest bardzo mała (rzędu mikro do miliwoltów). Nie wszystkie wzmocniacze operacyjne tolerują duże różnice pomiędzy sygnałami wejściowymi i nie pewnie pracują w takich warunkach. Impedancja wejściowa wzmacniaczy operacyjnych może znacząco zmaleć dla sygnałów wejściowych o dużej różnicy, gdyż zaczynają wtedy przewodzić diody zabezpieczające. W konsekwencji sterowanie pracą wzmacniacza jako komparatora zostaje zakłócone. W skrajnych wypadkach mogą nawet ulec uszkodzeniu. Komparatory pracują zwykle z sygnałami wejściowymi o wiele bardziej różniącymi się między sobą. Powszechnie używa się ich do porównywania napięć, których wartość jest niższa niż połowa napięcia zasilania. Wzmocniacze posiadają zwykle szeroki zakres CMVIN (Common-Mode Input Voltage, tj. zakres zmian wejściowego sygnału wspólnego, czyli identycznego dla obu wejść, w którym wzmacniacz pracuje prawidłowo), lecz ponownie uczulamy, iż nie wszystkie wzmocniacze dobrze sobie radzą z pracą w takich warunkach. Wzmocnienie i offset to parametry wspólne dla wzmacniaczy i komparatorów. Niemniej jednak praca przełączająca komparatora oznacza, że posiada on charakterystyki, które nie odnoszą się do standardowego użycia analogowego. Charakterystyki przełączania pokazane



Rysunek 5. Opóźnienie propagacji komparatora



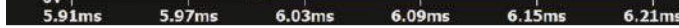
1

1

1

1

czami zwykle redukuje szybkość narastania



1

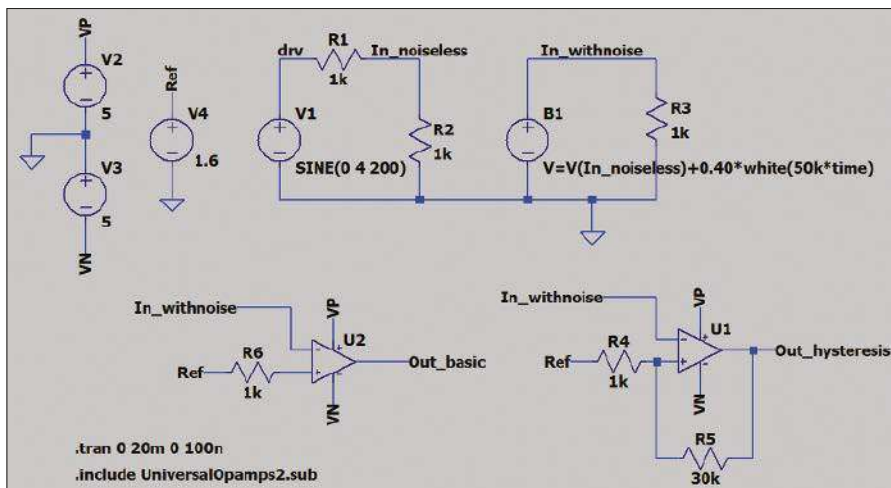
przełączać się szybciej.

Symulacje LM393 i LM741

problemem w tej roli gwarant, a jako naprzeciw



Rysunek 8. Charakterystyka przełączania komparatora z histerezą



Rysunek 9. Schemat symulacji LTSpice używający uproszczonego modelu wzmacniacza jako komparatora do pokazania różnic pomiędzy układami z histerezą i bez

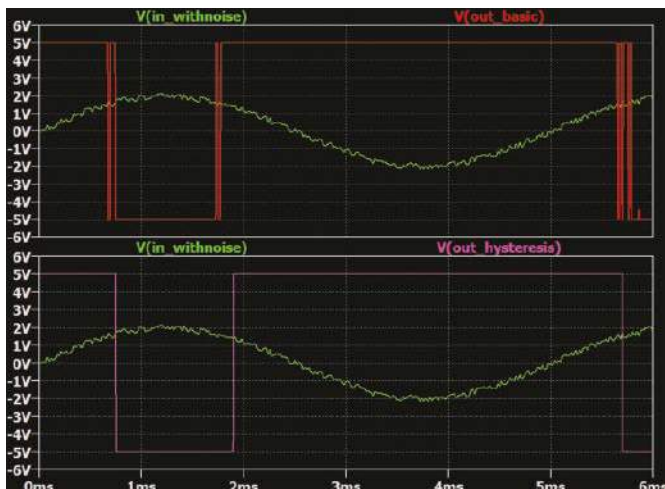
referencyjne użyty jest środek zakresu napięcia zasilania. Układ LM393 posiada wyjście z otwartym kolektorem więc konieczny jest rezystor podciągający R1. Wynik symulacji pokazano na rysunku 7. Widzimy wyraźnie, że napięcie wyjściowe układu LM741 wynosi ± 1 V względem linii zasilania, jak na forum słusznie stwierdził to Qasim. Jest to naturalna właściwość LM741, lecz nie jest to typowe dla wszystkich modeli wzmacniaczy pracujących w roli komparatora – dostępne są wzmacniacze z wyjściem pracującym w pełnym zakresie zasilania (tzw. rail-to-rail) i zgodnie z nazwą potrafią zbliżyć napięcie wyjściowe do skrajnych wartości zakresu zasilania. Dla porównania napięcie wyjściowe układu z LM393 wynosi 20 V bez obciążenia, i schodzi do 0,15 V (napięcie nasycenia kolektor-emiter dla tranzystora o otwartym kolektorze). Symulacja udowadnia także, iż wzmacniacz LM741 odpowiada znacznie wolniej niż komparator LM393, jak wspomniano już o tym wcześniej. Podsumowując,

wzmacniacze operacyjne mogą być używane jako komparatory, jednak nie bez kłopotów i w zastosowaniach o relatywnie małej szybkości. Wzmacniacze również mogą nie radzić sobie z sygnałami wejściowymi o dużej różnicy, co jest powszechnym wymaganiem stawianym układom stosującym dedykowane komparatory. Informacje na temat obsługiwanych sygnałów wejściowych wzmacniaczy znajdziemy oczywiście w ich kartach katalogowych. Warto ponadto wspomnieć, iż uzyskanie z wzmacniacza odpowiedniego sygnału wyjściowego dostosowanego do używanej w naszym systemie logiki może być również trudniejsze niż z komparatora dedykowanego do sterowania wymaganymi wejściami logicznymi.

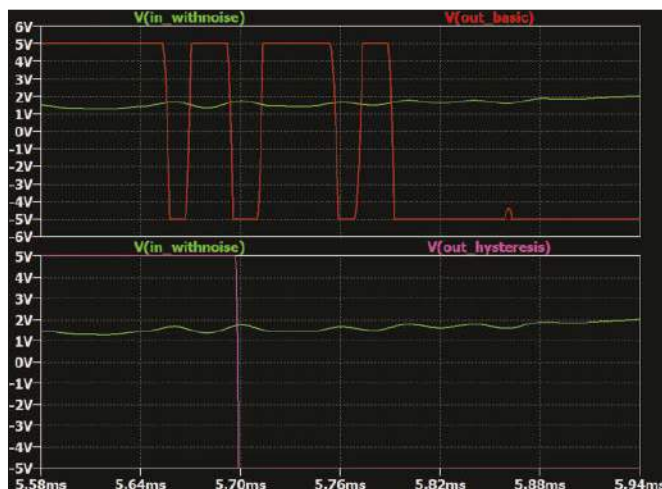
Histereza

Komparator używany z pojedynczym progiem (wartością referencyjną), jak to było omawiane na poprzednich przykładach, może przełączać się wielokrotnie w sytuacji, gdy

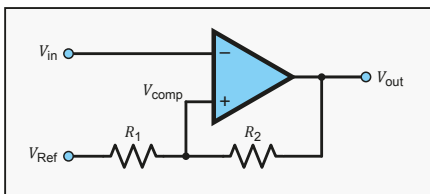
zastumiony, powolny sygnał przekracza próg. Zliczanie każdego przełączenia progu czy też „drżenie” sygnału wyjściowego gdy sygnał wejściowy znajduje się bardzo blisko tego progu jest zjawiskiem niepożądanym. Problem ten może zostać ominięty poprzez użycie dwóch progów V_{TH} i V_{TL} . Różnicę pomiędzy V_{TH} a V_{TL} nazywamy histerezą. Komparator z histerezą może zostać zbudowany poprzez zastosowanie dodatniego sprzężenia zwrotnego do lekkiego przesunięcia progu, w zależności od tego, w jakim stanie aktualnie znajduje się komparator. Komparatory z histerezą są również zwane przerzutnikami Schmitta. Odpowiedź wejścia-wyjścia komparatora z histerezą jest pokazana na rysunku 8. Gdy sygnał wejściowy przekroczy V_{TH} komparator się przełączy, lecz zachowa swój stan, gdy sygnał zmaleje poniżej V_{TH} . Do ponownego przełączania komparatora sygnał wejściowy musi przejść przez V_{TL} . Jeśli poziom szumów wejściowych jest nam znany, histereza może zostać ustawiona nieco powyżej tego poziomu a komparator nie przełączy się pod ich wpływem. Rysunek 9 prezentuje schemat symulacji LTSpice dla porównania obwodów komparatorów z histerezą i bez. Komparatory używają ogólnych, standardowych modeli wzmacniaczy (UniversalOpamps2), które są dostarczane z programem LTSpice, nie zaś dedykowanych modeli elementów LM393 czy LM741. Jest to w zupełności wystarczające do stworzenia przebiegów ilustrujących ogólną zasadę. Rysunki 10 i 11 prezentują rezultat podania tego samego, zastumionego sygnału do prostego komparatora (układ z U2 na rysunku 9) oraz do takiego z zaimplementowaną histerezą (układ z U1). Jak wcześniej wspomniano, podstawowy komparator przełączy się wiele razy – za każdym razem, gdy zastumiony sygnał przekroczy próg. Komparator z histerezą przełącza się „czysto”. Rysunek 10



Rysunek 10. Odpowiedź dwóch komparatorów ze schematu symulacyjnego z rysunku 9



Rysunek 11. Zoom, czyli powiększony na osi czasu widok jednego z miejsc przejścia przez próg z rysunku 10



Rysunek 12. Użycie dodatniego sprzężenia w celu dodania histerezy do komparatora

ilustruje zachowanie prostego komparatora, który przełącza się w wielu różnych momentach, pod wpływem losowości szumu. Na rysunku 11 możemy szczegółowo zaobserwować moment przejścia (lub tak naprawdę przejść, w przypadku prostego komparatora) progów.

Budowa układu

Jak pokazano na rysunku 12, komparator z histerezą może zostać wykonany przy użyciu komparatora z dodatnim sprzężeniem zwrotnym. Jako że progi zależą od napięć wyjściowych komparatora, powinny one być bardzo dokładnie kontrolowane. Analizując dalej układ z rysunku 12 widzimy, iż punkt przełączenia V_{comp} zależy od V_{ref} oraz V_{out} . V_{ref} zwykle jest stałe lecz V_{out} zależy od stanu w jakim znajduje się komparator. V_{out} może mieć więc jedną z dwóch wartości. Załóżmy, że oznaczmy je jako $\pm V_0$ (sygnał dodatni i ujemny ma zastosowanie do obwodów z zasilaniem symetrycznym). Ponadto załóżmy, że $V_{in} < V_{comp}$ więc $V_{out} = +V_0$ (zauważ, że V_{in} podawane jest na wejście odwracające). Jeśli V_{in} powoli narasta, warunek jest spełniony do momentu, gdy $V_{in} = V_{comp} = V_{TH}$ (patrz rysunek 7), gdzie:

$$V_{TH} = \frac{R_2}{R_1 + R_2} V_{ref} + \frac{R_1}{R_1 + R_2} V_0$$

Równanie jest wyprowadzone z zsumowania dwóch równań dzielników napięć. Najpierw należy określić wpływ V_{ref} na V_{comp} z $V_{out} = 0$ a następnie wpływ V_{out} na V_{comp} z $V_{ref} = 0$. Ma tu zastosowanie zasada superpozycji. Przy przełączeniu w punkcie $V_{comp} = V_A$ wyjście zmienia się w $V_{out} = -V_0$, zmieniając punkt przełączenia, $V_{comp} = V_B$.

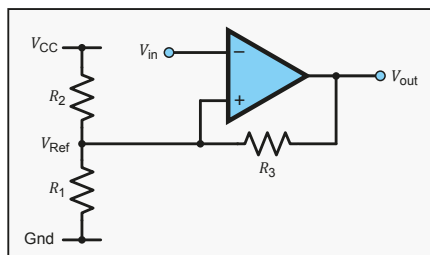
$$V_{TL} = \frac{R_2}{R_1 + R_2} V_{ref} - \frac{R_1}{R_1 + R_2} V_0$$

V_{out} pozostanie na poziomie $-V_0$ dopóki wejście nie spadnie ponownie poniżej V_{comp} . Różnica polega na zmianie progów przełączenia. Histereza (V_H) jest określona więc przez:

$$V_H = V_{TH} - V_{TL} = \frac{2R_1}{R_1 + R_2} V_0$$

Przełączanie jest symetryczne wokół średniej z V_A i V_B , która wynosi:

$$\frac{R_2}{R_1 + R_2} V_{ref}$$



Rysunek 13. Dzielnik potencjału może być użyty w połączeniu z rezystorem sprzężenia zwrotnego determinując wartości progów komparatora (układ dla zasilania pojedynczego)

Jeśli wartość R_2 jest dużo większa niż R_1 , średnia z V_A i V_B jest w przybliżeniu równa V_{ref} . W tych warunkach komparator przełącza się w punktach $V_H/2$ powyżej i poniżej V_{ref} . Dla obwodu z rysunku 9 progi przełączania to $1,548 \pm 0,161$ V (1,71 V i 1,39 V). Zauważ, że nie jest to dokładnie wartość referencyjna 1,6 V. Jeśli komparator jest zasilany napięciem pojedynczym, wyjście (w idealnym przypadku) przełącza się pomiędzy $+V_0$ i 0 V zamiast $+V_0$ i $-V_0$. Formuła na V_{TH} jest taka sama, lecz $-V_0$ zamieniamy na zero w równaniu na V_{TL} , które przyjmuje postać:

$$V_{TL} = \frac{R_2}{R_1 + R_2} V_{ref}$$

Jeśli wartość R_2 jest dużo większa niż R_1 to dolny próg (V_{TL}) jest blisko V_{ref} , a nie w przybliżeniu pomiędzy progami, jak to ma miejsce w obwodach o zasilaniu symetrycznym. Może to stanowić źródło wątpliwości przy układach zasilanych napięciem pojedynczym.

Dzielniki napięcia

V_{ref} dla układu z rysunku 12 może zostać uzyskany z dzielnika napięcia wpiętego pomiędzy linie zasilania. Aby zapobiec wpływowi elementów determinujących wartość progów (R_1 i R_2 na rysunku 12) na dzielnik napięcia i wpływania na napięcie referencyjne, należy zadbać gdyby rezystory te miały relatywnie niską rezystancję. Alternatywnie, można wykorzystać interakcję pomiędzy dzielnikiem napięcia a rezystorem sprzężenia zwrotnego do ustawienia dwóch progów za pomocą układu z rysunku 13. Jeśli zasilamy obwód z rysunku 13 z pojedynczego napięcia i założymy, że napięcie wyjściowe pokrywa cały zakres napięcia wejściowego (rail-to-rail, od V_{CC} do 0 V), łatwo możemy wyznaczyć te dwa progi. Gdy na wyjściu komparatora mamy 0 V, R_3 jest podłączony do masy równolegle do R_1 . Wartość równolegle połączonych tych dwóch rezystorów R_{P13} wynosi, stosując równanie dla dwóch rezystorów równoległych:

$$R_{P13} = \frac{R_1 R_3}{R_1 + R_3}$$

Pliki LTSpice omawiane w tym artykule są dostępne do pobrania ze strony EPE (www.epemag3.com/proj/0519.html).

