

Nowa rubryka - Edukacja w EdW dla szkół i uczelni
Wykład 1 - Od mikroprocesorów do mikrokontrolerów

ELEKTRONIKA

dla wszystkich

nr 12/2022 (323) • grudzień • www.elportal.pl

DIY PLUS
tylko dla prenumeratorów

Przedwzmacniacz o ultraniskich zniekształceniach

PROJEKTY dla elektroników

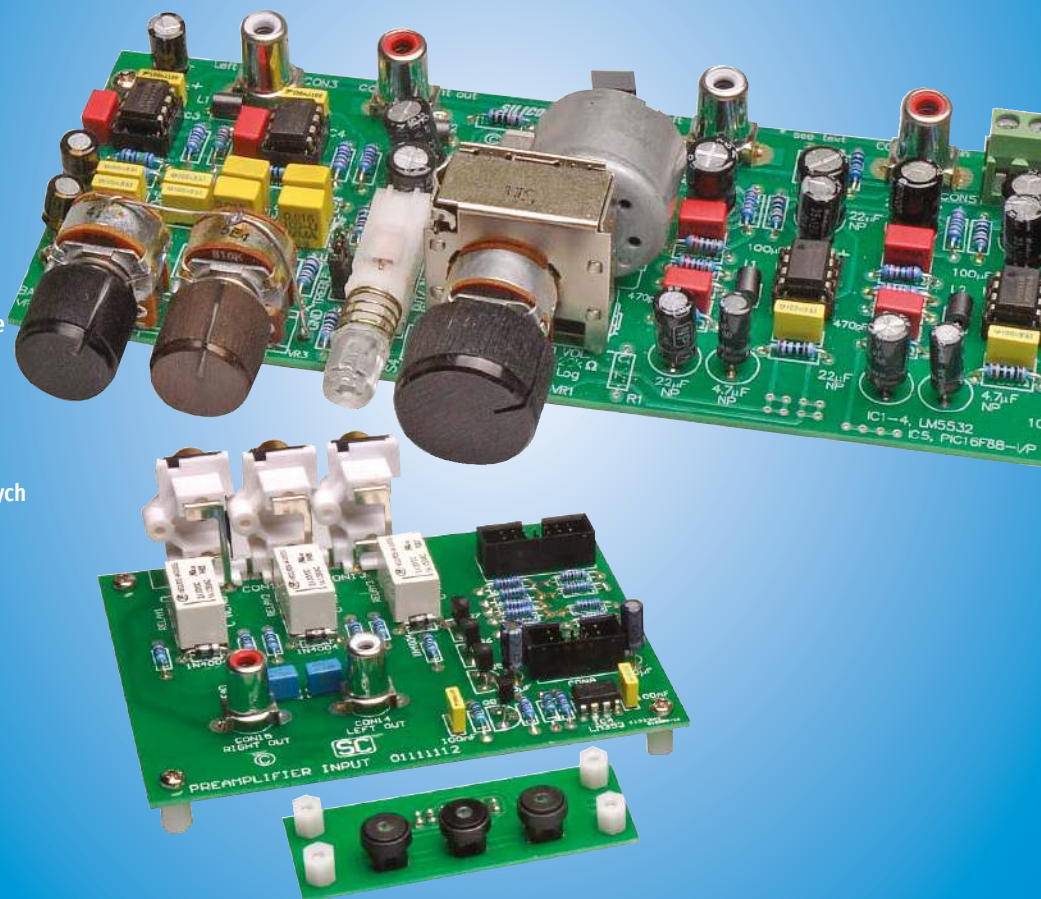
- ▶ Bezprzewodowy retransmitter danych w paśmie 433 MHz
- ▶ Stylowy - nasz nowy ściemniacz - dotykowy i zdalnie sterowany

DIY dla wszystkich

- ▶ Inteligentny sterownik oświetlenia
- ▶ System ostrzegania o przekroczeniu temperatury z sygnalizacją na smartfonie
- ▶ Zegar analogowy i termometr z wyświetlaczem TFT LCD
- ▶ Rejestrator danych pomiarowych z pamięcią USB
- ▶ Inteligentny sterownik urządzeń domowych z wykorzystaniem platformy Blynk

TUTORIALE

- ▶ Szkoła Konstruktorów
- ▶ Nakrętki i śruby w głośniku, część 2
- ▶ Silniki krokowe w praktyce, część 1: Wprowadzenie do technologii silników krokowych
- ▶ Tensometry i sygnały różnicowe



EP.com.pl

Największy
portal dla
elektroników
konstruktorów

FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

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

www.piekarz.pl



Król automatyki
jest w Tobie
AutomatykaB2B.pl



16,90 zł (w tym 8% VAT)
ISSN 1425-1698 Indeks 33362X
9 771425 169221
12 >
QR code

Innowacyjna seria zestawów do nauki lutowania:

- większe pady
- duże odstępy między punktami lutowniczymi
- atrakcyjna grafika
- praktyczne zastosowanie

Zestawy dostępne pojedynczo i w pakietach.



Choinka LED RGB

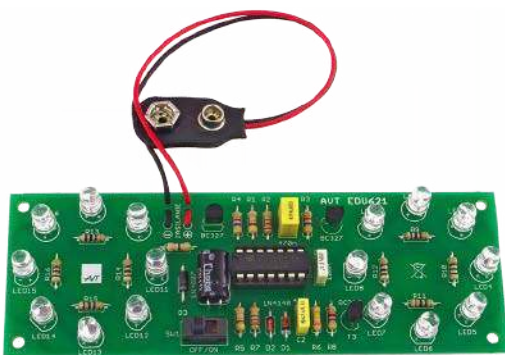
Świąteczny nastrój buduje mnóstwo pozornie drobnych szczegółów – migoczące lampki, klimatyczne melodie, zapach herbaty z pomarańczą i goździkami. Czego jeszcze brakuje? Może choinki! W te Święta spraw sobie prezent, który pozwoli Ci na rozwijanie swoich umiejętności w lutowaniu!

SPECYFIKACJA:

- źródło światła – płynnie zmieniające kolor diody LED RGB,
- bardzo prosty montaż,
- zasilanie: 3 VDC [2×AA] – zestaw nie zawiera baterii,
- wymiary płytki: 68×83 mm

kod handlowy: **AVTEDU640**

cena: **24zł**



Stroboskop policyjny LED

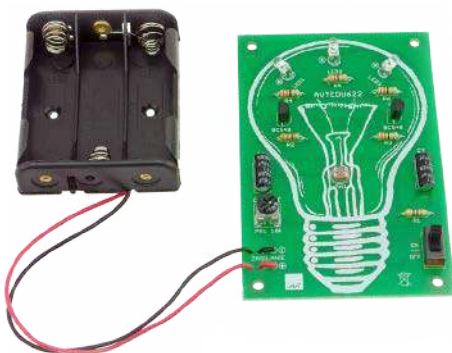
Efektom wizualnym generowanym przez moduł jest imitacja świateł pojazdu uprzywilejowanego.

SPECYFIKACJA:

- 2 pola świetlne z diodami LED (czerwone i niebieskie),
- 8 diod LED w każdym polu,
- wymiary płytki: 125×44 mm,
- napięcie zasilania: 9 VDC [6F22] – zestaw nie zawiera baterii

kod handlowy: **AVTEDU621**

cena: **24zł**



Zmierzchowa lampka LED

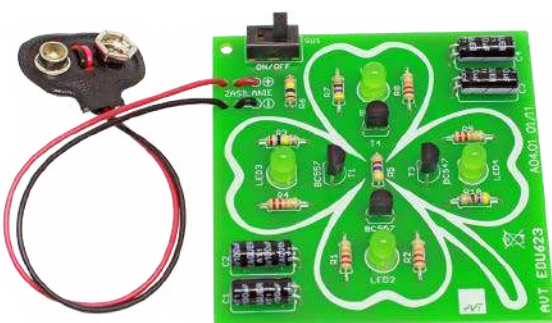
Praktyczna, wyróżniająca się designem lampka nocna z czujnikiem zmierzchu, która po zapadnięciu zmroku, rozbłyśnie jasnym światłem diod LED.

SPECYFIKACJA:

- źródło światła: 3 białe diody LED,
- płynna regulacja czułości zadziałania,
- napięcie zasilania: 5 VDC [3×AA] – zestaw nie zawiera baterii,
- wymiary płytki: 92×60 mm

kod handlowy: **AVTEDU622**

cena: **24zł**



Czterolistna koniczynka LED

Urokliwy i prosty w montażu gadżet będzie cieszył oko dzięki dwóm parom diod LED, które migają w zmiennym rytmie.

SPECYFIKACJA:

- 4 diody LED wysokiej jasności (pure green),
- automatyczna regulacja częstotliwości błysków,
- napięcie zasilania: 9 VDC [6F22] – zestaw nie zawiera baterii,
- mały pobór prądu (ok. 9 mA dla zasilania 9 V),
- wymiary płytki 65×65 mm

kod handlowy: **AVTEDU623**

cena: **22zł**



**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**



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

* 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 przyślemy 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 d96e5f8e8c



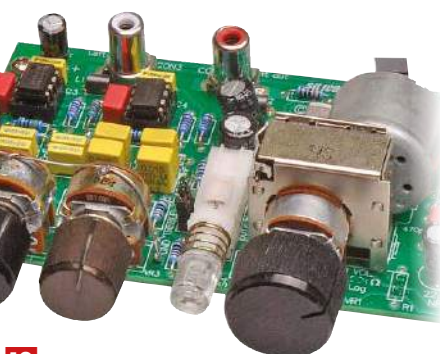
8

Projekty dla elektroników:

- Stylowy – nasz nowy ściemniacz – dotykowy i zdalnie sterowany 8
- Przedwzmacniacz o ultraniskich zniekształceniach, część 2 18
- Bezprzewodowy retransmitter danych w paśmie 433 MHz 25

Tutoriale:

- Szkoła Konstruktorów 34
- Nakrętki i śruby w głośniku, część 2 48
- Silniki krokowe w praktyce, część 1:
- Wprowadzenie do technologii silników krokowych 54
- Tensometry i sygnały różnicowe 58
- Edukacja w EdW dla szkół i uczelni:
- Wykład 1 „Od mikroprocesorów do mikrokontrolerów” 64



18



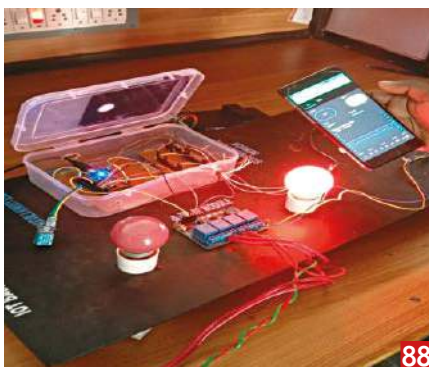
25

DIY dla wszystkich:

- System ostrzegania o przekroczeniu temperatury
- z sygnalizacją na smartfonie 78
- Rejestrator danych pomiarowych z pamięcią USB 80
- Zegar analogowy i termometr z wyświetlaczem TFT LCD 84
- Inteligentny sterownik oświetlenia 86
- Inteligentny sterownik urządzeń domowych
- z wykorzystaniem platformy Blynk 88

DIY PLUS

- Przetwornik częstotliwość/napięcie – obrotomierz
- z czujnikiem magnetycznym o zmiennej reluktancji 91
- Sterownik termoelektrycznej chłodziarki 91



88

Rubryki stałe:

- Prenumerata 3
- Od wydawcy 5
- Pocztą 6

A za miesiąc w styczniowym EdW



- * Gitarowy pedał efektów przesterowania i zniekształceń na lampach Nutube**
 Czy marzyło Ci się prawdziwe „brzmienie lampowe” Twojej gitary z pedałem efektów przesterowania i zniekształceń? Wśród wielu różnych pedałów z efektami gitarowymi najbardziej popularny jest efekt przesterowania i zniekształceń. Układy budowane na półprzewodnikach tylko imitują najbardziej pożądane „brzmienie lampowe”. Zamiast imitacji możesz mieć efekty z prawdziwych lamp elektronowych, dzięki dostępnym na rynku kompaktowym, niskonapięciowym podwójnym triodom 6P1.
- * Stacja zdalnego monitoringu**
 Któż nie doświadczył niepokoju o swoje wartościowe dobra (samochód, domek letniskowy, dom, łódź, itp.), pozostawione bez nadzoru? Na przykład awaria hydrauliczna może spowodować, że po powrocie do domu znajdziemy staw w piwnicy. Wystarczy kilka czujników i płytka rozszerzająca Arduino i możemy być pewni, że na końcu świata zostaniemy natychmiast zaalarmowani. Można nawet zaprojektować zdalne podjęcie interwencji, na przykład wyłączenie pompy wodnej. Przedstawiamy projekt w postaci otwartej na modyfikacje. Każdy może dostosować hardware i program do własnych potrzeb.
- * Uniwersalny ściemniacz światła, część 2**
 Już sama zapowiedź publikacji w tym wydaniu EdW pierwszej części projektu uniwersalnego ściemniacza wywołała spore zainteresowanie Czytelników. Zaproponowaliśmy dotykowy ściemniacz uniwersalny, działający dla wszystkich rodzajów źródeł światła, zarówno LED i halogenowych, jak też żarowych. Układ może być sterowany zdalnie pilotem na podczerwień. W części 2. artykułu przedstawiamy szczegółowo montaż i uruchomienie ściemniacza.
- * Plus zwykła porcja intrygujących projektów DIY.**
- * Plus wiele artykułów w Twoich ulubionych cyklach Tutoriali, w tym polecamy szczególnie nową rubrykę „Edukacja w EdW”.**

**W kioskach
od 30 grudnia**

Edukacja w EdW

Coś się kończy, coś się zaczyna – to tytuł opowiadania Andrzeja Sapkowskiego, czasami interpretowany jako przyznanie przez autora „wiedźmińskiej sagi”, że zakończył ją, gdy mu się znudziła. W tym wydaniu EdW kończy się wieloletni cykl Piotra Góreckiego – „Szkoła Konstruktorów”. Ogromna wiedza i wybitny talent belferski Autora nigdzie tak nie błyszczały jak w „Szkoła Konstruktorów”.

Jednak wszystko ma swój koniec, po którym zawsze coś nowego się zaczyna. W EdW już od kilku miesięcy publikujemy cykle dydaktyczne trzech charyzmatycznych popularyzatorów elektroniki, autorów w brytyjskim „Practical Electronics”. Pora przedstawić ich sylwetki Czytelnikom EdW. Z króciutkich biogramów, które prezentuję niżej, można się dowiedzieć, że każdy z nich ma kilkadziesiąt lat praktyki i jest uznanym ekspertem w swoim segmencie elektroniki. Trudno przecenić klasę tych autorów. Gdy czytam tekst o thereminach napisany przez najwybitniejszego obecnie na świecie konstruktora tych instrumentów, to wiem, że Jake Rothman dzieli się ze mną unikalną wiedzą, opartą na własnym przeogromnym doświadczeniu konstruktora. Gdy czytam tutorial Mike’a Tooleya, to wiem, że napisał to dla mnie autor wielu podręczników, uznanych w świecie bestsellerów. Ian Bell, doktor nauk technicznych, ponad 30 lat pracy wykładowcy i naukowca mówi samo za siebie.

Tutoriale w EdW to liga światowa, a jednak potrzebna jest też rubryka oparta na bliższym kontakcie z Czytelnikami. Odpowiedzią na tę potrzebę ma być nowa rubryka „Edukacja w EdW”. Składa się z dwóch części – Wykład i Konsultacje. Nie jest to serial. W każdym wydaniu EdW proponujemy autonomiczny temat wykładu. W tym wydaniu jest temat „Od mikroprocesorów do mikrokontrolerów”. Wiadomo, temat przeogromny, więc wymaga ociosania według pewnych kryteriów. Kryteria wynikają z założonego celu wykładu. Otóż nie tworzymy podręcznika elektroniki, tylko próbujemy „posklejać” różne elementy wiedzy o określonym temacie, by Czytelnika zainteresować i dać minimum wiedzy przydatnej w praktyce konstruktorom. A druga część – Konsultacje, to otwarte drzwi dla każdego Studenta z każdym pytaniem. Tu pojawia się sprawa najważniejsza. Potrzebujemy interakcji, czyli sprzężenia zwrotnego. Piszcie do nas listy z pytaniami do Konsultanta, a także z propozycjami tematów wykładów. Bez interakcji nie damy rady, bo nie da się na dłuższą metę wyklądać do pustej sali. Dlatego wdzięczność za Waszą aktywność będziemy wyrażali konkretnie. Za każdy list o długości nie krótszej niż tweet wysyłamy darmowe e-wydanie kolejnego numeru EdW (prenumeratorom EdW proponujemy e-wydanie dowolnego tytułu czasopisma na www.ulubionykiosk.pl). Proszę też zwrócić uwagę na Quizy – to nowość w EdW. Quizy są równolegle publikowane na www.Elportal.pl, gdzie również podajemy ich rozwiązania.

Wiesław Marciniak

PS Prezentacja autorów cykli tematycznych w dziale Tutoriale



Jake Rothman

Światowy autorytet w konstrukcjach sprzętu muzycznego i audio – od theremina do kolumn głośnikowych. Na sprzęcie jego konstrukcji grali wybitni muzycy, tacy jak Goldfrapp, The Beastie Boys, Super Furry Animals, Mark Almond, Supergrass. Wykładowca w University of South Wales. Jeszcze w 1995 roku publikowaliśmy w „Elektronice Praktycznej” jego artykuł o wzmacniaczu lampowym – konstrukcji teraz legendarnej.



Mike Tooley

Od ponad 30 lat wykłada elektronikę i awionikę dla inżynierów i techników. Był dziekanem wydziału Engineering w Brooklands College, Surrey, UK. Autor kilkunastu podręczników elektroniki i awioniki. Jego książka „Electronic Circuits: Fundamentals and Applications” miała wiele wydań i stała się światowym bestsellerem.



Ian Bell

Doktor nauk technicznych. Wybitny naukowiec i wykładowca od 35 lat w University of Hull – School of Engineering, UK. Autor 60 publikacji naukowych, dotyczących modelowania układów analogowych, architektury chipów multiprocessorowych i generowania testów układów analogowych.

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

Szanowna Redakcja

Z opóźnieniem dotarł do mnie numer EdW 9/22, w którym znalazłem zadanie przedstawione przez pana Władysława Borgiela.

Jestem czytelnikiem EP i EdW od pierwszych wydawanych numerów i bardzo cenię obydwa miesięczniki.

W załączniku przedstawiam propozycję rozwiązania zadania przy następujących założeniach:

- odcięcie jednego ogniwa z nieskończonego łańcucha nie zmienia rezystancji R (pojemności C) widzianej od strony zacisków wejściowych,
- schemat do obliczeń upraszcza się do postaci wg rysunków w załączniku.

Klasyczne rozwiązanie zadania było przedstawione w Zbiorze zadań z elektrotechniki P. Ciechanowicza (wydanie z lat 60.) i opierało się na wykorzystaniu sumy szeregu nieskończonego.

Serdecznie pozdrawiam i życzę dalszych sukcesów w zakresie popularyzacji i wszechstronnej edukacji elektronicznej.

Michał

a)

$$R = R_1 + \frac{R_2 R}{R_2 + R} \quad | : (R_2 + R)$$

$$R(R_2 + R) = R_1(R_2 + R) + R_2 R$$

$$R R_2 + R^2 = R_1 R_2 + R R_1 + R_2 R$$

$$R^2 - R_1 R - R_1 R_2 = 0$$

$$\Delta = R_1^2 + 4 R_1 R_2 \quad | \sqrt{\Delta} = \sqrt{R_1^2 + 4 R_1 R_2} > R_1$$

$$R = \frac{R_1 + \sqrt{\Delta}}{2} \quad | \text{ lub } R = \frac{R_1 - \sqrt{\Delta}}{2}$$

$R < 0$ - nie spełnia warunków zadania

b)

$$C = \frac{C_1 (C_2 + C)}{C_1 + C_2 + C} \quad | : (C_1 + C_2 + C)$$

$$C(C_1 + C_2 + C) = C_1 (C_2 + C)$$

$$C C_1 + C C_2 + C^2 = C_1 C_2 + C C_1$$

$$C^2 + C C_2 - C_1 C_2 = 0$$

$$\Delta = C_2^2 + 4 C_1 C_2 \quad | \sqrt{\Delta} = \sqrt{C_2^2 + 4 C_1 C_2} > C_2$$

$$C = \frac{-C_2 + \sqrt{\Delta}}{2} \quad | \text{ lub } C = \frac{-C_2 - \sqrt{\Delta}}{2}$$

$C < 0$ - nie spełnia warunków zadania

Prośba o rozszerzenie tematu BMS'ów

Dzień dobry

Czytając listopadowe wydanie EdW w kilku miejscach padła wzmianka o możliwości i konieczności uporządkowania wiedzy nt. układów zabezpieczających akumulatory litowe. Przez jednych określano BMS, przez innych balanser'em, a na chińskim portalu aukcyjnym można zobaczyć np. „3S 40A Enhanced” itp. W tym temacie nieco zostało już powiedziane w rozwiązaniu zadania 303, ale to ciągle niewiele i tak niejako przy okazji. Chciałbym zatem prosić (o ile to możliwe) o uporządkowanie wiedzy w tym temacie na łamach EdW.

pozdrawiam serdecznie całą redakcję
Andrzej

Patronat AVT

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

- Górnoląskie Centrum Edukacyjne im. Marii Skłodowskiej-Curie w Gliwicach, 44-100 Gliwice, Okrzei 20
- Regionalne Centrum Edukacji Zawodowej w Biłgoraju, 23-400 Biłgoraj, Kościuszki 98
- Regionalne Centrum Edukacji Zawodowej w Lubartowie, 21-100 Lubartów, 1 Maja 82
- Techniczne Zakłady Naukowe w Dąbrowie Górniczej, 41-300 Dąbrowa Górnicza, Zawidzkiej 10
- Technikum nr 4 im. Marii Skłodowskiej-Curie, 41-902 Bytom, Katowicka 35
- Zespół Placówek Edukacyjno-Wychowawczych w Gołdapi, 19-500 Gołdap, Wojska Polskiego 18
- Zespół Placówek Oświatowych w Rudniku, 32-440 Sułkowice, Rudnik, Szkolna 55
- Zespół Szkolno-Przedszkolny nr 2 w Wiśle, 43-460 Wiśła, Malinka 53
- Zespół Szkolno-Przedszkolny w Ostroźnicy, 47-280 Pawłowiczki, Ostroźnica, Kościelna 42
- Zespół Szkół Budowlano-Elektrycznych im. Jana III Sobieskiego w Świdnicy, 58-100 Świdnica Śląska, Wąbrzyska 35-37
- Zespół Szkół Elektronicznych i Telekomunikacyjnych w Olsztynie, 10-144 Olsztyn, Bałtycka 37a
- Zespół Szkół Elektronicznych im. I. Domeyki w Bolestawcu, 59-700 Bolestawiec, Tyrankiewiczów 2
- Zespół Szkół Elektronicznych w Rzeszowie, 35-078 Rzeszów, Hetmańska 120
- Zespół Szkół Elektronicznych, Elektrycznych i Mechanicznych, 43-300 Bielsko-Biała, Stowackiego 24
- Zespół Szkół Elektrycznych w Kielcach, 25-317 Kielce, Kaczorowskiego 8
- Zespół Szkół im. Bolesława Prusa, 42-207 Częstochowa, Prusa 20
- Zespół Szkół im. Ks. Dra Jana Zwierza w Ropczycach, 39-100 Ropczyce, Mickiewicza 14
- Zespół Szkół im. Ks. Stanisława Staszica, 39-400 Tarnobrzeg, Kopernika 1
- Zespół Szkół nr 10 im. Prof. Janusza Groszkowskiego w Zabrze, 41-807 Zabrze, Chopina 26
- Zespół Szkół nr 2 im. Gen. Józefa Bema, 05-822 Milanówek, Wójtowska 3
- Zespół Szkół nr 2 im. Ks. Prof. Józefa Tischnera w Żorach, 44-240 Żory, Boryńska 2
- Zespół Szkół nr 2 w Pabianicach im. Prof. Janusza Groszkowskiego, 95-200 Pabianice, Św. Jana 27
- Zespół Szkół nr 4 w Nowym Sączu, 33-300 Nowy Sącz, Św. Ducha 6
- Zespół Szkół nr 40 im. Stefana Starzyńskiego, 03-771 Warszawa, Objazdowa 3
- Zespół Szkół Politechnicznych im. Bohaterów Monte Cassino we Wrześni, 62-300 Września, Wojska Polskiego 1
- Zespół Szkół Ponadgimnazjalnych nr 1 w Jarocinie, 63-200 Jarocin, Franciszkańska 1
- Zespół Szkół Ponadpodstawowych nr 2 im. E. Kwiatkowskiego w Jarocinie, 63-200 Jarocin, Franciszkańska 2
- Zespół Szkół Ponadpodstawowych nr 3 im. Armii Krajowej w Zamościu, 22-400 Zamość, Zamoyskiego 62
- Zespół Szkół Powiatowych im. Stanisława Staszica w Opocznie, 26-300 Opoczno, Kossaka 1a
- Zespół Szkół Publicznych w Szewnie, 27-400 Ostrowiec Świętokrzyski, Szewna, Langiewicza 3
- Zespół Szkół Spożywczych i Hotelarskich w Radomiu, 26-600 Radom, Św. Brata Alberta 1
- Zespół Szkół Technicznych i Licealnych w Piechowicach, 58-573 Piechowice, Przemysłowa 21
- Zespół Szkół Technicznych i Ogólnokształcących nr 3 im. E. Abramowskiego, 40-659 Katowice, Harczerzy Września 1939 2
- Zespół Szkół Technicznych im. Armii Krajowej w Skarżysku-Kamiennej, 26-110 Skarżysko-Kamienna, Tyśiąclecia 22
- Zespół Szkół Technicznych w Kolbuszowej, 36-100 Kolbuszowa, Bytnara 2
- Zespół Szkół w Białowieży, 36-030 Białowieża, Kowala 3
- Zespół Szkół w Zarzeczu, 37-205 Zarzecze, Św. Jana Pawła II 7
- Zespół Szkół Zawodowych nr 1 im. Gen. F. Kleeberga w Dęblinie, 08-530 Dęblin, Tyśiąclecia 3
- Centrum Edukacji Zawodowej, 82-200 Malbork, De Gaulle'a 75a
- Centrum Edukacji Zawodowej i Biznesu, 66-400 Gorzów Wielkopolski, Pomorska 67
- Gminny Zespół Szkolno-Przedszkolny nr 4 w Włocławku, 42-110 Popów, Więcki, Szkolna 1
- Noworudzka Szkoła Techniczna w Nowej Rudzie, 57-401 Nowa Ruda, Stara Droga 4



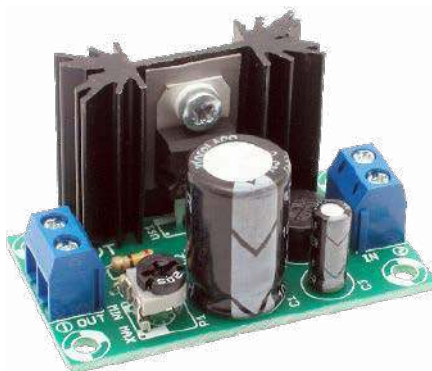
Najbardziej popularne kity AVT



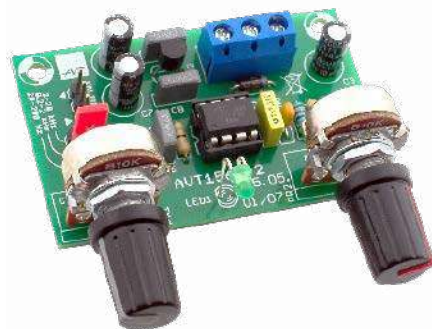
AVT1476 Automatyczny włącznik zmierzchowy
<https://sklep.avt.pl/avt1476.html>



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



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



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



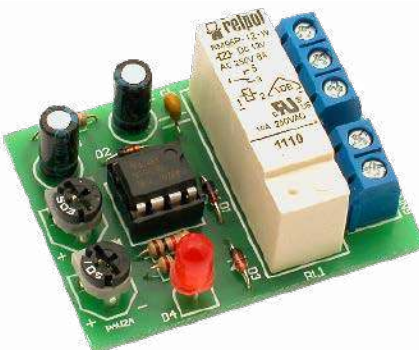
AVT1023 Przedwzmacniacz gramofonowy o charakterystyce RIAA
<https://sklep.avt.pl/avt1023.html>



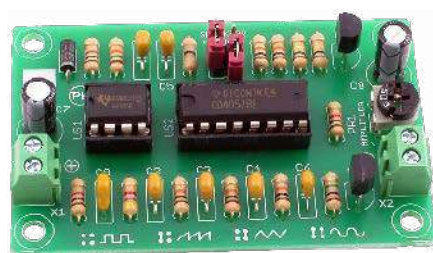
AVT5540 Radio FM z RDS
<https://sklep.avt.pl/avt5540.html>



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



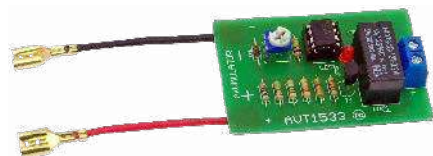
AVT1459 Uniwersalny układ czasowy
<https://sklep.avt.pl/avt1459.html>



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



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



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



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



Pełna oferta na: sklep.avt.pl



Stylowy – nasz nowy ściemniacz – dotykowy i zdalnie sterowany

Nasz nowy ściemniacz współpracuje z większością nowoczesnego oświetlenia, w tym ze ściemnialnymi diodami LED, ściemnialnymi świetlówkami i ściemnialnymi lampami halogenowymi z oprawami ustawianymi („downlight”), jak również ze staromodnymi żarówkami. Ma również naprawdę łatwe w użyciu sterowanie dotykowe, a nawet pilota na podczerwień, co zapewnia najwyższą wygodę! Jest ultranowoczesny, łatwy w budowie i prosty w podłączeniu.



Smukły wygląd, płynne ściemnianie w szerokim zakresie, sterowanie dotykowe i pilot na podczerwień to tylko niektóre z wyjątkowych cech tego nowego ściemniacza dotykowego i na podczerwień Trailing Edge Light Dimmer projektu Silicon Chip, działającego z obcinaniem zbocza opadającego prądu zmiennego 230 V.

Jest idealny do ściemniania nowoczesnych lamp LED i nie posiada pokrętki regulacyjnej w stylu „ubiegłego wieku”. Nie sterujesz swoim telefonem czy tabletem za pomocą pokrętki, prawda? Korzystasz z ekranu dotykowego.

Więc może również chcesz interfejsu dotykowego dla swojego oświetlenia?

A żebyś nawet nie musiał wstawać z fotela i chodzić po pokoju, do sterowania oświetleniem możesz użyć również stylowego, płaskiego pilota na podczerwień. Zapewnia on nawet wybór wstępnych ustawień, dzięki którym szybko uzyskasz pożądaną nastrój!

Praktycznie całe oświetlenie w nowych lub odnowionych mieszkaniach czy domach wykorzystuje obecnie diody LED, co często oznacza, że w domach tych brakuje ściemniaczy. Jeśli lampy są ściemnialne (lub mogą być łatwo zastąpione wersjami ściemnialnymi), to taki jak ten ściemniacz jest świetny jako wyposażenie, ponieważ są chwile, kiedy nie potrzebujesz pełnego światła.

Na przykład, gdy właśnie obudziłeś się rano! Jeśli masz nowoczesny dom, będziesz chciał mieć również nowoczesny ściemniacz, więc ten projekt jest świetnym rozwiązaniem.

Wizualnie, jego minimalistyczny styl z płytką ze szrotowanego aluminium oznacza, że będzie pasował do nowoczesnego mieszkania – chociaż wygląda świetnie również w bardziej tradycyjnym otoczeniu.

A dopełnia całości opcja zdalnego sterowania na podczerwień. Możesz trzymać pilota na stoliku nocnym, w jadalni, w salonie... gdziekolwiek spędzasz dużo czasu.

Oglądasz film? Nie musisz wstawać z kanapy, możesz przyciemnić światła tak jak w kinie. Dziecko wymaga przewijania w nocy? Nie ma potrzeby używania jasnego oświetlenia, które może zakłócić rytm snu. Dopiero co spałeś? Dostosuj się do dnia, powoli zwiększając oświetlenie sypialni.

Nie rzuca się w oczy, ponieważ jedyną widoczną częścią ściemniacza jest płytka ścienna.



Używamy dostępnej w handlu pustej płytki Clipsal Classic 2000, więc wygląda bardzo profesjonalnie i nowocześnie. Jest dodana mała soczewka, aby umożliwić odbiór transmisji podczerwieni z pilota ręcznego. W razie potrzeby można dodać dodatkowe ściennie płytki dotykowe w innych miejscach.

Nie musisz wcale budować samemu ręcznego pilota na podczerwień. Zamiast tego zastosuj małą, taną, dostępną w handlu jednostkę, która wygląda atrakcyjnie i profesjonalnie.

Mamy nadzieję, że udało się nam przekonać Cię do idei tego ściemniacza. Więc czytaj dalej, aby dowiedzieć się, jak działa i co może.

Wymagania dotyczące ściemniania świateł LED

Do ściemniania diod LED lub świetlówek kompaktowych potrzebny jest ściemniacz uniwersalny lub obcinający zbocze opadające zasilania prądem zmiennym z sieci (patrz



Cechy:

- Obcinanie zbocza opadającego – współpracuje z diodami LED
- Elegancki, minimalistyczny wygląd
- Regulacja za pomocą płytki dotykowej – bez pokrętła
- Opcjonalny pilot na podczerwień
- Miękkie włączanie/wyłączanie (jako szybkie zwiększanie lub zmniejszanie jasności)
- Obsługa wielu płytek dotykowych
- Szeroki zakres przyciemniania
- Niski poziom zakłóceń elektromagnetycznych (EMI)
- Możliwość pracy bez podłączenia przewodu neutralnego (N)

ramkę o rodzajach ściemniaczy). Ale trzeba się też upewnić, że lampy są przystosowane do współpracy ze ściemniaczami. Jeśli tak, będzie to napisane na opakowaniu i prawdopodobnie odpowiednie oznaczenie będzie wydrukowane na samych lampach.

Wiele lamp LED i fluorescencyjnych nie jest przystosowanych do pracy w trybie regulacji siły świecenia. Ale my przekonaliśmy się, że nawet te, które dopuszczają pracę ze ściemniaczami, nie zawsze funkcjonują prawidłowo z niektórymi!

Tak więc opłaca się przed instalacją przetestować oświetlenie ze ściemniaczem, którego zamierzasz użyć.

Nasz ściemniacz został przetestowany z kilkoma różnymi diodami LED i okazało się, że działa dobrze (tak jak powinien), ale mogą istnieć pewne lampy LED, które nie będą działać, gdy są z niego zasilane, więc zachowaj ostrożność.

To samo dotyczy halogenów z transformatorem elektronicznymi.

Niektóre są wyraźnie oznaczone jako umożliwiające ściemnianie i większość z nich z tym ściemniaczem będzie działać. Halogeny zasilane przez tradycyjne transformatory z rdzeniem żelaznym mogą również być regulowane. W przypadku zasilania kilku lamp halogenowych lub żarowych za pomocą tego ściemniacza należy uważać, aby nie przekroczyć jego maksymalnego obciążenia 250 W.

Sterowanie ściemnianiem

Lampą lub lampami podłączonymi do ściemniacza można sterować na dwa sposoby: za pomocą płytki dotykowej lub za pomocą pilota na podczerwień.

W przypadku płytki dotykowej ściemnianie jest inicjowane poprzez proste przytrzymanie dłoni na płytce dotykowej. Jasność światła będzie się płynnie zmniejszać lub zwiększać. Chwilowe podniesienie ręki, a następnie

ponowne przyłożenie jej do płytki dotykowej spowoduje przełączenie między zmniejszaniem i zwiększaniem jasności.

Przejsie od całkowitego wyłączenia do całkowitego włączenia lub odwrotnie trwa trzy sekundy. Ściemnianie zatrzymuje się po osiągnięciu minimalnej lub pełnej jasności.

Chcesz mieć natychmiast pełne światło? Szybkie dotknięcie płytki dotykowej powoduje włączenie światła, a kolejne szybkie dotknięcie – jego wyłączenie. Po natychmiastowym włączeniu lampa osiąga pełną jasność przez krótki okres około 0,4 s (400 ms). Daje to efekt płynnego włączania/wyłączania, a nie nagłej zmiany poziomu światła.

Należy pamiętać, że szybkie stuknięcie to każde dotknięcie o długości od 140 do 600 ms.



Spód naszego nowego dotykowego ściemniacza z opcją zdalnego sterowania. Montuje się go na standardowej płytce Clipsal, która z kolei akceptuje standardowy aluminiowy panel zewnętrzny.

Dotknięcia krótsze niż 140 ms są ignorowane (aby zapobiec niepożądanemu przełączaniu światła z powodu zakłóceń elektrycznych itp.), natomiast naciśnięcia dłuższe niż 600 ms inicjują funkcję ściemniania w górę/w dół.

Pilot na podczerwień

Podczas gdy płytka dotykowa posiada efektywnie tylko jeden element sterujący, który musi spełniać kilka funkcji, ręczny pilot IR, jak pokazano poniżej, posiada dziewięć przycisków.

Wszystkie te przyciski sterują w jakiś sposób ściemniaczem.

Przycisk „Operate” lub „On/off” na górze włącza lub wyłącza całkowicie oświetlenie, podobnie jak szybkie dotknięcie płytki dotykowej.

Przycisk w kształcie koła znajdujący się w środku strzałek kierunkowych również włącza lub wyłącza światło, jednak działa on nieco inaczej.

Po jego naciśnięciu w celu włączenia światła, powróci ono do tego samego poziomu jasności, jaki lampa miała przed ostatnim wyłączeniem.

Przytrzymanie przycisków strzałek w górę lub w dół powoduje odpowiednio powolne zwiększanie lub zmniejszanie jasności, przy czym przejście od całkowitego wyłączenia do całkowitego włączenia lub odwrotnie wymaga dziewięciu sekund. Przyciski strzałek w lewo i w prawo również zmniejszają lub zwiększają jasność, ale robią to szybciej, potrzebując tylko dwóch sekund od jednego skrajnego punktu do drugiego.

Przyciski A, B i C zapewniają trzy różne stałe poziomy jasności. Są to odpowiednio ustawienie lampy przyciemnionej, średniej i jasnej. Podobnie jak w przypadku regulacji on/off, zamiast natychmiastowego przejścia do nowego poziomu jasności, urządzenie szybko zwiększa lub zmniejsza jasność w zależności od potrzeb, zapewniając płynne przejście.

Podczas pierwszego uruchomienia ściemniacza, lampa pozostaje wyłączona. Moc w stanie czuwania pobierana przez ściemniacz z sieci wynosi niewiele ponad 1 W.

Co zrobić, gdy nie ma przewodu neutralnego?

W większości instalacji domowych przewód neutralny (zerowy) sieci nie jest doprowadzony do wyłącznika światła. Podłączenie przewodu

Specyfikacja:

Zakres napięcia sieciowego:	200–255 VAC
Częstotliwość sieci:	50 Hz (opcja 60 Hz)
Minimalne obciążenie:	8 W
Obciążenie maksymalne:	250 W
Minimalna jasność:	0% (całkowicie wyłączone)
Jasność maksymalna:	100%, gdy dostępny jest przewód neutralny; gdy przewód „N” nie jest dostępny, regulowana do około 95%
Kroki jasności:	2% (50 kroków od wyłączenia do pełnej jasności)
Czas ściemniania dotykowego:	trzy sekundy od pełnego włączenia do pełnego wyłączenia lub odwrotnie
Polecenia dotykowe:	włącz/wyłącz, jaśniej/ciemniej
Polecenia pilota na podczerwień:	włączanie/wyłączanie, zwiększanie/zmniejszanie jasności szybko (2 s) lub wolno (9 s), plus trzy ustawienia dedykowane
Kroki ściemniania:	50 dla sterowania dotykowego, 100/450 dla podczerwieni (szybkie/wolne ściemnianie)
Czas łagodnego włączenia/wyłączenia:	400 ms
Moc pobierana w stanie spoczynku:	około 1 W
Czas martwy sterowania dotykowego:	dotknięcie przez <140 ms – brak reakcji
Dotknięcie przez 140–600 ms:	naprzemienne działanie włączania/wyłączania
Dotknięcie przez >600 ms:	rozpoczyna zmianę jasności w górę lub w dół; przytrzymaj, aby kontynuować (działanie naprzemienne)

zerowego do lampy jest zwykle wykonane w suficie z puszkii instalacyjnej; tylko przewód fazowy lampy i przewód fazowy zasilający muszą być prowadzone przez puszkę ścienną do wyłącznika lub ściemniacza (to oszczędza kabel!).

Stanowi to problem przy zasilaniu obwodu ściemniacza. Kiedy lampa jest włączona na pełną jasność, teoretycznie nie ma różnicy potencjałów pomiędzy tymi dwoma przewodami, a więc nie ma dostępnej mocy do uruchomienia samego ściemniacza.

Ponieważ taka instalacja jest powszechnie stosowana, jakoś musieliśmy sobie z tym poradzić.

Rozwiązaniem jest ograniczenie maksymalnej jasności lampy tak, aby była tylko

trochę mniejsza od tej osiągananej, gdy otrzymuje ona pełne zasilanie 230 VAC. Ponieważ ściemnianie odbywa się poprzez wyłączenie napięcia sieci przed końcem każdego cyklu, pozostawia to małe okno, w którym napięcie sieciowe jest nadal obecne, ale lampa jest wyłączona. To właśnie w tym czasie ściemniacz pobiera z sieci moc potrzebną do działania.

Jeśli dostępny jest przewód neutralny, to ściemniacz jest zasilany niezależnie od tego, czy lampa jest cały czas włączona, a więc dostępna jest maksymalna jasność.

Nasz ściemniacz uwzględnia obie możliwości okablowania. W przypadku braku przewodu neutralnego maksymalną jasność lampy należy ustawić tak, aby napięcie sieciowe było wystarczające do zasilenia ściemniacza bez migotania lampy.

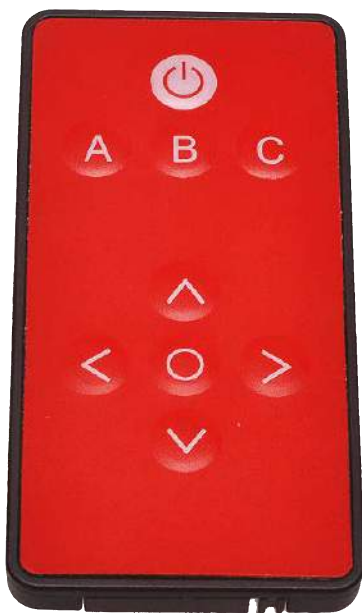
Jak to zrobić, opiszemy w przyszłym mieście w artykule opisującym budowę.

Efekt zatrzaśnięcia lampy LED

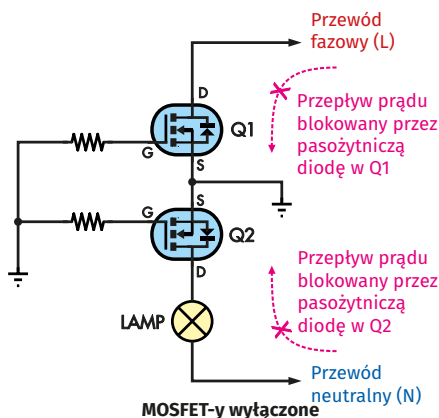
Wiele zasilanych z sieci lamp LED „włącza się skokowo”, gdy regulacja ściemniania jest zwiększana od wyłączenia do niskiego poziomu jasności. Oznacza to, że jasność lampy może nie wzrastać powoli, jak się tego oczekuje; zamiast tego lampa pozostaje całkowicie wyłączona, a następnie nagle ożywa po osiągnięciu określonego ustawienia jasności, z wyższą jasnością niż można by się spodziewać.

Jest to spowodowane tym, że sterownik LED wymaga pewnego napięcia i prądu do uruchomienia. Po uruchomieniu można zwykle ustawić jasność na niższym poziomie i światło pozostanie włączone.

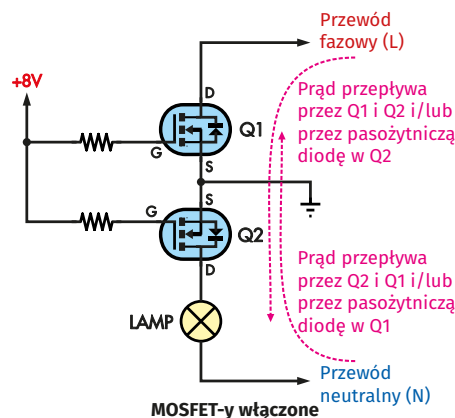
Innymi słowy, nie można sprawić, by lampa świeciła słabo, gdy zwiększa się jej jasność ze stanu wyłączonego. Zamiast tego musisz włączyć ją na pośrednią jasność, a następnie zmniejszyć jej jasność, aby uzyskać pożądane



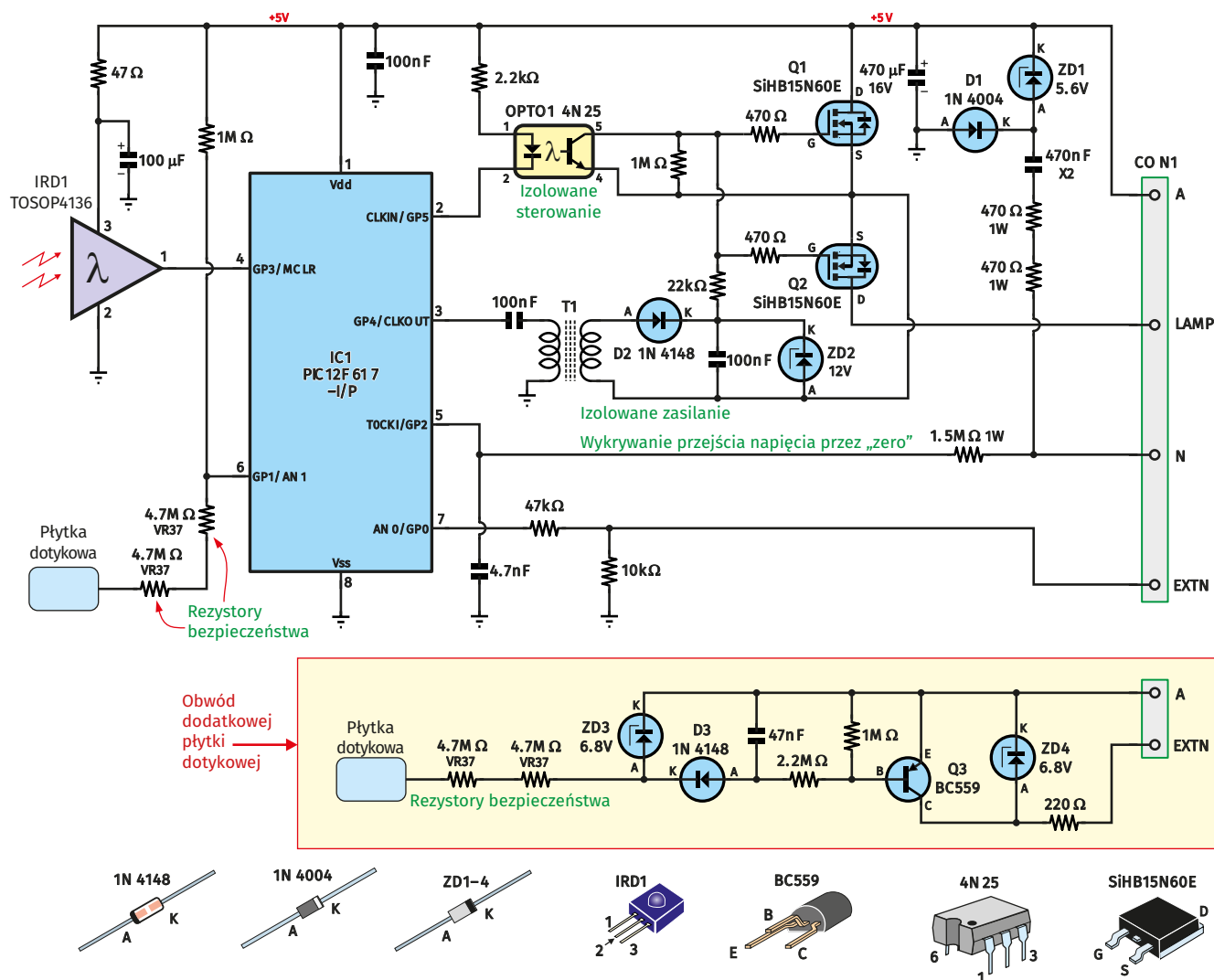
Dziesięcioprzyciskowy pilot, którego użyliśmy do tego projektu. Bez wątpienia istnieje wiele innych, które spełnią swoje zadanie. Nasz pochodzi z firmy Little Bird Electronics (www.littlebird.com.au), zakupiony za książeczką sumę 5,87 \$.



Rysunek 1 a). Kiedy MOSFET-y Q1 & Q2 są wyłączone, prąd nie może płynąć przez lampę niezależnie od polaryzacji napięcia, ponieważ jedna z dwóch diod w strukturze MOSFET-a będzie zawsze spolaryzowana wstecznie i zablokuje przepływ prądu. Gdyby zastosowano pojedynczy MOSFET, zawsze przewodziłby on przez co najmniej połowę czasu, co poważnie ogranicza możliwy zakres ściemniania



Rysunek 1 b). Kiedy bramki MOSFET-ów Q1 i Q2 są na potencjale co najmniej 8 V powyżej napięcia na ich źródłach (połączone tutaj z masą obwodu), oba MOSFET-y przewodzą, a więc prąd może płynąć przez lampę niezależnie od polaryzacji napięcia czy fazy sinusoidalnego napięcia sieci. Pasożytnicza dioda może przewodzić pewien prąd w zależności od napięcia na MOSFET-ach



DOTYKOWY I ZDALNIE STEROWANY ŚCIEMNIACZ OBCIĄJĄCY ZBOCZE OPADAJĄCE

Rysunek 2. Schemat ideowy uniwersalnego ściemniacza sterowanego dotykowo bądź zdalnie. Schemat na żółtym tle pokazuje opcjonalny, dodatkowy obwód, wymagany tylko wtedy, gdy do sterowania potrzebujesz dwóch lub więcej płytek dotykowych. Mikroukład IC1 wykonuje większość operacji, sterując MOSFET-ami Q1 i Q2 poprzez transoptor OPTO1 i izolowane zasilanie wykorzystujące transformator T1. Układ monitoruje na końcówce 5 fазę sieci i określa moment przełączania dwóch MOSFET-ów, aby osiągnąć pożądany poziom jasności lampy

przyciemnienie. Efekt ten jest bardziej zauważalny, gdy ściemniacz działa bez podłączenia oddzielnego przewodu neutralnego.

Opis obwodu

Schemat ideowy uniwersalnego ściemniacza pokazany jest na **rysunku 2**. Pomimo wielu przydatnych funkcji, obwód jest dość prosty, ponieważ większość czynności jest wykonywana przez oprogramowanie mikrokontrolera PIC12F617 (IC1).

MOSFET-y Q1 i Q2 przełączają napięcie sieciowe do lampy (lamp), aby regulować ich jasność. Sposób, w jaki sterują one obciążeniem lampy, pokazano na **rysunku 1**. Dzięki takiej konfiguracji możemy sterować mocą w całym przebiegu sieciowym, włączając lub wyłączając zasilanie sieciowe lampy w dowolnym momencie.

Powodem, dla którego są wymagane w tym celu dwa MOSFET-y, jest to, że MOSFET mocy

zawiera wewnętrzną (zintegrowaną lub „pasozżytniczą”) diodę, która nie może być usunięta; jest ona nieodłączna od struktury MOSFET-a.

Ponieważ prąd zmienia kierunek przepływu w każdej połowie sinusoidalnego przebiegu sieciowego, gdybyśmy użyli pojedynczego MOSFET-a, jego wewnętrzna dioda przewodziłaby przez połowę czasu i przykładała w tym czasie pełne napięcie do obciążenia, niezależnie od tego, czy MOSFET byłby włączony, czy nie.

Łącząc dwa MOSFET-y szeregowo przeciwnie (a więc z wewnętrznymi diodami w przeciwnych kierunkach), w momencie wyłączonych MOSFET-ów, przepływ prądu jest zablokowany w obu kierunkach, jak pokazano na **rysunku 1a**.

Gdy MOSFET-y są włączone poprzez podniesienie napięcia na ich bramkach, jak na **rysunku 1b**, prąd może płynąć w każdym

kierunku przez kanały MOSFET-ów, i w niewielkim stopniu, przez pasozżytnicze diody struktury.

Diody te będą przewodzić tylko wtedy, gdy prąd płynący przez kanał będzie na tyle duży, że wytworzy różnicę napięć na MOSFET-ach (ze względu na rezystancję kanału), która będzie wyższa niż napięcie przewodzenia diody.

Sterowanie bramkami MOSFET-ów

MOSFET-y: Q1 i Q2 włączają się, gdy napięcie na ich bramkach jest wyższe niż napięcie na wspólnym zacisku ich źródeł. Dla tych konkretnych MOSFET-ów, aby przewodziły z minimalnymi stratami, różnica ta musi wynosić co najmniej 8 V.

Ale napięcie bramki nie może być zbyt wysokie, ponieważ powyżej 20 V może uszkodzić MOSFET-y. To sprawia, że zapewnienie

odpowiedniego napięcia, aby utrzymać je włączone w razie potrzeby, jest dość trudne.

Najprostszym rozwiązaniem jest galwaniczne odizolowanie źródła napięcia bramki od reszty układu. Wynika to z faktu, że szyna +5 V jest podłączona bezpośrednio do przewodu fazowego („Active”), który jest niezbędny do działania sterowania dotykowego.

Problem polega na tym, że nawet gdybyśmy mogli wygenerować wymagane zasilanie na poziomie 8...20 V, a następnie przyłożyć je do bramek MOSFET-ów, z ich zaciskami źródeł podłączonymi do masy obwodu, to gdy tylko Q1 się włączy, połączy przewód fazowy (+5 V) z masą; efektywnie zwarłoby to zasilanie 5 V i tym samym wyłączyłoby cały układ.

Eliminujemy ten problem poprzez zastosowanie „pływającego” („floating”) zasilania bramki, zaciski źródeł MOSFET-ów nie muszą już być podłączone do masy obwodu, abyśmy mogli ustawiać napięcie na ich bramkach.

Transformator T1 zapewnia zarówno tę izolację, jak również zwiększa napięcie sterujące 5 V, aby uzyskać na bramkach napięcie powyżej 8 V. Transformator ten składa się z toroidalnego rdzenia ferrytowego do pracy z wysoką częstotliwością, z dwoma uzwojeniami DNE. Uzwojenie pierwotne (12 zwojów) jest zasilane, poprzez kondensator sprzęgający AC 100 nF, napięciem prostokątnym o częstotliwości 2 MHz generowanym na wyjściu zegarowym IC1 (końcówka 3),.

Uzwojenie wtórne ma cztery razy więcej zwojów niż pierwotne (tzn. 48 zwojów), i jest od niego izolowane. Napięcie zmienne z uzwojenia wtórnego jest prostowane półokresowo przez diodę D2 i filtrowany za pomocą kondensatora 100 nF. Dioda Zenera ZD2 12 V ogranicza napięcie maksymalne. W rezultacie otrzymujemy napięcie 10 VDC, którego ujemny biegun jest podłączony do połączonych źródeł MOSFET-ów Q1 i Q2, a dodatni do ich bramek poprzez rezystor 22 kΩ i dwa rezystory 470 Ω.

Napięcie bramek jest sterowane za pomocą transoptora OPTO1. Konieczne jest zachowanie izolacji pomiędzy IC1, którego szyna 5 V jest podłączona do przewodu fazowego („Active”), a bramkami MOSFET-ów Q1 i Q2. Gdy wyjście GP5 IC1 (końcówka 2) przechodzi w stan wysoki, wewnętrzna dioda podczerwieni OPTO1 jest wyłączona. Gdy to wyjście jest w stanie niskim, przez diodę przepływa prąd około 2 mA z zasilania 5 V, ograniczony rezystorem 2,2 kΩ.

Gdy dioda ta oświetla wewnętrzny fototranzystor OPTO1, zwiera napięcie zasilania 10 V bramek MOSFET-ów Q1 i Q2, wyłączając je. Gdy fototranzystor wyłączy się, zasilanie 10 V może ponownie podciągnąć bramki MOSFET-ów do wysokiego poziomu, dzięki czemu włączają się one ponownie.

Bramki MOSFET-ów, aby zapobiec oscylacjom przy włączaniu, są odizolowane od siebie rezystorami 470 Ω. Rezystor 1 MΩ pomiędzy kolektorem a emitorem tranzystora wyjściowego OPTO1 zapewnia, że Q1 i Q2 pozostają wyłączone, gdy układ mikrokontrolera IC1 nie jest zasilany.

Wykrywanie przejścia przez zero w sieci

Aby prawidłowo ustawić moment przełączania Q1 i Q2, tak, aby uzyskać pożądaną poziom ściemniania, IC1 posiada zegar, który jest zsynchronizowany z przejściem przez zero w sieci, tj. stanem, gdy napięcia przewodu fazowego i neutralnego są sobie równe (co zdarza się w naszej sieci o częstotliwości 50 Hz 100 razy na sekundę.). Potrzebny jest zatem sposób na wykrycie tego stanu, aby zsynchronizować zegar IC1 z częstotliwością sieci.

Detekcja następuje na końcówce 5 IC1, poprzez rezystor 1,5 MΩ podłączony do przewodu neutralnego (połączenie może iść przez lampę(y), w przypadkach gdy nie ma osobnego przewodu neutralnego). Wykrycie przejścia przez zero następuje tylko przy przejściu ujemnym,

Ściemniacze obcinania zbocza narastającego vs obcinania zbocza opadającego

Nasza sieć elektryczna (nominalnie 230 VAC) to sinusoida o częstotliwości 50 Hz. Aby zapewnić funkcję ściemniania, jest ona zwykle „siekana” w jakiś sposób przez urządzenie przetwarzające, które przerwa zasilanie sieciowe lampy, aby zmniejszyć jej jasność. Im dłużej to urządzenie przetwarzające jest włączone, tym jaśniej świeci lampa.

Najczęstszą metodą cięcia kształtu fali sieci jest „kontrola fazy”, gdzie moc jest dostarczana w sposób ciągły przez część każdej półokresu cyklu sieci. Każda półokresu cyklu sieci trwa 10 ms i przez cały ten okres napięcie przewodu fazowego jest albo wyższe albo niższe od napięcia przewodu neutralnego.

Pomiędzy każdą półokresu sinusoidy występuje „przejście przez zero”, gdzie napięcia: fazowe i neutralne są sobie równe. Przyjmuje się umownie, że każdy pełny przebieg sieciowy (trwający 20 ms) ma fazę od 0° do 360°, przy czym dwa przejścia przez zero mają kąty fazowe 0° i 180°, a szczyty napięcia znajdują się przy 90° i 270°; patrz **rysunek 3**.

Terminy: przyciemnianie poprzez obcinanie zbocza narastającego („leading-edge dimming”) lub obcinanie zbocza opadającego („trailing-edge dimming”) odnoszą się do faktu, że istnieją dwa główne sposoby zapewnienia sterowania obciążenia fazy. Działają one podobnie, ale są zazwyczaj stosowane w różnych okolicznościach.

Jeśli opóźnimy podawanie napięcia sieciowego do obciążenia do momentu osiągnięcia określonego kąta fazowego – powiedzmy 45° – a następnie pozwalamy na dalsze jego podawanie do momentu rozpoczęcia następnego półokresu, zmniejszymy napięcie RMS (RMS – Root Mean Square, zwane też napięciem skutecznym lub wartością średnią kwadratową) na obciążeniu, a tym samym zmniejszamy moc pobieraną przez obciążenie. Jest to znane jako przyciemnianie przez obcinanie zbocza narastającego, ponieważ opóźnia się początkowe, obcięte, zbocze napięcia sieciowego „widziane” przez obciążenie; patrz **rysunek 4**.

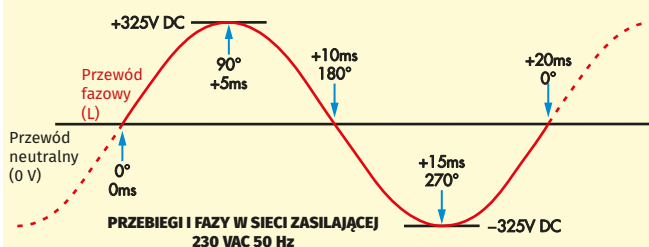
Alternatywnie, jeśli podamy moc do obciążenia od początku sinusoidy napięcia (tj. przy 0°), a następnie odetniemy sinusoidę przed końcem półokresu – powiedzmy przy 315° – wtedy przesuwasz krawędź obciętego zbocza przebiegu sieciowego widzianego przez obciążenie i to jest znane jako ściemnianie przez obcięcie zbocza opadającego (trailing edge dimming); patrz **rysunek 5**.

W obu tych przykładach napięcie RMS przyłożone do obciążenia jest takie samo – około 219 V RMS w systemie o nominalnym napięciu 230 VAC.

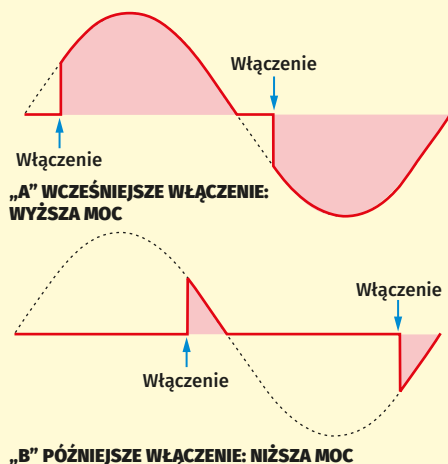
Ściemniacz pierwszego typu jest stosowany od około 50 lat, głównie do ściemniania lamp żarowych. Dzieje się tak dlatego, że można go zrealizować za pomocą prostego obwodu wykorzystującego triak, jak pokazano na **rysunku 6**.

Triak jest pięciowarstwowym elementem półprzewodnikowym, który włącza się, gdy jego bramka jestysterowana. Nie można go jednak wyłączyć poprzez bramkę; zamiast tego wyłącza się samorzutnie, gdy przepływ prądu przez niego spadnie do wartości bliskiej zeru.

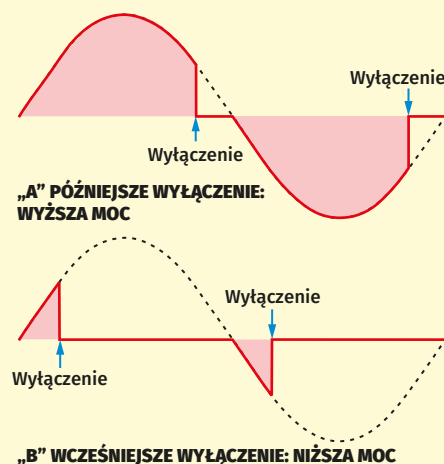
W praktyce, podczas zasilania obciążenia rezystancyjnego, takiego jak lampa żarowa, triak wyłącza się, gdy napięcie sieciowe jest



Rysunek 3: australijskie napięcie sieci, jak w Polsce, jest w przybliżeniu sinusoidalne i ma częstotliwość 50 Hz (tj. okres 20 ms), przy nominalnym napięciu 230 VAC. Przejście od ujemnego do dodatniego napięcia względem potencjału przewodu neutralnego jest uważane za początek każdego cyklu; odpowiada to kątowi fazowemu 0°. Drugie przejście przez zero jest przy kącie fazowym 180°, a dwa szczyty napięcia są przy 90° i 270°. Podczas regulacji fazowej moc dostarczana do obciążenia jest przetwarzana w statym, ustalonym punkcie cyklu.



Rysunek 4. Ściemniacz obcinania (włączania, sterowania) zbocza narastającego („leading edge”) zmienia punkt załączenia w cyklu sieciowym, ale zawsze wyłącza się przy przejściu przez zero. Im wcześniej się włączy, tym więcej mocy jest dostarczane do obciążenia i tym jaśniejsze jest światło. Z diodami LED, ani z innymi lampami, które mają sterowniki elektroniczne, taki ściemniacz nie działa dobrze.



Rysunek 5. Ściemniacz typu obcinania (wyłączania, sterowania) zbocza opadającego („trailing edge”) osiąga podobny rezultat, ale zamiast tego włącza lampę przy przejściu przez zero, a następnie wyłącza ją w pewnym momencie późniejszego cyklu sieciowego. Im późniejsze jest wyłączenie, tym jaśniej świeci lampka. Schemat ten jest kompatybilny z lampami posiadającymi sterowniki elektroniczne, w tym z większością ściemniaków LED.

bliskie 0 V. Stąd, w prosty sposób, można zapewnić sterowanie fazy na zboczu narastającym.

Ściemnianie diod LED

Ściemniacze obcinania zbocza narastającego nie nadają się do stosowania z lampami LED lub lampami halogenowymi z transformatorami elektronicznymi. To dlatego, że w obu przypadkach obwód sterowania prostuje napięcie sieciowe, a następnie filtruje je za pomocą kondensatora. To właśnie ładunek na tym kondensatorze zasilą następnie pozostałe obwody, w tym lampę.

Jeśli do tego typu obwodu zostanie nagle podłączone napięcie, diody w prostowniku natychmiast przewodzą i pobierają duży prąd, aby szybko naładować kondensator.

Taki wysoki prąd rozruchowy jest do opanowania, jeśli występuje rzadko, np. po włączeniu światła, ale jeśli dzieje się to w każdym cyklu sieciowym (kiedy triak w ściemniaczu się włącza), może to doprowadzić do przegrzania i awarii.