Rezystancja ta tworzy dzielnik napięcia z rezystorem R_2 i ustawia dolny próg:

$$V_{TL} = \frac{R_{P13}}{R_{P13} + R_2} V_{CC}$$

Gdy napięcie wyjściowe wynosi V_{CC} , R_3 jest połączony równolegle z R_2 . Wartość równolegle połączonych rezystorów R_{P23} wynosi:

$$R_{P23} = \frac{R_2 R_3}{R_2 + R_3}$$

Rezystancja ta buduje dzielnik napięcia z rezystorem R_1 i determinuje górny próg przełączania:

$$V_{TH} = \frac{R_1}{R_1 + R_{P23}} V_{CC}$$

Jeśli zakres napięć wyjściowych nie pokrywa się z zakresem napięcia zasilania, wtedy potrzebne są bardziej skomplikowane równania, które uwzględniają napięcia wyjściowe dwóch komparatorów oraz V_{CC} . Progi przełączania dla obwodu z rysunku 13 niekoniecznie rozłożone są w równej odległości od napięcia otwartego obwodu dzielnika R_1 – R_2 . Przykładowo, jeśli $V_{CC} = 5$ V i dobierzemy rezystory $R_1 = 5$ kΩ oraz $R_2 = 25$ kΩ to otrzymamy dzielnik napięcia dający wartość 833 mV bez dodatkowych obciążeń. Dokładając rezystor sprzężenia zwrotnego $R_3 = 100$ kΩ, otrzymujemy R_{P13} wynoszące 4,762 kΩ zaś R_{P23} równe jest 20 kΩ. Napięcia progowe to $V_{TL} = 0,800$ V, a $V_{TH} = 1,00$ V (co wynika z powyższej formuły). Mamy więc całkowitą histerezę o wartości 200 mV, lecz dolny próg (w tym przykładzie) jest znacznie bliżej wartości dzielnika napięcia otwartego obwodu niż górny próg. ■

Ian Bell

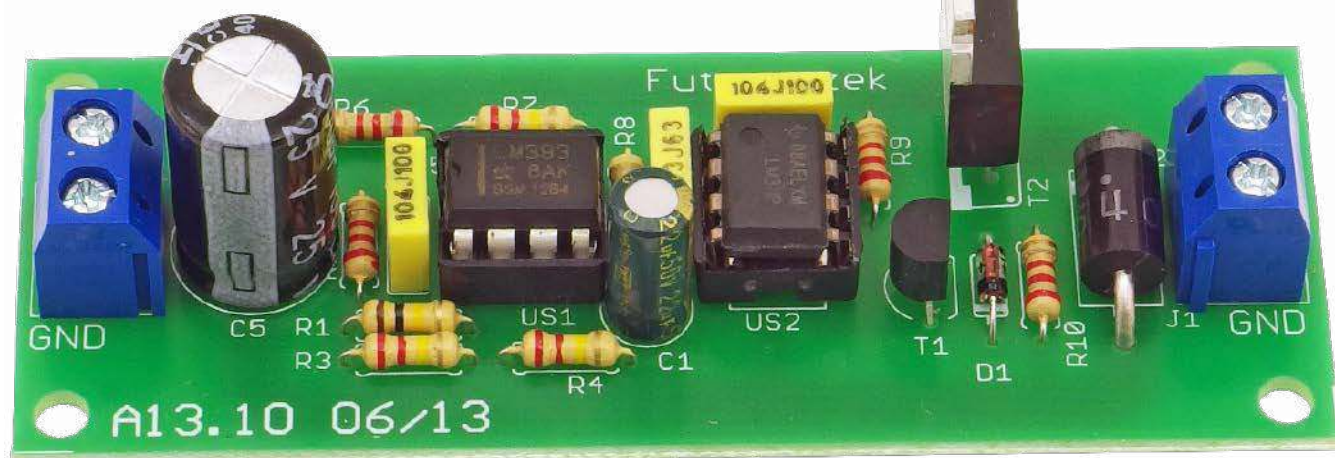
Ten artykuł był wcześniej opublikowany na łamach „Everyday Practical Electronics”, maj 2019 (www.epemag3.com).

REKLAMA

AVTEDU625

PIPEK DRĘCZYCIEL

sklep.avt.pl



Sterownik pulsującej taśmy LED



Taśmy z diodami LED służą nie tylko do oświetlania, lecz również do dekoracji. Na przykład wolno pulsujące światło może dać naprawdę ciekawy efekt wizualny, chociażby w witrynach sklepowych. A to wszystko bez użycia mikrokontrolera!

Do czego to służy?

Zadaniem tego układu jest cykliczne ściemnianie i rozjaśnianie odcinka taśmy LED, przystosowanej do zasilania napięciem 12 V. Obciążeniem może być również coś innego, na przykład żarówka albo dioda LED wysokiej mocy z odpowiednim rezystorem ograniczającym jej prąd. Regulacja odbywa

się poprzez zmianę wypełnienia impulsów, czyli PWM. Na płytce umieszczono odpowiedni element wykonawczy, jak i odpowiednie zabezpieczenia dla niego w wypadku podłączenia obciążenia indukcyjnego, na przykład silnika prądu stałego

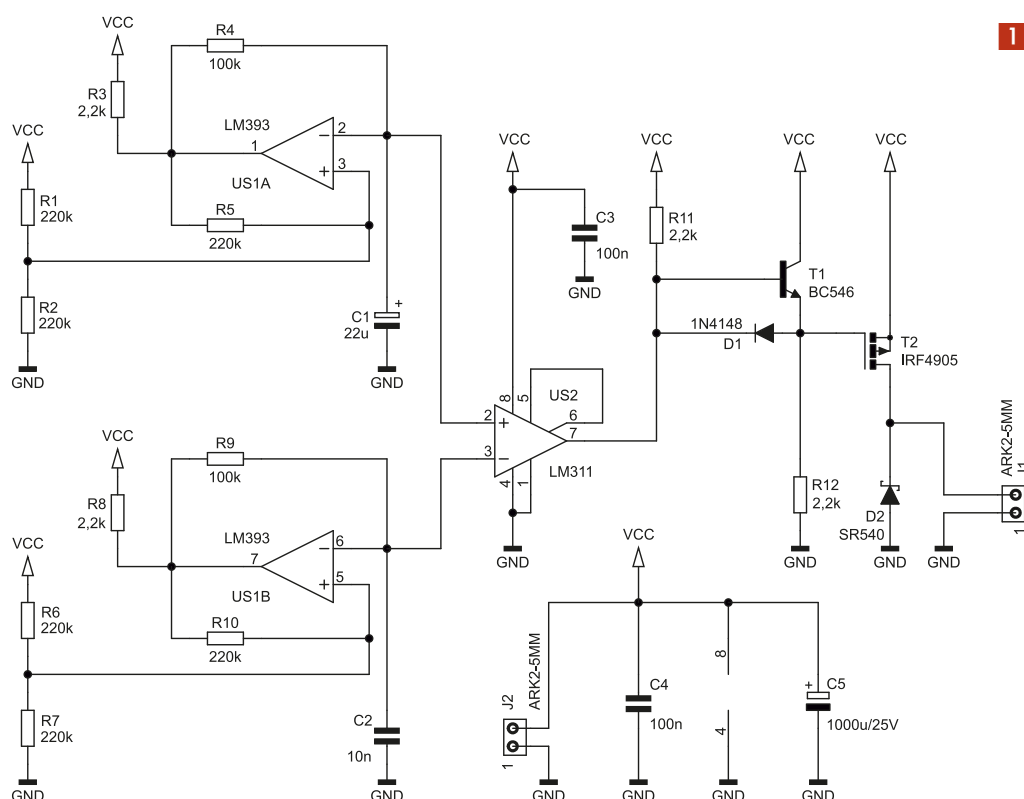
Układ jest przystosowany do zasilania napięciem 12 V i takie samo napięcie podaje

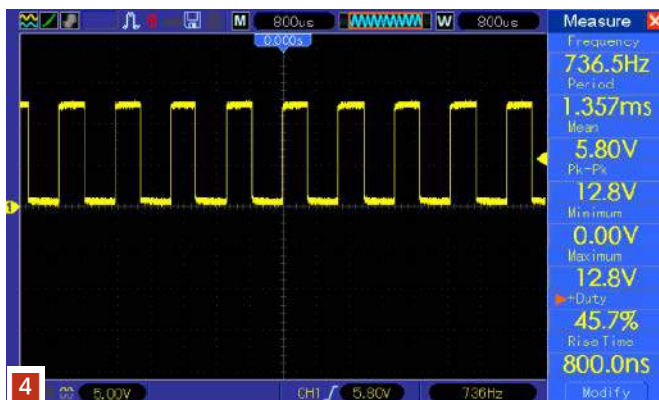
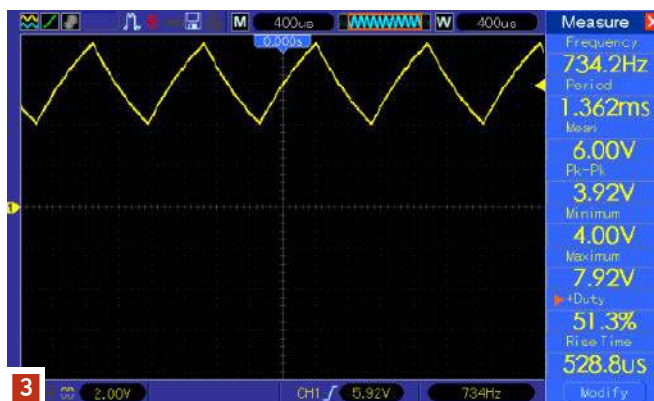
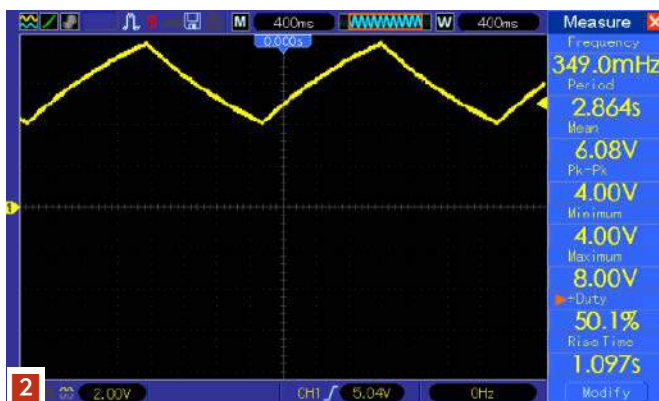
odbiornikowi. Pojedynczy cykl trwa około 2,8 s i w tym czasie sterowana taśma w pełni rozjaśnia się, a potem gaśnie. Nie ma przy tym nieestetycznego efektu skokowego rozjaśniania, jaki znamy z prostych sterowników wykorzystujących mikrokontrolery z ośmiobitowym generatorem PWM.

Jak to działa?

1 Schemat układu można zobaczyć na **rysunku 1**. Składa się z czterech bloków, a trzy z nich zostały zrealizowane na komparatorach. Wszystkie bloki są zasilane niestabilizowanym napięciem stałym o wartości około 12 V, zaś podczas wykonywania oscylogramów na złącze zasilające układ było podane równe 12 V. Pierwsze dwa bloki mają identyczną topologię, różnią się tylko wartością jednego elementu. Dokładniej, są to generatory sygnału prostokątnego, lecz są tutaj wykorzystywane w nieco inny sposób. Dla uproszczenia, szczegółowo omówię tylko górny.

Wyjście komparatora LM393 jest typu otwarty kolektor, więc potrafi jedynie „wciągać” prąd. Rezystor R3 umożliwia mu przyjmowanie wysokiego stanu





ponownie zmodyfikowany za sprawą rezystora R5.

Jednak w tym układzie jest dla nas ciekawy przebieg napięcia na kondensatorze C1, którego oscylogram widnieje na **rysunku 2**. Ma on postać wolno-zmiennego sygnału trójkątnego (dokład-

w rzeczywiście, ponieważ wartości minimalne i maksymalne tych sygnałów nie są idealnie jednakowe, jak również sam komparator ma pewien offset napięciowy, to uzyskuje się nieco węższy zakres. W układzie prototypowym zaobserwowano od 2% do 98%. Dobranie rezystorów w generatorach o mniejszej tolerancji (np. 1%) oraz komparatorów o mniejszym offsecie napięciowym pozwoliłoby na nieznaczną poprawę tego wyniku, lecz byłoby to okupione znaczącym wzrostem ceny układu – a nie o to chodzi w prostym, analogowym sterowniku.

Ostatnim blokiem jest tranzystor wykonawczy T2, który jest typu MOSFET z kanałem P, wraz z układem sterowania bramką. Otwieranie tranzystora realizuje tranzystor wyjściowy komparatora US2, który „wyciąga” prąd poprzez diodę D1. Z kolei zatykanie tranzystora T2 jest przyspieszane przez tranzystor T1, który pracuje tu w roli wtórnika napięciowego. Jego rolą jest zmniejszenie rezystancji przełączającej pojemność bramki – gdyby robił to tylko rezystor R11, ten proces trwałby zdecydowanie dłużej, a to wiązałoby się z wydzielaniem dużej ilości ciepła w tranzystorze T2. R12 stanowi obciążenie emitera T1, aby mógł przez niego płynąć prąd polaryzujący go. Dioda D2 zabezpiecza T2 przed uszkodzeniem, jakie mogłoby

logicznego (czyli potencjału zbliżonego do napięcia zasilania) w chwilach, kiedy tranzystor wyjściowy jest zatkany. Dostarcza on prądu, który rozplywa się po pozostałych fragmentach tego obwodu.

Komparator US1A jest zasilany pojedynczym napięciem, więc do poprawnej pracy potrzebuje „sztucznej masy”, którą wytwarza dzielnik oporowy złożony z rezystorów R1 i R2. Jego rezystancja wewnętrzna jest równa równoległemu połączeniu tych rezystorów, czyli wynosi 110 kΩ. Rezystor R5 wprowadza dodatnie sprzężenie zwrotne, ponieważ dostarcza prąd do wyjścia tego dzielnika lub ten prąd zeń odbiera, co ma wpływ na potencjał wejścia nieodwracającego. Dzieje się tak, ponieważ rezystancja wewnętrzna dzielnika jest stosunkowo wysoka. Dlatego potencjał wspomnianego węzła rośnie, kiedy napięcie na wyjściu jest wysokie oraz spada, kiedy jest niskie.

Z kolei rezystor R4 przeładowuje kondensator C1, a to z kolei ustala częstotliwość generowanego sygnału. Podniesienie potencjału wyjścia komparatora powoduje powolne ładowanie C1 tak długo, aż nie nastąpi przekroczenie progu przełączenia, ustalonego przez dzielnik R1+R2 z potencjałem zmodyfikowanym przez R5. Kiedy komparator przełączy się i potencjał jego wyjścia gwałtownie spadnie, C1 zacznie być rozładowywany. Jak długo? Tutaj mechanizm jest identyczny: dopóki nie zostanie przekroczony próg przełączenia,

niej: quasi-trójkątnego, gdyż przeładowywanie kondensatora ma przebieg wykładniczy). Jego częstotliwość wynosi około 0,35 Hz, czyli okres jest równy 2,857 s. Ponadto, ma równe zbocza (narastające trwa tyle samo co opadające) i dobrze ustaloną wartość minimalną oraz maksymalną. To bardzo ważne w tym układzie, o czym dalej. Ten sygnał odpowiada za zmianę poziomu jasności sterowanej taśmy.

Drugi generator, który różni się wyłączeniem pojemnością kondensatora, generuje przebieg trójkątny o przebiegu widocznym na **rysunku 3**. Jego częstotliwość wynosi około 730 Hz, również ma równe zbocza narastające i opadające. Ma również bardzo zbliżone wartości minimalne i maksymalne co przebieg z **rysunku 2**. Ten sygnał przekształca się w sygnał PWM o zmieniającym się wypełnieniu.

Oba te sygnały dostaje na swoje wejścia komparator US2, którego zadaniem jest wytworzenie sygnału prostokątnego o zmieniającym się cyklicznie wypełnieniu – jego fragment jest widoczny na **rysunku 4**. Szybszy sygnał zostaje przez komparator przekształcony na prostokątny, a moment przełączenia tegoż ustala sygnał wolniejszy. Ponieważ ten drugi również ulega zmianom, toteż moment przełączenia zmienia się i sygnał ma rosnące lub malejące wypełnienie.

W teorii, zakres zmian tego wypełnienia powinien wynosić od 0 do 100%. Jednak

Wykaz elementów:

Rezystory:

R1, R2, R5...R7, R10 220 kΩ
R3, R8, R11, R12 2,2 kΩ
R4, R9 100 kΩ

Kondensatory:

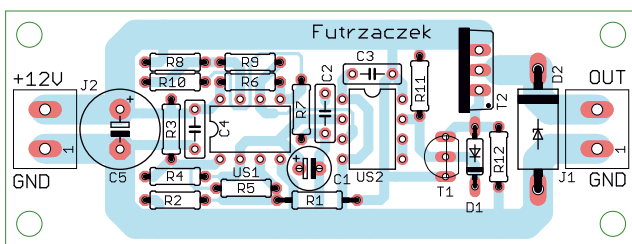
C1 22 μF/25V THT raster 2,5 mm
C2 10 nF raster 5 mm
C3, C4 100 nF raster 5 mm
C5 1000 μF/25 V THT raster 5 mm

Półprzewodniki:

D1 1N4148
D2 SR540
T1 BC546 lub podobny
T2 IRF4905
US1 LM393 DIP8
US2 LM311 DIP8

Inne:

J1, J2 ARK2 5 mm
Dwie podstawki DIP8



spowodować przełączanie obciążenia o charakterze indukcyjnym.

Uzyskano w ten sposób relatywnie szybkie przełączanie tranzystora T2. W układzie prototypowym zmierzono, że czas narastania napięcia na bramce T2 wynosi 440 ns, zaś czas opadania około 500 ns.

Montaż i uruchomienie

Układ prototypowy został zmontowany na jednostronnej płytce drukowanej

o wymiarach 80×30 mm. Wzór jej ścieżek i schemat montażowy przedstawia **rysunek 5**. W odległości 3 mm od krawędzi płytki znajdują się otwory montażowe o średnicy 3,2 mm każdy.

Montaż elementów polecam wykonać w właściwej kolejności, czyli zaczynając od elementów w obudowach o najniższej wysokości, takich jak rezystory i dioda D1. Pod układy scalone warto zastosować podstawki by nie uległy one przegrzaniu podczas lutowania.

Prawidłowo zmontowany układ zaczyna działać od razu po włączeniu zasilania, którym powinno być napięcie stałe o wartości około 12 V, podłączone do zacisków złącza J2. Należy pamiętać o zachowaniu prawidłowej polaryzacji. Z kolei sterowana taśma LED (lub inny odbiornik) powinna być podłączona do złącza J1. Pobór prądu przez sam układ, bez podłączonego obciążenia, zmienia się w granicach 10...16 mA.

Możliwa jest mała modyfikacja układu, która polega na zmianie częstotliwości powtarzania cykli rozjaśniania i ściemniania, co można uczynić wymieniając kondensator C1 na egzemplarz o innej pojemności. Na przykład 100 µF/25 V powinien dać pięciokrotne wydłużenie całego cyklu. Z kolei zmianę częstotliwości PWM można uzyskać poprzez dobór pojemności C2. ■

Michał Kurzela

michal.kurzela@ep.com.pl

REKLAMA

Sięgnij po archiwalne wydania ELEKTRONIKI dla WSZYSTKICH

Przesyłka
GRATIS

Zamów wygodnie na
www.UlubionyKiosk.pl

Przesyłka
GRATIS

Zamów wygodnie na
www.UlubionyKiosk.pl

Przesyłka
GRATIS

Czterokanałowa bezdotykowa tablica rozdzielcza dla świata po Covidzie



Covid-19 wpoił nam zasady higieny, które warto zachować.

W ramach środków ostrożności należy w miarę możliwości unikać dotykania jakichkolwiek przedmiotów, w tym przełączników elektrycznych. W miejscach publicznych, szkołach i urzędach wszyscy dotykają przełączników, co może powodować przenoszenie zarazków z jednej osoby na wiele innych.

Oto rozwiązanie w postaci czterokanałowej bezdotykowej tablicy rozdzielczej. Arduino Uno jako serce projektu, współpracujące z czterema ultradźwiękowymi czujnikami. To proste, tanie i nowatorskie rozwiązanie. Przełączniki można włączać i wyłączać naprzemiennie, umieszczając dłoń obok ultradźwiękowych czujników w urządzeniu, bez dotykania.

Układ i działanie

Na rysunku 1 przedstawiono schemat czterokanałowej centrali bezdotykowej. Składa się ona z czterech ultradźwiękowych czujników HC-SR04 (SEN1-SEN4), czterokanałowej płytki przekaźnikowej (RM1) i płytki Arduino Uno (Board1).

Mikrokontroler na płycie Arduino Uno odbiera sygnały przełączające z czujników, a kod programu wysyła je do czterokanałowego modułu przekaźnika i w ten sposób włącza/wyłącza urządzenia. Ultradźwiękowy czujnik HC-SR04 jest wykorzystywany do wykrywania czynności przełączania.

Szczegóły połączeń czujnika ultradźwiękowego HC-SR04 pokazano na rysunku 2. Posiada następujące cztery styki:

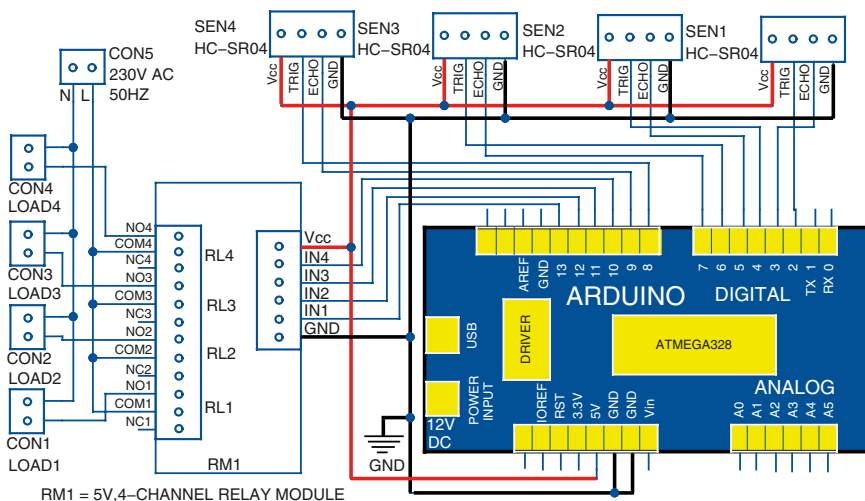
VCC: +5 V DC

Trig: Wyzwalanie (wejście)

Echo: Echo (wyjście)

GND: GND

Czujniki można łatwo połączyć z Arduino Uno. Gdy ręka znajdzie się w odległości ok. 5 cm od czujnika, program w Arduino odczyta sygnał i przekaźnik zostanie zasilony,



Rysunek 1. Schemat układu bezdotykowej tablicy rozdzielczej

Rodzaje tablic rozdzielczych

Tablice rozdzielcze	Szczegóły
Tradycyjne tablice rozdzielcze	- Wysoka wytrzymałość - W pełni ręczna obsługa
Moduł przekaźników sterowanych pilotem na podczerwień	- Obsługa za pomocą pilota na podczerwień - Sterowanie urządzeniem jest realizowane przez przekaźnik - Płytki przekaźników działa przy napięciu 12 V prądu stałego
Moduł przekaźników sterowanych pilotem radiowym	- Działa na zasadzie pilota bezprzewodowego RF (częstotliwość radiowa) 433 MHz - Płytki przekaźników działa przy napięciu 12 V prądu stałego - Sterowanie urządzeniem jest realizowane przez przekaźnik
Bezprzewodowy przełącznik Wi-Fi	- Bezprzewodowy przełącznik Wi-Fi współpracuje z urządzeniami Alexa, Google Home, Nexa - Aplikacja mobilna - Standard bezprzewodowy: bezpieczeństwo 802.11 b/g/n - Przeważnie posiada wewnętrzny zasilacz prądu stałego

Wykaz elementów:

Półprzewodniki:

Board1: Arduino Uno

Pozostałe:

SEN1-SEN4: czujnik HC-SR04

RM1: moduł czterokanałowego przekaźnika 1CO, 5 V

LOAD1-LOAD4: urządzenia 230 V prądu przemiennego

- Przewody/zworki



Rysunek 2. Szczegóły styków czujnika HC-SR04

aby wykonać operację przełączania. Gdy ręka zostanie ponownie umieszczona w pobliżu czujnika, Arduino ponownie odczyta sygnał i przełącznik przekaźnika wróci do poprzedniego położenia. Wszystkie inne czujniki ultradźwiękowe działają podobnie.

W tabeli 1 przedstawiono porównanie różnych typów ogólnie dostępnych tablic rozdzielczych. W tabeli 2 przedstawiono połączenia pinów Arduino Uno z komponentami.

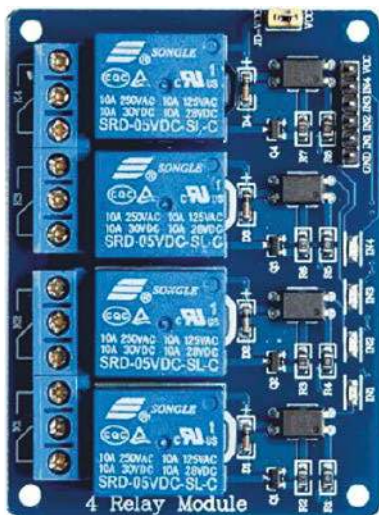
Płytki przekaźnika czterokanałowego

Na rysunku 3 przedstawiono płytkę przekaźnika czterokanałowego. Jest to płytka o niskim poziomie napięcia (5 V), której każdy kanał pobiera prąd sterujący 15–20 mA. Może być stosowany do zasilania różnych urządzeń pobierających duży prąd, ponieważ jest wyposażony w przekaźniki wysokoprądowe pracujące pod napięciem 250 V, 10 A prądu przemianowego lub 30 V, 10 A prądu stałego. Posiada on standardowy interfejs, który może być sterowany bezpośrednio przez mikrokontroler.

Ten moduł jest optycznie odizolowany od wysokiego napięcia dla zapewnienia bezpieczeństwa i ma cztery przekaźniki SPDT.

Oprogramowanie

Program fourchan.ino jest ładowany do pamięci wewnętrznej Arduino Uno. Program ten jest prosty i łatwy do zrozumie-



Rysunek 3. Płytki przekaźnika czterokanałowego

Połączenia pinów Arduino Uno z komponentami

Piny Arduino Uno	Komponenty	Czujniki HC-SR04
pin 5 v	- Podłączany do styków Vcc wszystkich czterech czujników HC-SR04 - Podłączany do styku Vcc na płytce przekaźnika	
pin 2	HC-SR04 -1 Pin wyzwalający (Trigger)	SEN1
pin 3	HC-SR04 -1 Pin Echo	
pin 4	HC-SR04 -2 Pin wyzwalający (Trigger)	SEN2
pin 5	HC-SR04 -2 Pin Echo	
pin 6	HC-SR04 -3 Pin wyzwalający (Trigger)	SEN3
pin 7	HC-SR04 -3 Pin Echo	
pin 8	HC-SR04 -4 Pin wyzwalający (Trigger)	SEN4
pin 9	HC-SR04 -5 Pin Echo	
pin 13	Wejście 1 płytki przekaźnika	
pin 12	Wejście 2 płytki przekaźnika	
pin 11	Wejście 3 płytki przekaźnika	
pin 10	Wejście 4 płytki przekaźnika	
GND	- Podłączany do styków GND wszystkich HC-SR04 - Podłączany pin GND płytki przekaźnika	

nia. Komentarze są umieszczone na końcu każdego wiersza poleceń.

W pierwszej kolejności należy zdefiniować w kodzie następujące elementy:

- styki echo i trigger czujnika ultradźwiękowego,
- cztery styki przekaźników,
- zmienne pośrednie dla działania przekaźnika,
- zmienne dla wszystkich czterech czujników.

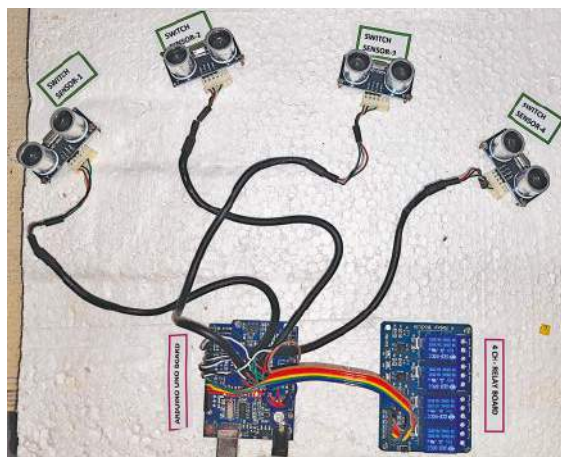
Następnie należy utworzyć funkcję void setup(). Odpowiednie piny należy zainicjować jako wejścia i wyjścia. Następnie należy zainicjować komunikację szeregową. Komunikacja szeregową nie ma nic wspólnego z projektem, służy jedynie do analizowania programu na monitorze.

Funkcja loop() pozwala programowi zmieniać się i reagować na zmiany. Służy do aktywnego sterowania płytką Arduino Uno. Aktywowanie czujnika Echo. Obliczania odległości i wykonywanie operacji włączania/wyłączania. Przekazywanie stanu przełącznika do przekaźnika. Należy napisać powyższy kod również dla wszystkich pozostałych trzech czujników i przekaźników.

Montaż i testowanie

Układ wymaga zasilania prądem stałym o napięciu 12 V i natężeniu 1 A. Podłączmy moduł tak, jak pokazano na schemacie. Należy zachować pewną odległość między czterema czujnikami Echo, tak aby ich działanie nie miało wpływu na sąsiednie czujniki.

Do osłony Arduino Uno i płytki z przekaźnikami można użyć puszek rozdzielczej



Rysunek 4. Prototyp autora

z PCW. Podłączmy urządzenia na prąd przemieniony zgodnie ze specyfikacją przekaźników. Połączenie przewodów fazowego (L) i neutralnego (N) pokazano na schemacie. Prototyp opracowany przez autora pokazano na rysunku 4.

Uwaga: Płytki Arduino Uno i czujniki HC-SR04 są bardzo czułe, dlatego należy obchodzić się z nimi ostrożnie. Należy powoli zbliżać dłoń do czujnika i powoli ją odsuwać.

Uwaga: Należy zachować odpowiednie środki ostrożności podczas pracy z siecią prądu przemianowego. Początkującym eksperymentatorom zaleca się konsultację z ekspertem przed wykonaniem połączeń między przekaźnikami a siecią prądu przemianowego. ■

K. Murali Krishna

Ten artykuł był wcześniej opublikowany na łamach „EFY”, wrzesień 2021 (efymag.com)

Automatyczny dozownik wody dla zwierząt domowych



Budujemy urządzenie, które pomaga w automatycznym dostarczaniu czystej i świeżej wody pitnej dla zwierzęcia, gdy jest to potrzebne. Poprzez zmianę sensorów przetwarzających wielkości nieelektryczne w elektryczne można to rozwiązanie zmodyfikować do wielu innych zastosowań.

Zwierzę domowe jest kochane jak dziecko w rodzinie i potrzebuje dużo uwagi oraz opieki. Ale czasami ze względu na napięty harmonogram, zapominamy o codziennych potrzebach zwierząt. Urządzenie to pomaga w automatycznym dostarczaniu czystej i świeżej wody pitnej dla zwierzęcia, gdy jest to potrzebne.

Schemat ideowy automatycznego dystrybutora wody pokazano na rysunku 1. Składa się on z transformatora X1 obniżającego napięcie do 15 V AC, mostka prostowniczego BR1, stabilizatora napięcia IC1 o wartości 12 V, płytki Arduino Uno Board1, czujnika ultradźwiękowego HC-SR04, czujnika poziomu wody SEN18, dwóch LEDów 5 mm (LED1 i LED2), tranzystora T1 typu 2N2222, przekaźnika RL1 o napięciu 12 V oraz pompy płynu P1 o zasilaniu 12 V DC. Ich funkcje zostały opisane poniżej.

Arduino Uno. Płytką Arduino Uno jest podłączona do czujnika poziomu wody i czujnika ultradźwiękowego, które wykrywają odpowiednio poziom wody i obecność zwierzęcia w pobliżu miski z wodą. Gdy czujnik ultradźwiękowy wykryje obecność zwierzęcia w pobliżu, przekazuje sygnał do Arduino.