Nawet jeśli ściemniacz i lampka tolerują taką sytuację, to i tak można by się spodziewać brzęczenia, skoków napięcia, nadmiernych zakłóceń elektromagnetycznych (EMI) i migotania lampy zamiast ściemniania. Tak więc wyraźnie nie jest to wykonalne.

Rozwiązaniem jest zastosowanie ściemniacza obcinającego zbocze opadające. Urządzenie przełączające włącza się teraz przy przejściu napięcia sieci przez zero, gdy nie ma różnicy potencjałów między przewodem fazowym i neutralnym. Napięcie lampy

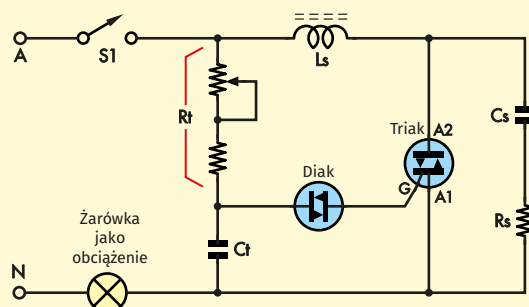
wzrasta stosunkowo powoli, a diody prostownicze przewodzą, gdy napięcie sieciowe przekroczy napięcie na (częściowo naładowanym w poprzednim cyklu) kondensatora. Prąd jest pobierany z sieci w znacznie mniejszych i lepiej tolerowanych impulsach.

Ponieważ diody LED w zasadzie opanowały rynek oświetleniowy, ściemniacze starego typu („leading edge”) ustępują miejsca ściemniaczom typu obcinania zbocza opadającego („trailing edge”) lub uniwersalnym (które mogą pracować w obu trybach).

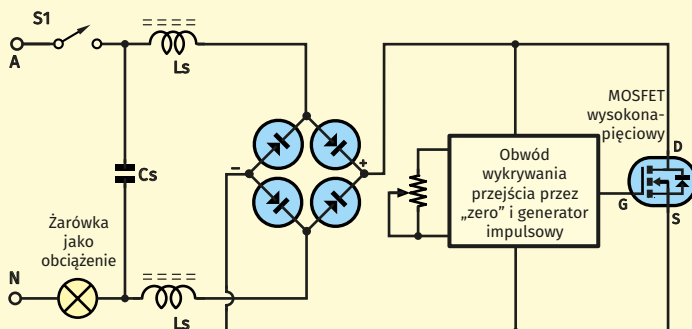
W przypadku nowoczesnych ściemniaczy konieczne jest zastosowanie innego elementu przełączającego niż triak, takiego, który może być wyłączony za pomocą sterowania bramką w dowolnej części przebiegu sieciowego. **Rysunek 7** pokazuje uproszczony obwód typowego ściemniacza z obcinaniem zbocza opadającego. Elementem przełączającym jest zwykle jeden lub dwa MOSFET-y lub IGBT (transystory bipolarnie z izolowaną bramką).

W przedstawionym obwodzie używamy dwóch MOSFET-ów, połączonych antyszeregowo (źródło do źródła). Zjrzyj do opisu obwodu, aby dowiedzieć się, dlaczego zastosowaliśmy taką konfigurację. Pozwala nam ona na włączanie lub wyłączanie zasilania lampy w dowolnym momencie cyklu sieciowego.

Więcej informacji na temat ściemniaczy oraz ich zastosowania z lampami LED można znaleźć w artykule zatytułowanym „LED lights/ downlights and dimmers” w numerze Silicon Chip z lipca 2017 roku (www.siliconchip.com.au/Article/10712).



Rysunek 6. Pokazuje jak prosty może być ściemniacz zbudowany na triaku. Choć wygląda to na uproszczony układ, rzeczywisty ściemniacz jest niewiele bardziej skomplikowany. Rt i Ct zapewniają zmienną stałą czasową, która ustala, jak późno w cyklu diak podaje impuls wyzwalający triak. Ten z kolei zasilą lampę prądem. Triak wyłącza się samorzutnie przy najbliższym przejściu przez zero.



Rysunek 7. Obwód ściemniacza obcinającego zbocze opadające jest nieco bardziej skomplikowany. Ten uproszczony schemat ukrywa większość złożonej struktury i jej funkcji wewnątrz żółtej ramki po prawej stronie. Zasilanie sieciowe jest prostowane, aby zapewnić zasilanie układowi sterowania, a także po to, aby można było zastosować pojedynczy MOSFET, ponieważ musi on wtedy przełączać tylko napięcie o jednej polaryzacji. Jest wymagany (nie pokazany) kondensator podtrzymujący zasilanie obwodu sterowania, gdy MOSFET jest włączony.



Ściemniacz jest zbudowany z dwóch płytek PCB w układzie „kanapki”, jedna na drugiej. Zespół jest zamontowany na płytce Clipsal z płytką dotykową po przeciwnej stronie.

oczywiste jest, że przejście dodatnie następuje 10 ms później.

Rezystor ograniczający prąd 1,5 MΩ tworzy w połączeniu z kondensatorem 4,7 nF filtr dolnoprzepustowy RC, który jest niezbędny do zmniejszenia efektów różnych sygnałów i zakłóceń, które mogą być nałożone na sieć 50 Hz. W przeciwnym razie spowodowałyby one zauważalne migotanie lampy z powodu przypadkowego wykrywania napięcia zerowego.

Opóźnia to jednak wykrycie przejścia przez zero. IC1 kompensuje to znane opóźnienie programowo, aby określić rzeczywisty moment przejścia przez zero.

Należy zauważyć, że tylko jedno z dwóch przejść przez zero jest faktycznie wykrywane. Drugie jest obliczane na jego podstawie w oparciu o oczekiwane opóźnienie dla częstotliwości sieci 50 Hz (lub 60 Hz, co nas nie interesuje). Poprawia to stabilność ściemniacza, zwłaszcza gdy pracuje bez podłączonego przewodu neutralnego.

Ponadto oprogramowanie sprawdza stan wejścia 5 tylko w okolicach spodziewanego czasu przejścia przez zero.

Gdyby detekcja napięcia zerowego była aktywna przez cały cykl, to włączenie i wyłączenie lampy przez ściemniacz powodowałoby fałszywą detekcję ze względu na zmianę napięcia na wejściu przewodu neutralnego ściemniacza przy przełączaniu lampy. Jest to ważne zwłaszcza wtedy, gdy detekcja przejścia przez zero odbywa się poprzez lampę/y, a nie bezpośrednio z przewodu neutralnego.

Zasilanie

Jak mogłeś się zorientować z powyższego wyjaśnienia, konfiguracja zasilania dla tego układu jest ściśle związana z jego działaniem.

Dzieje się tak dlatego, że ściemniacz zasilany jest z tej samej sieci, którą zarówno monitoruje (dla zdarzeń związanych z przejściem przez zero), jak i przełącza. A w przypadku, gdy do urządzenia nie jest podłączony przewód neutralny, sprawa staje się naprawdę trudna.

Poza opisaną powyżej izolowaną galwanicznie od IC1 sekcją sterownika bramek MOSFET-ów, reszta układu także „pływa” z nominalnie 230 VAC przebiegiem sieciowym na przewodzie fazowym („Active”). W rzeczywistości, przewód fazowy i szyna zasilająca +5 V są połączone bezpośrednio. Można by więc myśleć, że prąd zasilający obwód jest pobierany z przewodu neutralnego; w praktyce, płynie on pomiędzy przewodem fazowym i neutralnym, z kierunkiem prądu zmieniającym się 100 razy na sekundę.

Ten prąd płynie z/do przewodu neutralnego przez dwa szeregowo połączone rezystory 470 Ω/1 W i kondensator 470 nF typu X2.

Gdy napięcie na przewodzie fazowym jest niższe od napięcia na przewodzie neutralnym, kondensator 470 nF ładuje się przez dwa szeregowo rezystory 470 Ω i diodę Zenera (ZD1), która wtedy przewodzi i działa jak standardowa dioda.

Gdy napięcie na przewodzie fazowym jest wyższe od napięcia na przewodzie neutralnym, kondensator 470 nF rozładowuje się przez diodę D1, ładując kondensator elektrolityczny

470 mF (obniżając potencjał jego ujemnej okładki), który następnie zasila resztę obwodu.

Gdy napięcie na kondensatorze 470 nF osiągnie 5 V, każdy dodatkowy prąd pobierany przez obwód zasilania płynie przez ZD1, aby zapobiec dalszemu wzrostowi napięcia. Ogranicza to napięcie zasilania do 5 V (a nie do 5,6 V), ze względu na 0,6 V napięcia przewodzenia diody D1. Inaczej mówiąc, napięcie na kondensatorze 470 nF jest równe napięciu diody Zenera ZD1 pomniejszonemu o napięcie przewodzenia diody D1.

Gdy ZD1 przewodzi, to impedancja dwóch rezystorów 470 Ω i kondensatora 470 nF zapobiega pobieraniu nadmiernego prądu z sieci.

Rezystory 470 Ω ograniczają również prąd rozruchowy przy każdym włączeniu włącznika światła, ponieważ chwilowe przyłożone napięcie może wynosić nawet 350 V DC (typowo napięcie szczytowe pomiędzy przewodem fazowym a neutralnym przy zasilaniu sieciowym 230 V wynosi 325 V (patrz rysunek 3), ale w niektórych rejonach z anormalnie wysokim napięciem sieci zasilającej może być znacznie wyższe). Nawiasem mówiąc sytuacja taka spotykana jest coraz częściej, ze względu na domowe instalacje fotowoltaiczne przesyłające w ciągu dnia nadmiar wytworzonej energii do sieci publicznej.

Jeśli w miejscu, w którym zainstalowany jest ściemniacz, nie ma dostępnego przewodu neutralnego, podłączenie do zacisku neutralnego jest wykonywane poprzez obciążenie lampy.

W tym przypadku zasilanie obwodu jest dostępne tylko wtedy, gdy lampa jest wyłączona. Gdy lampa jest włączona, napięcie na MOSFET-ach Q1 i Q2 jest mniejsze niż 1 V i jest ono niewystarczające do wygenerowania napięcia zasilania 5 VDC.

Zatem, gdy nie ma oddzielnego przewodu neutralnego, zakres regulacji fazy musi być ograniczony do mniej niż pełnego cyklu sieci.

W ten sposób, aby zapewnić wystarczającą ilość mocy do uruchomienia reszty obwodu, lampa nie świeci przez cały cykl. Oznacza to, że bez podłączenia przewodu neutralnego maksymalna jasność lampy jest ograniczona.

Kontrola ściemniania

Sterująca płytka dotykowa jest podłączona do końcówki 6 IC1 poprzez dwa szeregowo wysokonapięciowe rezystory bezpieczeństwa 4,7 MΩ, natomiast opcjonalny moduł dodatkowej płytki sterującej (drugiej płytki dotykowej – lub kolejnej) jest podłączony do końcówki 7 poprzez standardowy rezystor 47 kΩ.

Istotne jest, aby używać dedykowanych rezystorów (tj. Vishay VR37 seria 4,7 MΩ lub analogicznych). Oprócz ograniczenia przepływu



Oto dodatkowa dotykowa płytki sterująca, która jest podobna do głównej płytki ściemniacza i montowana w ten sam sposób (patrz poniżej)...

prądu przez osobę dotykającą płytki dotykowej do poniżej 36 mA, te konkretne rezystory zapewniają wysoki margines bezpieczeństwa, ponieważ ich napięcie znamionowe wynosi 2,5 kVAC dla każdego. Dwa rezystory połączone szeregowo mają napięcie znamionowe >5 kV, zapewniając dodatkowe bezpieczeństwo.

Zaufaj nam; na pewno nie chcesz ryzykować bezpośredniego podłączenia się do przewodu fazowego – to co najmniej boli! I dotyczy to wszystkich innych osób, które będą używać ściemniacza, nie tylko Ciebie.

Zazwyczaj wejście z płytki dotykowej na końcówce 6 układu IC1 jest utrzymywane na poziomie 5 V (tj. na potencjale przewodu fazowego) przez rezystor podciągający 1 MΩ,



...a tutaj dodatkowa płytki zamontowana na panelu Clipsal.

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

- 1 szt. dwustronna płytka drukowana o kodzie 10111191, 66×104 mm
- 1 szt. płytka drukowana o kodzie 10111193, 58,5×104 mm
- 1 szt. standardowa płytka z serii C2000 z pustą pokrywą aluminiową Clipsal CLOPTO1031VXBA
- 1 szt. blok montażowy CL1449AWE (opcja; patrz tekst artykułu konstrukcyjnego w przyszłym miesiącu)
- 1 szt. soczewka Fresnela do czujnika podczerwieni (Murata IML0688) [RS components Cat 124-5980]
- 1 szt. pilot na podczerwień [LitT1e Bird Electronics SF-COM-14865]
- 1 szt. pastylkowa bateria litowa CR2025 3,2 V (do pilota IR)
- 1 szt. 4-drożna listwa zaciskowa, 25 A/300 VAC, moduł 9,5 mm [Jaycar HM-3162] (CON1)
- 1 szt. rdzeń toroidalny, materiał w.cz. L8, 18×10×6 mm [Jaycar LO1230] (T1)
- 3 szt. poliamidowe kablowe opaski zaciskowe długości 100 mm
- 4 szt. śrubki M3×6
- 8 szt. nakrętek sześciokątnych M3
- 1 podstawka układu scalonego DIL-8 (do IC1)
- 1 odcinek długości 25 mm rurki termokurczliwej o średnicy 16 mm
- 1 odcinek długości 1,26 m emaliowanego drutu miedzianego o średnicy 0,25 mm (do nawinięcia T1)
- 1 odcinek długości 15 mm ocynowanego drutu miedzianego o średnicy 0,71 mm

Półprzewodniki:

- 1 szt. mikrokontroler PIC12F617-L/P zaprogramowany wsadem 1011119A.hex (IC1)
- 1 szt. transoptor 4N25 (OPTO1)
- 1 szt. odbiornik podczerwieni TSOP4136 (IRD1)
- 2 szt. N-kanalowe MOSFET-y SIHB15N60E, 15 A/600 V (Q1, Q2)
- 1 szt. dioda 1N4004 1A/400 V (D1)
- 1 szt. dioda 1N4148 (D2)
- 1 szt. dioda Zenera 5,6 V/1 W (ZD1)
- 1 szt. dioda Zenera 12 V/1 W (ZD2)

Kondensatory:

- 1 szt. kondensator elektrolityczny 470 mF/16 V PC MB
- 1 szt. kondensator elektrolityczny 100 mF/16 V PC MB
- 1 szt. przeciwzakłóceńowy kondensator polipropylenowy MKP-X2 470 nF/275 VAC, moduł 22,5 mm
- 3 szt. kondensatory poliestrowe MKT 100 nF/63/100 V
- 1 szt. kondensator poliestrowy MKT 4,7 nF/63/100 V

Rezystory: (0.25W, 1%, chyba że podano inaczej)

- 2 szt. rezystory bezpieczeństwa 4,7 MΩ/0,5 W Vishay VR37 3,5 kV [RS Components 484-4400]
- 1 szt. 1,5 MΩ 1W 5%
- 2 szt. 1 MΩ
- 1 szt. 47 kΩ
- 1 szt. 22 kΩ
- 1 szt. 10 kΩ
- 1 szt. 2,2 kΩ
- 2 szt. 470 Ω/1 W 5%
- 2 szt. 470 Ω
- 1 szt. 47 Ω

Części dla każdej dodatkowej płytki dotykowej

- 1 szt. dwustronna płytka drukowana kod 10111192, 58,5×104 mm
- 1 szt. płytka drukowana kod 10111193, 58,5×104 mm
- 1 szt. standardowa płytka z serii C2000 z pustą pokrywą aluminiową Clipsal CLOPTO1031VXBA
- 1 szt. blok montażowy CL1449AWE (opcja; patrz tekst)
- 1 szt. 4-drożna listwa zaciskowa, 25 A 300 VAC moduł 9,5 mm [Jaycar HM-3162] (CON1)
- 4 szt. śrubki M3×6 z tłem walcowym
- 8 szt. nakrętek sześciokątnych M3
- 1 odcinek długości 15 mm ocynowanego przewodu miedzianego o średnicy 0,71 mm

Półprzewodniki:

- 1 szt. tranzystor BC559 PNP (Q3)
- 1 szt. dioda 1N4148 (D3)
- 2 szt. diody Zenera 6,8 V/1 W (ZD3, ZD4)

Kondensatory:

- 1 szt. kondensator poliestrowy MKT 47 nF

Rezystory: (0.25 W, metalizowane 1% jeśli nie podano inaczej)

- 2 szt. rezystory bezpieczeństwa 4,7 MΩ/0,5 W Vishay VR37 3,5 kV [RS Components 484-4400]
- 1 szt. 2,2 MΩ
- 1 szt. 1 MΩ
- 1 szt. 220 Ω

Dodatkowe części do zewnętrznego sterowania przełącznikiem

- 1 szt. chwilowy zwierający włącznik sieciowy Clipsal 30MBPR i pasująca ramka lub standardowa płytka pojedynczego przełącznika przyciskowego (np. od dzwonka)

Pilot na podczerwień wykorzystujący protokół PDP (Pulse Distance Protocol – Kodowanie Długości i Odstępów Ramek Impulsów)

Większość kontrolerów podczerwieni wykorzystuje częstotliwość modulacji 36...40 kHz, typowo 38 kHz, gdzie dioda nadajnika podczerwieni jest włączana i wyłączana z tą częstotliwością. Odbywa się to w seriach (ramkach, pakietach), przy czym długość i odstępy między ramkami (pauzami) wskazują, który przycisk został wciśnięty.

Seria impulsów i pauz jest zwykle kodowana w określonym formacie (lub protokole) i istnieje kilka różnych stosowanych powszechnie protokołów. Obejmuje to protokoły RC5 i RC6 wykorzystujące kodowanie przy zastosowaniu formatu Manchester, których pomysłodawcą jest firma Philips. Istnieje również protokół Pulse Width Protocol używany przez firmę Sony. Ręczny pilot IR użyty w tym projekcie wykorzystuje protokół Pulse Distance Protocol, pochodzący od firmy NEC.

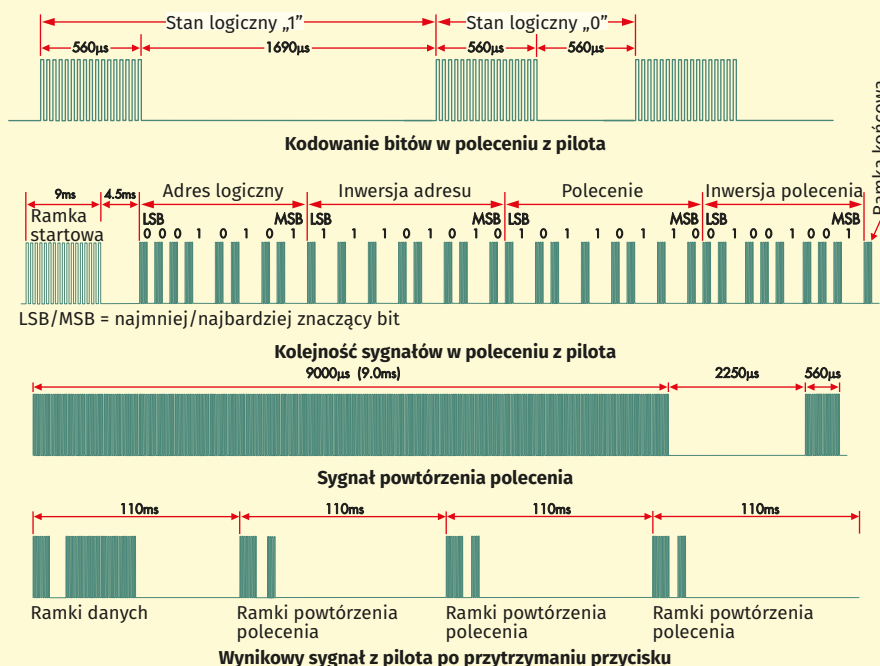
Jeśli jesteś zainteresowany szczegółami dotyczącymi wszystkich tych protokołów, a także innych, zobacz notę aplikacyjną AN3053 firmy Freescale Semiconductors (dawniej Motorola) pod adresem: <https://bit.ly/3ksdVT>.

Szczegóły tego protokołu pokazane są na sąsiednim schemacie (**rysunek 8**). Jest on rozbity na cztery panele.

Górny panel pokazuje, jak transmitowane są bity „zero” i „jeden”. Oba zaczynają się od ramki 560 ms impulsów modulowanych z częstotliwością 38 kHz. Logiczna „1” ma następnie 1690 ms pauzy, podczas gdy logiczne „0” ma krótszą, 560 ms, pauzę.

Drugi panel pokazuje strukturę transmisji pojedynczego polecenia. Zaczyna się ona od ramki 9 ms i 4,5 ms przerwy. Po tym następuje osiem bitów adresowych, potem kolejne osiem bitów, które są „dopełnieniem” tych samych ośmiu bitów adresowych (inwersja adresu, bit „0” jest zamieniony na bit „1”, a bit „1” zamienia się na bit „0”). Bity adresu identyfikują sprzęt sterowany przez pilota (TV, DVD, radio, itp.).

Po nich następuje osiem bitów polecenia, plus ich dopełnienie (inwersja), wskazujących, która funkcja powinna być aktywowana, a następnie finalna ramka 560 ms, kończąca transmisję.



Rysunek 8. Szczegóły czasowe protokołu PDP zdalnego sterowania podczerwienią. Pierwszy panel pokazuje czas trwania ramek logicznych stanów „0” i „1” (składających się z pakietów-ramek impulsów IR o częstotliwości 38 kHz; jeden impuls w ramce trwa ok. 26 ms). Drugi panel pokazuje, jak te bity danych są połączone z ramką startową i ramką końcową („tail burst”), aby zakodować naciśnięcie przycisku pilota. Trzeci panel pokazuje sygnał powtórzenia nadawany po przytrzymaniu przycisku pilota, a czwarty panel pokazuje serię sygnałów, które wynikają z naciśnięcia, a następnie przytrzymania przycisku.

Należy zauważyć, że dane adresowe i polecenia są wysyłane z najmniej znaczącym bitem jako pierwszym.

Komplementarne bajty adresu i polecenia są wysyłane jako sposób wykrywania błędów. Jeśli otrzymana wartość danych komplementarnych nie jest uzupełnieniem otrzymanych wcześniej danych, oznacza to, że jedna lub druga wartość została nieprawidłowo wykryta i zdekodowana.

Brak danych uzupełniających sugeruje, że odebrane dane nie odpowiadają protokołowi PDP, a więc sygnał jest wysyłany przez inny pilot ręczny.

Po wciśnięciu przycisku, jeśli jest on dalej przytrzymywany, pilot wysyła powtarzające się sygnały. Składają się one z ramki 9 ms; 2,25 ms pauzy; i ramki 560 ms. Ten sygnał jest powtarzany w odstępach 110 ms, aż do zwolnienia przycisku.

Sygnał powtórzenia jest używana do poinformowania odbiornika, aby ewentualnie powtórzył tę konkretną czynność, w zależności od tego, co jest wybrane.

Na przykład, polecenie „zwiększenie głośności” lub „przeskocz do następnego kanału”; ale już polecenie „Wycisz” („Mute”) czy „Kanał nr 3” może się nie powtarzać.

ale jeśli płytką dotykową zostanie dotknięta, pojemność osoby dotykającej względem masy (uziemienia) otoczenia sprawdza potencjał płytki dotykowej do potencjału masy otoczenia, gdy napięcie fazowe wzrasta.

To efektywnie ściąga końcówkę 6 do potencjału masy zasilania na wystarczająco długo, aby IC1 mógł wykryć ten stan.

Wejście rozszerzeń na końcówce 7 jest normalnie utrzymywane w stanie niskim przez rezystor 10 kΩ. Jest ono podciągane do stanu wysokiego, czyli do zasilania 5 V, w momencie dotknięcia płytki dodatkowego modułu sterowania. Rezystor 47 kΩ zabezpiecza końcówkę wejściową 7 przed stanami nieustalonymi lub nieprawidłowym podłączeniem.

Zauważ, że musimy użyć oddzielnego wejścia dla dodatkowych płytek dotykowych. Gdybyśmy wykorzystywali wejście na końcówce 6 do podłączenia innej płytki sterującej, dodatkowa pojemność i oddziaływanie dodatkowego przewodu połączeniowego prowadziłyby na tym wejściu 6 o wysokiej impedancji do fałszywego wyzwania.

Pilot na podczerwień

Jeśli jest zamontowany, moduł odbiornika podczerwieni IRD1 odbiera i demoduluje kody z ręcznego pilota na podczerwień. IRD1 zawiera wzmacniacz i automatyczną regulację wzmocnienia oraz filtr pasmowy 38 kHz umożliwiający odbiór wyłącznie sygnałów z pilota. Następnie

wykrywa i usuwa nośną 38 kHz. Powstały w ten sposób sygnał jest podawany na wejście 4 układu IC1, gotowe do wykrywania kodu.

Ręczny pilot IR jest niewielkim urządzeniem o wymiarach zaledwie 80×40×7 mm. Jest on zasilany litowym ogniwo pastylkowym CR2025 3,2 V. Na górnym panelu ma dziewięć membranowych przycisków zwiniętych. Przyciski obejmują: przycisk włączania/wyłączania zasilania („Operate”), trzy przyciski oznaczone jako przyciski A, B i C oraz zestaw 5 przycisków: góra, dół, lewo, prawo i centralny przycisk akceptacji lub OK.

Układ 5 przycisków jest powszechnie używany do wyboru głośności i kanałów lub funkcji przód, tył, lewo i prawo.

Nie ma zbyt wielu informacji na temat elektroniki w tym pilocie, poza tym, że wykorzystuje on 16-stykowy układ scalony SMD do zdalnego sterowania, oznaczony HB8101P.

Za każdym razem, gdy przycisk jest wciśnięty, wysyła on poprzez pulsowanie podczerwonej diody LED (IR) unikalny kod. Sygnał podczerwieni jest wysyłany w postaci impulsów o częstotliwości 38 kHz, przy użyciu tak zwanego protokołu PDP (Pulse Distance Protocol). Protokół ten jest opisany w sąsiednim panelu.

IC1 odbiera ten sygnał i dekoduje go. Jeśli sygnał zostanie rozpoznany jako prawidłowy kod skojarzony z przyciskiem na pilocie IR, aktywowana zostanie wymagana funkcja ściemniacza.

Układ rozszerzający

Układ płytki rozszerzającej, wymaganej do dodania drugiej (lub trzeciej...) płytki dotykowej do sterowania tym samym zestawem świateł, jest pokazany na dole rysunku 2. Jest on dość prosty, ponieważ jest to tylko układ wysyłania sygnału do mikrokontrolera IC1 na płytce głównej, który następnie traktuje to zdarzenie identycznie jak dotknięcie własnej lokalnej płytki.

Gdy dodatkowa płytka dotykowa jest swobodna, tranzystor PNP Q3 jest wyłączony poprzez rezystor 1 MQ pomiędzy jego bazą i emitery. Gdy płytka dotykowa zostanie dotknięta, a napięcie fazy jest powyżej potencjału masy, potencjał bazy Q3 zostaje obniżony poprzez dwa rezystory zabezpieczające, diodę D3 i rezystor 2,2 MΩ.

To włącza Q3 i złącze EXTN jest podciągane do potencjału przewodu fazowego, który jest również zasilaniem +5 V dla IC1 na płytce głównej. Zwiększa to potencjał końcówki 7 (wejście cyfrowe GPO) układu IC1 do wysokiego poziomu, wysyłając mikrokontrolerowi sygnał, że płytka została dotknięta.

Kondensator 47 nF działa jako filtr i zapobiega przed włączeniem Q3 impulsami czy zakłóceniami elektrycznymi (np. piorunami lub EMI). Kondensator ten utrzymuje również tranzystor Q3 włączony wystarczająco długo, aby IC1 na głównym ściemniaczu mógł wykryć ten fakt, nawet przy bardzo szybkim dotknięciu płytki.

Dioda Zenera ZD3 chroni przed nadmiernym napięciem na katodzie diody D3 podczas dotykania płytki, ponieważ różnica potencjałów, zależnie od fazy napięcia sieciowego, może wynosić nawet setki woltów. Prąd płytki stykowej jest ograniczony do bardzo niskiego poziomu przez szeregowo rezystory zabezpieczające.

Dioda Zenera ZD4 oraz rezystor 220 Ω podłączony do kolektora Q3 to typowy



Nawet jeśli w salonie itp. znajduje się wiele lampek LED podłączonych do jednego wyłącznika, nasz ściemniacz będzie je obsługiwał – maksymalnie do 250 W obciążenia. A ponieważ większość domowych lamp LED ma moc rzędu 8-20 W każda, umożliwia to sterowanie naprawdę dużą ilością diod LED.

układ „idiotoodporny” – zapewnia ochronę w przypadku odwrotnego podłączenia do obwodu głównego. W takim przypadku ZD4 będzie spolaryzowana w kierunku przewodu zasilania, chroniąc Q3, natomiast rezystor 220 Ω ograniczy prąd zwarcia. Cienka ścieżka na płytce drukowanej stopi się, jeśli takie błędne połączenie będzie trwało dłużej. W przypadku takiej spowodowanej awarii musiałbyś wtedy naprawić PCB, oczywiście po odłączeniu od sieci i po usunięciu błędów okablowania!

Masz również możliwość użycia zwiernego chwilowego przycisku sieciowego (np. Clipsal 30MBPR i płytka przełącznika) zamiast dodatkowej płytki dotykowej, jako wtórnego przełącznika ściemniacza. Należy go tylko okablować, aby

połączyć po naciśnięciu zaciski przewodu fazowego i szynę rozszerzeń („Extension” – EXTN).

Może być podłączonych równolegle wiele płytek dodatkowych, pomiędzy zaciskami fazy (L) i EXTN, jeśli potrzebne są więcej niż dwie płytki sterowania ściemniaczem.

W przyszłym miesiącu

W części 2 w przyszłym miesiącu podamy wszystkie szczegóły budowy i okablowania, etapy testowania i regulacji oraz kilka wskazówek dotyczących użytkowania. ■

John Clarke

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

REKLAMA

Certyfikat Underwriters Laboratories
UL 94V-0 E480148 TYPE 1

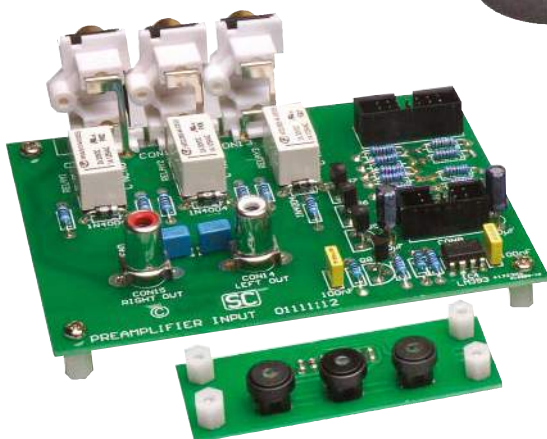
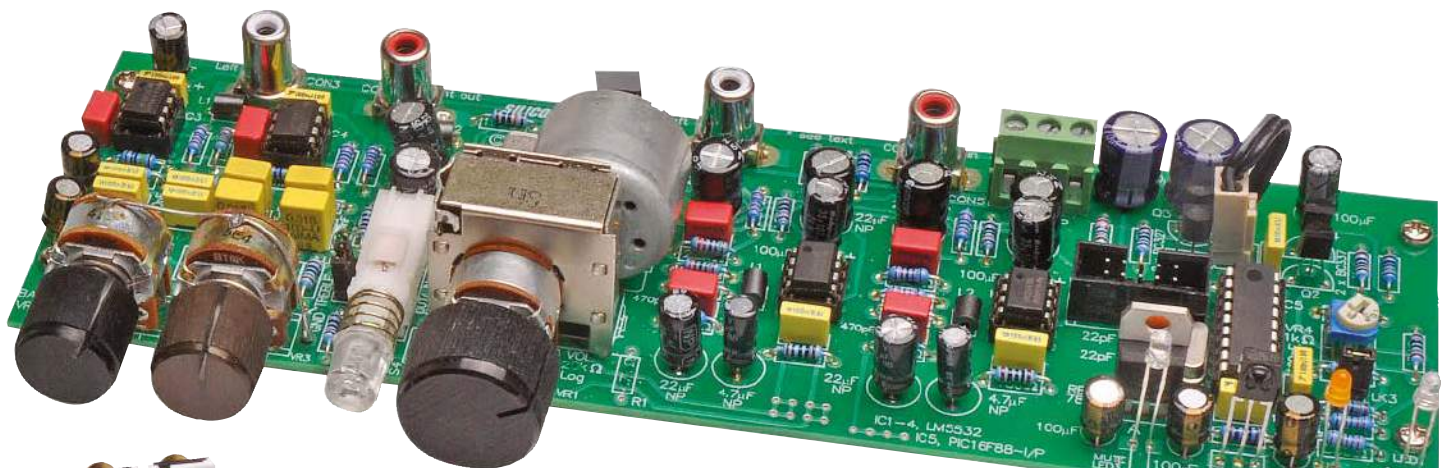
Zakład produkcyjny:
05-660 Warka
ul. M. Rópelewskiej 17
tel. 22 781 63 95
22 761 95 80
fax. 22 781 63 95 w 23
www.elmax.com.pl
elmax@elmax.com.pl

OBWODY DRUKOWANE

Produkcja, Projektowanie, Montaż

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

Przedwzmacniacz o ultraniskich zniekształceniach, część 2



Z opcjami:

- ☑ Regulacja tonów niskich i wysokich
- ☑ Zdalna regulacja głośności
- ☑ Zdalne sterowanie lub regulacja ręczna
- ☑ Przetwarzanie i separacja wejść przekaźnikami
- ☑ Współpracuje praktycznie z KAŻDYM wzmacniaczem!