Program Arduino porównuje otrzymany sygnał z zaprogramowaną wartością i podaje sygnał na swój cyfrowy pin wyjściowy 13, aby uruchomić (załączyć) przekaźnik RL1, jeśli spełnione są zaprogramowane warunki. W przeciwnym razie przekaźnik pozostaje niewzbudzony.

Zasilacz 12 V. Zasilą Arduino Uno i pompę wodną oraz pomaga sterować przekaźnikiem. Składa się on z transformatora X1 zmniejszającego napięcie z 230 V na 15 V, 2 A, pełnego mostka prostowniczego BR1, kondensatorów C1 i C2, układu scalonego 7812 stabilizującego napięcie 12 V, rezystora R1 o wartości 2,2 k oraz diody LED1 jako wskaźnika. Napięcie prądu przemiennego 15 V z transformatora jest przekształcane na napięcie prądu stałego przez prostownik i filtrowane przez

dwa kondensatory. Układ scalony 7812 stabilizuje napięcie do 12 V prądu stałego.

Czujnik poziomu wody SEN18. Czujnik ten ma dziesięć odsłoniętych ścieżek przewodzących, z których pięć naprzemiennych pasków to ścieżki zasilające, a pozostałe pięć to ścieżki detekcyjne. Gdy czujnik jest zanurzony w wodzie (pionowo), między pięcioma parami ścieżek zasilających i detekcyjnych tworzą się połączenia elektryczne spowodowane przewodnością wody. Czujnik działa więc jak zmienny opornik, którego wartość zmienia się wraz ze zmianą poziomu wody. Czujnik może być zasilany z wbudowanego zasilacza prądu stałego 3,3 V lub 5 V.

Przekaźnik. Ten elektromagnetyczny przełącznik SPDT 12 V zapewnia całkowitą izolację elektryczną między obwodem sterującym a wyjściowym. Jest on sterowany niskim napięciem prądu stałego, ale jego styki mogą wytrzymać duży przepływ prądu przy całkowitej izolacji elektrycznej od delikatnego układu elektronicznego o niskim poborze mocy. Przekaźnik włącza i wyłącza pompę wodną po otrzymaniu impulsu elektrycznego.

Czujnik ultradźwiękowy HC-SR04 Urządzenie wyposażone w nadajnik i odbiornik ultradźwiękowy wykrywa obecność obiektu znajdującego się w odległości od 2 cm do 400 cm. Jego nadajnik ultradźwiękowy wytwarza sygnał o częstotliwości 40 kHz,

a odbiornik odbiera sygnał odbity, jeśli obiekt znajduje się w zasięgu.

Oprogramowanie

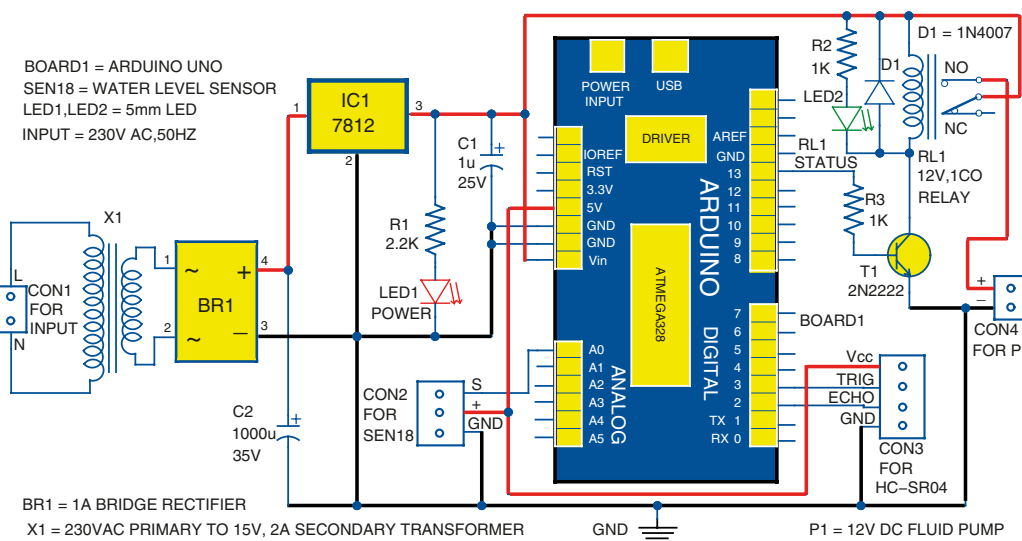
Program Arduino `dispenser.ino`, użyty w prototypie, należy wgrać do Arduino Uno. Do tego celu potrzebne jest oprogramowanie typu open source Arduino IDE.

W programie piny 3 i 13 Arduino są identyfikowane jako wyjście cyfrowe, a pin 2 jako wejście cyfrowe. Pin A0 służy do odczytu sygnału analogowego. Poniżej opisano główne funkcje kodu Arduino.

PinMode(). Konfiguruje określony pin tak, aby zachowywał się jak wejście lub wyjście. Informacje o funkcjach pinów cyfrowych można znaleźć w ich opisie. Można je również wykorzystywać jako piny PWM.

DigitalWrite(). Jeśli pin został skonfigurowany jako wyjście za pomocą funkcji `pinMode( )`, jego napięcie zostanie ustawione na odpowiednią wartość: 5 V (lub 3,3 V na płytce 3,3 V) dla stanu wysokiego, 0 V (masa) dla stanu niskiego.

Port Szeregowy. Służy do komunikacji między płytką Arduino a komputerem lub innymi urządzeniami. Wszystkie płytki Arduino mają co najmniej jeden port szeregowy (znany także jako UART lub USART), a niektóre mają ich kilka.



Rysunek 1. Schemat układu automatycznego dozownika wody

begin(). Określa szybkość transmisji danych w bitach na sekundę (bod) dla szeregowego przesyłania danych. Do komunikacji z komputerem należy używać jednej z tych szybkości: 300, 600, 1200, 2400, 4800, 9600, 14400, 19200, 28800, 38400, 57600 lub 115200.

Można jednak określić inne szybkości transmisji, np. w celu komunikacji przez piny 0 i 1 z komponentem wymagającym określonej szybkości transmisji.

AnalogRead(). Odczytuje wartość z określonego pinu analogowego. Płytki Arduino posiadają wielokanałowy, 10-bitowy przetwornik analogowo-cyfrowy. Oznacza to, że odwzorowuje on napięcia wejściowe z przedziału od 0 do napięcia roboczego (5 V lub 3,3 V) na wartości całkowite z przedziału od 0 do 1023.

println(). Wyświetla dane do portu szeregowego w postaci czytelnej dla człowieka tekstu ASCII, po którym następuje znak powrotu karetki.

Opóźnienie(). Służy do wstrzymania programu na czas (w milisekundach) określony jako parametr.

delayMicroseconds(). Powoduje wstrzymanie programu na czas (w mikrosekundach) określony przez parametr.

Wykaz elementów:

Półprzewodniki:

Board1: płytka Arduino Uno
IC1: 7812, stabilizator napięcia 12 V
BR1: prostownik mostkowy 1 A
D1: dioda prostownicza 1N4007
T1: tranzystor 2N2222 NPN
LED1, LED2: diody LED 5 mm

Rezystory:

(wszystkie 1/4-watowe, ±5% węglowe)
R1: 2,2 kΩ
R2, R3: 1 kΩ

Kondensatory:

C1: 1 μF, 25 V elektrolityczny
C2: 1000 μF, 25 V elektrolityczny

Pozostałe:

RL1: przekaźnik przełączający 12 V
P1: pompa płynów, 12 V prądu stałego
SEN18: 3-pinowy czujnik poziomu wody – czujnik ultradźwiękowy HC-SR04
X1: transformator prądu przemiennego z uzwojeniem pierwotnym 230 V i wtórnym 15 V, 2 A
CON2: 3-stykowe złącze żeńskie Berg dla SEN18
CON3: 4-stykowe złącze żeńskie Berg dla HC-SR04
– 8-stykowe złącze żeńskie Berg dla Arduino
– 2-stykowe złącze dla wtórnego uzwojenia transformatora X1
– uniwersalna płytka drukowana
– przewód/zworka

Obliczanie odległości

Sygnał przechodzi od nadajnika do obiektu, a następnie wraca do odbiornika. Zatem, gdy obliczamy odległość obiektu od czujnika, mnożąc czas i prędkość, musimy podzielić tę odległość przez dwa. Odległość jest obliczana na podstawie zależności:

$$\text{odległość} = (\text{czas} \times \text{prędkość}) / 2$$

W tym przypadku prędkość fali wynosi 340 m/s, czyli 0,0340 cm/mikrosekundę.

Zatem odległość w centymetrach = $(\text{czas} \times 0,0340) / 2$

Obecnie największa wartość, która pozwoli uzyskać dokładne opóźnienie, wynosi 16 383 mikrosekundy. Może się to zmienić w kolejnych wersjach Arduino. W przypadku opóźnień większych niż kilka tysięcy mikrosekund można użyć funkcji delay().

Montaż i testowanie

Na rysunku 2 pokazano rzeczywistej wielkości płytkę drukowaną układu dozownika wody, a na rysunku 3 rozmieszczenie jego elementów. Do zmontowania układu można także użyć uniwersalnej płytki drukowanej.

Po zmontowaniu układu podłączmy napięcie 15 V z uzwojenia wtórnego transformatora do płytki drukowanej przez X1 zaznaczony na rysunku 3. Podłączmy uzwojenie pierwotne transformatora do sieci 230 V prądu przemiennego przez złącze CON1 (nie pokazane na schemacie urządzenia). Nie zapomnijmy przesłać kodu dispenser.ino do Arduino przed włączeniem układu.

Jak pokazano na rysunku 1, pin sygnałowy (S) czujnika poziomu wody jest podłączony do pinu analogowego A0 płytki Arduino Uno, a piny GND i +Vcc są podłączone odpowiednio do pinów masy i 5 V. Czujnik poziomu wody wykrywa poziom wody w misce z wodą dla zwierząt.

Po stronie wyjściowej zastosowano przekaźnik SPDT RL1 o napięciu 12 V. Wysokie napięcie na pinie 13 Arduino wyzwala przekaźnik przez tranzystor przełączający T1 w celu uruchomienia pompy wody P1 na prąd stały. D1 jest używana jako dioda zabezpieczająca (flyback diode), chroniąca tranzystor przed wsteczną SEM (siłą elektromotoryczną). Wskaźnik LED2 służy do sygnalizacji stanu

przełącznika. Pompa P1 jest podłączona jako obciążenie do przełącznika RL1 pomiędzy stykami normalnie otwartym (NO) i wspólnym (C) z szeregowym podłączeniem źródła zasilania 12 V prądu stałego.

Za pomocą czujnika SEN18 należy ustawić progowy poziom wody w misce zwierzęcia. Dlatego przed ostateczną konfiguracją należy najpierw skalibrować czujnik poziomu wody zgodnie z potrzebami użytkownika. Po włączeniu zasilania i wgraniu do Arduino programu calibration.ino można to zrobić, wykonując dwie czynności wymienione poniżej.

Krok 1. Zaznaczyć poziom wody w misce w miejscu, które ma być uznane za pełne.

Krok 2. Sprawdzić wartość na monitorze szeregowym i zanotować wartość pełnego poziomu.

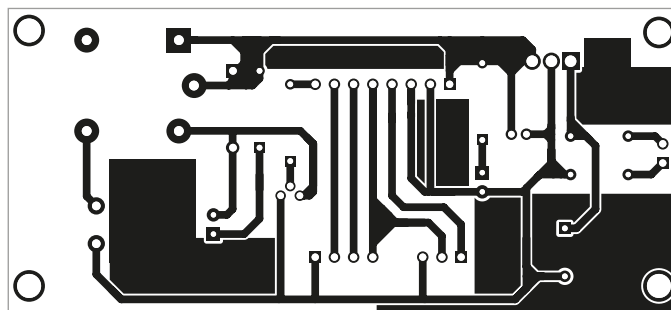
Teraz należy umieścić tę wartość w kodzie dispenser.ino. Skompilujmy kod i prześlijmy go do Arduino. Dzięki temu system będzie działał precyzyjnie.

Po skonfigurowaniu całego układu, włączmy zasilanie. Czujnik ultradźwiękowy stale sprawdzi, czy zwierzę znajduje się w pobliżu miski, aby napić się wody, czy też nie. Gdy zwierzę zbliży się do miski, program sprawdza poziom wody na podstawie sygnału zwrotnego z czujnika. Jeśli czujnik poziomu wody wykryje niski poziom, program Arduino poda sygnał wysoki na cyfrowy pin 13 We/Wy Arduino, co spowoduje włączenie przekaźnika RL1 przez tranzystor T1. Pompa P1 podłączona do styku normalnie otwartego (NO) przekaźnika włącza się i zaczyna tłoczyć wodę do miski. Gdy poziom wody osiągnie żądany poziom, program wysła niski sygnał do wyjścia cyfrowego na pinie 13 w Arduino, co powoduje wyłączenie przekaźnika RL1, a tym samym odłączenie zasilania i zatrzymanie pompy P1.

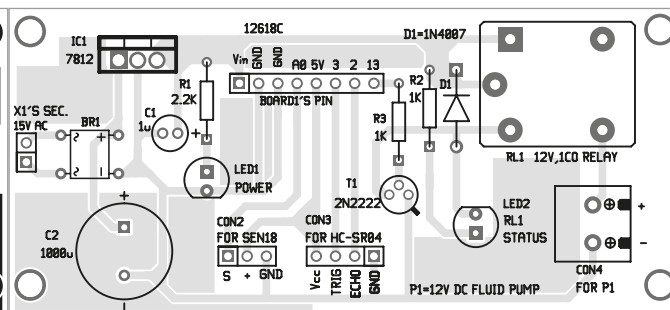
W ten sposób można zapewnić zwierzętom domowym czystą i świeżą wodę w sposób automatyczny, bez konieczności ręcznej interwencji. ■

Tej Vijaykumar Patel

Ten artykuł był wcześniej opublikowany na łamach „EFY”, wrzesień 2021 (efymag.com)



Rysunek 2. Schemat płytki drukowanej dozownika wody



Rysunek 3. Rozmieszczenie elementów na płycie drukowanej

Zapisywanie danych przy użyciu microSD i Raspberry Pi Pico



Opisany tu system rejestracji danych można wykorzystać do połączenia zewnętrznego lub wewnętrznego czujnika z płytką Raspberry Pi Pico.

Projekt ma wiele zalet, między innymi wykorzystuje wewnętrzny czujnik temperatury płytki Raspberry Pi Pico, dzięki czemu nie jest wymagany zewnętrzny czujnik tego typu. Po drugie, łatwo wyjmowana karta microSD może być używana do rejestrowania różnych danych dotyczących temperatury w danym miejscu lub instalacji przemysłowej. Kartę SD można wyjąć i przenieść do laboratorium lub dyspozytorni w celu przeprowadzenia analizy lub wykorzystania w przyszłości.

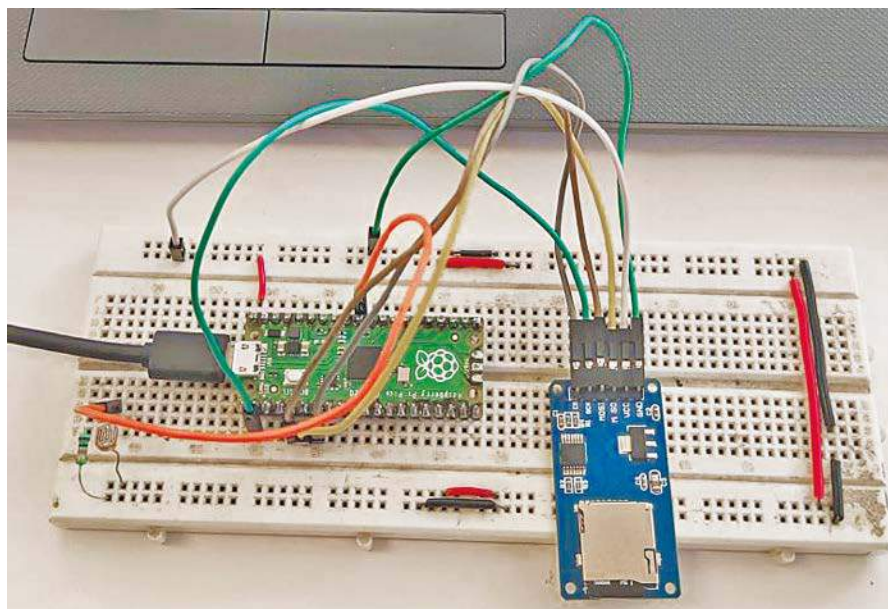
W projekcie uwzględniono również opcjonalne czujniki zewnętrzne, takie jak czujniki siły i światła.

Komponenty wymagane do realizacji projektu przedstawiono w tabeli poniżej.

Na rysunku 1 przedstawiono autorski prototyp zamocowany na płytce prototypowej, a na rysunku 2 – połączenie systemu rejestracji danych z kartą microSD i Raspberry Pi Pico. Można zastosować fotorezystor lub inny odpowiedni czujnik analogowy, taki jak przedstawiony czujnik siły.

W konfiguracji wykorzystano piny szeregowego interfejsu peryferyjnego (SPI) MOSI, MISO, CS i CLK płytki Raspberry Pi Pico Board1 (Płytki#1). Czujnik zewnętrzny (czujnik światła, fotorezystor lub czujnik siły) jest podłączony do pinu GP26 Płytki#1.

Czujnik oparty na fotorezystorze służy do rejestrowania stanu oświetlenia w obszarze, w którym jest zainstalowany. Moduł karty microSD jest połączony z pinami GP1, GP3, GP4 i GP5 płytki Raspberry Pi Pico. VBUS z płytki



Rysunek 1. Prototyp autora

służy do zasilania karty microSD i fotorezystora. Rezystor 10 k pull-down jest podłączony do pinu GP26 Płytki#1 w celu podłączenia czujników zewnętrznych.

Oprogramowanie

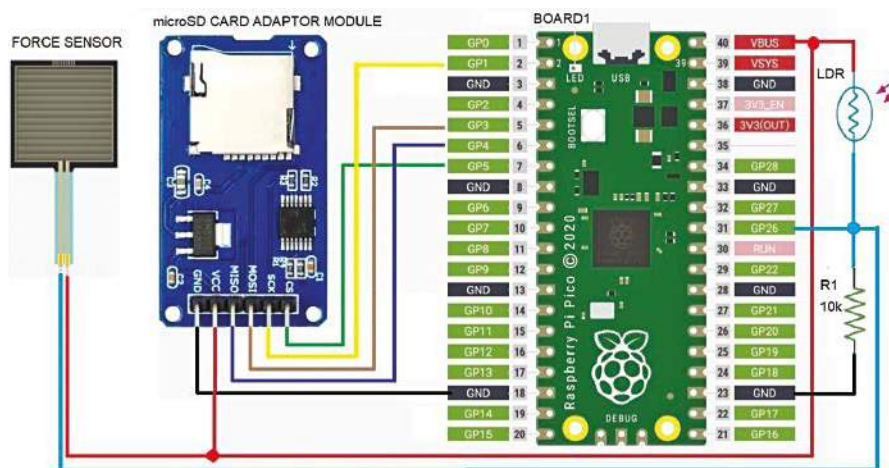
Projekt jest kontrolowany przez program w języku Python opracowany w edytorze Mu. Mu editor jest prostym edytorem kodu, który działa z płytkami Adafruit CircuitPython,

w tym Raspberry Pi Pico, i na większości systemów operacyjnych, w tym Windows i Linux. Program CircuitPython jest oparty o język programowania Python i jest łatwy do zrozumienia.

Do tego projektu wymagane są zarówno Mu, jak i CircuitPython. Edytor Mu można pobrać ze strony [https:// codewith.mu](https://codewith.mu), a następnie zainstalować, korzystając z instrukcji podanych na stronie [https:// learn.adafruit.com](https://learn.adafruit.com).

Wymagane komponenty

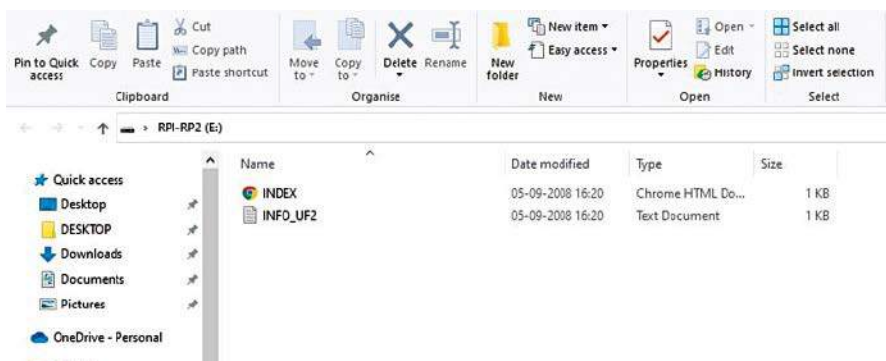
Komponenty	Liczba	Przeznaczenie
Raspberry Pi Pico	1	obliczenia
karta microSD (16 GB) z adapterem	1	przechowywanie danych
fotorezystor (LDR)	1	czujnik zewnętrzny
czujnik siły	1	czujnik zewnętrzny
rezystor 10 kΩ	1	ogranicznik prądu
przewody/zworki	9	połączenia



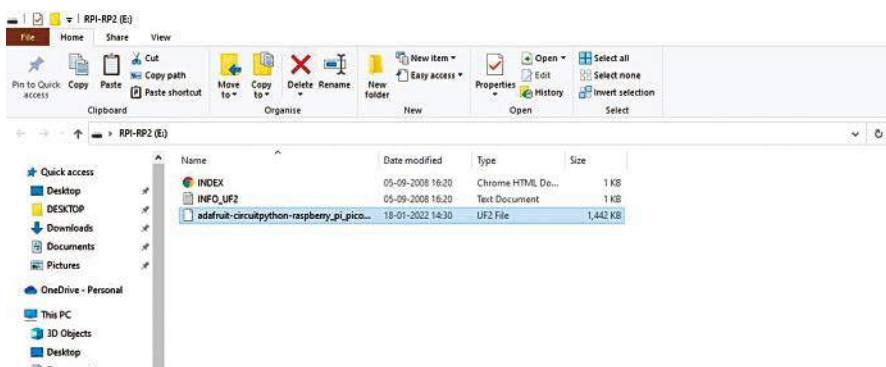
Rysunek 2. Połączenie systemu rejestracji danych z kartą microSD i Raspberry Pi Pico



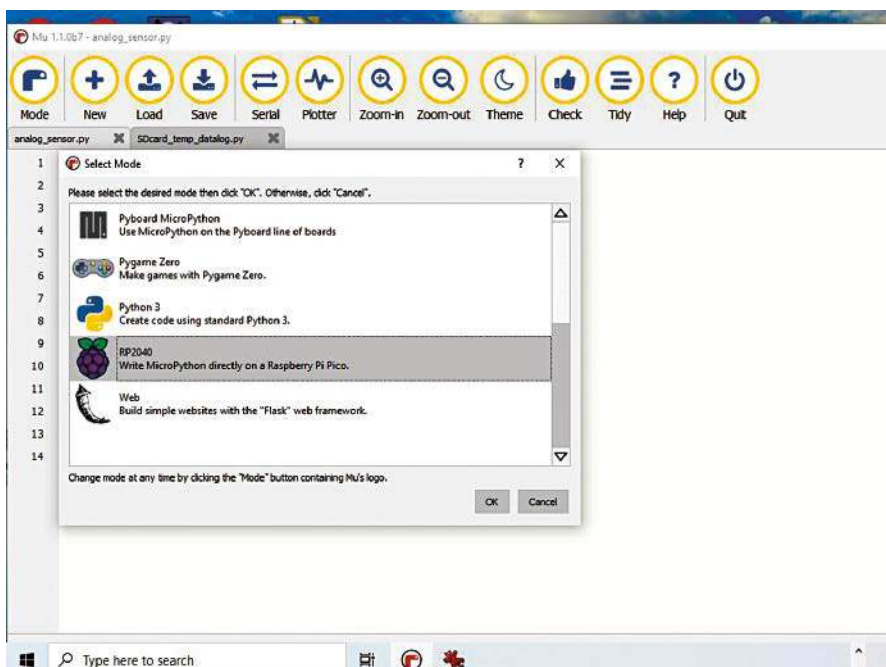
Rysunek 3. Okno edytora Mu



Rysunek 4. Nowe okno napędu



Rysunek 5. Plik CircuitPython w napędzie



Rysunek 6. Wybór trybu pracy

com/getting-started-with-raspberry-pi-pico-circuitpython/ installing-mu-editor. Po zainstalowaniu edytora Mu uruchomimy go, a zobaczymy okno pokazane na rysunku 3.

Podczas testowania tego projektu użyto programu CircuitPython w wersji 7.1.1.1. Płytką CircuitPython wymaga bootloadera o nazwie UF2 (format pamięci USB), który sprawia, że instalacja i aktualizacja CircuitPython jest szybka i łatwa. Pobierzmy najnowszą wersję pliku bootloadera CircuitPython .UF2 ze strony https://circuitpython.org/board/raspberry_pi_pico/.

Instalacja CircuitPython na płytce jest prosta. Wystarczy skopiować i wkleić plik w sposób opisany tutaj. Najpierw naciśniemy przycisk BOOTSEL na płytce Raspberry Pi Pico, wkładając jednocześnie przewód USB płytki RPI do laptopa/komputera PC, i zwolnimy go po pojawieniu się okna nowego napędu (RPI-RP2), jak pokazano na rysunku 4.

Następnie skopiujemy plik adafruit-circuitpython-raspberry_pi_pico-en_US-7.1.1.uf2 pobrany w poprzednim kroku. Wklejmy go do nowego okna, jak pokazano na Rys. 5. Teraz otworzymy okno edytora Mu i zamknijmy wszystkie inne okna.

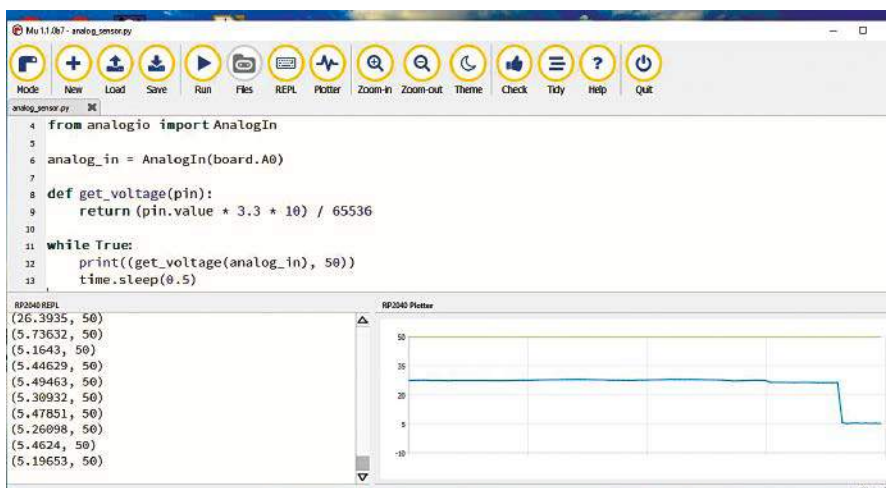
Następnie, pierwszą rzeczą, którą należy zrobić w programie Mu, jest wybranie z menu głównego opcji „Tryb”. W tym projekcie należy wybrać opcję RP2040 w oknie Select Mode, jak pokazano na rysunku 6.

Rozpocznijmy kodowanie w oknie edytora Mu lub wczytajmy istniejący kod. Można załadować kod analog_sensor.py, który jest dołączony do tego projektu. Kod ten będzie działał zarówno z fotorezystorem, jak i z czujnikiem siły, ale nie z obydwojema jednocześnie.

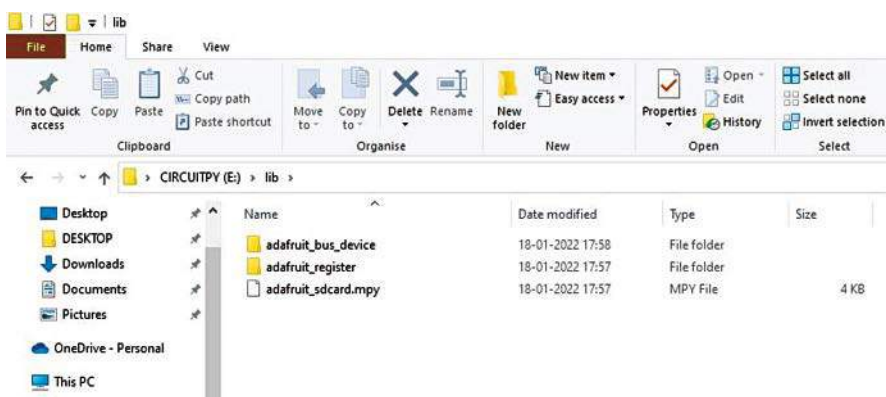
Podłączmy fotorezystor do układu, jak pokazano na schemacie, i uruchomimy kod, klikając przycisk Run. Po przykryciu fotorezystora nieprzezroczystym obiektem i ponownym odsłonięciu go można zaobserwować, że arbitralne wartości natężenia światła zmieniają się w zakresie od 0 do 32, gdzie 0 oznacza najciemniejsze, a 32 najjaśniejsze światło.

Dodatkową cechą programu Mu jest możliwość wyświetlania przebiegu danych. Po wybraniu opcji Plotter z menu głównego zostanie wyświetlony przebieg wyjściowy, jak pokazano na rysunku 7. Podobnie można zastąpić fotorezystor dowolnym innym czujnikiem analogowym, np. czujnikiem siły. Dane z czujnika zostaną wyświetlone na wyświetlaczu wraz z przebiegiem.

Zbadajmy teraz wbudowany czujnik temperatury na płytce Raspberry Pi Pico i zapiszmy



Rysunek 7. Przebieg wyjściowy czujnika (fotorezystora)



Rysunek 8. Napęd CIRCUITPY

dane o temperaturze na karcie microSD, często używanej w telefonach komórkowych. Tutaj używamy bibliotek CircuitPython wraz z głównym kodem (Sdcard_temp_data.py). Musimy zaimportować biblioteki adafruit dla karty SD, aby połączyć ją z płytką Raspberry Pi Pico. W tym celu należy pobrać biblioteki (Bundle for Version 7.x) ze strony <https://circuitpython.org/libraries>.

Rozpakujmy folder Bundle i skopiuje my foldery adafruit_bus_device i adafruit_register. Wklejmy je do folderu 'lib' w napędzie CIRCUITPY(E). Skopiuje my również plik adafruit_sdcard.mpy z folderu Bundle i wklejmy do folderu 'lib', jak pokazano na rysunku 8.

Otwórzmy edytor Mu i załadujmy główny kod w pliku Code.py. Zapiszmy plik i uruchommy go. Jeśli coś pójdzie nie tak, w dolnej części okna edytora Mu pojawi się komunikat o błędzie. Na przykład, jeśli karta SD nie jest prawidłowo podłączona do układu, pojawi się komunikat o błędzie „OSError: no SD card” (brak karty SD), jak pokazano na rysunku 9.

Jeśli wszystko jest w porządku, zamiast komunikatu o błędzie pod oknem edytora Mu pojawi się następujący tekst (zob. rysunek 10): „The first line of text on SD card, Check

temperature data on SD card” (pierwszy wiersz tekstu na karcie SD, sprawdź dane temperatury na karcie SD).

Teraz zamknijmy edytor Mu i odłączmy kartę SD od adaptera. Włóżmy kartę SD do gniazda kart microSD w laptopie/komputerze PC i sprawdźmy różne wartości temperatury dostępne w formacie pliku tekstowego.

Montaż i testowanie

Cały układ można zmontować na płytce prototypowej lub uniwersalnej płytce drukowanej. Większość płytek Raspberry Pi Pico nie ma podłączonych lub przylutowanych pinów złącz.

Bez tych pinów nie można połączyć płytki z urządzeniami zewnętrznymi, takimi jak czujniki, mierniki lub inne urządzenia wymagane w projekcie. Przewody można przylutować bezpośrednio do otworów na płytce, ale ze względu na stosunkowo niewielki raster wyprowadzeń może to być kłopotliwe.