W zeszłym miesiącu (EdW 11/2022) przedstawiliśmy nasz najnowocześniejszy przedwzmacniacz stereofoniczny. Oprócz niemal niemierzalnych szumów i zniekształceń (typowo 0,0003% THD+N!), ma on zdalną regulację głośności, wybór wejścia i wyciszenie oraz – na panelu przednim – pokrętki regulacji tonów niskich i wysokich. Teraz zbudujemy płytki wyboru wejść i zasilacz.

Schematy ideowe opcjonalnych płytek: wyboru wejścia oraz płytki przycisków na panelu przednim zostały pokazane w zeszłym miesiącu na rysunkach 8 i 9.

W tamtym artykule wymieniliśmy również części potrzebne do zmontowania tych dwóch płytek.

Rysunki 10 i 11 pokazują schematy montażowe płytek drukowanych dla tych dwóch modułów, więc możesz zobaczyć, jak są rozplanowane te części.

Przy okazji, nie musisz budować żadnej z tych płytek, jeśli używasz tylko jednego zestawu wejść stereo.

W takim przypadku podłączysz gniazda wejściowe montowane w obudowie

bezpośrednio do CON1 i CON2 na głównej płycie przedwzmacniacza.

Jeśli potrzebujesz tylko funkcji zdalnego sterowania i zdalnego wyboru wejść, ale nie chcesz przycisków/wskaźników wejść na panelu przednim, możesz zbudować płytkę selektora wejść (rysunek 10), ale nie potrzebujesz montować płytki przycisków na panelu przednim (rysunek 11).

Możesz wtedy używać pilota do wyboru pomiędzy trzema wejściami, chociaż nie będzie sygnalizacji wybranego wejścia – będziesz musiał zapamiętać swój ostatni wybór.

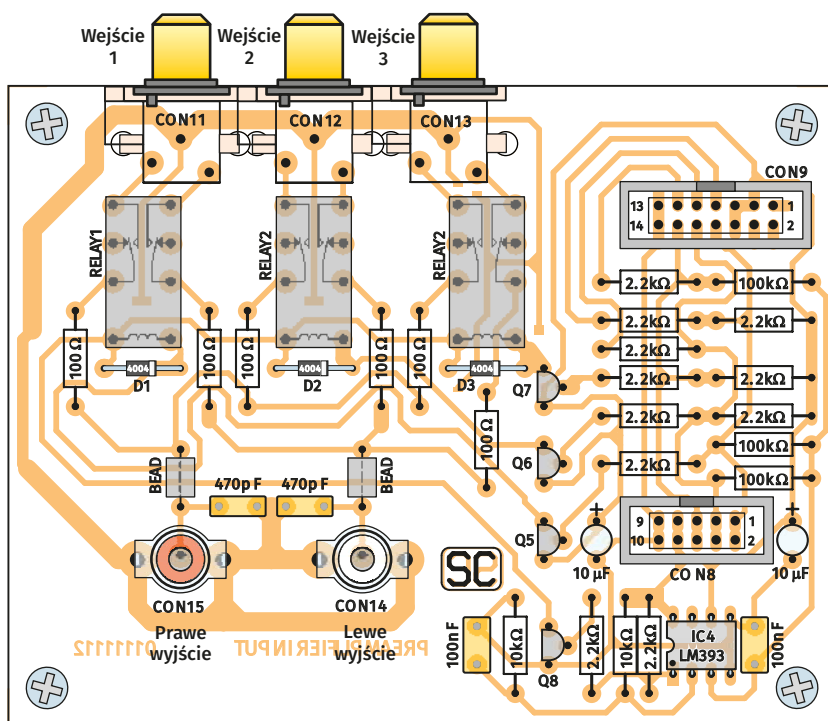
Nawiasem mówiąc, nie wymienialiśmy teraz ponownie wszystkich funkcji i parametrów – zapoznaj się z nimi oraz z wykresami

charakterystyk w wydaniu listopadowym EdW. Zgodzisz się chyba, że jest to znakomity projekt!

Budowa selektora wejść

Płytkę wyboru wejścia jest łatwa w montażu. Jest ona zbudowana na dwustronnej płytce drukowanej o wymiarach 110×85 mm oznaczonej kodem 0111112.

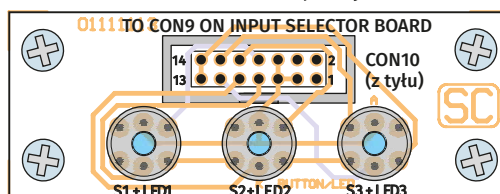
Zacznij od zamontowania rezystorów w miejscach pokazanych na rysunku. Opublikowaliśmy kody rezystorów w zeszłym miesiącu, ale zawsze najlepiej jest sprawdzić ich wartości za pomocą DMM ustawionego na pomiar rezystancji, aby upewnić się, że trafią we właściwe miejsca.



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

Rysunek 10. (po lewej) Postępuj zgodnie z tym schematem montażowym, aby zbudować płytkę wyboru wejść. Upewnij się, że dwa wtyki IDC, CON8 i CON9, są prawidłowo ustawione i zauważ, że Q5-Q7 to tranzystory PNP BC327, a Q8 to tranzystor NPN BC337.

Rysunek 11. (poniżej) Z przodu płytki przycisków są zamontowane trzy przełączniki, natomiast wtyk IDC Z-WS14 umieszczony jest pod spodem (kierunek klucza-wycięcia w stronę S2). Zwróć uwagę na orientację przełączników (patrz tekst): z sześciu końcówek dla każdego przełącznika cztery są dla samych styków przełącznika plus dwa zasilania wbudowanej diody LED.



Zajmij się diodami D1...D3, upewniając się, że ich paski katodowe są usytuowane jak na rysunku, a następnie przełóż kawałki odciętych wcześniej końcówek rezystorów przez koraliki ferrytowe i przylutuj je we wskazanych miejscach.

Zalecamy wlutowanie układu IC4 bezpośrednio do płytki, chociaż możesz użyć podstawki, jeśli Ci to odpowiada.

Tak czy inaczej, upewnij się, że znacznik przy końcówce 1 jest skierowany w lewo, jak pokazano na rysunku.

Następnie zamontuj kondensatory (do wyboru: MKT/MKP/ ceramiczne).

Wyjaśniliśmy szczegółowo w zeszłym miesiącu, z wykresami

charakterystyki, dlaczego istnieją trzy różne opcje wyboru kondensatorów 470 pF i że jeśli chcesz zastosować kondensatory ceramiczne, to aby uzyskać zamierzone parametry, muszą to być kondensatory typu NPO/CoG.

W naszym prototypie użyliśmy kondensatorów polipropylenowych MKP (np. popularne czerwone WIMA, jak na głównej płytce). Zamontuj teraz, plus dwa kondensatory poliestrowe 100 nF MKT.

Następnie przylutuj cztery tranzystory, pamiętając, że Q5...Q7 to BC327, a Q8 to BC337. Kolejno można wlutować dwa kondensatory elektrolityczne, z dłuższymi dodatkami wyprowadzeniami przełożonymi przez otwory oznaczone „+”

na PCB, a potem 10- i 14-szpilkowe wtyki IDC Z-WS10 i Z-WS14, CON8 i CON9.

Red. EDW: w zeszłym numerze wyjaśniliśmy, dlaczego stosujemy taką nomenklaturę elementów złączy.

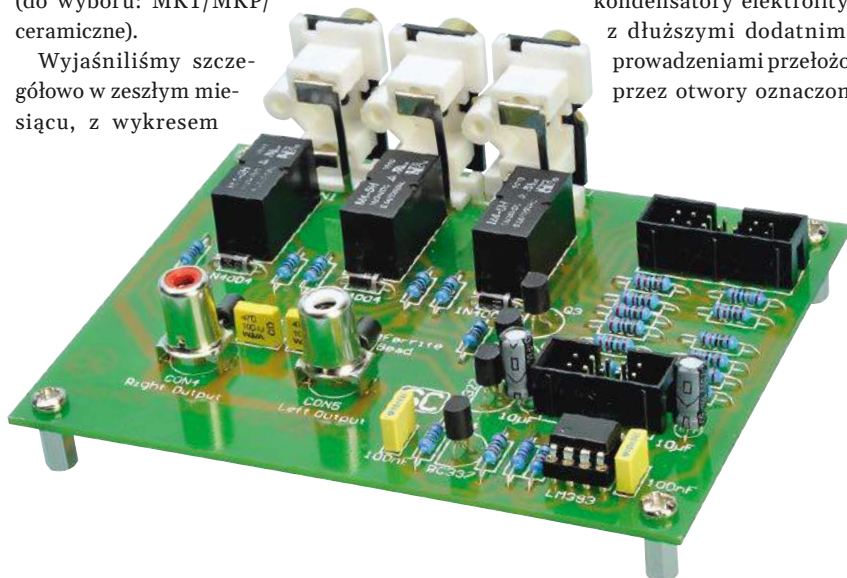
Wtyki te muszą być zamontowane kluczami-wycięciami skierowanymi do góry PCB.

Na końcu wlutuj przełączniki, trzy stereo-foniczne, podwójne gniazda wejściowe RCA i dwa gniazda wyjściowe RCA do montażu pionowego na PCB.

Zwróć uwagę na oznaczenie lewego i prawego gniazda wyjściowego RCA – nie jest to błąd, a ułożenie ich w ten nietypowy sposób daje optymalne prowadzenie ścieżek na płycie drukowanej.

Montaż płytki panelu przedniego z przyciskami

Na płytce przycisków znajdują się tylko cztery części – trzy przyciski przełączników z jednej strony i 14-szpilkowy wtyk



Te zdjęcia pokazują ukończoną płytkę wyboru wejść oraz (po prawej) obie strony płytki przycisków. Zwróć uwagę na orientację wtyków IDC na obu modułach – sprawdź, przed przylutowaniem ich wyprowadzeń, czy te złącza, przełączniki, gniazda RCA i przyciski przełączników umieszczone są zgodnie ze schematami montażowymi płytek PCB.



IDC z drugiej (patrz rysunek 11). Płytkę jest oznaczona kodem 01111113 i ma wymiary 66×25 mm.

Trzy przyciski można założyć jako pierwsze, ale należy pamiętać, że muszą być one zamontowane we właściwy sposób.

Mają one „zagięte” styki w każdym rogu oraz dwie proste końcówki dla wbudowanej niebieskiej diody LED. Końcówka anodowa jest z tych dwóch dłuższa i musi ona trafić w otwór oznaczony „A” na PCB (w kierunku nagłówka).

Gdy końcówki są już włożone, przyciski należy wcisnąć do końca, tak, aby przylegały do płytki przed przylutowaniem ich wyprowadzeń.

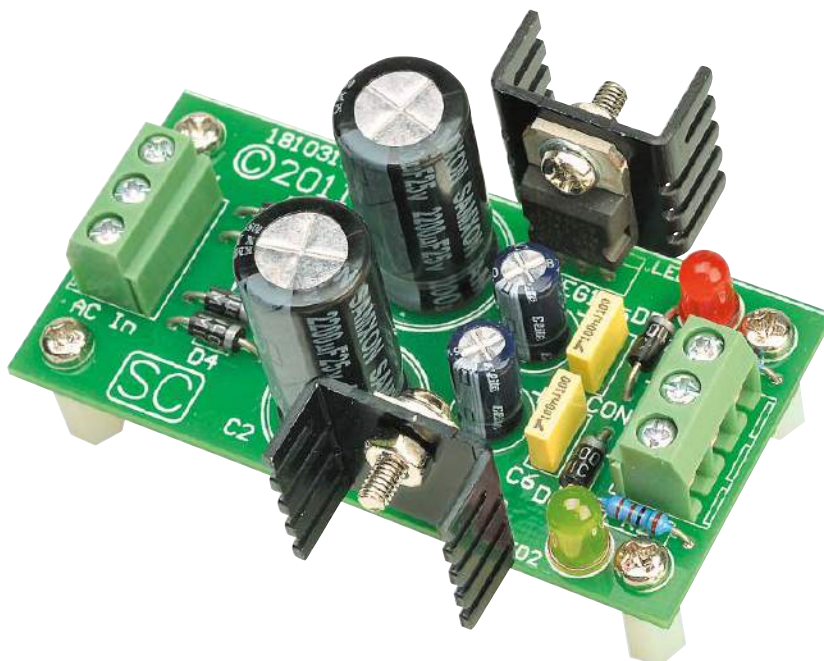
Następnie można po drugiej stronie PCB zamontować wtyk IDC, wycięciem-kłuczem skierowanym w stronę przycisków.

Wybór zasilacza

Jeśli budujesz ten przedwzmacniacz jako część kompletnego wzmacniacza, istnieje szansa, że masz już odpowiedni zasilacz, który posiada wymagane zasilanie ± 15 V. W przeciwnym razie, w zeszłym miesiącu wspomnieliśmy o kilku różnych odpowiednich płytkach zasilających.

Można tu wymienić: Universal Regulator z marca 2011 (<https://siliconchip.com.au/Article/930>) [dostępny jako zestaw Jaycar, nr katal. KC5501], zestaw AVT – Modułowy zasilacz symetryczny ± 15 V, kit AVT3140/15 oraz płytkę zasilającą Ultra-LD Mk.2/3/4, opisaną w numerze z lipca 2022 EdW. Na wypadek gdybyś nie miał tych numerów, szybko omówimy tutaj budowę obu tych zasilaczy.

Zasilacz uniwersalny jest dobrym wyborem, jeśli budujesz samodzielny przedwzmacniacz,



Płytkę zasilacza uniwersalnego może obsługiwać szeroki zakres napięć wejść i wyjść. Z transformatorem 15-0-15 VAC uzyskasz stabilizowane zasilanie ± 15 VDC, idealne dla Twojego przedwzmacniacza o ultraniskich zniekształceniach (i wielu innych projektów!).

lub wbudowujesz przedwzmacniacz do wzmacniacza, który ma już zasilacz, ale nie ma szyn ± 15 V.

Zasilacz Ultra-LD jest najlepszy, jeśli wbudowujesz przedwzmacniacz do kompletnego wzmacniacza, który wykonujesz od podstaw.

Budowa zasilacza uniwersalnego

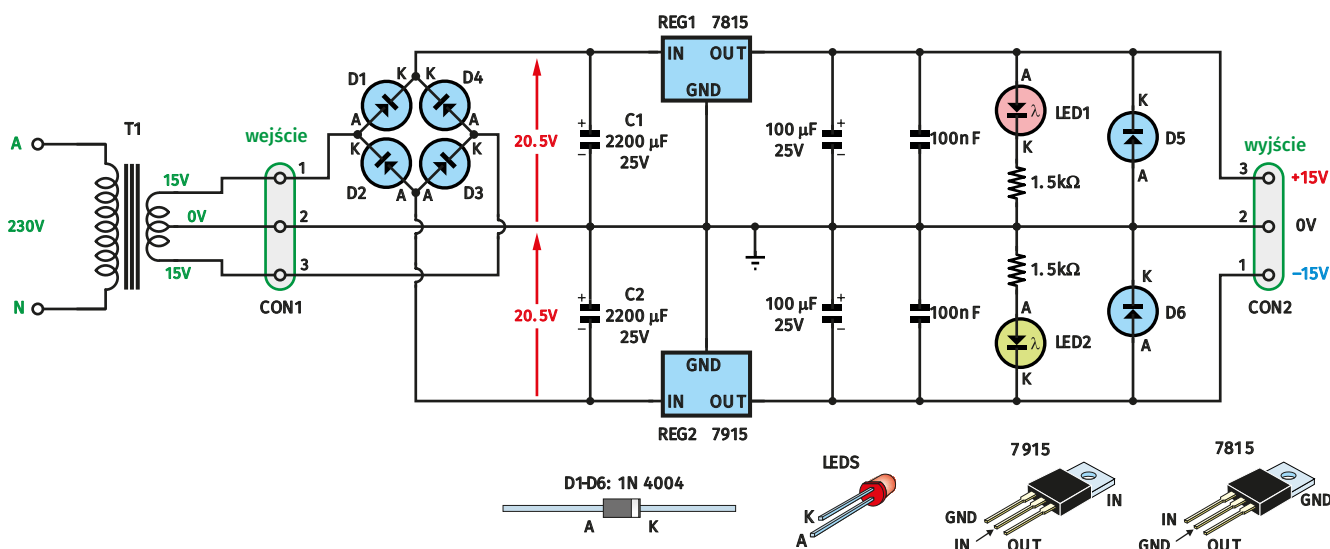
Rysunek 12 przedstawia schemat ideowy zasilacza uniwersalnego, natomiast

rysunek 13 to schemat montażowy płytki drukowanej.

Możesz go zasilać napięciem AC 30 V (15-0-15 V) z wtórnego uzwojenia transformatora lub z pojedynczego uzwojenia 15 V.

Opcja ze środkowym odczepem uzwojenia wtórnego, jeśli masz taką możliwość, jest lepsza, ponieważ powoduje ona mniejsze tętnienia na wejściach stabilizatorów.

Wyjście AC z transformatora jest prostowane przez mostek z diod D1-D4 i filtrowane przez parę kondensatorów po 2200 mF.



UNIERSALNY ZASILACZ STABILIZOWANY W KONFIGURACJI SYMETRYCZNEJ, ZASILANY Z TRANSFORMATORA O UZWOJENIU WTÓRNYM Z ODCZEPEM

Rysunek 12. Schemat ideowy uniwersalnego zasilacza symetrycznego ± 15 V. Diody D1...D4 tworzą mostek prostowniczy, podczas gdy kondensatory C1 i C2 filtrują wyprostowany prąd zmienny. Stabilizatory REG1 i REG2 zapewniają stałe napięcie wyjściowe, podczas gdy diody LED1 i LED2 sygnalizują obecność napięć na wyjściu. Można również zastosować transformator z pojedynczym uzwojeniem wtórnym (lub zasilacz wtyczkowy AC) podłączony pomiędzy zaciski 1 i 2 lub 2 i 3 CON1.

Następnie jest stabilizowane na poziomie +15 V przez REG1 i -15 V przez REG2. Te stabilizowane napięcia dostępne są na bloku zacisków CON2, skąd podłączone są do przedwzmacniacza.

Zasilacz zbudowany jest na płytce o kodzie 18103111 i wymiarach 71×35,5 mm. Możesz ją dostać w sklepie internetowym Silicon Chip (SC0782) lub, jeśli kupisz zestaw od Jaycar, płytka będzie dołączona (płytka PCB ze wszystkimi elementami, także w wersji zmontowanej, kupisz również w sklepie AVT).

Poniżej znajduje się lista części potrzebnych do budowy.

Zacznij montaż od zamontowania dwóch rezystorów, a następnie sześciu diod (z polaryzacją pokazaną na rysunku 13).

Następnie zamontuj diody LED dłuższymi przewodami (anodowymi) w kierunku dołu płytki. Potem zamontuj dwa kondensatory MKT.

Następnie można zamontować dwa 3-drożne bloki zacisków śrubowych, przy czym otwory na przewody powinny być skierowane na zewnątrz krawędzi płytki.

Teraz przylutuj REG1 i REG2, metalowymi radiatorami w kierunku krawędzi płytki, jak pokazano na rysunku, uważając, aby nie zamienić wzajemnie tych dwóch elementów. Na koniec przylutuj cztery kondensatory elektrolityczne, upewniając się, że ich dłuższe (dodatnie) wyprowadzenia trafiają na pola lutownicze oznaczone symbolem „+”.

Na zdjęciu powyżej widać dwa radiatory w kształcie litery „U” (i są one wymienione na liście części). Nie zaszkodzi je zamontować, ale jeśli zamierzasz zasilac tylko przedwzmacniacz, a napięcie wtórne Twojego transformatora ma zalecaną wartość, to nie powinny być one potrzebne, gdyż przedwzmacniacz nie pobiera dużego prądu.

Budowa pełnego zasilacza

Schemat ideowy zasilacza Ultra-LD pokazany jest na rysunku 14. Dolna część (na schemacie montażowym prawa) jest podobna do opisanego wyżej zasilacza uniwersalnego i działa w ten sam sposób. Moduł prostownika mostkowego montowany w obudowie jest wykorzystywany do prostowania wysokich napięć AC z wtórnych uzwojeń transformatora zasilającego. Są tu one pokazane jako 40-0-40 V, ale z tą płytką mogą być stosowane również niższe napięcia.

Powstałe w ten sposób na szynach zasilających napięcia DC są następnie filtrowane przez sześć kondensatorów po 4700 mF każdy, po trzy na szynę, i udostępniane na CON1 i CON2, w celu doprowadzenia ich do modułów wzmacniacza mocy.

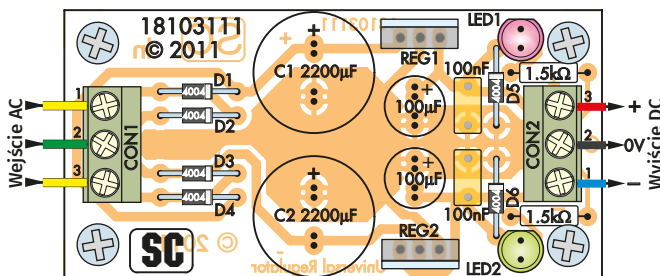
Rysunek 15 to schemat montażowy płytki drukowanej dla tego zasilania. Możesz nabyć tę płytkę w sklepie internetowym Silicon Chip (nr SC0716). Dwa druciane mostki nie powinny być potrzebne, ponieważ nasze płytki są dwustronne i mają miedziane ścieżki na górnej stronie łączące te punkty, ale jeśli wytrawiasz płytkę jednostronną, będziesz musiał zamontować te dwa połączenia używając pocynowanego drutu miedzianego lub DNE o średnicy 1 mm.

Następnie zamontuj diody w pokazanej pozycji, a potem diody LED, dłuższymi (anodowymi) przewodami w kierunku góry płytki. Potem możesz wygiąć przewody regulatorów/stabilizatorów, aby dopasować je do układu otworów na płytce i przymocować do płytki ich metalowe radiatory za pomocą śrubek M3 i nakrętek.

Po sprawdzeniu, że końcówki są proste, przylutuj je i przytnij.

W następnej kolejności zamontuj bloki zacisków. Przed przylutowaniem należy połączyć CON4 z CON5 oraz CON3 z CON6 na wypustki „jaskółczy ogon” i wlutować do PCB, tak, aby otwory na przewody były skierowane na zewnątrz krawędzi płytki. Następnie możesz zamontować dwa rezystory 5 W, umieszczając ich korpusy kilka milimetrów nad powierzchnią płytki, aby umożliwić obieg powietrza chłodzącego.

Po tym zamontuj męskie złącza konektorowe. Jeśli używasz typu pionowego z dwoma stykami lutowniczymi, wcisnij je w otwory w płytce



Rysunek 13. Ten schemat montażowy płytki drukowanej odpowiada schematowi ideowemu z rysunku 12. Możesz zamontować U-kształtne radiatory na REG1 i REG2, ale nie są one absolutnie niezbędne do pracy z przedwzmacniaczem, ponieważ nie pobiera on dużego prądu.

Wykaz elementów, kupuj w sklepie sklep.avt.pl (W-wa, ul. Leszczynowa 11,

tel. +48222578451, e-mail: handlowy@avt.pl);

Lista części do zasilacza uniwersalnego (wyjścia ±15 V)

- 1 szt. płytka drukowana, kod 18103111, 71×35,5 mm
- 1 szt. transformator sieciowy, uzwojenie pierwotne 230 VAC, uzwojenie wtórne z odczepem 15-0-15 VAC lub zasilacz wtyczkowy 15 VAC (patrz tekst)
- 2 szt. trójdrożne bloki zacisków śrubowych, raster 5,08 mm (CON1, CON2)
- 4 szt. gwintowane tulejki dystansowe
- 4 szt. śrubki M3×6
- 2 szt. radiatory profil FK239 do obudowy TO-220 (opcjonalnie)
- 2 szt. śrubki M3×10, nakrętki i podkładki ząbkowane dla radiatorów (opcjonalnie)

Półprzewodniki:

- 1 szt. regulator/stabilizator 7815 TO-220 +15 V (REG1)
- 1 szt. regulator/stabilizator 7915 TO-220 -15 V (REG2)
- 6 szt. diod 1N4004
- 1 szt. czerwona dioda LED 5 mm (LED1)
- 1 szt. zielona dioda LED 5 mm (LED2)

Kondensatory:

- 2 szt. kondensatory elektrolityczne 2200 µF/25 V
- 2 szt. kondensatory elektrolityczne 100 µF/25 V
- 2 szt. kondensatory poliestrowe MKT 100 nF

Rezystory: (wszystkie 0,25 W 1% metalizowane)

- 2 szt. 1,5 kΩ

Lista części zasilacza do wzmacniacza Ultra-LD i przedwzmacniacza (wyjścia ±57 V i ±15 V)

- 1 szt. płytka drukowana, kod 01109111, 141×80 mm
- 1 szt. transformator sieciowy, z uzwojeniami wtórnymi z odczepami 40-0-40 VAC i 15-0-15 VAC (patrz tekst)
- 4 szt. trójdrożne bloki zacisków śrubowych do montażu na płytce drukowanej, raster 5,08 mm (Altronics P2035A lub równoważne) (CON1-4)
- 2 szt. dwudrożne bloki zacisków śrubowych do montażu na płytce drukowanej, raster 5,08 mm (Altronics P2034A) (CON5-6)
- 3 szt. pojedyncze złącza konektorowe męskie 6,3 mm do montażu na płytce drukowanej (lub przykręcane) [Altronics H2094]
- 3 szt. śrubki M4×10, nakrętki, podkładki płaskie i ząbkowane (w przypadku korzystania z konektorów przykręcanych)
- 4 szt. gwintowane poliamidowe tulejki dystansowe M3×9
- 6 szt. wkrętów M3×6
- 2 szt. podkładki sprężyste lub ząbkowane i nakrętki M3
- 1 odcinek długości 150 mm drutu miedzianego ocynowanego o średnicy 0,7 mm

Półprzewodniki:

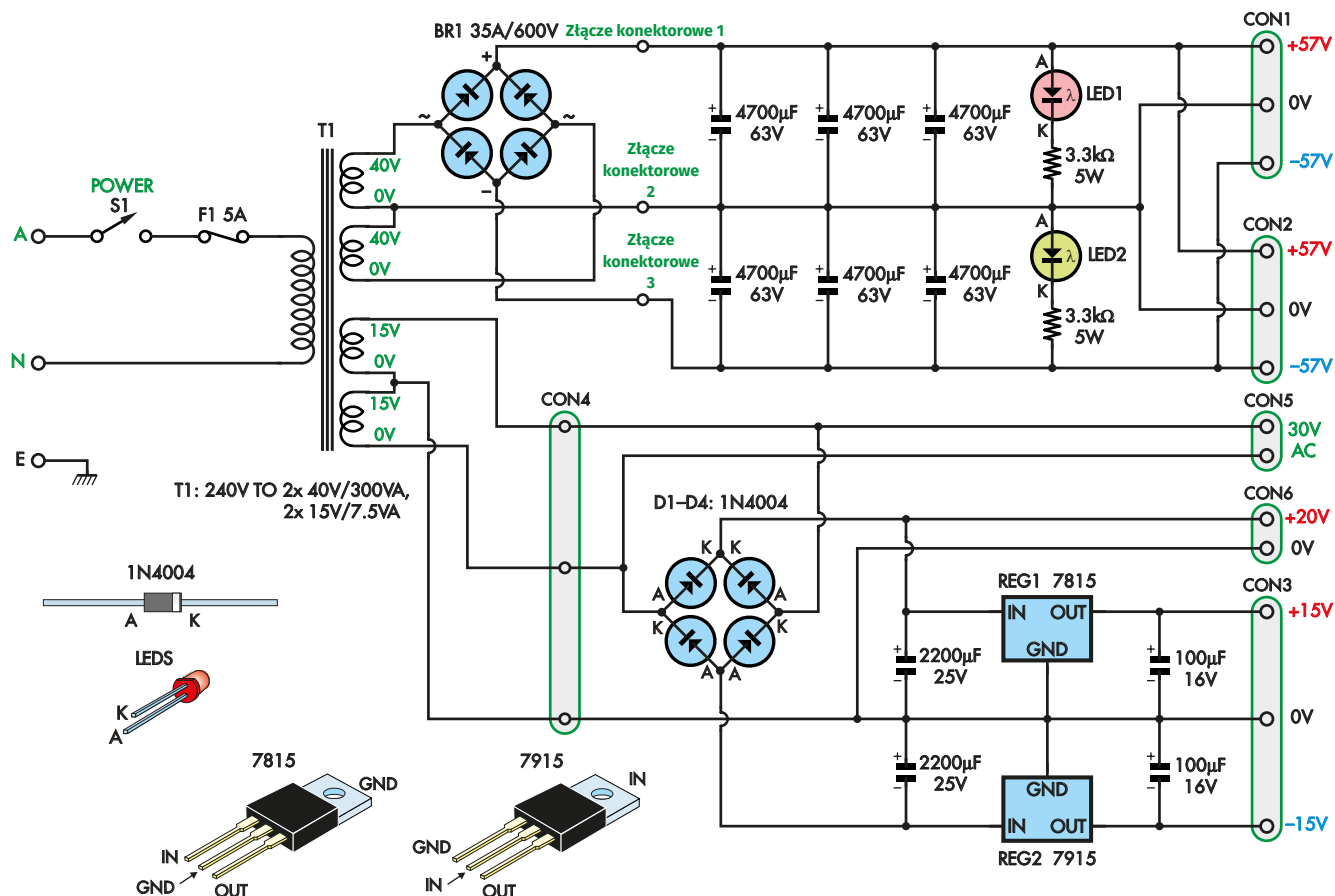
- 1 szt. mostek prostowniczy do montażu w obudowie 35 A/400 V (BR1)
- 1 szt. regulator/stabilizator 7815 TO-220 +15 V (REG1)
- 1 szt. regulator/stabilizator 7915 TO-220 -15 V (REG2)
- 1 szt. czerwona dioda LED 5 mm (LED1)
- 1 szt. zielona dioda LED 5 mm (LED2)

Kondensatory:

- 6 szt. kondensatory elektrolityczne 4700 µF/63 V
- 2 szt. kondensatory elektrolityczne 2200 µF/25 V
- 2 szt. kondensatory elektrolityczne 220 µF/16 V

Rezystory:

- 2 szt. 3,3 kΩ 5 W



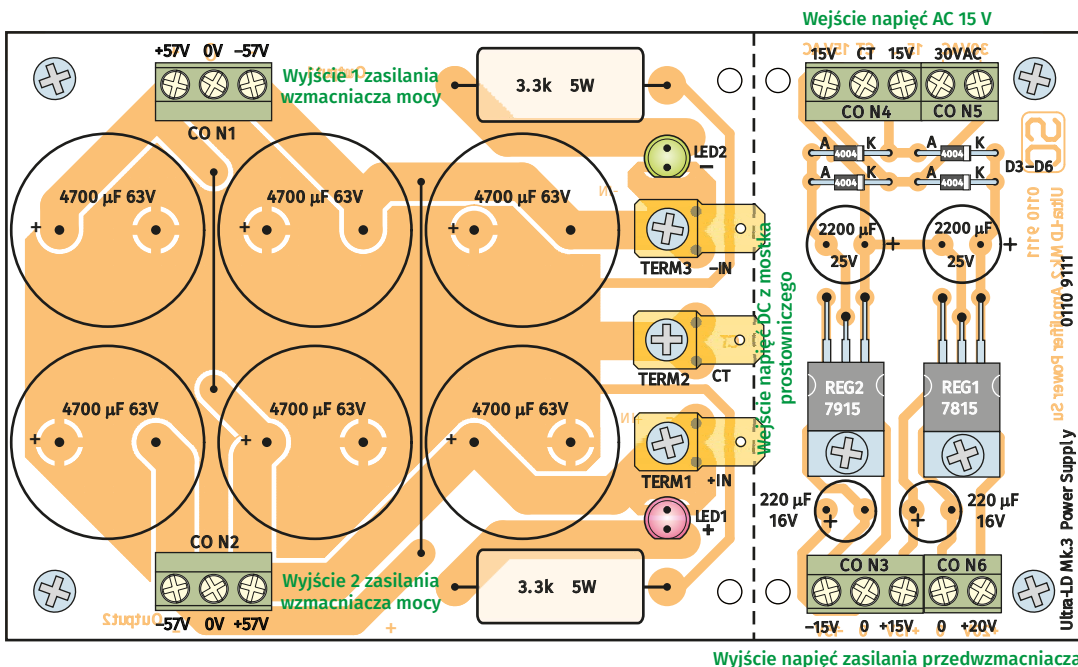
ZASILACZ WZMACNIACZA MOCY I PRZEDWZMACNIACZA

Rysunek 14. Zasilacz tego wzmacniacza wykorzystuje transformator toroidalny (T1) z dwoma uzwojeniami wtórnymi: 2×40 V i 2×15 V, ale w razie potrzeby można użyć dwóch oddzielnych transformatorów. Można użyć transformatora z uzwojeniem wtórnym o niższym napięciu, niż 2×40 V, aby uzyskać niższe napięcie zasilania wzmacniacza mocy, bez dokonywania jakichkolwiek zmian na płycie.

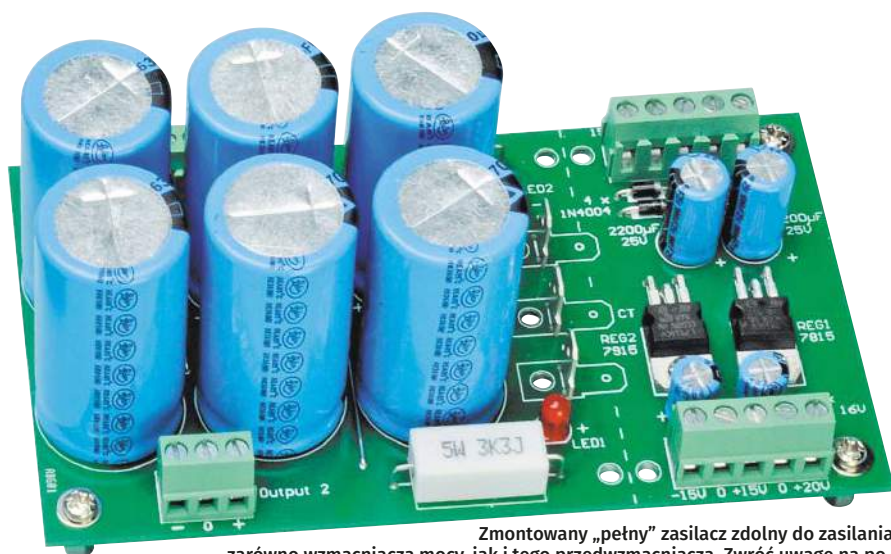
i przylutuj – ze względu na ich dużą masę będziesz potrzebował do tego gorącej lutownicy. Jeśli używasz złączy pojedynczych typu

przykręcanego, przymocuj je do płytki za pomocą śrubek M4, podkładek płaskich i ząbkowanych oraz nakrętek, jak pokazano na rysunku 16.

Teraz pozostało już tylko przylutować małe, a następnie duże, kondensatory elektrolityczne. W obu przypadkach dłuższe (dodatnie)



Rysunek 15. Zamontuj elementy na płycie zasilacza zgodnie z rysunkiem, zwracając uwagę na to, aby wszystkie kondensatory elektrolityczne były zamontowane z właściwą polaryzacją. Należy również pamiętać o zastosowaniu właściwego stabilizatora w każdym miejscu. Dwie diody LED sygnalizują podłączenie zasilania i świecą po wyłączeniu aż do rozładowania kondensatorów 4700 mF.



Zmontowany „pełny” zasilacz zdolny do zasilania zarówno wzmacniacza mocy, jak i tego przedwzmacniacza. Zwróć uwagę na polaryzację kondensatorów elektrolitycznych, diod, LED-ów i stabilizatorów.

wyprowadzenia muszą trafić do pól lutowniczych oznaczonych „+” na płycie.

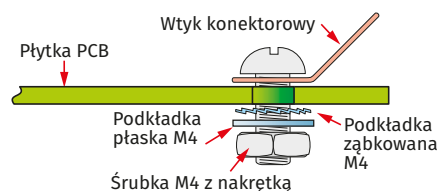
Wstępne testowanie

Przed zainstalowaniem na płycie przedwzmacniacza trzech układów scalonych, warto sprawdzić napięcia zasilające. Będziesz musiał podłączyć do swojego zasilacza transformator, następnie podłączyć wyjścia +15 V, 0 V i -15 V zasilacza do odpowiednich wejść na płycie głównej przedwzmacniacza.

Bezpieczniej jest użyć do testów zasilacza wtyczkowego 15 VAC (Uwaga!), jeśli nie

masz jeszcze transformatora i zasilacza zainstalowanych w uziemionej przewodem ochronnym metalowej obudowie. Wystarczy podłączyć jeden przewód z wyjścia tymczasowego zasilacza do któregoś z zacisków wejściowych niskiego napięcia AC na płycie zasilacza, a drugi do punktu podłączenia środkowego odczepu transformatora.

Podłącz wtyczkę do sieci. Teraz należy sprawdzić napięcia na końcówkach „8” i „4” czterech 8-stykowych podstawek pod układy scalone IC1...IC4 na płycie przedwzmacniacza, tzn. pomiędzy każdym



Rysunek 16. Oto sposób mocowania męskich pojedynczych końcówek konektorowych do płytki drukowanej zasilacza. Można również zastosować konektory pionowe z lutowanymi stykami.

z tych styków a zaciskiem 0 V (środkowym) złącza CON6. Powinieneś uzyskać odczyty odpowiednio +15 V oraz -15 V.

Podobnie sprawdź napięcie na styku „14” podstawki dla IC5. Powinno ono zawierać się w przedziale od +4,8 V do +5,2 V.

Jeśli te napięcia są prawidłowe, wyłącz zasilanie i zamontuj układy scalone. Zwróć uwagę, że układy IC1...IC4 są skierowane w jedną stronę, natomiast mikrokontroler IC5 w drugą.

Testowanie pilota/przetłacznika

Można teraz przetestować funkcje pilota używając odpowiedniego pilota uniwersalnego, np. Altronics A1012. Jak wspomniano wcześniej, domyślnym trybem urządzenia zaprogramowanym w mikrokontrolerze jest TV, ale jeśli koliduje to z innym sprzętem, można wybrać jako urządzenie SAT1 lub SAT2.

Niezależnie od wybranego trybu, należy również zaprogramować odpowiedni kod w pi-locie (patrz panel).

Wykonanie przewodów połączeniowych

Aby połączyć trzy płytki, musisz wykonać dwa kable IDC. Poniższe schematy przedstawiają, w jaki sposób to zrobić.

Red. EdW: być może znajdziesz u siebie albo u znajomego stary kabel komputerowy od podłączenia stacji dyskie-tek FDD 3,5". Obejrzyj go dokładnie, aby zorientować się w sposobie zaciskania gniazd IDC na taśmie. Jeśli kabel nie jest już do niczego potrzebny, masz gratis potrzebne Ci odcinki taśmy 10- i 14-żyłowej.

Strona zacisku 1 na gniazdach IDC jest zwykle oznaczony małym trójkątem na plastikowej listwie i do niego musi zawsze trafić przewód taśmy z czerwonym paskiem.

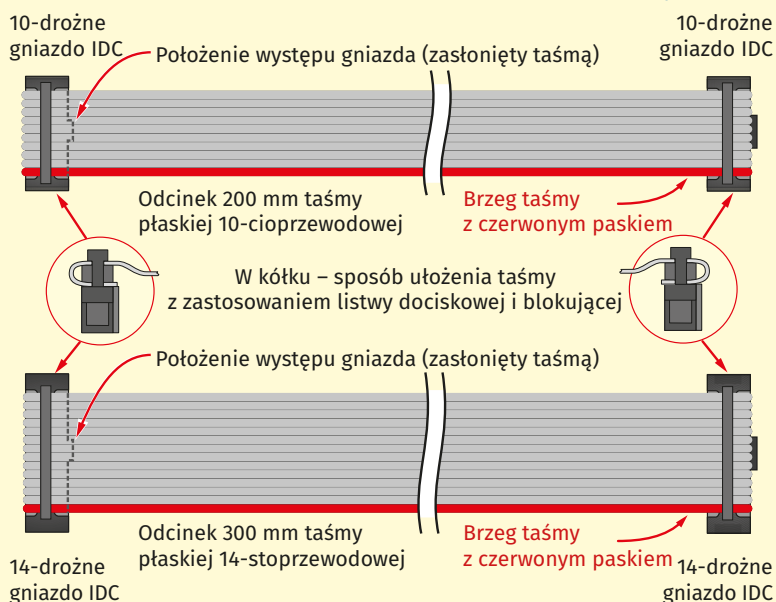
Gniazda IDC można zacisnąć na kablu (tzn. gniazdo z taśmą i listwą dociskową) w imadle lub użyć narzędzia do zaciskania IDC (np. Altronics T1540 lub Jaycar TH1941). Zwróć uwagę, że złącze składa się z 3-ch elementów: właściwego gniazda z zaciskami nożowymi do przewodów, listwy dociskowej i listwy blokującej. Dwa ostatnie elementy mają po bokach delikatne zatrzaski, uważaj, aby ich nie wyjąć. Jeśli nie masz wprawy w zaciskaniu gniazd IDC, np. taśma układa Ci się krzywo, możesz przed tym przykleić taśmę do listwy zaciskowej minimalną ilością kleju momentalnego. Uważaj, aby przy przycinaniu taśmy na długość nie zostawić poszarpanych drucików, które mogłyby zwrzeć sąsiednie przewody.

Po zacisnięciu nie zapomnijcie o założeniu listew zabezpieczających (blokujących).

Po wykonaniu kabli warto sprawdzić, czy zostały

one prawidłowo zaciśnięte. Najlepszym sposobem jest włożenie ich do odpowiednich wtyków IDC na płytkach PCB, a następnie sprawdzenie połączenia pomiędzy odpowiednimi stykami na obu końcach za pomocą multimetru.

Red. EdW: zalecamy wykonanie tej czynności z użyciem wtyków IDC (złącze ze szpilkami) przed lutowaniem ich do PCB, ponieważ płytka PCB zwiera niektóre przewody ze sobą).



Wybór trybu i programowanie pilota



Jak wspomniano w tekście, konieczne jest prawidłowe zaprogramowanie pilota uniwersalnego. Domyślnie kod RC5 oprogramowania mikrokontrolera rozpoznaje polecenia TV, ale można też wybrać SAT1 lub SAT2. Wystarczy podczas włączania zasilania nacisnąć i przytrzymać przycisk S1 na płytce przycisków, aby wybrać SAT1 lub przycisk S2, aby wybrać SAT2. Naciśnięcie przycisku S3 podczas włączania zasilania powoduje powrót do trybu TV.

Po wybraniu trybu lub „urządzenia” należy zaprogramować w pilocie odpowiedni kod. Polega to na wybraniu przyciskami na pilocie jako urządzenia: TV, SAT1 lub SAT2 (zgodnie się z ustawieniami mikrokontrolera), a następnie zaprogramowaniu trzy- lub czterocyfrowym numerem urządzenia firmy Philips. Wynika to z faktu, że większość urządzeń Philipsa (ale nie wszystkie) używa standardu kodu RC5, którego oczekuje przedwzmacniacz.

Można zastosować większość pilotów uniwersalnych, w tym model pokazany powyżej, Altronics A1012 (29,95\$) oraz Jaycar AR1955 (29,95\$) lub AR1954 (39,95\$). Dla Altronics A1012 należy użyć kodu 023 lub 089 dla trybu TV, 242 dla SAT1 lub 245 dla SAT2.

Podobnie w przypadku pilotów Jaycar, należy użyć kodu 1506 dla TV, 0200 dla SAT1 lub 1100 dla SAT2.

W przypadku innych pilotów uniwersalnych, jest to tylko kwestia testowania różnych kodów, aż znajdziesz taki, który działa poprawnie. Zwykle dla każdego urządzenia Philipsa podanych jest nie więcej niż 15 kodów (a zwykle mniej), więc znalezienie właściwego nie powinno zająć dużo czasu.

Należy pamiętać, że niektóre kody mogą działać tylko częściowo, np. mogą sterować głośnością, ale nie wyborem wejścia. W takim przypadku należy spróbować innego kodu. Ponadto, niektóre piloty mogą działać tylko w jednym trybie (np. TV, ale nie SAT).

Uwaga: jeśli nie masz jeszcze gotowego zasilacza symetrycznego, to możesz sprawdzić działanie pilota używając pojedynczego zasilania DC 9...15 V podłączonego pomiędzy zaciski +15 V i 0 V na zaciskach CON5 przedwzmacniacza (zwróć uwagę na polaryzację).

Tak jak poprzednio sprawdź napięcie na styku „14” podstawki IC5 (musi być pomiędzy +4,8 V a +5,2 V), następnie wyłącz zasilanie i zamontuj IC5 (styk 1 w kierunku IRD1). Załóż także zwórkę na LK3, aby wyłączyć funkcję powrotu wyciszenia.

Teraz połącz trzy płytki za pomocą przewodów taśmowych. Wszystkie złącza są oznaczone kluczem – wycięciem na występ, więc wszystko powinno być prawidłowo podłączone.

Następnie obróć VR2 całkowicie w lewo i użyj pilota do sprawdzenia różnych funkcji. Najpierw należy sprawdzić czy wejścia mogą być wybierane za pomocą przycisków 1, 2 i 3 na pilocie oraz przycisków S1...S3 na płytce przycisków. Po każdym naciśnięciu przycisku powinieneś usłyszeć „kliknięcie”, gdy jego przełącznik się włączy, a niebieska dioda LED w odpowiednim przycisku powinna się zaświecić.

Również pomarańczowa dioda LED potwierdzenia (ACK) powinna migać po każdym naciśnięciu przycisku na pilocie. Jeżeli dioda ACK nie miga, sprawdź czy kod zaprogramowany w pilocie odpowiada trybowi pracy urządzenia (np. TV, SAT1 lub SAT2). Jeśli kod nie jest poprawny, dioda ACK nie będzie migać.

Teraz sprawdź, czy potencjometr głośności obraca się w prawo po naciśnięciu przycisków „Głośność +” (Volume Up) i „Kanał +” (Channel Up) oraz w przeciwnym kierunku po naciśnięciu przycisków „Głośność -” (Volume Down) i „Kanał -” (Channel Down). Powinien obracać się dość szybko, gdy wciskane są przyciski „Głośność +/-”, a w wolniejszym tempie, gdy używane są przyciski „Kanał +/-”.

Jeśli potencjometr obraca się w złym kierunku, należy zamienić przewody do silnika.

Regulacja potencjometru nastawnego VR2

Następnie ustawiamy potencjometr głośności w pozycji środkowej, skręcamy VR2 całkowicie w lewo i naciskamy przycisk „Wyciszenie” (Mute). Potencjometr będzie obracał się w lewo i gdy tylko trafi na ogranicznik, sprzęgło zacznie się ślizgać.

W tym czasie powoli obracaj VR2 w prawo, aż silnik się zatrzyma. Teraz naciśnij przez kilka sekund „Głośność +”, by obrócić potencjometr w prawo, i ponownie naciśnij

„Wyciszenie”. Tym razem silnik powinien zatrzymać się, gdy tylko potencjometr osiągnie granicę obrotu w lewo.

Program zatrzyma również silnik, jeśli będzie się on nadal obracał przez 13 sekund po włączeniu funkcji „Wyciszenie”. Oznacza to, że należy wyregulować VR2 w ciągu tego okresu 13 sekund. Jeżeli silnik zatrzymuje się przedwcześnie lub pracuje przez pełne 13 sekund po osiągnięciu krańcowego położenia, należy spróbować ponownie przeprowadzić regulację.

Rozwiązywanie problemów

Jeżeli urządzenie nie reaguje na sygnały z pilota, należy sprawdzić, czy pilot znajduje się w odpowiednim trybie (TV, SAT1 lub SAT2) i czy został prawidłowo zaprogramowany.

Jeśli używasz pilota innego niż wymienione w panelu, przetestuj różne kody, aż znajdziesz ten, który działa. Zaczynaj od kodów wymienionych pod marką Philips, ponieważ mają one największe szanse na działanie.

Jeśli urządzenie reaguje na przyciski 1, 2 i 3 na pilocie, ale przełączniki przycisków nie działają, sprawdź, czy taśma do płytki przycisków została prawidłowo zaciśnięta. Podobnie, jeśli funkcja zdalnej regulacji głośności działa, ale nie działa zdalny wybór wejścia, sprawdź przewód od płytki przedwzmacniacza do płytki wyboru wejścia.

Zauważ, że kabel z płytki przedwzmacniacza dostarcza również zasilanie do pozostałych dwóch płytek.

Warto więc sprawdzić, czy pomiędzy stykami „8” i „4” podstawki IC4 na płytce wyboru wejść jest napięcie 5 V i ponownie sprawdzić kabel, jeśli brakuje tego zasilania.

Testy audio

Jeśli do testowania przedwzmacniacza używasz zasilania ± 15 V, możesz dalej testować przedwzmacniacz podłączając jego wyjścia do wzmacniacza stereo i podając sygnały audio z telefonu komórkowego, tabletu, iPod-a, odtwarzacza CD/DVD/Blu-Ray lub dowolnego innego źródła.

W zależności od urządzenia, do wykonania połączenia może być potrzebny kabel z wtykiem Jack stereo 3,5 mm na jednym końcu i czerwono-białymi wtykami RCA na drugim końcu. Są one powszechnie dostępne. ■

John Clarke

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

Bezprzewodowy retransmitter danych w paśmie 433 MHz



Istnieje wiele urządzeń „zdalnego sterowania”, które wykorzystują pasmo 433 MHz (a dokładnie: 433,050...434,790 MHz) do przesyłania danych.

Być może posiadasz jedno z nich i nawet nie zdajesz sobie z tego sprawy – pilot alarmu samochodowego, sterownik drzwi garażu/bramy wjazdowej, a nawet zewnętrzna stacja pogodowa – to tylko niektóre przykłady.

Ale czy Twoje urządzenia są na pewno w 100% niezawodne? A może zasięg jest nieco mniejszy niż byś chciał? Być może pilot jest zbyt daleko od odbiornika, albo na drodze sygnału stoją wzgórza, drzewa lub inne przeszkody? Oto rozwiązanie: mały, zasilany energią słoneczną wzmacniacz – retransmitter sygnału radiowego, który umieszcza się, z zachowaniem dobrej widoczności, między nadajnikiem a odbiornikiem. W ten sposób uzyskasz niezawodność i dodatkowy zasięg, którego potrzebujesz.

Istnieje sporo urządzeń, które przesyłają okresowo pakiety danych w paśmie 433 MHz UHF „LIPD” (Low Interference Potential Devices – Urządzenia o potencjalnie niskich zakłóceniach), w tym kilka naszych konstrukcji, takich jak nasz Driveway Monitor – Detektor samochodów (SC: lipiec i sierpień 2015; <https://siliconchip.com.au/Series/288>).

Dotyczy to również niektórych urządzeń komercyjnych, takich jak zdalne stacje pogodowe. Niestety, nie zawsze jest możliwe uzyskanie niezawodnego odbioru.

Czasami jest to spowodowane tym, że pomiędzy umiejscowieniem nadajnika i odbiornika znajdują się wzgórza, drzewa, budynki itp.

Może to być też spowodowane ograniczonymi rozmiarami anten lub prawnym limitem mocy emisji wynoszącym 25 mW (Red. EdW: [patrz dodatkowe informacje w ramce „Czy tego retransmitera można używać legalnie bez licencji?”](#)). Aby nie komplikować opisu, pozostawiamy tekst wg norm australijskich, natomiast Czytelnik musi pamiętać o przepisach europejskich, wymienionych w ramce) dla nielicencjonowanych urządzeń działających w tym paśmie (wiele nadajników 433 MHz ma znacznie mniejszą moc).

Wpływ może mieć nawet pogoda: krzew lub drzewo, które w niewielkim stopniu ogranicza zasięg podczas suchej pogody, może przy wilgotnym powietrzu praktycznie całkowicie stłumić sygnał UHF.

Wprawdzie sygnały 433 MHz nie są tak bardzo tłumione jak wyższe częstotliwości (np. pasmo 2,4 GHz, które jest również używane do przesyłania danych, przykładowo Wi-Fi), jest to niewielkim pocieszeniem, gdy masz kłopoty z powodu wygłuszonego sygnału, co może uniemożliwić przesyłanie danych.

Ten retransmitter może być umieszczony w miejscu, gdzie ma wyraźny i niezawodny odbiór sygnałów z nadajnika, a które to miejsce jest również lepsze do odbioru przez jednostkę

Parametry

- * Zwiększa zasięg nadajników 433 MHz
- * Pokonuje ograniczenia „brak widoczności” spowodowane przez drzewa, przeszkody itp.
- * Odbiera sygnał w paśmie 433 MHz i ponownie nadaje z częstotliwością 433 MHz, z krótkim opóźnieniem
- * Nadaje się do stosowania w projektach z modulacją ASK
- * Odróżnia prawidłowy sygnał od szumu
- * Możliwe użycie kilku retransmiterów w łańcuchu
- * Regulowany okres opóźnienia
- * Regulowana maksymalna szybkość wykrywania poprawnych danych
- * Zasilanie słoneczne z buforem – akumulatorkiem LiFePO₄
- * Zasięg do 200 m na otwartej przestrzeni dzięki zoptymalizowanej antenie



Dzięki panelowi słonecznemu, który utrzymuje wewnętrzny akumulator naładowany, urządzenie jest bezobsługowe. Zwiększ zasięg Twojej transmisji danych nawet dwukrotnie!

odbiorczą (tj. może być umieszczony gdzieś pomiędzy tymi dwoma urządzeniami).

Przechowuje on w pamięci odebrane dane, a następnie, z krótkim opóźnieniem, transmituje ten sam sygnał w tym samym paśmie częstotliwości. Pozwala to nawet na dwukrotne zwiększenie zasięgu sieci bezprzewodowej, gdy możliwa jest transmisja w linii widoczności.

Jest również niezwykle skuteczny w poprawianiu integralności sygnału, gdy między dwoma urządzeniami znajdują się przeszkody, w tym budynki, drzewa i przeszkody terenowe.

Co powinienś wypróbować w pierwszej kolejności?

Zanim zbudujesz retransmitter, istnieje kilka prostych sposobów na poprawę zasięgu, które mogą dać ci zasięg, jakiego potrzebujesz. Po pierwsze, wypróbuj efektywniejszą antenę. (Stare powiedzenie radioamatorów – najlepszym wzmacniaczem sygnału jest antena)

W naszych projektach używamy jako anteny zazwyczaj krótkiego odcinka drutu, o rozmiarze jednej czwartej długości fali. Jest

to około 170 mm dla nadajnika lub odbiornika 433 MHz.

Użycie komercyjnej anteny typu „whip” („bicz”, albo prętowa, przykładem anteny teleskopowej do odbiorników przenośnych) do nadajnika i/lub odbiornika może poprawić zasięg w porównaniu z prostą anteną drucianą, podobnie jak zastosowanie dłuższej anteny półfalowej (340 mm dla 433 MHz).

Ale musimy Cię ostrzec, że jeśli Twój nadajnik jest bliski prawnej granicy 25 mW mocy promieniowania, użycie lepszej anteny (o większym zysku) może być nielegalne. Dzieje się tak dlatego, że 25 mW jest limitem efektywnie emitowanej mocy, więc obejmuje zysk anteny. Każdy wzrost zysku anteny o 3 dB jest równoznaczny z podwojeniem mocy emisji.

Nie można więc legalnie używać anteny +3 dB z urządzeniami, których moc nadawcza przekracza 12,5 mW.

Ważna jest również orientacja anteny. Posiadanie anteny nadawczej i odbiorczej o tej samej orientacji, np. obie zorientowane

Czy tego retransmitera można używać legalnie bez licencji?

Jednym słowem – tak.

Możesz zobaczyć licencję klasy „LIPD” (Low-Interference-Potential-Devices – urządzenia o potencjalnie niskim oddziaływaniu/zakłóceniach) na pasmo 433MHz „ISM” (Industrial, Scientific, Medical – „przemysłowe, naukowe, medyczne”), która dotyczy wszystkich w Australii, na stronie: <https://bit.ly/3CZNt8b>

Odpowiednik tego dokumentu dla Nowej Zelandii jest dostępny pod adresem: <https://bit.ly/3N6BTnq>.

Zauważ, że nowozelandzki limit EIRP (Effective Isotropic Radiated Power – efektywna moc promieniowana izotropowo) wynoszący -16 dBm jest taki sam jak australijski limit 25 mW. Jest on po prostu określony w innych jednostkach.

Żaden z tych dokumentów nie nakłada żadnych ograniczeń na wykorzystanie pasma 433/434 MHz LIPD innych niż maksymalna efektywna moc promieniowana. Nic nie ogranicza tego jak często można nadawać w tym paśmie, ani jak długie mogą być te pakiety danych. Nie ma też żadnej wzmianki o retransmiterach.

Ponieważ nasz retransmitter używa komercyjnie dostępnych nadajników 433 MHz, które są zgodne z limitem mocy, i ponieważ nadaje tylko po zakończeniu oryginalnej transmisji, może całkowicie legalnie pracować w Australii i Nowej Zelandii.

Jednakże nie zalecamy używania tego retransmitera z jakimikolwiek sygnałami, które nadawane są często. Zazwyczaj należy go używać w połączeniu z urządzeniem, które wysyła krótki pakiet danych (w czasie znacznie poniżej jednej sekundy) nie częściej niż, powiedzmy, raz na 30 sekund. Jeśli używałbyś go z urządzeniem nadającym za często, mógłbyś pokryć pasmo 433 MHz swoimi transmisjami w promieniu 100–200 m.

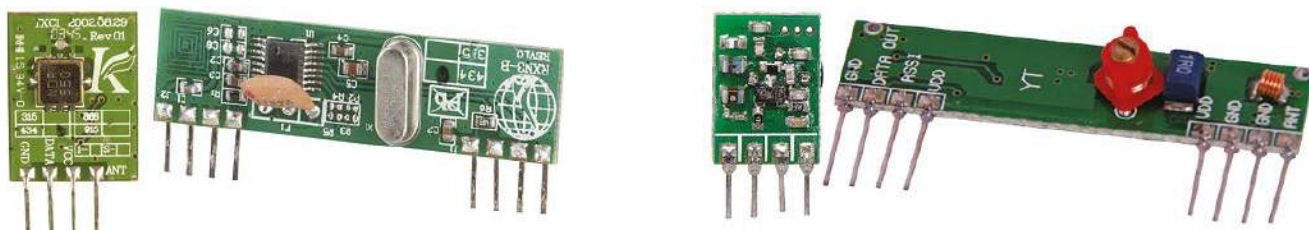
Class License stwierdza, że: „W przypadku wystąpienia zakłóceń, na użytkownika LIPD spoczywa obowiązek podjęcia środków w celu ich usunięcia, na przykład poprzez przestrojenie lub zaprzestanie pracy urządzenia.” (Przestrojenie tych urządzeń byłoby trudne, jeśli nie niemożliwe, bez specjalistycznego sprzętu). Należy więc o tym pamiętać i zachować zdrowy rozsądek przy ustawianiu urządzenia nadawczego i retransmitera(ów).

pionowo lub obie zorientowane równolegle poziomo, może poprawić zasięg.

Jeśli te zmiany okażą się niepraktyczne lub nie będą wystarczająco skuteczne, wtedy sensowne będzie zbudowanie retransmitera.

Retransmitter umieszcza się w torze propagacji sygnału pomiędzy nadajnikiem a odbiornikiem.

Krótko mówiąc, retransmitter składa się z odbiornika i nadajnika UHF,



Retransmitter danych 433MHz korzysta z komercyjnych modułów nadawczo-odbiorczych, jak te na fotografii. Po lewej stronie: nadajnik Jaycar ZW3100 i odbiornik ZW3102, po prawej nadajnik Altronics Z6900 i odbiornik Z6905. Są one w zasadzie identyczne, każdy z nich będzie pasował bezpośrednio do naszej płytki.

a także mikrokontrolera, paru kB pamięci i zasilacza. Gdy retransmitter odbierze przeznaczony dlań sygnał, jest on przechowywany w pamięci, a następnie, z określonym opóźnieniem, retransmitowany, aby mógł być odebrany przez odbiornik.

Zwiększa to skutecznie zasięg transmisji, ponieważ retransmitter może być umieszczony bliżej zarówno nadajnika, jak i odbiornika niż ich wzajemna odległość, i ewentualnie w bardziej korzystnej lokalizacji (np. wyżej), gdzie na drodze sygnału będzie mniej przeszkód.

Istnieją inne typy retransmiterów, które działają nieco inaczej niż ten. Na przykład, wiele retransmiterów przesyła odebrany sygnał dalej na innej częstotliwości.

Zapobiega to konfliktom pomiędzy nadajnikiem a retransmitterem i pozwala retransmitterowi pracować efektywnie bez opóźnień. Ale odbiornik końcowy musi być w stanie odbierać na nowej częstotliwości, więc ten typ retransmitera nie jest „przezroczysty” (tzn. obecność retransmitera jest wtedy konieczna, w przeciwieństwie do niniejszego projektu).

Ten retransmitter przesyła dalej sygnał w tym samym paśmie częstotliwości, co odbierana transmisja. Oznacza to, że odbiornik końcowy nie musi być w żaden sposób modyfikowany.

Jednak retransmitter musi poczekać na zakończenie transmisji przed dalszym wysłaniem. W przeciwnym razie sygnały: odbierany i nadawany będą się wzajemnie zakłócać, a retransmitter może się nawet zapętlić, nieustannie wysyłając te same dane, odbierając je i znów wysyłając.

Projekty kompatybilne

Niektóre z projektów opublikowanych wcześniej w SC, które mogą skorzystać z tego retransmitera:

- zdalny przełącznik UHF (styczeń 2009; <https://bit.ly/3N6CTBa>),
- uniwersalny 10-kanalowy odbiornik zdalnego sterowania (czerwiec 2013; <https://bit.ly/3MWub8t>),
- wspomniany wcześniej Detektor samochodów oraz Transmitter sygnału IR w paśmie 433 MHz UHF (czerwiec 2013; <https://bit.ly/3U6UQSV>).

Wszystkie te wymienione projekty wykorzystywały standardowe nadajniki i odbiorniki 433 MHz UHF sprzedawane przez firmy Jaycar i Altronics (jak pokazano powyżej; ze względów prawnych odradzamy zakupy na pewnym portalu).

Kody katalogowe Jaycar to ZW3100 dla nadajnika i ZW3102 dla odbiornika, natomiast kody

katalogowe Altronics to Z6900 dla nadajnika i Z6905 dla odbiornika.

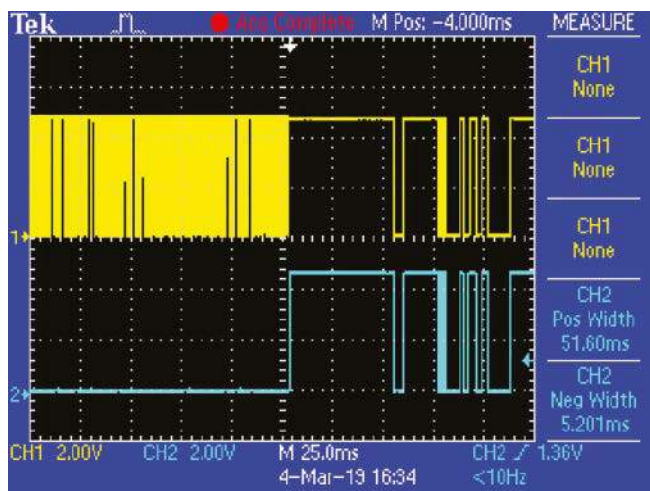
Wzmacniacz ten może współpracować z niektórymi innymi urządzeniami komercyjnymi transmitującymi dane w paśmie 433 MHz, jednak to, czy będzie działał, zależy od protokołów tych transmisji, więc trudno powiedzieć, że dane urządzenie będzie lub nie będzie działało, dopóki tego nie sprawdzisz.

Należy pamiętać, że retransmitera należy używać w sytuacjach, w których nie ma znaczenia, czy odbiornik mógłby odebrać w krótkim odstępie czasu dwa identyczne pakiety.

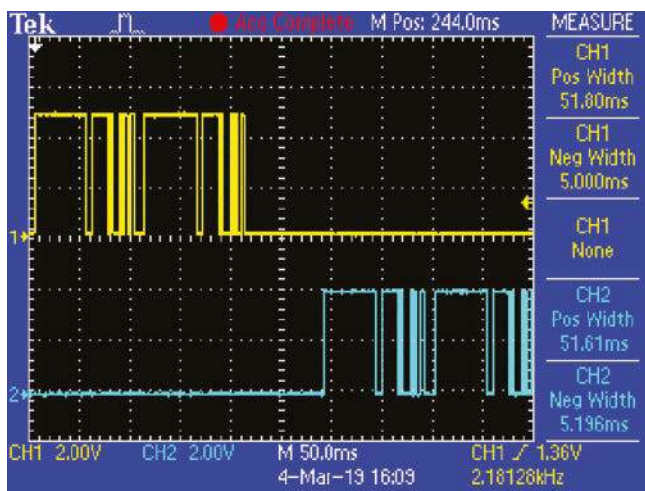
Wynika to z tego, że w niektórych przypadkach odbiornik może odebrać zarówno bezpośrednią transmisję, jak i tę z retransmitera.