Dlatego, aby uzyskać dostęp do pinów GPIO, należy ostrożnie przylutować każdy pin za pomocą listwy kołkowej i odpowiedniej lutownicy. Następnie można zamontować piny na płytce drukowanej lub podłączyć

```

7 import microcontroller
8 from time import sleep
9
10 SD_CS = board.GP5
11
12 spi = io.SPI(board.GP2, board.GP3, board.GP4)
13 cs = digitalio.DigitalInOut(SD_CS)
    
```

RP2040 REPL
>OK
soft reboot
raw REPL; CTRL-B to exit
>OK

Traceback (most recent call last):
File "<stdin>", line 14, in <module>
File "adafruit_sdcard.py", line 108, in __init__
File "adafruit_sdcard.py", line 121, in __init_card
OSError: no SD card

Rysunek 9. Komunikat o błędzie

```

4 import digitalio
5 import storage
6 import adafruit_sdcard
7 import microcontroller
8 from time import sleep
9
10 SD_CS = board.GP5
11
12 spi = io.SPI(board.GP2, board.GP3, board.GP4)
13 cs = digitalio.DigitalInOut(SD_CS)
14 sdcard = adafruit_sdcard.SDCard(spi, cs)
15 vfat = storage.VfatFat(sdcard)
    
```

RP2040 REPL
soft reboot
raw REPL; CTRL-B to exit
>OK

The first line of text on SD Card
Check temperature data on SD Card

Rysunek 10. Wynik końcowy

je za pomocą zworki do innych elementów lub modułów.

Po wykonaniu wszystkich połączeń należy ponownie sprawdzić poprawność polaryzacji, zwłaszcza styków Vcc i masy. Należy pamiętać, że napięcie VBUS na styku 40 wynosi zwykle 5 V, co stanowi główne zasilanie Vcc w układzie. Przed włożeniem karty microSD należy ją najpierw sformatować w systemie FAT32, a następnie włożyć do gniazda w adapterze.

Po pierwszym podłączeniu Raspberry Pi Pico do laptopa edytor Mu próbuje automatycznie wykryć płytkę i wyszukuje w niej napęd CIRCUITPY. Przedtem należy się upewnić, że program Circuit-Python został prawidłowo zainstalowany w napędzie, jak wyjaśniono powyżej. Następnie należy wybrać tryb i uruchomić kod analog_sensor.py.

Jeśli uzyskamy prawidłowe dane wyjściowe z czujnika i przebieg, otwieramy główny kod, klikając opcję Load na pasku menu głównego. Uruchommy kod i upewnijmy się, że w dolnej części okna edytora Mu wyświetlany jest komunikat wyjściowy. Program można zamknąć, a obwód wyłączyć w dowolnym momencie, aby sprawdzić dane z czujnika na karcie SD z dowolnego komputera wyposażonego w gniazdo kart microSD. ■

Sani Theo

Ten artykuł był wcześniej opublikowany na łamach „EFY”, marzec 2022 (efymag.com)

AVTEDU

Poznaj całą serię

Zupełnie nowa edukacyjna seria kitów AVTEDU. Wypróbuj je wszystkie i zostań mistrzem lutownicy, poznaj świat elektroniki i zgłębiaj go razem z nami

#AVTEDU #NaukaLutowania #KityAVT

Zestaw umożliwiający rozpoczęcie nauki techniki lutowania elementów elektronicznych. Wraz z serią kitów AVTEDU tworzy idealne uzupełnienie zagadnienia montażu prostych urządzeń elektronicznych.

Zestaw zawiera **lutownicę**, wysokiej jakości **podstawkę** z czyszcikiem, **cynę** z topnikiem, **kalafonię**, **pęsety**, **odsysacz** do cyny oraz **szczypce** tnące boczne.

W komplecie na dobry początek znajduje się również **zestaw AVTEDU do złutowania**.



AVTEDUSTART - zestaw narzędzi do nauki lutowania



sklep.avt.pl

AVT SPV Sp. z o.o., 03-197 Warszawa, ul. Leszcynowa 11
tel.: (22) 257 84 51 e-mail: handlowy@avt.pl

Arduino programowane „ręcznie”

Interesujące rozwiązanie umożliwiające programowanie Arduino z wykorzystaniem wyświetlacza OLED oraz klawiatury 5-przyciskowej.

Do czego to służy?

Arduino chyba nie trzeba nikomu przedstawiać. Za pomocą tego układu można zrealizować mnóstwo ciekawych projektów. Wystarczy kupić odpowiednią płytkę, podłączyć ją do komputera za pomocą kabla USB i można zaprogramować ją przy użyciu darmowego oprogramowania, pisząc program w języku C/C++. I w tym cały problem. Otóż do zaprogramowania prostego układu potrzebujemy „potężnego komputera” i dodatkowego oprogramowania. W większości przypadków nasz układ realizuje bardzo proste zadanie i czasami chcemy wykonać drobne zmiany w jego działaniu. Czy nie da się zrobić tak, aby programowanie było realizowane już przez samo Arduino? Przedstawione rozwiązanie jest dowodem na to, że jest to możliwe. Prezentowany układ to nic innego jak płytka rozszerzająca podłączona do popularnego Arduino. Moduł wyposażony jest w wyświetlacz OLED oraz 5-przyciskową klawiaturę. Warto zwrócić uwagę na prostotę tego rozwiązania, użycie popularnych elementów oraz niski koszt.

Jak to działa?

Istotą działania tego urządzenia jest program wgrany jednorazowo do Arduino. Program ten zajmuje się obsługą klawiatury, wyświetlacza oraz wykonywaniem zadanych przez użytkownika instrukcji. Przy pomocy 5-przyciskowej klawiatury poruszamy się po menu wyświetlanym na wyświetlaczu OLED. Wybierając odpowiednie pozycje w menu jesteśmy w stanie zaprogramować lub edytować listę instrukcji jakie ma wykonać nasze Arduino. Instrukcje te budową zbliżone są do języka assemblerowego. **Rysunek 1** przedstawia poszczególne elementy instrukcji.

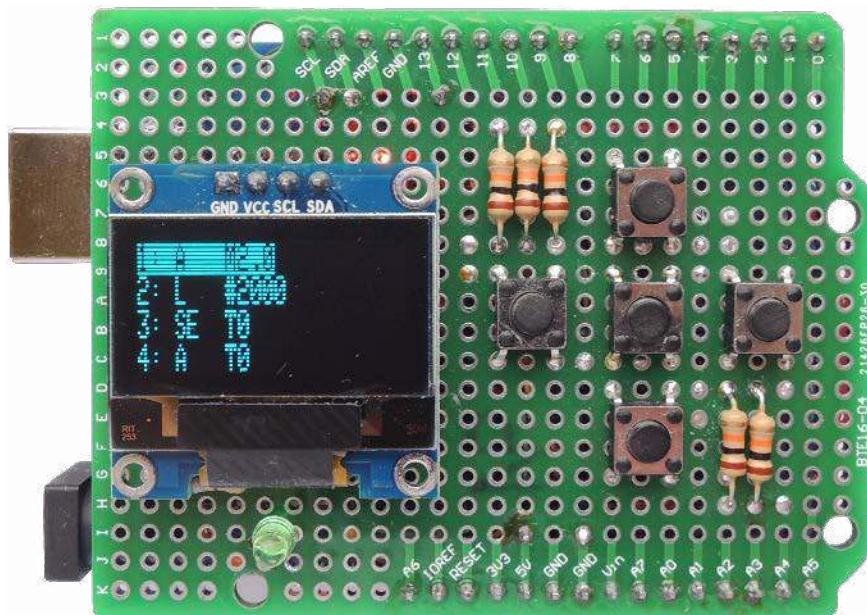
1: A M2.0

2: = D2

3: L #5

Nr bitu
Nr komórki pamięci
Obszar pamięci
Rozkaz
Nr linii

Parametr



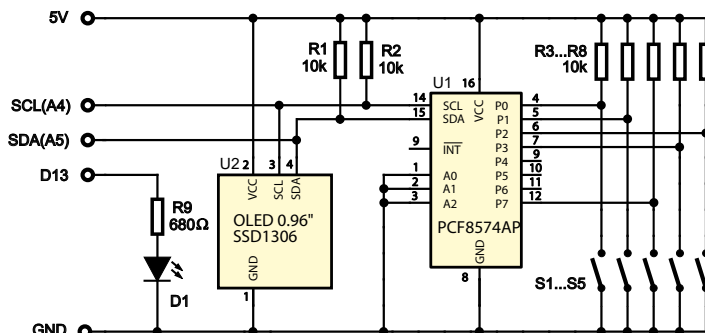
Opis budowy

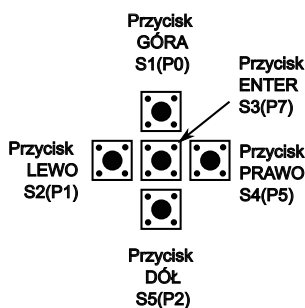
Elementem niezbędnym do prawidłowego funkcjonowania urządzenia jest wspomniany moduł rozszerzeń. Dla lepszego zrozumienia budowy i działania urządzenia, warto przeanalizować schemat ideowy, pokazany na **rysunku 2**. Układ zasadniczo składa się z dwóch części. Pierwszą z nich jest moduł wyświetlacza OLED oznaczony na schemacie jako U2. Obecnie na rynku występuje wiele różnych tanich modułów wyświetlaczy z obsługą magistrali I²C. W tym przypadku jest to moduł wyświetlacza OLED 0,96" 128×64 px oparty na sterowniku SSD1306. Drugim elementem modułu jest klawiatura składająca się z układu PCF8574AP (8-bitowego ekspandera wejść cyfrowych z obsługą komunikacji I²C) oznaczonego na schemacie jako U1, przycisków S1...S5 i rezystorów R1...R5. Dodatkowo na module umieszczono LED D1 podłączony do wyjścia cyfrowego D13 za pośrednictwem rezystora R9.

Klawiatura i wyświetlacz podłączone są do wspólnej magistrali I²C wyprowadzonej z modułu wyjściami SCL i SDA (na płytce Arduino oznaczone jako A4, A5). Cały moduł zasilany jest z Arduino. Wejścia A0, A1, A2 układu U1 podłączone są do GND w celu ustawienia domyślnego adresu magistrali I²C układu PCF8574AP (adres 0x38h). Rezystory R1, R2 pełnią funkcję podciągającą linie SDA i SCL do stanu wysokiego.

Montaż i uruchomienie

Wszystkie elementy modułu należy podłączyć zgodnie ze schematem ideowym, pokazanym na **rysunku 2**. **Należy zwrócić uwagę na odpowiednie podłączenie sygnałów przycisków klawiatury.** Zapewnia to prawidłowe funkcjonowanie układu. W przeciwnym wypadku niektóre przyciski nie będą działały lub będą działały nieprawidłowo. **Rysunek 3** przedstawia rozmieszczenie przycisków, ich funkcję oraz





3

odpowiednie wejście do układu PCF8574AP. **Należy zwrócić uwagę, aby był to dokładnie ten układ, ponieważ często zamiennie sprzedawany jest układ PCF8574P (bez literki A).** Jednak ten układ nie będzie działał, ponieważ posiada inną adresację. Wymagałoby to zmian w kodzie źródłowym.

Kolejną rzeczą, na którą należy zwrócić uwagę jest wyprowadzenie pinów z modułu OLED. Muszą być wyprowadzone jak na **fotografii 4** (kolejność pinów od lewej strony patrząc z góry: GND, VCC, SCL, SDA). Na rynku występują moduły z inną kolejnością wyprowadzeń.

Po podłączeniu modułu do płytki Arduino (w tym przypadku Arduino UNO), podłączeniu zasilania i wgraniu programu do mikrokontrolera za pomocą oprogramowania Arduino, na wyświetlaczu OLED powinien pojawić się napis „No program”. Po 3 sekundach napis zniknie i pojawią się menu główne z opcjami „Run” (uruchamia program), „Edit” (tryb edycji), „Program” (programowanie EEPROM) oraz „Clear” (czyszczenie EEPROM). Po menu poruszamy się przyciskami GÓRA, DÓŁ, a wybraną opcję aktywujemy przyciskiem ENTER (aby wrócić poziom wyżej wciskamy przycisk LEWO). Schemat na **rysunku 5** przedstawia listę kroków jakie należy wykonać, aby stworzyć prosty program mrugający diodą LED na płytce

Arduino (wyjście D13 – LED_BUILDIN) z częstotliwością 1Hz w momencie, gdy przyciśniemy przycisk DÓŁ na klawiaturze.

Każdy blocek na rysunku przedstawia widok z wyświetlacza w danym momencie edycji. Napis zakreślony na czarno to opcja wybrana przyciskami GÓRA/DÓŁ. Strzałki na rysunku wskazują na kolejny ekran jaki pojawia się po zatwierdzeniu wybranej opcji przyciskiem ENTER (z wyjątkiem opisanych strzałek w pierwszej i ostatniej linii). Na początku (patrząc od góry) widać jak wyświetla się wspomniany wcześniej komunikat „No program” i po 3 sekundach przeniesieni zostajemy do menu głównego. Wybieramy „Edit” i możemy dodawać kolejne linie programu klikając w „[+]”. Wybieramy wskazane na rysunku pozycje i zatwierdzamy przyciskiem ENTER. Na ekranach z napisami „Enter byte nr:”, „Enter bit nr:” możemy edytować wartości liczbowe parametrów. Pozycję cyfry w liczbie wybieramy przyciskami LEWO/PRAWO. Przyciskami GÓRA/DÓŁ zmieniamy wartość wybranej w ten sposób cyfry. Przyciskiem ENTER zatwierdzamy całą liczbę. Strzałka zawracająca na rysunku oznacza zakończenie edycji nowej linii i powrót do listingu całego programu. Kolejne linie programu dodajemy ponownie klikając w „[+]”. Na dole rysunku z lewej strony widać cały listing programu. Aby teraz



4

uruchomić stworzony program wychodzimy z edycji przyciskiem LEWO i przechodzimy do menu głównego. Wybieramy „Run”, a następnie „Save & Run”. Program zostaje zapisany do nieulotnej pamięci wewnętrznej EEPROM i rozpoczyna się jego wykonanie. W **tabeli 1** znajduje się opis zaprogramowanych instrukcji.

Uwagi końcowe

Arduino zapamiętuje program po wyłączeniu zasilania. Przy podaniu zasilania jeśli zostanie wykryty program automatycznie jest wykonywany (przechodzi w tryb „Running”). Aby wyjść z trybu „Running...” do menu głównego należy jednocześnie wcisnąć przyciski LEWO, PRAWO i trzymać przez 3 sekundy. W trybie „Running...” wyświetlacz wyłącza się automatycznie po 30 sekundach.

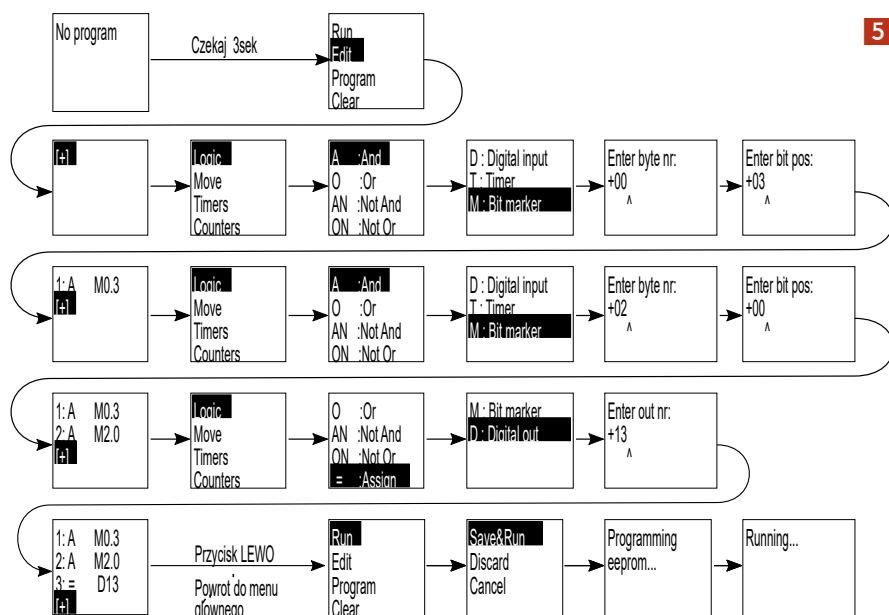
W menu „Edit”, gdzie widać listing instrukcji po wybraniu danej linii programu i naciśnięciu przycisku ENTER, pojawia się menu z opcjami „Insert”, „Edit” oraz „Remove”. Opcja „Insert” powoduje wstawienie nowej linii programu przed wybraną linią. „Edit” zmienia wybraną linię. Natomiast „Remove” usuwa wybraną linię.

Wszystkie dodatkowe materiały (szczegółowy opis, lista dostępnych instrukcji, schemat poruszania się po menu urządzenia oraz kod źródłowy programu) znajdują się wśród materiałów dodatkowych do tego numeru EdW. Większość przykładów w celu prezentacji rezultatów wykorzystuje wbudowaną w Arduino LED (na module jest to LED podłączony do wyjścia D13) oraz przyciski i wyświetlacz niniejszego modułu. Nie trzeba więc podłączać dodatkowych urządzeń, aby rozpocząć zabawę z urządzeniem. ■

Miłosz Gołębiowski
milgo@o2.pl

Wykaz elementów:

R1...R8: 10 kΩ
R9: 680 Ω
D1: LED zielona 3 mm
S1...S5: przycisk
U1: PCF8574AP
U2: OLED 0,96 oparty na SSD1306



5

IoT na oku: sterowanie urządzeniami domowymi



Prezentujemy tu prawdopodobnie pierwsze na świecie urządzenie, które umożliwia włączanie i wyłączanie dowolnego urządzenia elektrycznego/elektronicznego za pomocą mrugnięcia okiem.

Czy kiedykolwiek myśleliśmy o tym, aby sterować elektrycznymi/elektronicznymi urządzeniami domowymi, po prostu na nie patrząc? Może się to zdarzyć w filmach, ale wydaje się niemożliwe w prawdziwym życiu.

Cóż, już nie. Oto prawdopodobnie pierwsze na świecie urządzenie, które umożliwia włączanie i wyłączanie dowolnego urządzenia elektrycznego/elektronicznego za pomocą mrugnięcia okiem. To rozwiązanie IoT może także osobom niepełnosprawnym w samodzielnym sterowaniu takimi urządzeniami.

Kodowanie

Urządzenie musi rozpoznawać inne urządzenia, które mają być sterowane za pomocą poleceń wysyłanych z oka. Dlatego poniższy kod umożliwia przechwytywanie obrazu wideo w czasie rzeczywistym w celu wykrywania obiektów i włączania lub wyłączania urządzeń.

OpenCV jest wykorzystywany do przechwytywania obrazu z kamery, a TensorFlow (TF) do rozpoznawania oglądanych urządzeń domowych.

Zainstalujemy Pythona i wymagane moduły w Raspberry Pi za pomocą następujących poleceń:

```
sudo pip3 install python-opencv
sudo pip3 install tensorflow
sudo pip3 install keras
sudo pip3 install gpiozero
sudo pip3 install pyserial
git clone --depth 1 https://github.com/tensorflow/models.git
```

Jeśli wystąpi problem z instalacją TensorFlow i OpenCV lub pojawi się błąd związany z wersją TensorFlow, należy postępować

Komponenty		
Nazwa komponentu	Ilość	Opis
Raspberry Pi 4	1	2 GB pamięci RAM lub więcej
Kamera RPi	1	Moduł kamery
Raspberry Pi Pico	1	Mikrokontroler
Bluetooth HC05	1	Moduł Bluetooth
Moduł przekaźnikowy	1	Moduł przekaźnika 5 V



Rysunek 1. Prototyp autora

zgodnie z instrukcjami dotyczącymi instalacji TensorFlow i wykrywania obiektów podanymi na stronie internetowej TensorFlow.

Następnie należy sklonować bibliotekę wykrywania obiektów. Potem w folderze testowym utworzymy plik zawierający listę urządzeń i zapiszmy go jako eyeiot.pbtxt

```
import os
import cv2
import numpy as np
from picamera.array import PiRGBArray
from picamera import PiCamera
import tensorflow as tf
import argparse
import sys
import time
import serial

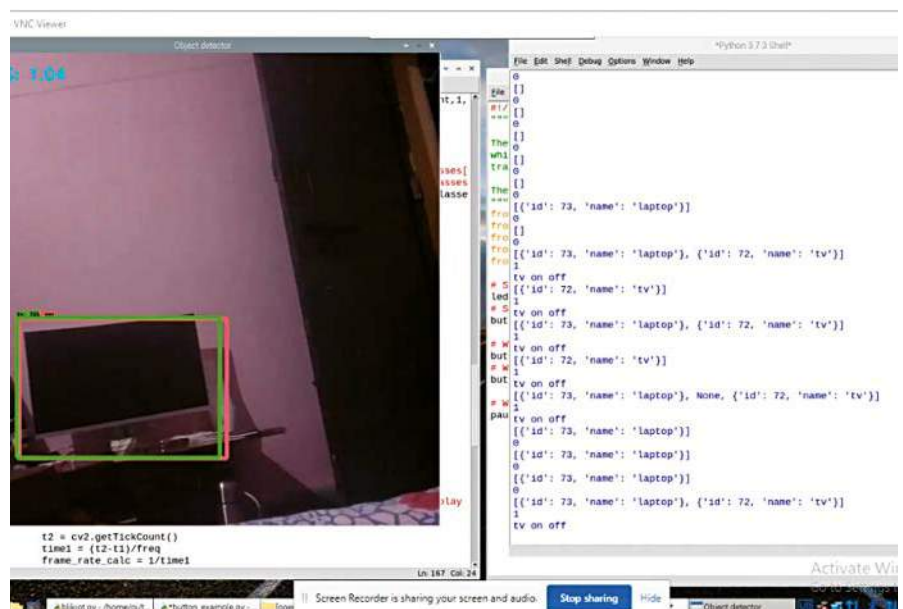
from signal import pause
import threading
ser = serial.Serial("/dev/rfcomm2", baudrate=9600)

Time=0
a=time.time()
#IM_WIDTH = 880
```

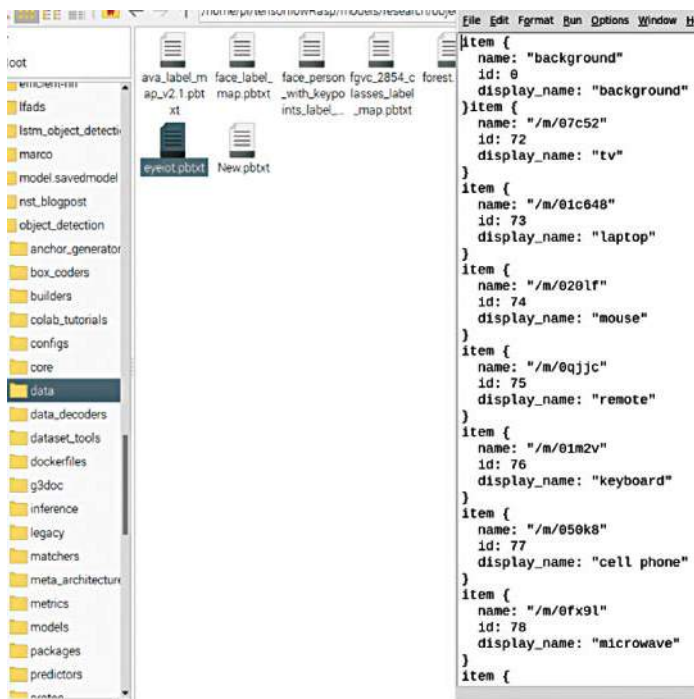
Rysunek 3. Kod ustawiający szybkość transmisji szeregową i nazwę portu

Aby utworzyć funkcje wykrywania i sterowania obiektami IoT, należy skopiować kod wykrywania obiektów do nowego pliku. Należy pamiętać o zaimportowaniu do kodu kilku dodatkowych modułów i bibliotek, takich jak gpiozero do sterowania wejściami i wyjściami pinów oraz pyserial do wysyłania poleceń do urządzenia Pico w celu sterowania przekaźnikiem, a tym samym urządzeniami.

Używamy modułu pyserial do połączenia z RPi Pico za pomocą Bluetooth HC05 dołączonego do RPi Pico oraz



Rysunek 2. Kod do testowania



Rysunek 4. Plik .pbtxt zawierający listę urządzeń do wykrycia



Rysunek 5. Kod wykrywający urządzenie, na które patrzy użytkownik

wbudowanego Bluetooth w Raspberry Pi 4. W kodzie należy więc zaimportować plik py serial i ustawić szybkość transmisji oraz nazwę portu szeregowego Bluetooth.

Na koniec należy ustawić w kodzie nazwę pliku eyiot.pbtxt, aby program mógł wykryć, które urządzenie ma być sterowane z podanej listy urządzeń (lodówka, telewizor, toster czy inne). Po wykonaniu tych czynności

otrzymujemy wynik, w którym stringi zawierają listę nazw obiektów wykrytych na granicy wideo.

Kod wyszukuje w wynikach formatu stringu różne nazwy podstringów. Te podstringi zawierają nazwy urządzeń, którymi chcemy sterować. Kilka instrukcji warunkowych sprawdzi dostępność tego stringu, czyli obecność urządzenia.

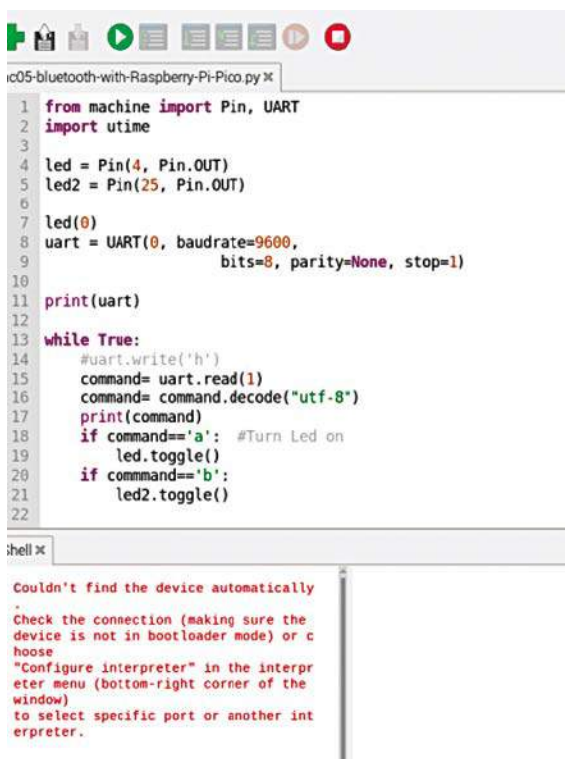
Zalóżmy, że użytkownik chce obsługiwać telewizor. Kod wykryje jego obecność, gdy użytkownik na niego spojrzy. Za pomocą warunku if będzie liczony czas, gdy użytkownik

będzie na niego patrzył. Po pięciu sekundach telewizor włączy się lub wyłączy.

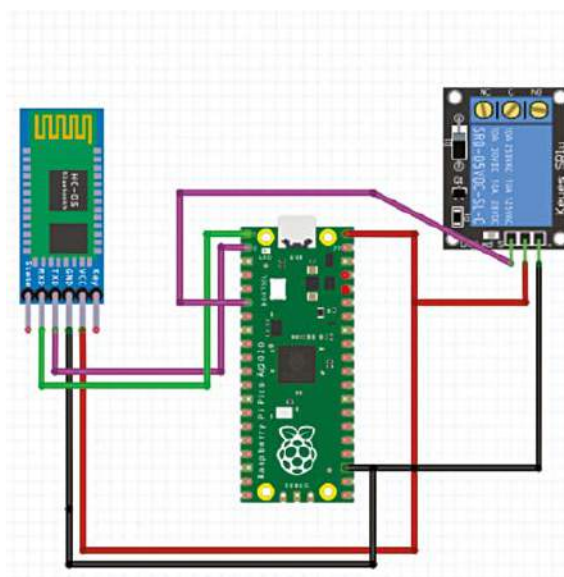
Kodowanie Raspberry Pi Pico

Musimy stworzyć kod dla Raspberry Pi Pico, tak aby mógł on połączyć się z głównym Raspberry Pi i sterować bezprzewodowo urządzeniem, które oglądamy.

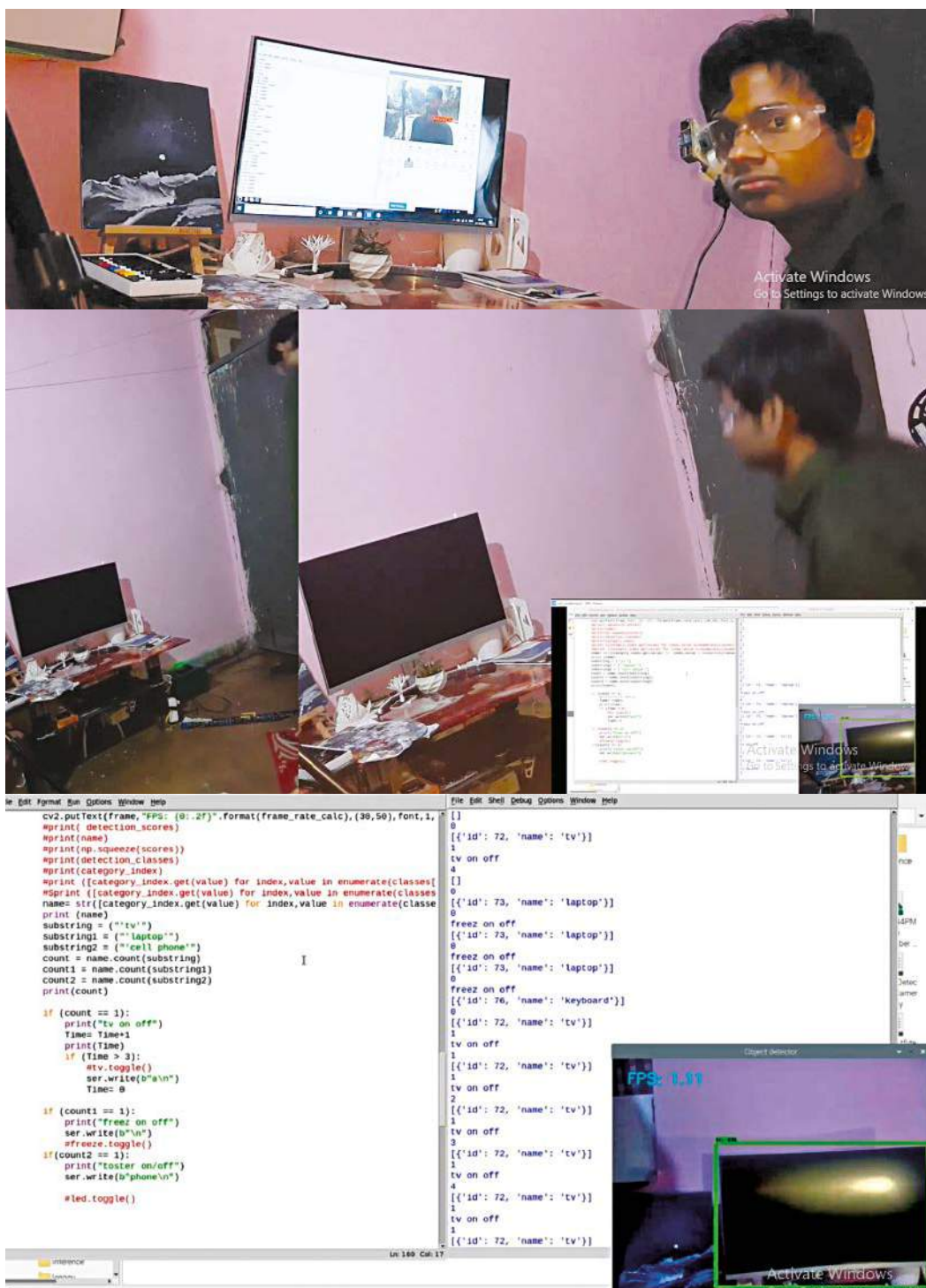
Będziemy używać języka micropython dla Raspberry Pi Pico. Do komunikacji z modulem Bluetooth HC 05 używamy USART, a następnie za pomocą warunku if sprawdzamy polecenia przełączania przełączników podłączonych urządzeń. Zapiszmy kod jako main.py na Raspberry Pi Pico.