We wszystkich wymienionych projektach nie powinno to mieć znaczenia, gdyż odbiorniki są efektywnie „bez zdefiniowanego stanu pracy” (nie mam pojęcia, czy istnieje sensowniejszy termin, posłużę się przykładem bramy wjazdowej – sygnał z pilota ją otwiera, powtórzony sygnał z retransmitera ją zamyka, czyli każdy sygnał ustawia odbiornik w określonej nowej konfiguracji. Jest to popularny gag w komediach).

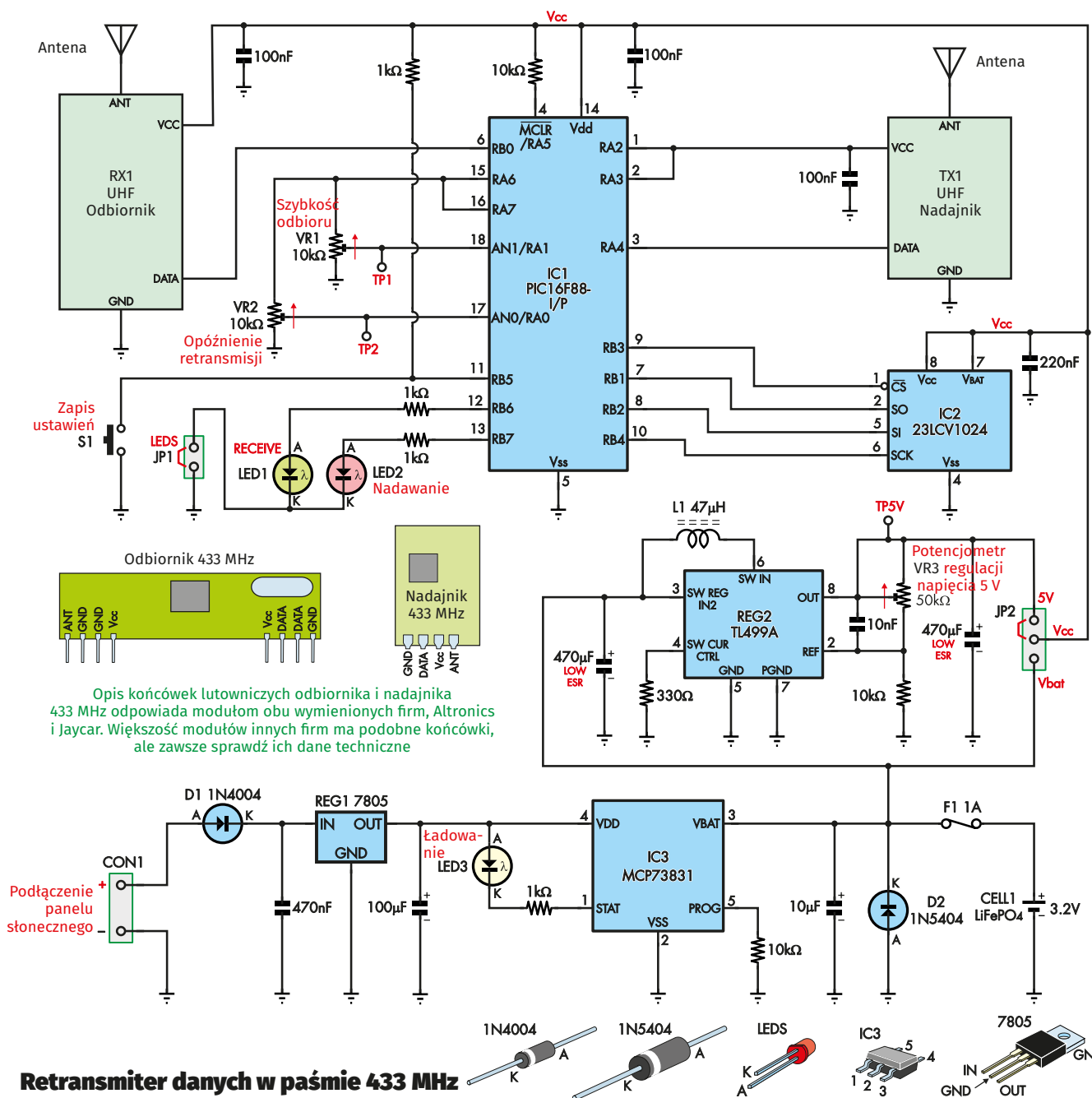
Powinno to być prawdziwe w przypadku wielu innych urządzeń, takich jak stacje pogodowe. Ale znowu trzeba będzie wypróbować, czy odbieranie



Oscylogram 1. Żółty przebieg na górze to sygnał na wyjściu odbiornika UHF, RX1. Widać szum o wysokiej częstotliwości przed odebraniem poprawnych danych. Gdy zostanie odebrany sygnał, losowy sygnał szumu jest stłumiony przez działanie ARW i rozpoczyna się detekcja nadawanego kodu. Układ IC1 ignoruje szum i akceptuje tylko prawidłowy kod, co widać na dolnym siniebieskim przebiegu.



Oscylogram 2. Żółty przebieg na górze pokazuje oryginalny sygnał odbierany z nadajnika źródłowego, natomiast przebieg siniebieski na dole pokazuje sygnał nadawany przez retransmitter. Widać, że nie rozpoczyna on transmisji, dopóki nie odbierze całego pakietu, a przed retransmisją występuje krótkie opóźnienie, w tym przypadku około 60 ms.



Retransmitter danych w paśmie 433 MHz

Rysunek 1. Schemat ideowy retransmitera. Transmisje danych są odbierane przez odbiornik UHF RX1 i podawane na wejście RB0 mikrokontrolera IC1. Są one następnie zapisywane w statycznej pamięci RAM IC2, a po zakończeniu transmisji odczytywane z powrotem z pamięci i przesyłane dalej do nadajnika UHF TX1. Następnie IC1 czeka przez zaprogramowany czas przed nasłuchiwaniami kolejnej transmisji. Napięcie zasilania z akumulatora LiFePO₄ jest zwiększane do 5 V przez REG2, natomiast akumulator jest ładowany z panelu słonecznego przy pomocy układu zarządzania ładowaniem IC3.

wielu identycznych pakietów nie ma negatywnego wpływu na działanie odbiornika.

Prezentacja

Nasz retransmitter jest umieszczony w szczelnej obudowie IP65, co oznacza, że nadaje się do użytku na zewnątrz, w miejscach, gdzie może być narażony na działanie czynników atmosferycznych.

Został zaprojektowany do zasilania z panelu słonecznego i wykorzystuje jeden akumulator LiFePO₄ o rozmiarze AA do przechowywania energii, więc może być

używany tam, gdzie nie ma dostępnego zasilania sieciowego.

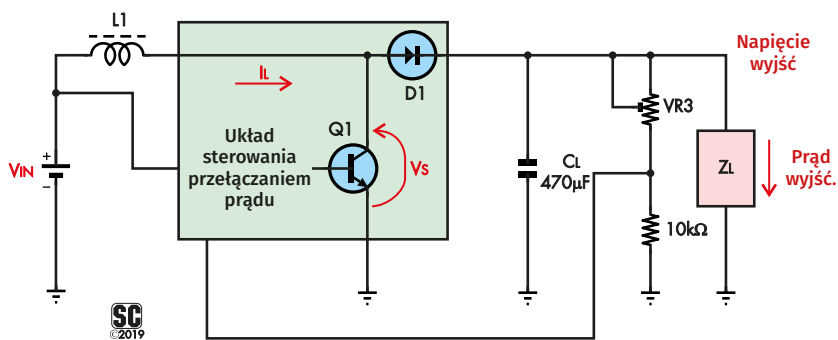
Jest to idealne rozwiązanie, ponieważ możesz zamontować go na przykład na słupie, gdzie będzie miał dobry „widok” na urządzenia: nadawcze i odbiorcze, a także powinien mieć na tyle dużo światła słonecznego, aby akumulator był naładowany.

Szczegóły układu

Schemat ideowy retransmitera jest pokazany na **rysunku 1**. Podstawowe podzespoły to: mikrokontroler IC1, wcześniej wspomniany

odbiornik (RX1) i nadajnik (TX1) 433 MHz, 1024 kbitów/128 kбайtów statycznej pamięci RAM (IC2), plus części zasilające takie jak ładowarka (IC3) i akumulator LiFePO₄ oraz przetwornica step-up 5 V (REG2).

Mikrokontroler IC1 monitoruje sygnał z odbiornika UHF (RX1), przechowuje odebrane dane w pamięci statycznej RAM (IC2), a następnie przekazuje te dane do nadajnika UHF (TX1), aby wysłać zapisany kod, który został wcześniej odebrany. IC1 posiada również dwa potencjometry nastawne (VR1 i VR2), które służą do ustawienia maksymalnej szybkości



Rysunek 2. Pokazuje jak przetwornica podwyższająca napięcie generuje 5 V do zasilania nadajnika i odbiornika UHF o napięciu 3,2...4,2 V z akumulatora. Układ przetaczania polaryzuje bazę wewnętrznego tranzystora Q1, który włącza przepływ prądu przez cewkę indukcyjną L1, co prowadzi do gromadzenia w niej energii pola magnetycznego. Gdy Q1 zostanie wyłączony, energia pola magnetycznego zgromadzona w dławiku L1 zamienia się na końcówkach L1 na energię impulsu prądowego o wysokim napięciu i takim samym kierunku przepływu, jak napięcie zasilania, a więc prąd ten poprzez diodę D1 ładuje kondensator wyjściowy CL do napięcia wyższego niż napięcie zasilania. Dioda D1 zabezpiecza jednocześnie kondensator CL przed rozładowywaniem poprzez tranzystor Q1, gdy jest on wyłączony. Końcowe napięcie na kondensatorze CL wynoszące 5 V jest regulowane przez porównanie napięcia odniesienia w układzie przetwarzającym z napięciem z dzielnika utworzonego przez potencjometr nastawny VR3 i rezystor 10 kΩ.

transmisji danych i minimalnego opóźnienia retransmisji; więcej o tym dalej.

Odbiornik RX1 jest zasilany w sposób ciągły napięciem 5 V, dzięki czemu w każdej chwili może odebrać sygnał. Gdy nie ma odbieranego sygnału, jego końcówka wyjściowa danych dostarcza losowy sygnał (szum) o wysokiej częstotliwości. Jest to spowodowane tym, że automatyczna regulacja wzmocnienia (ARW) w odbiorniku zwiększa wzmocnienie aż do momentu odbioru sygnału, nawet jeśli ten sygnał jest tylko wzmocnionym szumem.

Kiedy zostanie odebrany kodowany sygnał 433 MHz, ARW zmniejsza wzmocnienie odbiornika RX1, aby zapobiec wewnętrznemu przesterowaniu z powodu nadmiernego wzmocnienia. Ponieważ wzmocnienie ARW zmienia się stosunkowo wolno, po zatrzymaniu transmisji sygnału 433 MHz, wyjście odbiornika będzie w stanie niskim przez kilkadziesiąt milisekund, zanim działanie ARW zwiększy wzmocnienie wystarczająco, aby ponownie generować szum.

Nadajnik i odbiornik 433 MHz wykorzystują podstawowy system modulacji, znany jako kluczkowanie z przesunięciem amplitudy lub ASK. Gdy jego wejście jest w stanie wysokim (logiczne jeden), nadajnik wytwarza nośną 433 MHz. Kiedy jego wejście jest w stanie niskim (logiczne zero), transmisja nośnej 433 MHz zatrzymuje się.

Szybkość przesyłania danych jest zwykle na tyle duża, że wzmocnienie odbiornika RX1 nie zmienia się znacząco podczas transmisji stanów logicznych „0”, choć wówczas brak jest nośnej.

Istnieją różne schematy, stosowane w celu uniknięcia występowania długich ciągów stanów logicznych „0” lub „1”, niezależnie od przesyłanych danych, aby rozwiązać ten problem. Jedną z metod jest tzw. kodowanie Manchester, w którym każdy bit jest kodowany: albo jako

niski (0), a następnie wysoki (1); albo jako wysoki (1), a następnie niski (0), z ustaloną szybkością. Przy użyciu kodowania Manchester para nadajników i odbiorników UHF może przesyłać dane z prędkością do 5 kbit/s.

Odróżnianie sygnału od szumu

Działanie ARW odbiornika jest poważnym problemem dla naszego oprogramowania, ponieważ musi ono być w stanie odróżnić serię logicznych zer i jedynek, które stanowią treść poprawnej transmisji danych, od pseudo „zer” i „jedynek”, pojawiających się wskutek wzmocnienia w odbiorniku szumu, gdy nie ma żadnego sygnału.

IC1 monitoruje sygnał z odbiornika UHF na swoim wejściu cyfrowym RBO (końcówka 6). Przy każdej zmianie poziomu napięcia decyduje, czy jest to tylko szum, czy też prawidłowy sygnał danych. Poprawne dane są wykrywane przez porównanie prędkości odbierania danych z ustawieniem maksymalnej prędkości transmisji.

Ustawia się to za pomocą potencjometru nastawnego VR1, który zmienia również napięcie w punkcie testowym TP1. Przy napięciu na TP1 równym 0 V maksymalna szybkość transmisji danych wynosi 233 bps (bits per second, czyli bitów na sekundę), a przy napięciu 5 V maksymalna szybkość przesyłania danych wynosi 5 kbps (kbits per second, czyli kilobitów na sekundę). Pośrednie napięcia dają pośrednie wartości maksymalnej prędkości transmisji.

Jeśli szybkość przychodzących danych jest większa niż ustawienie szybkości VR1, dane są uważane za szum i odrzucane jako nieważne (zob. oscylogram 1).

Jeżeli prędkość transmisji odbieranych danych jest mniejsza niż ustawiona maksymalna prędkość ich przesyłania,

dane są uważane za ważne i są zapisywane w pamięci. Po tym, jak szybkość transmisji ponownie przekracza ustawioną szybkość maksymalną, oprogramowanie przyjmuje, że transmisja została zakończona i zapisane dane są następnie transmitowane.

Odbywa się to poprzez odczyt danych z pamięci RAM, realizowany za pomocą portów RB1, RB3, RB4 (końcówki 7, 9 i 10) układu IC1 i podanie ich na wyjście cyfrowe RA4 (końcówka 3) układu IC1 z taką samą szybkością, z jaką zostały odebrane. W tym samym czasie TX1, nadajnik UHF jest zasilany z portów RA2 i RA3 układu IC1 i zapisane dane zostają wysłane (patrz oscylogram 2).

IC2 jest pamięcią typu Static RAM, służącą do przechowywania danych. Jest to 1024-kbittowa pamięć o dostępie swobodnym i organizacji 128k x 8-bit. Odczyt i zapis pamięci odbywa się za pomocą interfejsu Serial Peripheral Interface (SPI).

Podczas zapisu dane są wysyłane do wejścia SI układu pamięci IC2 (końcówka 5) z wyjścia SDO RB2 (końcówka 8) układu IC1, po jednym bajcie na raz. Podczas odczytu dane są wysyłane z wyjścia SO układu IC2 (końcówka 2) do wejścia RB1 (końcówka 7) układu IC1; również po jednym bajcie. W obu przypadkach dane są taktowane sygnałem z portu RB4 (końcówka 10) IC1, który jest podawany na wejście SCK pamięci IC2 (końcówka 6).

Interfejs SPI pamięci jest włączany niskim poziomem na wejściu „chip select” (CSinv – wybór układu) (końcówka 1) układu IC2, sterowanym z wyjścia cyfrowego RB3 układu IC1 (końcówka 9).

Aby zapisać dane do pamięci, linia CSinv jest ustawiana na niskim poziomie, a następnie z IC1 do IC2 wysyłana jest instrukcja zapisu, po czym podawany jest adres komórki pamięci, do której dane mają być zapisane. W naszej aplikacji jest to zawsze pierwsza lokalizacja (adres zero). Jest to 24-bitowy adres wysyłany w postaci trzech (8-bitowych) bajtów. Siedem najbardziej znaczących bitów adresu ma zawsze wartość „zero”, ponieważ tylko 17 bitów jest potrzebne do zaadresowania 128 kbajtów (adresowane są komórki 8-bitowe) w pamięci RAM.

Następnie można zapisać dane, jeden bajt na raz. Domyślnie adres jest automatycznie zwiększany o „1” po każdym zapisanym bajcie danych, więc bajty są zapisywane do pamięci RAM sekwencyjnie.

Otrzymane dane są przechowywane jako wartości 16-bitowe. Najbardziej znaczący bit (bit Fhex) wskazuje na poziom odebranego sygnału: niski (0) lub wysoki (1). Pozostałe 15 bitów służy do zapisania czasu, przez jaki dane pozostawały na tym poziomie. Okres ten jest przechowywany z przyrostami po 4 µs, co daje minimalny czas trwania 4 µs i maksymalny 131 ms.

Odczyt danych z pamięci jest procesem sekwencyjnym podobnym do zapisu, z tą różnicą, że używana jest inna instrukcja, a dane wysyłane są w przeciwnym kierunku, od IC2 (port SO, końcówka 2) do IC1 (port RB1, końcówka 7).

Funkcje oszczędzania energii

Ponieważ zasilamy retransmitter za pomocą panelu słonecznego i małego akumulatora buforowego, zużycie energii przez moduł musi być zminimalizowane, zwłaszcza w stanie bezczynności i oczekiwania na dane. Odbywa się to poprzez wyłączenie zasilania niepotrzebnych w danej chwili podzespołów.

Dwa potencjometry nastawne, VR1 i VR2, są podłączane do zasilania 5 V tylko wtedy, gdy ich pozycje są odczytywane. Dzieje się tak tylko podczas początkowego procesu włączania

zasilania oraz w momencie wciśnięcia przełącznika S1. W każdym innym przypadku końcówki RA6 i RA7, które dostarczają napięcie 5 V do potencjometrów, są w stanie niskim (0 V), co ma zapobiec przepływowi prądu przez te potencjometry. Pozwala to zaoszczędzić 1 mA, co daje 24 mAh na dobę.

Podobnie nadajnik (TX1) jest wyłączony, dopóki nie jest wymagane wysłanie sygnału UHF. TX1 jest zasilany bezpośrednio z wyjść cyfrowych IC1: RA2 i RA3 (końcówki 1 i 2); przechodzą one w stan wysoki (do 5 V), aby zasilić TX1. Oszczędność energii jest znaczna, ponieważ TX1 pobiera około 10 mA podczas zasilania i nadawania. To oszczędza 240 mAh/dobę.

Pamięć IC2, jeśli nie jest używana, pozostaje w stanie czuwania. Tak więc gdy nie ma ważnych danych do odebrania/wysłania, pobiera tylko 10 µA zamiast 10 mA wymaganych

w stanie aktywnym. Zazwyczaj pamięć jest zasilana tylko dwa razy dłużej niż nadajnik; pierwsza połowa to okres odbioru, a druga połowa to okres nadawania. To również pozwala zaoszczędzić nieco mniej jak 240 mAh/dobę.

Mikrokontroler IC1 pobiera typowo 1,7 mA, a odbiornik UHF RX1 zwykle 2,9 mA. Diody LED nadawania i odbioru są zasilane, gdy aktywne są odpowiednio TX1 i RX1 i pobierają po ok. 3 mA. Diody LED, jeśli nie są potrzebne, mogą być odłączone za pomocą zwory (JPL) w celu oszczędzania energii. Są one przewidziane głównie do celów testowych.

Układ jest zasilany z jednego akumulatora 600 mAh typu LiFePO₄ o rozmiarze AA. Pobór prądu w stanie spoczynkowym wynosi około 9,4 mA, tzn. kiedy nadajnik, pamięć i diody LED są wyłączone. Oznacza to, że ogniwo będzie działać przez około 60 godzin lub 2,5 dnia po pełnym naładowaniu.

Obwód ładowania akumulatora

Akumulator LiFePO₄ jest ładowany przez układ IC3, zasilany z regulatora/stabilizatora 5 V (REG1), a ten z kolei jest zasilany z panelu słonecznego. Zauważ, że akumulator jest podłączony do reszty układu poprzez bezpiecznik (F1), który zapobiega uszkodzeniu w przypadku odwrotnego włożenia ogniwa. Jeśli ogniwo zostanie źle włożone, prąd popłynie przez diodę D2 i przepali bezpiecznik. Dioda D1 zapobiega uszkodzeniu układu, jeśli panel słoneczny zostanie podłączony z niewłaściwą polaryzacją.

IC3 to miniaturowy układ zarządzania ładowaniem pojedynczych ogniw Li-ion/LiFePO₄. Ładuje on ogniwo stałym prądem, aż do napięcia końcowego równego 4,2 V. Prąd ładowania jest ustawiany przez rezystor przy końcówce 5, i w naszym układzie jest ustawiony na 100 mA przez rezystor 10 kΩ. Dioda (LED3) świeci, gdy ogniwo jest ładowane.

Odbiornik (RX1) i nadajnik (TX1) 433 MHz UHF mogą pracować zasilane napięciem 2,5...5 V. Ponieważ nadajnik będzie miał większą moc wyjściową, a tym samym lepszy zasięg, gdy będzie zasilany napięciem 5 V, a nie napięciem 3,2...4,2 V bezpośrednio z akumulatora LiFePO₄, używamy przetwornicy podwyższającej napięcie (step-up boost) do generowania napięcia 5 V zasilającego te moduły.

Jednakże moduł może być zbudowany bez tej przetwornicy, jeśli nie jest wymagany maksymalny zasięg. Pozwala to zaoszczędzić czas i pieniądze. Reszta układu będzie wtedy zasilana bezpośrednio z akumulatora LiFePO₄. To również wydłuży żywotność tego ogniwa, ponieważ przetwornica „step-up” ma sprawność tylko około 70%, a niższe napięcie zasilania będzie również oznaczać, że przez IC1, IC2, TX1 i RX1 jest pobierany mniejszy prąd.

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

- 1 szt. dwustronna płytka drukowana o kodzie 15004191, 103,5×78 mm
- 1 szt. obudowa IP65, 115×90×55 mm z tworzywa sztucznego (ABS) [Jaycar HB6216 lub podobna]
- 1 szt. akumulator 600 mAh LiFePO₄ (rozmiar AA (!!!), o średnicy 14 mm i długości 50 mm) [Jaycar SB2305]
- 1 szt. panel słoneczny 12 V/5 W [Jaycar ZM9050 lub podobny]
- 1 szt. toroid ze sproszkowanego żelaza w zalewie epoksydowej 15×8×6,5 mm [Jaycar L01242] (L1)
- 1 szt. nadajnik 433 MHz ASK [Altronics Z6900, Jaycar ZW3100] (TX1)
- 1 szt. odbiornik 433 MHz ASK [Altronics Z6905, Jaycar ZW3102] (RX1)
- 1 szt. chwilowy przełącznik przyciskowy SPST do montażu na PCB [Altronics S1120, Jaycar SP0600] (S1)
- 1 szt. zacisk śrubowy 2-stykowy, raster 5,08 mm (CON1)
- 1 szt. 2-szpilkowy odcinek listwy kołkowej prostej do druku, raster 2,54 mm (JP1)
- 1 szt. 3-szpilkowy odcinek listwy kołkowej prostej do druku, raster 2,54 mm (JP2)
- 2 szt. zworki (do JP1, JP2)
- 1 szt. podstawka pod układ scalony DIL18 wąska (dla IC1)
- 1–2 szt. podstawka pod układ scalony DIL8 (opcjonalnie; dla IC2 i REG2)
- 1 szt. pojemnik na ogniwa AA montowany na płytce drukowanej
- 1 szt. radiator, 19×19×9,5 mm [Altronics H0630, Jaycar HH8502 lub inny, profil „C” FK239]
- 1 szt. wodoszczelny przepust kablowy IP65 pasujący do kabla o średnicy 3...6,5 mm
- 6 szt. kołków PC (opcjonalnie)
- 4 szt. śrubki M3×5
- 1 szt. śrubka M3×6
- 1 szt. nakrętka sześciokątna M3
- 2 szt. wkrety samogwintujące #0 o długości 4,75 mm
- 2 szt. opaski zaciskowe 100 mm
- 1 etykieta panelu (patrz tekst)
- 1 bezpiecznik 1 A M5×20 (F1)
- 1 komplet oprawek bezpiecznika M5×20 do montażu na płytce drukowanej (F1)
- 1 odcinek o długości 500 mm emaliowanego drutu miedzianego (DNE) o średnicy 1 mm (na cewkę L1)
- 2 odcinki o długości 175 mm izolowanego drutu potężeniowego o średniej obciążalności lub:
- 2 odcinki o długości 175 mm emaliowanego drutu miedzianego (DNE) o średnicy 1 mm (patrz tekst)

Półprzewodniki:

- 1 szt. 8-bitowy mikrokontroler PIC16F88-I/P zaprogramowany za pomocą wsadu 1500419A.hex (IC1)
- 1 szt. 23LCV1024-I/P 128kB Static RAM obudowa PDIP [Mouser, Digi-Key] (IC2) LUB:
- 1 szt. 23LCV1024-I/SN 128kB Static RAM obudowa SOIC [Mouser, Digi-Key] (IC2)
- 1 szt. MCP73831T-2ACI/OT układ ładowania pojedynczych ogniw Li-ion/LiFePO₄, obudowa SOT-23-5 [Mouser, Digi-Key] (IC3)
- 1 szt. układ przetwornicy „step-up” TL499A [Jaycar ZV1644] (REG2)
- 1 szt. regulator/stabilizator 5 V 7805 (REG1)
- 1 szt. dioda 1N4004 1 A (D1)
- 1 szt. dioda 1N5404 3 A (D2)
- 1 szt. zielona dioda LED 3 mm o wysokiej jasności (LED1)
- 1 szt. czerwona dioda LED 3 mm o wysokiej jasności (LED2)
- 1 szt. żółta dioda LED 3 mm o wysokiej jasności (LED3)

Kondensatory:

- 2 szt. kondensator elektrolityczny low-ESR 470 µF/16 V
- 1 szt. kondensator elektrolityczny 100 µF/16V
- 1 szt. kondensator elektrolityczny 10 µF/16V
- 1 szt. kondensator poliestrowy 470 nF/63V MKT (kod 0.47, 474 lub 470n)
- 1 szt. kondensator poliestrowy 220 nF/63V MKT (kod 0.22, 224 lub 220n)
- 2 szt. kondensator poliestrowy 100 nF/63V MKT (kod 0.1, 104 lub 100n)
- 1 szt. kondensator ceramiczny MLC 100 nF
- 1 szt. kondensator poliestrowy 10 nF/63V MKT (kod 0.01, 103 lub 10n)

Rezystory: (wszystkie 0,25 W, 1%, metalizowane, ew. 0,6 W)

- 3 szt. 10 kΩ (brązowy czarny pomarańczowy brązowy lub brązowy czarny czerwony brązowy)
- 4 szt. 1 kΩ (brązowy czarny czerwony brązowy lub brązowy czarny czarny brązowy)
- 1 szt. 330 Ω (pomarańczowy pomarańczowy brązowy brązowy lub pomarańczowy pomarańczowy czarny czarny brązowy)
- 2 szt. 10 kΩ miniaturowe potencjometry nastawne do montażu poziomego (VR1, VR2)
- 1 szt. 50 kΩ miniaturowy potencjometr nastawny do montażu poziomego (VR3)

Zworka JP2 służy do wyboru trybu zasilania: z przetwornicy 5 V, albo bezpośrednio z akumulatora.

Podwyższenie napięcia realizowane jest przez układ przełączający TL499A – REG2. Składa się on z układu sterującego przełączaniem, tranzystora i szeregowej diody. Wymaga on dławika L1 do realizacji funkcji podwyższania napięcia oraz kondensatora wyjściowego 470 µF „low-ESR” do magazynowania energii i filtrowania impulsów prądu.

Uproszczony układ objaśniający działanie przetwornicy podwyższającej napięcie przedstawiono na rysunku 2. Początkowo wewnętrzny tranzystor Q1 jest wyłączony i prąd płynie przez cewkę L1 (z szybkością wzrostu natężenia ograniczoną jej indukcyjnością), aż do osiągnięcia określonego natężenia. Ten maksymalny prąd jest ustawiany przez rezystor podłączony do końcówki 4 REG2.

Gdy Q1 zostanie wyłączony, energia pola magnetycznego zmagazynowana w dławiku L1 zamienia się na końcówkach L1 na energię impulsu prądowego o wysokim napięciu i takim samym kierunku przepływu, jak napięcie zasilania, a więc prąd ten poprzez diodę D1 ładuje kondensator wyjściowy CL do napięcia wyższego niż napięcie zasilania. Dioda D1 zabezpiecza jednocześnie kondensator CL przed rozładowywaniem poprzez tranzystor Q1, gdy jest on wyłączony.

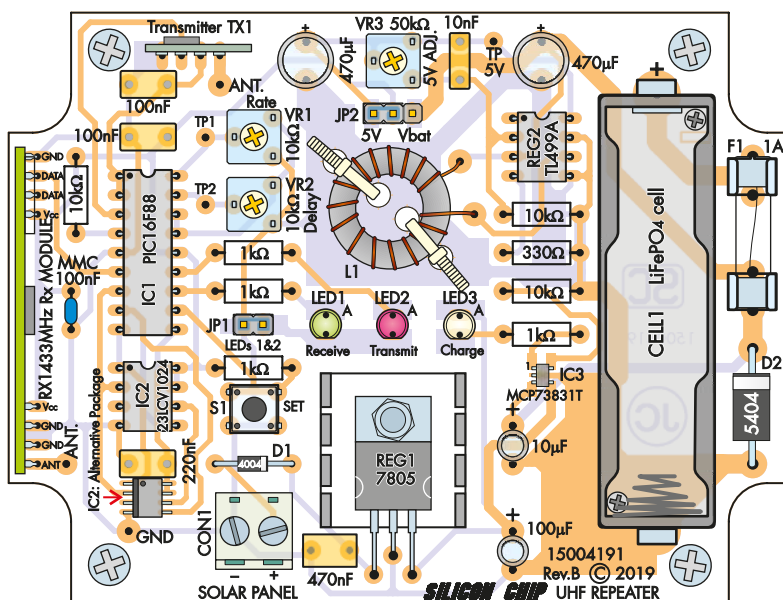
Proces ten jest powtarzany po ponownym włączeniu Q1 w następnym cyklu ładowania CL.

Napięcie wyjściowe jest próbkiwane poprzez dzielnik napięcia składający się z potencjometru nastawnego VR3 i rezystora 10 kΩ. Napięcie z dzielnika podawane na końcówkę 2 układu REG2 jest porównywane z wewnętrznym napięciem odniesienia równym 1,26 V. Cykl pracy Q1 jest uruchamiany i zatrzymywany tak, aby utrzymać napięcie 1,26 V na wejściu 2 układu REG2.

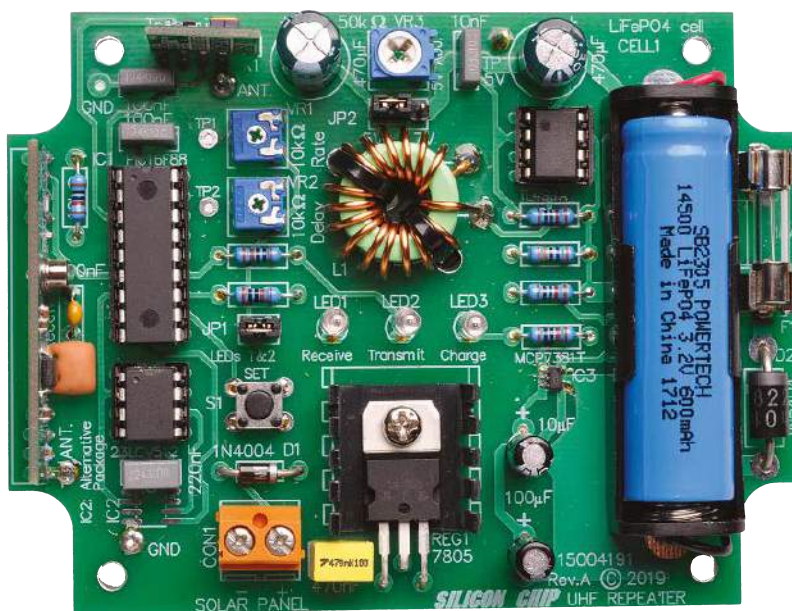
Dlatego zmieniając rezystancję VR3, możemy zmieniać napięcie wyjściowe przetwornicy. Im większe jest tłumienie tego dzielnika rezystancyjnego, tym większe musi być napięcie wyjściowe, aby utrzymać potencjał 1,26 V na końcówce 2. Jeśli VR3 ustawimy na 29,7 kΩ, to dzielnik utworzony z rezystorem 10 kΩ zmniejsza napięcie wyjściowe 3,97 razy. Oznacza to, że na wyjściu będzie $3,97 \times 1,26 \text{ V} = 5 \text{ V}$.

Jeśli napięcie wyjściowe wzrośnie nieznacznie powyżej 5 V, układ TL499A wyłączy tranzystor Q1, aż do momentu, gdy napięcie spadnie nieco poniżej poziomu 5 V. Jeśli napięcie spadnie poniżej 5 V, tranzystor Q1 rozpocznie cykliczną pracę aż do przywrócenia napięcia do poziomu 5 V.

Należy pamiętać, że napięcie referencyjne 1,26 V jest tylko wartością nominalną



Rysunek 3. Ten schemat montażowy płytki drukowanej oraz zdjęcie poniżej pokazują sposób montażu elementów na PCB. Istnieją dwie możliwe lokalizacje dla IC2, w zależności od tego, czy używasz wersji do montażu przewlekane (DIP) czy SMD (SOIC). Uważaj, aby przed włączaniem prawidłowo zorientować diody, układy scalone, pojemnik akumulatora, nadajnik i odbiornik, jak pokazano na rysunku. Można pominąć niektóre elementy, jeśli nie jest potrzebna funkcja ładowania akumulatora z panelu słonecznego (szczegóły w tekście).



i może to być dowolne napięcie w zakresie od 1,20 V do 1,32 V, w zależności od konkretnego układu scalonego. Tak więc VR3 umożliwia dokładne ustawienie napięcia wyjściowego.

Łączenie wielu retransmiterów

Jak wspomniano w panelu parametrów, możliwe jest stosowanie więcej niż jednego retransmitera, aby jeszcze bardziej zwiększyć zasięg transmisji. Retransmiter najbliższy źródła (oryginalnego nadajnika) będzie wysyłał sygnał do drugiego retransmitera. Gdy drugi retransmiter wyśle swój sygnał, pierwszy

retransmiter musi go zignorować; w przeciwnym razie oba retransmitery będą w nieskończoność retransmitować ten sam pakiet.

Zapobiega temu regulowane opóźnienie pomiędzy końcem każdej transmisji a przyjęciem przez urządzenie nowego pakietu. Opóźnienie to wynosi od 50 ms do 12,5 s i jest ustawiane za pomocą VR2. Napięcie na TP2 wskazuje ustawienie opóźnienia, przy czym każdy wolt odpowiada 2,5 s opóźnienia. Tak więc, jeśli VR2 jest ustawiony w pozycji odpowiadającej napięciu 2 V na TP2, to opóźnienie wynosi $2,5 \times 2 = 5 \text{ s}$. Ustawienie 0 V daje opóźnienie 50 ms (minimum).

Budowa

Wzmacniacz jest lutowany na dwustronnej płycie PCB o kodzie 15004191, o wymiarach 103,5×78 mm. Mieści się ona w szczelnej obudowie klasy IP65 z tworzywa sztucznego o wymiarach 115×90×55 mm. Podczas wykonania należy korzystać ze schematu montażowego, pokazanego na **rysunku 3**, jako instrukcji.

Zacznij od wlutowania układu IC3 sterowania ładowaniem akumulatora. Znajduje się on w małej pięciostykowej obudowie SMD. Prawidłowa orientacja jest oczywista, ponieważ ma on dwa styki z jednej strony i trzy z drugiej. Przylutuj jeden ze styków (najlepiej w prawym górnym rogu), a następnie sprawdź jego położenie w stosunku do pól lutowniczych i przylutuj styk po przekątnej.

Następnie możesz przystąpić do lutowania pozostałych styków, a pierwszy przelutować odrobiną lutu lub żelu topnikowego.

Jeśli przypadkowo zewrzesz trzy styki, które są blisko siebie, dodaj odrobinę topnika, a następnie oczyść końcówki za pomocą miedzianej plecionki lutowniczej.

Płytkę ma możliwość zastosowania układu pamięci (IC2) w obudowie DIP (przewlekanej) lub SOIC (SMD). Należy wlutować tylko jeden typ! Jeśli używasz obudowy SOIC, przylutuj pamięć w następnej kolejności, używając podobnej procedury jak opisana powyżej. Najpierw jednak upewnij się, że kropka lub szczelina przy styku 1 znajduje się u góry po lewej stronie, jak pokazano na **rysunku 3**. Krawędź od strony styku 1 powinna być dodatkowo fabrycznie ścięta.

Obudowa SOIC dla IC2 jest większa niż IC3, więc powinno być Ci nieco łatwiej. Ponownie, wszelkie przypadkowe zwarcia pomiędzy stykami można usunąć topnikiem i plecionką lutowniczą.

Następnie wlutuj rezystory. Są one oznaczone kolorowym kodem paskowym odpowiadającym rezystancji, jak podano na liście części. Do sprawdzenia wartości rezystorów należy użyć multimetru cyfrowego, ponieważ kolorowe paski mogą być trudne do odczytania.

Kolejno polutuj diody, pamiętając o tym, aby włożyć je z właściwą polaryzacją, tzn. z paskami skierowanymi tak, jak pokazano na schemacie montażowym. D2 jest znacznie większa niż D1.

Zalecamy użycie podstawki dla IC1. Pozostałe układy scalone (w tym IC2, jeśli jest używana wersja w obudowie DIP) mogą być zamontowane w podstawkach lub wlutowane bezpośrednio do PCB, co daje lepszą długoterminową niezawodność. Zwróć uwagę na orientację podczas montażu podstawek i układów scalonych. Dodatkowo upewnij się, że IC2 i REG2 nie są zamienione.

Następnie do zainstalowania jest sześć opcjonalnych kołków typu PC. Ułatwiają one

podłączenie przewodów i sprawdzanie punktów testowych. Znajdują się one w miejscach oznaczonych: TP5V, GND, TP1, TP2 oraz po jednym dla podłączenia anten modułów RX1 i TX1.

W następnej kolejności należy zamontować kondensatory, zaczynając od wielowarstwowego kondensatora ceramicznego 100 nF obok odbiornika UHF RX1, a następnie kondensatory poliestrowe typu MKT, z dowolną orientacją. Po nich należy zamontować kondensatory elektrolityczne, które muszą być wlutowane zgodnie z pokazaną polaryzacją; dłuższy przewód przechodzi przez pole lutownicze oznaczone znakiem „+”, umieszczone bliżej góry płytki.

Teraz można zamontować REG1. Jest on umieszczony poziomo na radiatorze. Wygnij wyprowadzenia tak, aby pasowały do otworów w płycie, a otwór montażowy pokrywał się z otworem w PCB. Umieść radiator pomiędzy regulatorem a płytką drukowaną i przykręć śrubkę oraz nakrętkę przed przylutowaniem wyprowadzeń. Odrobina pasty termoprzewodzącej pomiędzy spodem REG1 i radiatorem zdecydowanie poprawi chłodzenie.

Potencjometry nastawne VR1 do VR3 są następne w kolejce. VR1 i VR2 mają po 10 kΩ i zazwyczaj są oznaczone kodem 103. VR3 ma wartość 50 kΩ i może być oznaczony jako 503. Następnie zamontuj diody LED, od LED1 do LED3. W każdym przypadku anoda (dłuższe wyprowadzenie) przechodzi przez otwór w polu lutowniczym na PCB oznaczonym literą „A”. Po wlutowaniu dolna część diod powinna znajdować się około 5 mm nad powierzchnią płytki. Następnie można zamontować przełącznik przyciskowy S1.

Zamontuj teraz krótkie, 3-szpilkowe i 2-szpilkowe odcinki pojedynczej prostej listwy kołkowej pod zworki JP1 i JP2. Następnie zamontuj 2-drożny zacisk śrubowy CON1, skierowany otworami na przewody na zewnątrz dolnej krawędzi PCB.

Dławik L1 ma nawinięte 17 zwojów emaliowanego drutu miedzianego o średnicy 1 mm na toroidalnym rdzeniu o średnicy 15 mm ze sprężynowanego żelaza w zalewie epoksydowej. DNE powinien być nawinięty starannie, równomiernie wokół pierścienia, jak pokazano na fotografii pod rysunkiem 3. Indukcyjność gotowego dławika powinna wynosić ok. 47 μH. Usuń przy pomocy noża emalię z końców drutu, aby można je było pocynować, a następnie przylutować w miejscach wskazanych na PCB. Cewka jest utrzymywana w miejscu za pomocą dwóch zaciskowych opasek kablowych, przechodzących przez otwory w PCB, jak na **rysunku**.

Pojemnik na akumulator musi być zorientowany wg schematu (czerwony przewód do „+”) i przymocowany do PCB za pomocą dwóch wkrętów samogwintujących przechodzących

przez pojemnik do otworów w PCB. Skrót przewody od pojemnika, odizoluj ich końcówki i przylutuj je do PCB.

Włóż gniazda bezpieczników do otworów oznaczonych F1 tak, aby ograniczniki w klipsach były skierowane na zewnątrz. Przed przylutowaniem włóż bezpiecznik, aby gniazda były prawidłowo ustawione, dla dobrego kontaktu z bezpiecznikami.

Na koniec można zamontować nadajnik i odbiornik UHF. Muszą być one również prawidłowo zorientowane. Oznaczenia kołków są wydrukowane na płycie nadajnika. Ustaw kołki anten nadajnika i odbiornika tak, aby trafiły do odpowiednich pól antenowych na PCB.

Masz dwie możliwości wykonania anten: albo użyj 170-milimetrowych odcinków drutu zwinętego w spiralę o średnicy 3...5 mm (np. drutu sprężynowego) wewnątrz obudowy, albo – dla lepszego zasięgu (>40 m) – 170-milimetrowych odcinków sztywnego emaliowanego drutu miedzianego na zewnątrz obudowy.

Dodatkowe 5 mm długości podane na liście części ma na celu zapewnienie wystarczającej ilości drutu do przylutowania do zacisków anteny (w przypadku drutu spiralnego) lub do zagięcia na końcówce (w przypadku emaliowanego drutu miedzianego).

Po wybraniu typu drutu na antenę, którego chcesz użyć, utnij odpowiednie odcinki i przylutuj je do kołków PC anteny lub bezpośrednio do pól lutowniczych anteny, jeśli nie używasz kołków PC.

Zwróć uwagę, że będziesz musiał zdrapać trochę izolacji z końca emaliowanego drutu miedzianego (np. nożem), aby móc go pocynować, a następnie przylutować do płytki.

Montaż w pudełku

Montaż płytki w pudełku z tworzywa sztucznego nie wymaga wiele pracy. Wywierciliśmy otwór z boku dla wodoszczelnego przepustu kablowego wymaganego do okablowania panelu słonecznego.

Otwór ten znajduje się 25 mm w górę od zewnętrznej podstawy obudowy naprzeciwko CON1. Jeśli wymagany jest tylko zasięg transmisji UHF mniejszy niż 40 m, przewody antenowe można zagiąć wokół wewnętrznej obudowy skrzynki.

Dla maksymalnego zasięgu transmisji (do 200 m), sztywny przewód antenowy odbiornika powinien przechodzić przez mały otwór w górnej krawędzi, a przewód odbiornika analogicznie przez mały otwór w dolnej krawędzi obudowy.

Po przełożeniu przez nie drutów, zagnij ich końcówki, aby utworzyć małe, 3 mm, pętle, albo wklej małe plastikowe kulki.

Zapobiega to wydlubaniu sobie oka ostrym końcem anteny.

Użyj drutu DNE o średnicy 1 mm, aby był wystarczająco sztywny i pozostał prosty. Otwory wylotowe drutu w obudowie powinny być następnie uszczelnione za pomocą neutralnego, samoutwardzalnego uszczelnacza silikonowego.

Płytkę PCB wzmacniacza jest utrzymywana wewnątrz obudowy za pomocą śrubek M3, które wchodzą w zintegrowane gwintowane tuleje w podstawie.

Neoprenowa uszczelka pokryw musi być umieszczona wewnątrz rowka dolnej części obudowy, a następnie przycięta na wymiar. Styk pomiędzy końcami uszczelki powinien znajdować się na dolnej długiej krawędzi pokryw.

Etykieta

Aby wykonać etykietę na panel przedni, masz kilka możliwości.

W celu uzyskania trwałej etykiety, należy wykonać odbicie lustrzane grafiki i wydrukować je na przezroczystej folii do rzutników (używając folii odpowiedniej dla danego typu drukarki). W ten sposób tusz znajdzie się na tylnej stronie folii, gdy etykieta zostanie przyklejona. Przymocuj ją za pomocą przezroczystego uszczelnacza silikonowego.

Istnieją alternatywne rozwiązania, takie jak etykiety „Dataflex” i „Datapol” do stosowania w drukarkach atramentowych i laserowych – więcej informacji i bezpośrednie linki do tych produktów znajdziesz na stronie: www.siliconchip.com.au/Help/FrontPanels

Panel słoneczny lub zasilanie sieciowe

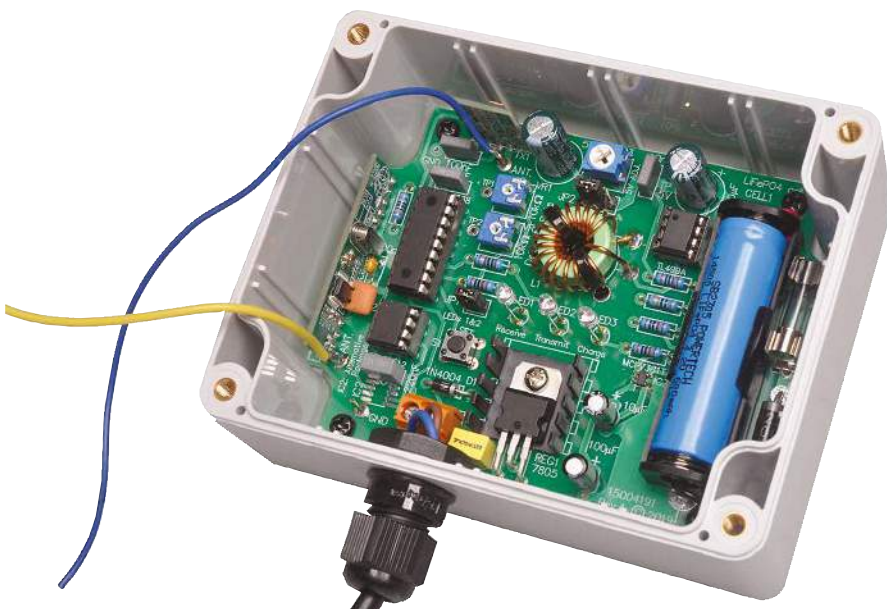
Do zasilania urządzenia użyliśmy panelu słonecznego 12 V 5 W. Panel 6 V byłby bardziej wydajny, ponieważ zmniejszamy napięcie do 5 V. Jednak panele 6 V nie są łatwe do znalezienia. Moc znamionowa panelu musi wynosić tylko 1 W.

Jeśli chcesz zasilac urządzenie z sieci, do CON1 można podłączyć zasilacz wtyczkowy 9 V.

Należy upewnić się, że zasilacz sieciowy nie jest narażony na działanie czynników atmosferycznych, a do retransmitera prowadzą tylko przewody niskiego napięcia.

W tym przypadku IC3 i akumulatory LiFePO4 nie są wymagane, chociaż można je pozostawić, aby urządzenie działało nawet podczas przerw w dostawie prądu (zakładając, że jednostki nadawcza i odbiorcza są również zasilane bateriami).

Jeśli zrezygnujesz z układu IC3, to możesz też pominąć F1, D2, LED3, a także układ IC4 i elementy z nim związane. Wyjście 5 V z REG1 mogłoby być użyte bezpośrednio do zasilania



Oto jak wygląda retransmitter zamontowany w wodoodpornej obudowie. Przewody w izolacji: niebieski i żółty to anteny: nadawcza i odbiorcza o długości po 170 mm – można je pozostawić luźno w obudowie, ale należy upewnić się, że nie mają odizolowanych końców, które mogłyby zewrzeć się z jakimiś elementami lub z płytką drukowaną.

układu poprzez podłączenie przewodu z wyjścia regulatora do zacisku 5 V na JP2.

Ustawienie

Ważne jest, aby na JP2 nie była umieszczona zworka, dopóki VR3 nie zostanie ustawiony na 5 V na wyjściu REG2. W tym celu należy włożyć akumulatory LiFePO4 do pojemnika i zmierzyć napięcie pomiędzy kołkami GND i TP5V. Wyreguluj VR3 tak, aby odczyt wynosił 5 V.

Instalacja

Retransmitter powinien być zamontowany w miejscu, które zapewni dobry odbiór oryginalnego sygnału UHF. Jeśli na JP1 jest założona zworka, wskaźniki LED (LED1 i LED2) zasygnalizują, czy sygnał jest odbierany i retransmitowany.

VR1 musi być tak wyregulowany, aby dioda odbiorcza nie migała w ogóle, a przynajmniej nie za często, gdy nie jest odbierany żaden sygnał. Ale jeśli zostanie skrócony za bardzo w lewo, retransmitter nie będzie działał, więc trzeba sprawdzić, czy jest w stanie przesyłać poprawne dane.

W tym celu należy początkowo ustawić VR1 do końca w prawo i nacisnąć S1, aby ustawienie VR1 zostało zaktualizowane. Następnie skrócić VR1 w lewo o kilka stopni i nacisnąć S1, aby ponownie zaktualizować ustawienie. Sprawdzić, czy retransmitter nadaje otrzymane sygnały.

Jeśli retransmitter działa prawidłowo, spróbuj dalszej regulacji w lewo. Ostateczna regulacja będzie kompromisem pomiędzy niezawodnym działaniem retransmitera i ignorowaniem szumów z odbiornika UHF.

Ustawienie VR1 zbyt mocno w lewo uniemożliwi poprawną pracę.

VR2 powinno być ustawione całkowicie w lewo, jeśli używasz pojedynczego retransmitera. Jeśli używasz wielu retransmiterów, ustaw VR2 na wszystkich retransmiterach do końca w prawo, co daje opóźnienie 12,5 s. Jeśli twój nadajnik może wysyłać sygnały w krótszych odstępach czasu, będziesz musiał poeksperymentować z maksymalnym obrotem VR2 w prawo, przy którym będą przekazywane wszystkie ważne pakiety.

Pamiętaj, że ustawienia potencjometrów nastawnych VR1 i VR2 są odczytywane przez IC1 tylko przy pierwszym włączeniu zasilania oraz po naciśnięciu S1. Po naciśnięciu przycisku S1 świecą diody LED1 i LED2, aby potwierdzić aktualizację ustawień.

Po zakończeniu regulacji VR1 i VR2 należy sprawdzić, czy odbiornik końcowy prawidłowo dekoduje kod przesyłany z retransmitera(ów). Jeśli nie, to być może trzeba będzie je przemieścić.

Następnie można zamontować na stałe retransmitter(y). W tym celu należy wykorzystać otwory montażowe przewidziane w obudowie.

Otwory te są dostępne po zdjęciu pokryw. Alternatywnie, można użyć wspornika i przymocować go do pudełka za pomocą otworów montażowych.

Unikaj wiercenia dodatkowych otworów w obudowie, ponieważ mogłoby to pogorszyć jej wodoszczelność. ■

John Clarke

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

Szkoła Konstruktorów



Rozwiązanie zadania głównego 316

Zaproponowany przez **Andrzeja Fryźlewicza** z Odrowąża, temat lipcowego zadania 316 brzmiał: **Zaproponuj układ elektryczny przydatny w ogrodzie.**

Interesujące rozwiązanie nadesłał **Jarosław Węgliński** z Warszawy. W e-mailu napisał (...) Przesyłam (...) **sterowanie zewnętrznym gniazdkiem zasilającym** przy użyciu transponderów RFID. Niestety przez brak czasu jest to wstępne rozwiązanie – planowana jest wymiana czytnika na obsługujący antenę zewnętrzną oraz przepisanie kodu na inny procesor niż ATmega (obecny znaczący wzrost cen procesorów Atmela może uzasadniać taką modyfikację). (...)

W opisie można przeczytać: (...) Co to jest działka rekreacyjna? Jest to teren zielony (oczywiście nie zawsze, jednak kwestia utrzymania go w stanie zieloności to temat na inny projekt, związany z automatycznym nawadnianiem), położony w oddaleniu od miejsca zamieszkania właściciela. Konsekwencją tego faktu są następujące:

- zieloność wymaga zabiegów porządkowych, w celu utrzymania jej na poziomie umożliwiającym swobodne przemieszczanie się,
- zabiegi porządkowe znacznie ułatwiają urządzenia techniczne (przykładowo – kosiarka do trawy),
- oddalenie od stałego miejsca zamieszkania naraża działkę na działalność przestępczą (w szczególności w zakresie kradzieży i dewastacji),
- każdorazowe przewożenie ze sobą wszystkich wartościowych rzeczy potrzebnych na działce jest niepraktyczne (zakładając skończoną pojemność bagażnika samochodowego).

Sytuacja jednak tylko pozornie jest beznadziejna – jednym z rozwiązań może być przechowywanie na działce urządzeń o niewielkiej wartości – mało atrakcyjnych dla złodzieja i w razie konieczności tanich w odkupieniu. W przypadku kosiarki do trawy

oznacza to urządzenie elektryczne. To z kolei wiąże się z koniecznością posiadania długich przedłużaczy – zajmujących miejsce w bagażniku lub będących atrakcyjnym łupem dla przestępcy, bądź instalacją gniazdek elektrycznych w pobliżu miejsc pracy kosiarki. Z kolei niezabezpieczone gniazdko nie dość, że narażają nas na kradzież prądu, to ułatwiają także prace ewentualnemu złodziejowi (który może użyć posiadanych elektronarzędzi celem pokonania innych zabezpieczeń zainstalowanych na terenie i tym samym zwiększenia wagi popełnionego przestępstwa).

Tak więc należy użyć gniazdko elektrycznego z mechanizmem kontroli dostępu, możliwego do instalacji na zewnątrz budynku i zabezpieczonego przed nieuprawnioną ingerencją. Pierwszym, narzucającym się rozwiązaniem są dostępne w handlu różnego rodzaju przekaźniki Wi-Fi. Najprostsze rozwiązanie – w postaci gniazdko sterowanego zdalnie należy od razu odrzucić – tutaj cała elektronika znajduje się w gniazdku, a ktoś zaopatrzony w śrubokręt może ominąć zabezpieczenie.

Drugim rozwiązaniem jest przekaźnik sterowany Wi-Fi, zainstalowany w pomieszczeniu i sterujący zewnętrznym gniazdkiem. Tu już jest lepiej, wyłączone gniazdko nie jest zasilane, więc trywialne próby ominięcia nie powiedzą się. Jednak należałoby pomyśleć o rodzaju obciążenia – typowa kosiarka to obciążenie typu indukcyjnego, więc bardziej wymagającego dla elementów stykowych (które często specyfikowane są na inne prądy maksymalne w przypadku obciążeń typowo rezystancyjnych, a dużo mniejsze dla obciążeń o charakterze indukcyjnym). Tak więc niezbędny będzie odpowiednio wyskalowany przekaźnik albo przekaźnik i dodatkowy stycznik.

Osobną kwestią jest wygoda sterowania – pomijając już przekaźniki sterowane przez system dostawcy zainstalowany w chmurze obliczeniowej (gdzie trzeba zapewnić

stabilne dołączenie do Internetu dla urządzeń na działce – co może być utrudnione i/lub generujące dodatkowe koszty), nawet sterowanie przez lokalną sieć bezprzewodową wymagać będzie noszenia ze sobą naładowanego smartfona.

Jeżeli nie chcemy być zmuszeni do używania smartfona celem włączenia kosiarki za garażem, trzeba wymyślić inny system sterowania. Możliwe jest użycie różnego rodzaju pilotów – czy to radiowych, czy na podczerwień, ale to niekoniecznie byłoby ułatwienie w stosunku do poprzedniego rozwiązania – zwykle potrzebny jest oddzielny pilot (integracja z posiadanym urządzeniem – np. od bramy garażowej może nie być trywialna).

Z tego względu zdecydowano się na użycie autentykacji przy zastosowaniu transponderów RFID – które często i tak nosi się ze sobą razem z kluczami, więc trudno o nich zapomnieć, nie wymagają baterii, a niewielki zasięg nie wydaje się być problemem, gdyż do gniazdko i tak trzeba podejść, żeby podłączyć przedłużacz.

Gdy już technologia została wybrana, pora zastanowić się nad szczegółami konstrukcyjnymi wynikającymi ze specyfiki projektu. Jak to wyżej wykazano, ważne jest aby układ był odporny nie tylko na warunki pogodowe, ale także próby dewastacji. W tym celu należy instalować go na ścianie budynku, gdzie większa część będzie znajdowała się wewnątrz, a tylko niezbędne elementy od strony zewnętrznej. Na pewno na zewnątrz należy zainstalować samo gniazdko, ale stycznik sterujący musi już być od strony wewnętrznej. Jeżeli chodzi o czytnik RFID to możliwe są różne warianty:

- jeżeli dysponujemy czytnikiem o dużym zasięgu to możliwe jest zainstalowanie go po stronie wewnętrznej i odczyt kart przez ścianę (jest to najbezpieczniejsze rozwiązanie jeżeli chodzi o zagrożenie wandalizmem),

- jeżeli nie, ale gdy czytnik ma antenę zewnętrzną, to na zewnątrz instalujemy tylko antenę (w obudowie odpornej na warunki atmosferyczne),
- w najgorszym przypadku na zewnątrz trzeba zainstalować cały moduł czytnika (dołączony do centralki za ścianą), jednak takie rozwiązanie naraża czytnik na uszkodzenie tak przez warunki pogodowe jak i celowe działanie osób trzecich.

Na rynku dostępne są bardzo tanie czytniki RFID oparte o kłony układu MFRC522 firmy NXP. Pracują one w paśmie 13,56 MHz i charakteryzują się zintegrowaną anteną i niewielkim zasięgiem pracy. Za to dostępna biblioteka do Arduino wygląda na dopracowaną, a nawet zawiera jako jeden z przykładów kompletną implementację zamka RFID z obsługą wielu kart i mechanizmem sterowania przy użyciu karty „master”, której dostosowanie do naszych potrzeb nie zajmie wiele czasu.

Niestety zintegrowana antena oraz niewielki zasięg pracy (dla posiadanego przez Autora modułu wynoszący poniżej 2 cm dla transponderów w formie breloka) wymusza albo instalację w płaskiej obudowie zewnętrznej albo dokonywanie przeróbek celem dodania zewnętrznej anteny. Wstępne eksperymenty z użyciem anteny zewnętrznej wykazały dużą czułość układu na parametry tejże anteny – już niewielka zmiana geometrii powodowała znaczące zmiany czułości (od około 3 cm do braku wykrywania transpondera z dowolnej odległości) a próba podłączenia anteny za pomocą dłuższego przewodu także zakończyła się brakiem komunikacji.

W związku z tym w kolejnym kroku zakupiono moduł czytnika z anteną zewnętrzną. Czytnik taki pracuje w paśmie 125 kHz i wraz z dostarczoną anteną zapewnia zasięg około 4 cm (dla transpondera w formie breloka). Taka konstrukcja wydaje się być bardziej dostosowana do planowanego zastosowania – na zewnątrz zainstalowana byłaby tylko antena.

W chwili obecnej projekt jest na etapie pierwszego prototypu – składa się z modułu

Arduino Uno z dołączonym czytnikiem RC522 i modułem przekaźnika. Arduino zaprogramowane zostało przykładem AccessControl dołączonym do biblioteki MFRC522-spi-i2c-uart-async (<https://github.com/miguelbalboa/rfid>), nieznacznie zmodyfikowanym na potrzeby bieżącego zastosowania. Zawiera elementy: D1 – diodę RGB, U1 – moduł Arduino Uno, U2 – moduł czytnika RFID RC522, U3 – moduł przekaźnika do Arduino.

Oryginalny przykład po przyłożeniu prawidłowej karty uruchamia przekaźnik na krótką chwilę (300 ms) i ponownie wyłącza, a sama operacja blokuje pracę programu. Jako że taki sposób działania nie jest odpowiedni do naszego zastosowania, modyfikacja polega na wyniesieniu stanu włączenia przekaźnika ponad główną pętlę programu. W ten sposób kolejne przyłożenia karty przełączają przekaźnik, który dodatkowo jest samoczynnie wyłączany po zadanim (długim – np. jednej godzinie) czasie, co ma zapobiec przypadkowemu pozostawieniu odblokowanego gniazdka.

Prototyp zmontowany jest zgodnie z proponowanym w użytym programie sposobem połączeń. Moduł RC522 podłączony jest przy użyciu SPI jak poniżej:

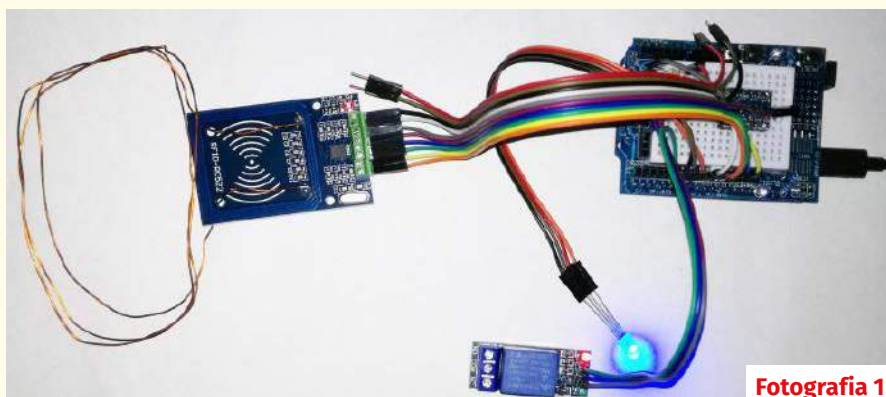
Sygnał	RC522	Arduino Uno
Reset	RST	9
SPI SS	SDA(SS)	10
SPI MOSI	MOSI	11
SPI MISO	MISO	12
SPI SCK	SCK	13

Osobną kwestią jest problem dostosowania poziomów napięć – moduł RC522 zasilany jest napięciem 3,3 V uzyskanym z Arduino i według dokumentacji NXP nie jest dozwolone bezpośrednio dołączanie go do elementów zasilanych napięciem 5 V. Jednak nie używamy oryginalnego układu tylko uproszczonej, dużo tańszej kopii (identyfikowanej przez bibliotekę jako „podróbka” – rozpoznawana poprzez odczyt wartości 0x12 jako wersji firmware

z modułu) i w wielu dostępnych w Internecie przykładach (przykładowo: <https://elportal.pl/kursy/arduino/1660-polaczenie-rfid-rc522-z-arduino>) widać bezpośrednie połączenie z Arduino Uno – jednak dla pewności warto zastosować tutaj konwerter poziomów 3,3 V ↔ 5 V (nawet jeżeli układ bez konwertera działa poprawnie, to takie połączenie może skutkować obniżoną niezawodnością modułu czytnika – a że w tym przypadku będzie on instalowany na zewnątrz, co implikuje zabezpieczoną obudowę i być może pokrycie lakierem lub zalewą – wymiana po ewentualnej awarii może nie być prosta). Dodatkowo do układu dołączony jest moduł przekaźnika (do końcówki 4 modułu Arduino) oraz trójkolorowa dioda LED sygnalizująca stan pracy (do końcówek 7, 6 i 5). W zależności od polaryzacji diody należy w kodzie włączyć lub zakomentować (dla diody ze wspólną katodą) definicję COMMON_ANODE. Dodatkowo do końcówki 3 można dołączyć przycisk zerowania – umożliwiający wykasowanie wszystkich zapamiętanych kart i wejście w tryb programowania nowej karty „master”.

Prowizoryczny model pokazany jest na **fotografii 1**.

Dobrze znany z łamów EdW Rafał Orodziński napisał tak: (...) Chciałbym pokazać dwa układy. Układ pokazany na **fotografii 2** to finalna wersja sterownika nawadniania. Układ umożliwia sterowanie nawadnianiem za pomocą sześciu niezależnych sekcji. Dla każdej z sekcji można ustawić maksymalny czas trwania nawadniania, zadaną wilgotność



Fotografia 1



Fotografia 2



Fotografia 3

i histerezę. Układ może pracować z dowolnymi czujnikami wilgotności z wyjściem prądowym 4...20 mA, do sterownika można wprowadzić za pomocą klawiatury równanie opisujące charakterystykę czujnika wilgotności w postaci wielomianu drugiego stopnia. Układ ma możliwość zdalnej blokady pracy za pomocą współpracującego z nim modułu włącznika internetowego działającego na zasadzie styku pozwolenia na pracę.

Układ pokazany na **fotografii 3** to „komputer ogrodniczy” – zostały do wykonania w nim już tylko drobne poprawki. Układ mierzy ciśnienie atmosferyczne, wilgotność powietrza, temperatury na zewnątrz i pod powierzchnią gruntu, nasłonecznienie oraz steruje pracą dwóch termostatów podgrzewających palmy uprawiane na zewnątrz (szorstkowiec i karlatka). Pierwsza z palm jest w stanie wytrzymać krótkotrwałe

spadki temperatur poniżej kilkunastu stopni Celsjusza, druga jest nieco mniej odporna, ale każdą z nich da się utrzymać w Polsce w ogrodzie za pomocą odpowiednich osłon i wspomagania ogrzewaniem. Dane pomiarowe rejestrowane są na karcie SD. Układ składa się z dwóch modułów zewnętrznego i wewnętrznego. Moduł zewnętrzny jest modułem wykonawczym sterującym pracą ogrzewania i zbierającym dane z czujników. Moduł wewnętrzny wizualizuje dane pomiarowe. Połączenie między modułami jest połączeniem bezprzewodowym. Moduł wyświetlacza zamocowany na module wykonawczym ułatwia testowanie układów (wersja robocza). Oba układy są autorstwa mego syna Mateusza Orodzińskiego a ja pełniłem funkcję doradczą oraz pomagałem rozwiązywać powstające ewentualnie problemy. Tyle o zadaniu 316.