Rysunek 6. Kod Raspberry Pi Pico do sterowania urządzeniem



Rysunek 7. Połączenie RPi Pico z modulem Bluetooth i przekaźnika



Rysunek 8. Autor testujący urządzenie

Połączenie

Połączmy HC05, Pico i moduł przekaźnika w sposób pokazany na rysunku 7. Podłączmy styk COM przekaźnika do jednego z przewodów urządzenia sterowanego, a jego styk normalnie otwarty do przewodu fazy napięcia sieci zasilającej sterowane urządzenie. Podłączmy drugi przewód urządzenia sterowanego do przewodu neutralnego sieci (uwaga: należy zachować ostrożność podczas pracy z zasilaniem sieciowym,

ponieważ nieprawidłowe podłączenie może spowodować porażenie prądem elektrycznym lub zwarcie).

Do złącza 25-pinowego można dodać jeszcze jeden przekaźnik do ładowania telefonu i laptopa.

Następnie należy podłączyć Raspberry Pi Pico do zasilacza prądu stałego 5 V i podłączyć Raspberry Pi 4, z kamerą zamontowaną na szkle okularowym, z baterią 5 V, a następnie uruchomić kod iotoneye.py

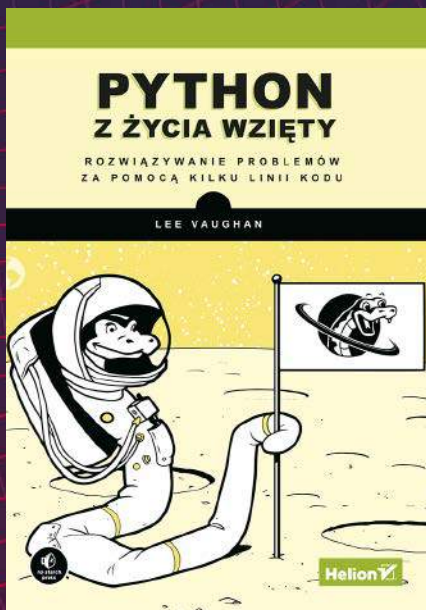
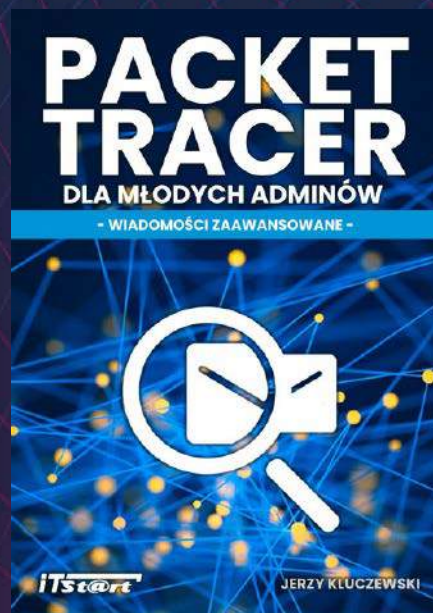
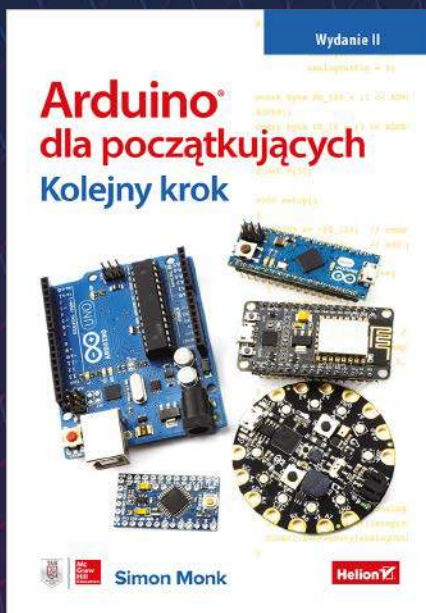
na Raspberry Pi 4. W EFY, do kodu testowego dołączyliśmy telewizor. Gdy patrzymy na telewizor przez pięć sekund, automatycznie włącza się lub wyłącza. W ten sam sposób można włączać i wyłączać dowolne urządzenie. ■

Ashwini Kumar Sinha

Ten artykuł był wcześniej opublikowany na łamach „EFY”, październik 2021 (efymag.com)

NOWOŚCI W ULUBIONYM KIOSKU

KSIĄŻKI Z RABATEM DO 30%

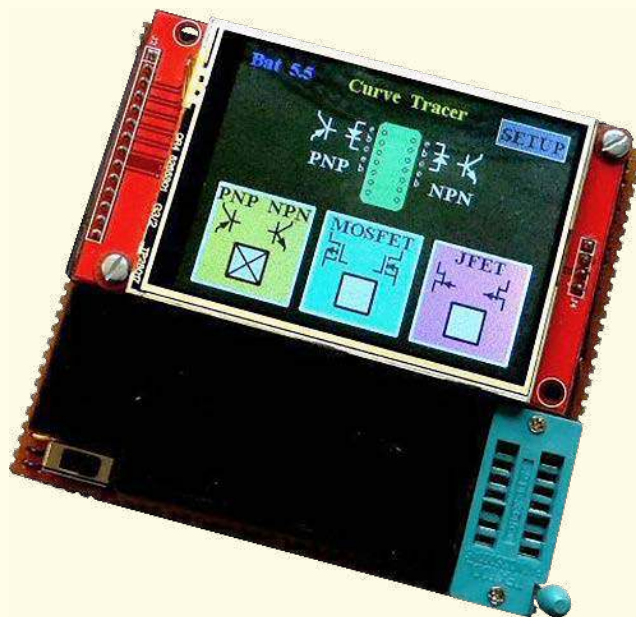
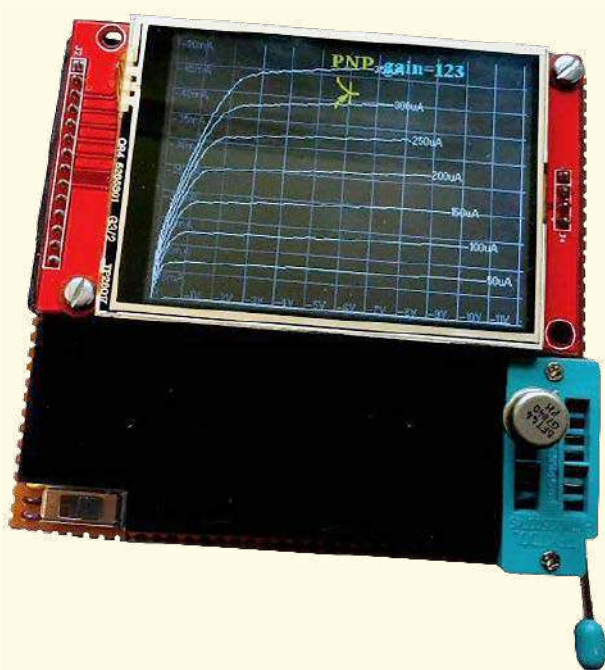


Zobacz pełną ofertę książek i zamów
wygodnie na UlubionyKiosk.pl

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**.

Obserwacja charakterystyk tranzystora

Autor opisuje jak zrobić urządzenie do obserwacji charakterystyk tranzystora. Takie urządzenie to najlepszy sposób, aby zrozumieć, jak działa element.



Zawsze chciałem mieć urządzenie do obserwacji charakterystyk tranzystora. To najlepszy sposób, aby zrozumieć, jak działa dany element. Po zbudowaniu i użyciu tego urządzenia wreszcie rozumiem różnicę między różnymi rodzajami FET-ów.

Przydaje się do:

- dopasowywania tranzystorów
- mierzenia wzmocnienia tranzystorów bipolarnych
- mierzenia progu zadziałania tranzystorów MOSFET

- mierzenia progu zadziałania tranzystorów JFET
- pomiaru napięcia przewodzenia diod
- pomiaru napięcia przebicia diod Zenera
- itd.

Byłem pod wielkim wrażeniem, kiedy kupiłem jeden ze wspaniałych testerów LCR-T4 Markusa Frejka i innych, ale chciałem, żeby mówił mi więcej o komponentach, więc zacząłem projektować własny tester.

Zacząłem od użycia tego samego ekranu, co w LCR-T4, ale nie miał on wystar-

czająco wysokiej rozdzielczości, więc zmieniłem go na 2,8-calowy wyświetlacz LCD 320×240. Tak się składa, że jest to kolorowy ekran dotykowy, co jest miłe. Układ do śledzenia krzywych działa na Arduino Pro Mini 5 V Atmega328p 16 MHz i jest zasilany przez 4 ogniwa AA.

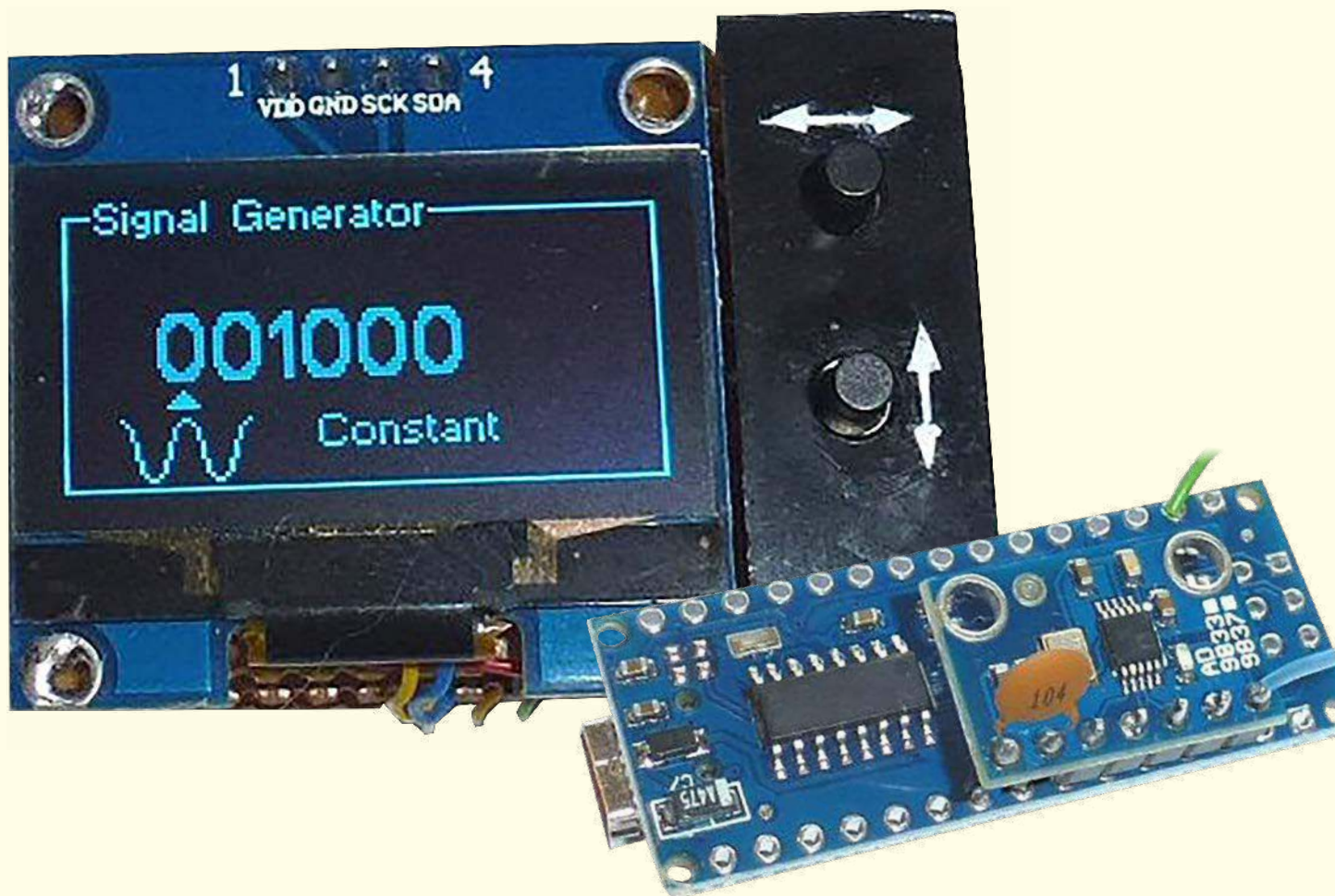
Dokończenie artykułu na stronie:
<https://bit.ly/3FvggIY>

Niektóre projekty aktualnie dostępne dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl:

1. Sonarowy theremin MIDI
2. Wyświetlacz EKG
3. Łatwy do zbudowania robot kroczący
4. Generator impulsów o zakresach zmiennych płynnie
5. Zamek elektroniczny na kod
6. Prosty tester tranzystorów
7. Zegar binarny
8. Miernik pojemności
9. Generator impulsów o zakresach zmiennych skokowo
10. Oscylator pierścieniowy
11. Generator sygnałów funkcji
12. Minutnik do gotowania jajek
13. Przerzutnik SR oparty na PLD
14. Programowalne bramki logiczne
15. Generator impulsów o zakresach krokowych

Generator sygnałów AD9833

Generator sygnału to bardzo przydatny element wyposażenia testowego. Ten w niniejszym artykule wykorzystuje moduł AD9833 i Arduino Nano – to wszystko, nie ma nawet płytki drukowanej. Opcjonalnie można dodać wyświetlacz OLED. Moduł AD9833 może generować fale sinusoidalne, trójkątne i prostokątne o częstotliwościach od 0,1 Hz do 12,5 MHz – oprogramowanie w tym projekcie jest ograniczone do zakresu od 1 Hz do 100 kHz.



Istnieją już inne projekty wykorzystujące Arduino i AD9833, sprawdź pod tymi linkami. Poniższy projekt jest prostszy i może być używany jako generator przemiatający. Generatory przemiatające pomagają testować charakterystykę częstotliwościową filtrów, wzmacniaczy itp. W przeciwieństwie do innych projektów, ten nie zawiera wzmacniacza ani

regulatora amplitudy, ale możesz je dodać, jeśli tylko chcesz.

Najprostszy generator sygnału

Aby uzyskać najprostszy generator sygnału, wystarczy przylutować moduł AD9833 do tylnej części Arduino Nano. Nie jest potrzebna płytka drukowana.

Moduł AD9833, który wybrałem, nie twierdę, że jest to najlepszy albo najtańszy, ale powinieneś kupić dokładnie taki, jak na zdjęciu powyżej.

Dokończenie artykułu na stronie:
<https://bit.ly/3sqYcUQ>

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

Dział Reklamy:
Katarzyna Gugała
katarzyna.gugała@elportal.pl
tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Grzegorz Becker
grzegorz.becker@elportal.pl

DTP, okładka, redakcja strony internetowej www.elportal.pl:
MAD Sp. z o.o.

Prenumerata:
W Wydawnictwie AVT, e-mail: prenumerata@avt.pl
tel. 22 257 84 22, (godz. 10:00–14:00)

W RUCH S.A., e-mail: prenumerata@ruch.com.pl
tel. 801 800 803, 22 717 59 59, www.prenumerata.ruch.com.pl

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”.

Przenośna stacja lutownicza KD862 na gorące powietrze



Cyfrowa stacja Hot Air KD 862 umieszczona w kolbie. Kompaktowa forma, łatwa do przenoszenia i transportu - wygodne rozwiązanie dla mobilnych serwisantów. Oprócz typowych zastosowań, nadaje się również do spawania tworzyw sztucznych, obkurczania, usuwania starych powłok z farb, itp.



Sterowanie umieszczone w kolbie: pokrętko do regulacji przepływu powietrza (**max 120l/min**) i przyciski do regulacji temperatury (**od 100°C do 480°C**)

- wyświetlacz LED
- zasilanie 230V
- pobór prądu 650W
- długość całkowita 30.5cm
- system schładzania grzałki
- mocna grzałka wykonana z grubego drutu (większa wytrzymałość i trwałość)
- źródło przepływu powietrza: wentylator z silnikiem bezszczotkowym

217zł

W zestawie:

- kolba
- uchwyt
- instrukcja
- 3 dysze okrągłe
- 1 dysza kwadratowa

kod handlowy: KD862



sklep.avt.pl



AVT SPV Sp. z o.o.
03-197 Warszawa, ul. Leszczynowa 11
Dział Handlowy tel.: (22) 257 84 51
e-mail: handlowy@avt.pl