Rozwiązanie zadania głównego 317

Temat sierpniowego zadania 317 brzmiał: **Zaproponuj układ elektroniczny w jakikolwiek sposób związany z pozyskiwaniem energii z otoczenia (energy harvesting).**

Temat jest naprawdę bardzo aktualny, ponieważ przeżywamy kryzys energetyczny i mnóstwo osób chciałoby się uniezależnić od klasycznych źródeł energii. Cel szczytny, jednak zbyt często wyobrażenia na temat takiego energetycznego uniezależnienia się okazują się mało realne.

Przede wszystkim należy przypomnieć, że według współczesnej wiedzy, **niemożliwe jest „wytworzenie energii z niczego”, natomiast możliwe są przemiany rozmaitych form energii.** To co nazywamy „pozyskiwaniem” czy też „wytworzeniem” energii, w rzeczywistości jest tylko jakimś rodzajem przemiany jednego rodzaju energii w inny. Najczęściej interesuje nas uzyskanie energii elektrycznej. Aby uzyskać dużą jej ilość, musimy dysponować odpowiednio dużą ilością innego rodzaju energii, którą przetworzymy na elektryczną. Czyli po pierwsze, trzeba mieć odpowiednio obfite źródło energii pierwotnej, po drugie trzeba tę energię pierwotną zamienić na energię elektryczną z jak najlepszą sprawnością, najlepiej 100-procentową. I z jednym, i z drugim jest kłopot.

Określenie **energy harvesting**, co dość zgrabnie można przetłumaczyć jako **pozyskiwanie energii**, dotyczy z reguły małych, a nawet bardzo małych ilości energii. Jednak zadanie główne 317 można śmiało rozciągnąć

na wszelkie sposoby pozyskiwania energii. Oto nadesłane rozwiązania.

Znany z wieloletniego udziału w Szkole Konstruktorów **Jacek Konieczny** z Poznania tym razem napisał: *Przesyłam propozycje dotyczące najpopularniejszych „odnawialnych” źródeł energii, tj. fotowoltaiki oraz wiatraków. Moje propozycje dotyczą zastosowania pewnego dodatkowego elementu włączanego przełącznikiem (lub stycznikiem), dlatego nie załączam żadnych szczególnych schematów. Zarówno panele fotowoltaiczne, jak i wiatraki mają pewną cechę wspólną. Jest nią wrażliwość na kapryśne zachowanie się naturalnych źródeł energii (...). W przypadku paneli fotowoltaicznych wystarczy, że panel ów zostanie chwilowo przestłonięty chmurką; wówczas od razu „siada” wartość prądu generowanego przez taki panel. Problemy z wiatrakami są bardziej rozbudowanej natury. Słynna już dokuczliwość wiatraków (dokuczliwość generowanego przez nie hałasu) dotyczy najczęściej spotykanych wiatraków „śmigłowych” o poziomej osi obrotu. Ale oprócz nich istnieją wiatraki o pionowej osi obrotu, tzw. **wiatraki VAWT** (najpopularniejsze wśród nich to tzw. **turbiny Savoniusa** oraz **turbiny Darrieusa**). Jednak powszechnie stosuje się tradycyjne wiatraki „śmigłowe”, ponieważ są one podobno sprawniejsze (...) Jednak wspomniane turbiny o pionowej osi obrotu mają inne zalety, które mogą kompensować ich mniejszą sprawność. Przede wszystkim turbiny (wiatraki) VAWT mogą*

pracować w dużo większym zakresie zmienności wiatrów niż wiatraki „śmigłowe”, tj. mogą pracować zarówno przy wiatrach bardzo słabych (przy których wiatraki „śmigłowe” w ogóle nie ruszają), jak i przy wiatrach bardzo silnych (przy których wiatraki „śmigłowe” wyłączają się i unieruchamiają, aby zapobiec ich uszkodzeniu). Zatem w ciągu np. roku wiatrak o osi pionowej może wyprodukować więcej prądu mimo niższej sprawności. (...) Proponowany przeze mnie wiatrak powinien mieć również poziomą oś obrotu. Nawet powinien przypominać tradycyjną turbinę „śmigłową”, ale taką, w której końcówki śmigieł umieszczone wewnątrz pierścienia wirującego w płaszczyźnie pionowej. W proponowanej przeze mnie konstrukcji wiatraka proponuję, aby tradycyjne śmigła zastąpić wirującymi szybko walcami. Innymi słowy, proponuję wykorzystać „efekt Magnusa” Siła Magnusa jest proporcjonalna zarówno do prędkości wiatru v opływającego wirnik (walec Magnusa), jak i do prędkości kątowej Ω , z jaką wiruje ów wirnik. Zatem w przypadku osłabienia „siły wiatru” wystarczy zwiększyć prędkość wirowania owych walców, aby zachować ten sam moment siły poruszający proponowany „wiatrak Magnusa”. W przypadku silnych wiatrów owe wirujące walce napędzane byłyby również wiatrem (np. rotorami Savoniusa). W przypadku słabych wiatrów przełącznik lub stycznik włączałby dodatkowe silniki elektryczne rozpędzające „walce Magnusa” do znacznie większych

prędkości wirowania. Czujnikiem mogłaby być nieduża prądnica prądu stałego (również napędzana pomocniczym wiatraczkiem), pełniąca funkcję prądnicy tachometrycznej i podłączona do wejścia komparatora okienkowego. Wyjście tego komparatora steruje (przez ewentualny wzmacniacz) przełącznikiem lub stycznikiem włączającym dodatkowy napęd elektryczny „walców Magnusa”.

Panel fotowoltaiczny. Proponuję wesprzeć panel fotowoltaiczny przez ogniwo termoelektryczne ogrzewane Słońcem (np. przez umieszczenie tego ogniwa w ognisku zwierciadła wklęsłego lub w ognisku soczewki Fresnela). Ogniwo termoelektryczne powinno być znacznie mniej wrażliwe na chwilowe przesłonięcie Słońca przez jakąś chmurkę (choćby ze względu na „pojemność cieplną” całego zestawu termoelektrycznego). Dlatego w przypadku chwilowego zasłonięcia Słońca przez chmurkę odpowiedni czujnik fotoelektryczny powinien uruchamiać przełącznik dołączający ogniwo termoelektryczne równolegle do pracującej baterii fotowoltaicznej. Ogniwo termoelektryczne ma znacznie mniejszą sprawność niż ogniwo fotoelektryczne, zatem przy tym samym napięciu znamionowym będzie wykazywać znacznie mniejszą wydajność prądową niż ogniwo fotoelektryczne. Zatem bezpośrednie dołączenie równolegle ogniwa termoelektrycznego do pracującego ogniwa fotoelektrycznego może nie być zdolne do skompensowania przesłonięcia tego ostatniego przez chmurę. Dlatego proponuję, aby ogniwo termoelektryczne stale ładowało baterię superkondensatorów. Wówczas awaryjne dołączenie ogniwa termoelektrycznego wraz z baterią superkondensatorów powinno skutecznie „buforować” krótkotrwały zanik natężenia prądu z ogniwa fotoelektrycznego.

Temat odnawialnych źródeł energii, a w szczególności „pozyskiwania darmowej energii” silnie rozbudza wyobraźnię i skłania

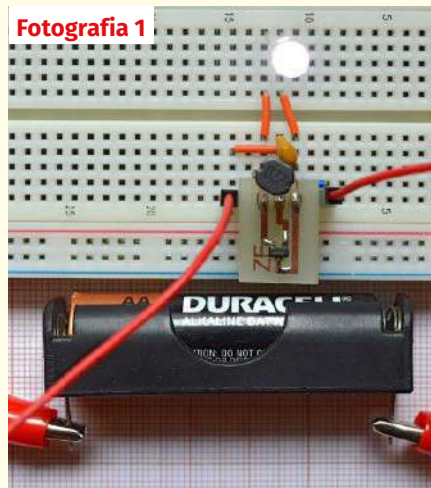
do poszukiwania rozmaitych pomysłów. Idea wspierania paneli fotowoltaicznych jest interesująca, ale nieprzypadkowo nie znajduje praktycznego zastosowania. W praktyce ilość pozyskiwanej energii należy zwiększać przez zwiększenie powierzchni paneli PV, a nie przez dodawanie modułów Peltiera. Po pierwsze energia pozyskiwana z modułów Peltiera okaże się w sumie wielokrotnie droższa od uzyskiwanej z paneli fotowoltaicznych. Termoelektryczne ogniwo Peltiera ma zupełnie inne właściwości niż panel fotowoltaiczny i musi być obsługiwane przez odpowiednio skonfigurowaną przetwornicę impulsową, a warunkiem pozyskiwania energii jest nie sama wysoka temperatura, tylko różnica temperatur obu stron modułu Peltiera.

Zygmunt Flisak z Opola poruszył atrakcyjny i ostatnio znów bardzo modny temat: Dzień dobry, popularne baterie AA, wbrew spadkowi popularności akumulatorów NiMH, nie odeszły jeszcze do lamusa. Nadal zasilają wiele urządzeń, takich jak zegary, piloty do odbiorników TV, wskaźniki laserowe itp. Akcesoria te kontrolują stan naładowania i odmawiają współpracy z bateriami, które można jeszcze z powodzeniem wykorzystać do zasilania źródeł światła z diodami LED. Przykładem takiego rozwiązania jest układ znany jako Joule thief.

Oryginalny układ bazuje na tranzystorze bipolarnym i obciążony jest szeregiem wad. Oscylator startuje przy napięciu przynajmniej 0,65 V, napięcie nasycenia tranzystora decyduje o stratach mocy, a ponadto przy odłączeniu diody świecącej (lub jej uszkodzeniu) pojawia się ryzyko zniszczenia tranzystora. Niektóre z tych wad udało się wyeliminować w scalonych przetwornicach do zasilania diod LED. Postanowiłem zbadać po tym kącie łatwo dostępny układ PR4401, umożliwiający budowę takiej przetwornicy z niewielkiej liczby elementów.

W prototypie wykonanym zgodnie z kartą katalogową zastosowałem

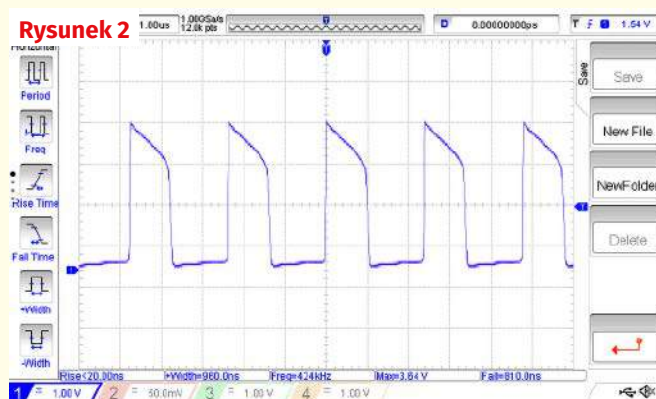
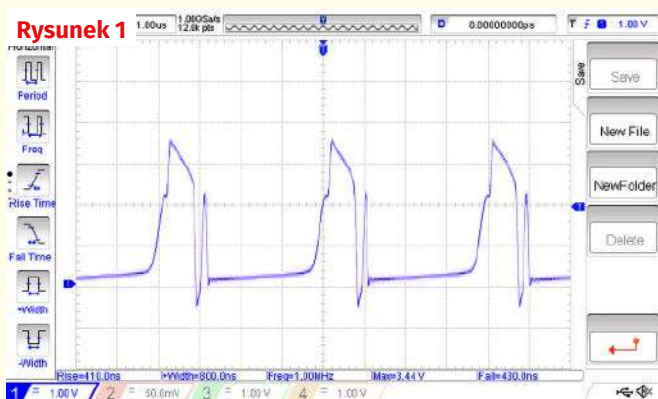
Fotografia 1



samodzielnie wykonaną przejściówkę umożliwiającą współpracę układu scalonego w obudowie SOT23 z płytką stykową, dławik do montażu powierzchniowego DLG-0504-220 o indukcyjności 22 μH i prądzie znamionowym 1,1 A z dolutowanymi wyprowadzeniami oraz kondensator ceramiczny o pojemności 1 μF do montażu przewlekane podłączony na wejściu (**fotografia 1**). Do testowania posłużyłem się zasilaczem laboratoryjnym. Pomiary wykazały, że mój egzemplarz przetwornicy uruchamia się przy napięciu wejściowym ok. 0,9 V, a drgania podtrzymywane są do minimalnego napięcia wynoszącego ok. 0,7 V. Dioda świeci wtedy słabo, a przebieg na jej wyprowadzeniach charakteryzuje się łagodnymi zboczami oraz oscylacjami (**rysunek 1**). Zwiększenie napięcia wejściowego do 1,3 V powoduje wzrost częstotliwości generowanych drgań i wyraźną poprawę jakości przebiegu (**rysunek 2**). Należy zwrócić uwagę, że wewnętrzny algorytm oscyloskopu błędnie obliczył częstotliwość przebiegu z rysunku 1; faktycznie wynosi ona ok. 250 kHz.

Najciekawszą część eksperymentu polegała na obserwacji przebiegu wyjściowego przy

REKLAMA



odłączonej diodzie LED i napięciu wejściowym wynoszącym 1,3 V. Okazało się, że tych warunkach przetwornica generuje wąskie impulsy o amplitudzie rzędu 19 V (**rysunek 3**). Można więc wnioskować, że układ posiada zabezpieczenie przed samozniszczeniem na skutek indukowania się wysokich napięć, którego pozbawiony jest oryginalny Joule thief.

Przy eksploatacji baterii aż do głębokiego rozładowania należy mieć na uwadze możliwość wycieku elektrolitu. Ponadto budowa przetwornicy impulsowej na płytce stykowej nie jest idealnym rozwiązaniem. Do tego typu doświadczeń przydatny byłby także układ MCP1643, lecz nie miałem do niego dostępu. Pozdrawiam

Temat układów typu Joule thief jest bardzo interesujący. Podjąłem już przygotowania do eksperymentów w tym zakresie. Najprostsze rozwiązania układowe z tranzystorem bipolarnym mają fatalnie małą sprawność i rzeczywiście minimalne napięcie zasilania to 0,65...0,7 V. Istnieją jednak bardzo proste sposoby realizacji podobnych układów o minimalnym napięciu zasilania rzędu 0,1...0,2 wolta, a nawet mniej. Będziemy się tym zajmować, ale nie tutaj.

A teraz kolejne interesujące rozwiązanie. **Andrzej Nowicki** z Warszawy na wstępie napisał: *Dzień dobry. Może trochę nie na temat, ale chciałbym dodać kilka uwag w kwestii magazynowania energii. Z energii odnawialną jest bowiem tak, że zwykle jest dostępna wtedy, gdy nie za bardzo jej potrzebujemy. Jakiś czas temu szacowaliśmy energię wyzwalaną przez ciężki wózek zjeżdżający z pagórka, i otrzymaliśmy całkiem spore wartości. Spróbuję policzyć, czy na tej zasadzie można zbudować mechaniczny „power bank” magazynujący znaczącą ilość energii.*

Przypuśćmy, że udało nam się w składnicy złomu zdobyć stary wagon kolejowy i jakieś 100...150 m toru kolejowego z podkładami, a obok naszego domu są jakieś niewielkie wzgórza, o wysokości, założmy, 20 m. Budujemy więc prostą (???) konstrukcję

– odcinek toru kolejowego, po którym przetaoczmy wagon z poziomu zerowego na szczyt 20-metrowego pagórka. Wagon wypełniamy kamieniami, złomem, gruzem i czym tam mamy, aby jego masa wyniosła 40 ton = 40000 kg.

Opisana konstrukcja może zmagazynować:

$$E = m \cdot h \cdot g$$

$$E = 40000 \text{ kg} \cdot 9,81 \frac{\text{m}}{\text{s}^2} \cdot 20 \text{ m} =$$

$$7848000 \frac{\text{m} \cdot \text{kg}}{\text{s}^2} = 7848 \text{ kJ} =$$

$$7848 \text{ kWs} = 2,18 \text{ kWh}$$

Porównam to z popularnym akumulatorem Li-ion typu 18650. Przeciętne parametry to:

- zgromadzony ładunek $Q = 2 \text{ A} \cdot \text{h}$
- średnie napięcie $U = 3,7 \text{ V}$

stąd zgromadzona energia to $E_{\text{akum}} = 7,4 \text{ Wh}$.

Dla porównania akumulatory kwasowo-ołowiowe: bezobsługowy akumulator AGM 100 Ah o napięciu znamionowym 12 V może zgromadzić $1200 \text{ Wh} = 1,2 \text{ kWh}$ przy masie 25 kg.

Potrzebowalibyśmy około 300 akumulatorów 18650 dla zmagazynowania takiej samej energii albo dwóch akumulatorów ołowiowych.

Oczywiście porównanie obu hipotetycznych rozwiązań jasno wskazuje, o jakim projekcie możemy myśleć w warunkach „domowych”, nawet uwzględniając potencjalne zagrożenie pożarowe ze strony ogniw Li-ion.

Z innych rozwiązań mechanicznych:

- raczej nikt nie ma warunków, aby zbudować mini elektrownię szczytowo-pompową

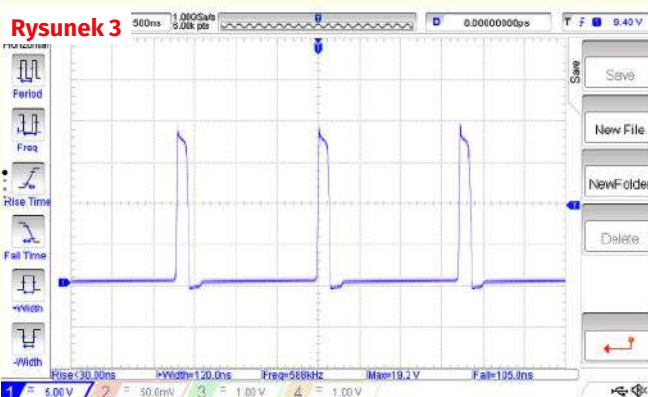
- instalacje ze sprężonym gazem są zbyt skomplikowane i niebezpieczne
- zostaje jeszcze magazynowanie energii w kole zamachowym. Rozwiązanie to zostało wdrożone praktycznie w „żyrobusech” – autobusach mogących przejechać dystans pomiędzy przystankami dzięki energii zgromadzonej w kole zamachowym. Podczas postoju na przystanku koło zamachowe było ponownie rozkręcane. Podaję link do pliku PDF z opisem takich konstrukcji: <http://bit.ly/3WYvzMb>.

Pokazane tam niewielkie koło zamachowe może zgromadzić około 500 kJ (139 kWh) energii.

Wniosek: Baterie akumulatorów pozostają w warunkach domowych praktycznie jedynym sposobem magazynowania energii pozyskiwanej ze „źródeł odnawialnych”. W przypadku akumulatorów Li-ion problemem jest jednak zapewnienie bezpieczeństwa przeciwpożarowego.

Temat pozyskiwania energii, „darmowej energii” budzi duże zainteresowanie, a nawet silne emocje. Warto interesować się tymi zagadnieniami. A jeżeli chodzi o praktyczne próby, nie należy się spodziewać spektakularnych sukcesów i niezależnienia od sieci energetycznej. Warto zaczynać od prób pozyskiwania małych ilości energii elektrycznej w niekonwencjonalny sposób, bo to też sprawia ogromną satysfakcję.

Piotr Górecki



Co tu nie gra? Rozwiązanie zadania 316

Na **rysunku A** pokazany jest zamieszczony w EdW 7/2022 schemat ocieplacza dźwięku na tranzystorze JFET z kanałem N.

Zadanie było łatwe, bo bez problemu można wskazać wymagany w tym konkursie jeden błąd. A błędów na schemacie jest kilka. Wiem, bo sam ten schemat rysowałem na potrzeby tego zadania – zawarłem na nim szereg usterek i dyskusyjnych rozwiązań.

Nie jest natomiast błędem sam pomysł zastosowania tranzystora polowego w „ocieplaczu dźwięku”. Jeden z uczestników niesłusznie stwierdził: *Chociaż obwody wejściowe są podobne do tych lampowych, to tranzystor JFET jest elementem półprzewodnikowym. Nie wprowadza on w sygnał audio nieparzystych harmonicznych, tak charakterystycznych dla lampowych ocieplaczy dźwięku.* Po pierwsze należałoby się zastanowić, czy chodzi o harmoniczne nieparzyste, czy może raczej parzyste, w szczególności drugą. Po drugie – powszechnie wiadomo, że tranzystory polowe pod wieloma względami mają właściwości podobne do lamp elektronowych.

Najbardziej dociekliwi mogą zainteresować się tematem „triodopodobnych” tranzystorów polowych SIT (Static Induction Transistor), spośród których wiele lat temu na rynku były dostępne najbardziej znane typy 2SK82 i 2SJ82. Do dziś można znaleźć oferty sprzedaży tranzystorów SIT nie tylko Sony i Tokin, ale też późniejszych, np. SemiSouth (FirstWatt). Tranzystory SIT mają prawie wszystkie charakterystyki podobne do triody. Natomiast popularne tranzystory JFET i MOSFET mają charakterystyki podobne raczej do pentody niż triody. Jednak charakterystyki wejściowe wszystkich tranzystorów polowych są podobne do (kwadratowych) charakterystyk wejściowych lamp próżniowych. Tematem lamp i ocieplaczy dźwięku będę się szeroko

zajmować, także od strony praktycznej, teoretycznej, ale nie tu.

A wracając do rysunku A trzeba zgodzić się z opiniami uczestników, którzy głośnym chórem stwierdzili, że głównym błędem jest to, że bramka tranzystora BF245 jest spolaryzowana dodatnio względem źródła. Złączone tranzystory polowe typu N wymagają polaryzacji zaporowej złącza bramka-źródło dla jego poprawnej pracy w zakresie liniowym i że tranzystory J-FET przy napięciu dodatnim na bramce względem źródła zachowują się jak zwykła dioda wpięta między bramkę a źródło, a poprawnie działają tylko przy napięciach ujemnych na bramce względem źródła.

Układ można mocno uprościć, ale jak słusznie zauważyło kilku uczestników, nie jest on całkowicie błędny i można byłoby go poprawić, dodając tylko jeden element: rezystor w obwodzie źródła, zaznaczony na **rysunku B** kolorem zielonym. Trzeba byłoby też sensowniej dobrać wartości elementów, o czym za chwilę.

Spośród drobnych, ale ważnych usterek, słusznie zgłosziliście błędną, odwrotną biegunowość elektrolitycznych kondensatorów wejściowego, a zwłaszcza wyjściowego. Słuszne wątpliwości budziły wartości wszystkich rezystorów. Jak napisał jeden z uczestników „zagwozdką” jest obecność kondensatora 10 nF. Rzeczywiście nie ma on sensu, niemniej rola podobnego kondensatora, tylko dającego filtr o dużo niższej częstotliwości jest uzasadniona, ale raczej tylko w wersji z MOSFET-em (klasycznym, wzbogacanym, a nie zubożanym), co jest tematem zadania konkursowego Policz316 i nie jest to zbieżność przypadkowa.

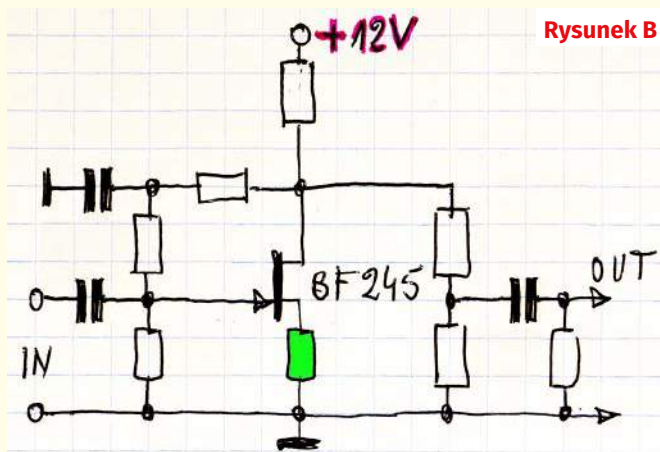
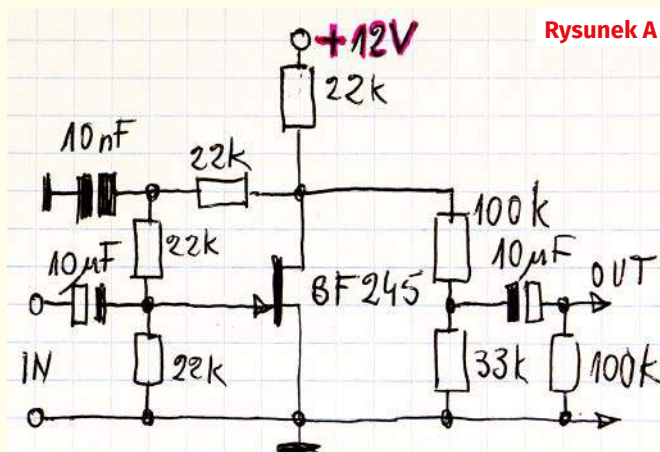
Dla wielu nieprzeniklioną zagadką był sens dodania wyjściowego dzielnika

napięcia. Jeden z młodych uczestników napisał: (...) poza tym dzielnik napięcia złożony z rezystorów 100 kΩ i 33 kΩ tylko obniża poziom napięcia wyjściowego 3-krotnie, co jest też niepotrzebne, chyba że zależy nam tylko na działaniu jako ocieplacz, a nie wzmacniacz.

W treści zadania była tylko króciutka informacja, że to ma być ocieplacz dźwięku. Bez żadnych szczegółów. Tymczasem temat lampowych i nielampowych ocieplaczy dźwięku jest bardzo szeroki i ma różne aspekty. Będziemy się nimi wspólnie zajmować – już mam przygotowany projekt do testowania układów audio. W „ocieplaczu dźwięku” wyjściowy dzielnik – tłumik może być jak najbardziej potrzebny, jeżeli ocieplacz ma być przystawką o wzmocnieniu 1× (0 dB). Tylko tu nasuwa się pytanie: czy tego tłumienia nie należałoby realizować inaczej – lepiej?

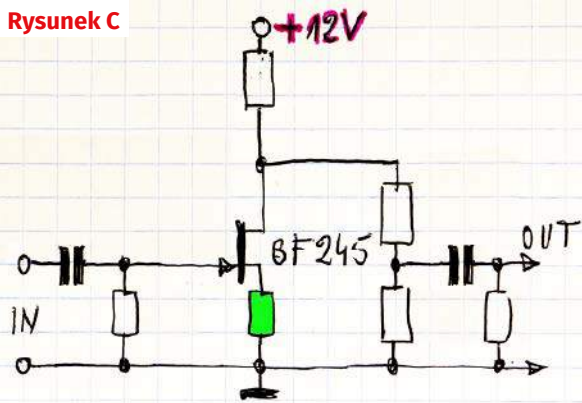
A tak w ogóle, to właśnie w przypadku ocieplaczy dźwięku, bazujących głównie na nieliniowości charakterystyki wejściowej, poziom sygnałów wejściowych i wyjściowych ma ogromne znaczenie – kwestii poziomów sygnału na pewno nie można pozostawić przypadkowi. To szeroki temat do oddzielnego omówienia.

Jeden ze stałych uczestników słusznie stwierdził: Schematów układów „pseudo-lampowych” na FET-ach jest w Internecie mnóstwo, więc każdy może sobie dobrać, co mu wygodnie. Przykładowa dyskusja: **Przedwzmacniacz „emulujący” lampowy dźwięk?** Niestety, w karcie katalogowej tranzystora BF245 są dość poważne braki, dlatego posłużę się charakterystyką mojego ulubionego BF861A. Charakterystyka (...) łopatologicznie tłumaczy, co to jest „napięcie odcięcia bramki”. Oczywiście „napięcie odcięcia bramki” jest dla tranzystora BF245B podane w notach katalogowych,





Rysunek C

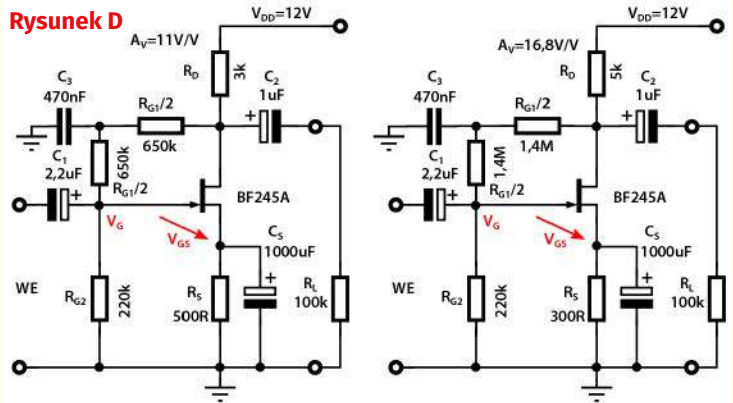


ale „dziwnym trafem” w niektórych brak przed nim znaku „-”. (...) podstawowy błąd proponowanego schematu: bramka próbuje pracować z napięciem dodatnim względem źródła, wymuszonym przez dzielnik rezystancyjny $4 \times 22 \text{ k}\Omega$ (...) względem potencjału źródła 0 V (na masie). Prawdłowo pomiędzy źródłem a masą powinien być rezystor ustalający polaryzację źródła na odpowiednim potencjale dodatnim (czyli inaczej potencjał ujemny bramki względem źródła – kłania się znajomość układów lampowych, gdzie w identyczny sposób polaryzuje się siatkę względem katody). Wartość tego rezystora powinna wynosić (...), zależnie od pożądanego prądu drenu. Natomiast bramka powinna być spolaryzowana na potencjale 0 V rezystorem pomiędzy bramką a masą. Wartość tego rezystora może być bardzo duża, zaczyna się zwykle od kilku $\text{M}\Omega$ i często sięga $\text{G}\Omega$. Z najlepszymi pozdrowieniami

Wspomniana znajomość układów lampowych prowadzi do uproszczonej wersji według **rysunku C**, a może jeszcze bardziej uproszczonej, takiej bez tłumika wyjściowego.

Inny stały uczestnik napisał: Przedstawiony układ jest konfiguracją układu wspólnego źródła. Tranzystor BF245 występuje w trzech grupach A, B, C. Brak jest na rysunku tego oznaczenia. Ma to znaczenie przy doborze punktu pracy. A nawet dla tej samej grupy występuje duży rozrzut parametrów, np. napięcia $V_{GS(OFF)}$, co skutkuje innym przebiegiem ch -ki przejściowej $I_D = f(U_{GS})$. Konfiguracja układu sugeruje, że ma to być układ ze sprzężeniem dren-bramka stabilizujący prąd DC drenu. (...) W obwodzie źródła tranzystora brak jest rezystora R_S . Tylko jego obecność zapewni nam właściwą (ujemną) polaryzację bramki. Równolegle do R_S możemy włączyć kondensator. (...) Zakładając pasmo przenoszenia 20 Hz...100 kHz, wartości kondensatorów sprzęgających mogą być mniejsze. Ale to możemy określić tylko na podstawie prawidłowego układu i prawidłowych wartości użytych rezystorów. (...) [na **rysunku D**]

Rysunek D



przedstawiam dwa schematy. Pominięcie kondensatora C_S skutkuje zmniejszeniem wzmacnienia.

Istnieją inne układy wzmacniaczy na JFET, ale nie zamieszczam ich rysunków bo wydaje mi się, że nie o to chodziło w zadaniu. A ich schematy są na tyle popularne, że „wszędzie” można znaleźć ich rysunki. Również nie analizuję przedstawionych układów. Analizę zostawiłem na zadanie Policz316, które jest tematycznie zbliżone. Dotyczy układu z tranzystorem MOSFET N. Pozdrawiam.

Kolejny częsty uczestnik konkursów napisał: Rysunek przedstawia schemat bufora z tranzystorem polowym złączowym (JFET). Charakterystyka tego elementu jest zbliżona kształtem do charakterystyki triody lub pentody. W typowej aplikacji przy napięciu na bramce równym zero względem źródła prąd dren-źródło osiąga maksymalną dopuszczalną wartość katalogową i maleje praktycznie do zera wraz ze spadkiem tego napięcia do wartości zwanej napięciem odcięcia. Złącze bramka-źródło jest więc cały czas spolaryzowane zaporowo, a prąd bramki (podobnie jak prąd siatki w lampie elektrowej) nie płynie.

I tutaj ujawnia się pierwsza wątpliwość co do prezentowanego układu: otóż dzielnik złożony z czterech rezystorów $22 \text{ k}\Omega$ ustalałby na bramce dodatnie napięcie o wartości $+3 \text{ V}$ względem źródła [co jednak zmieniłby przewodzący tranzystor]. Aby usunąć tę usterkę, pomiędzy źródło a masę należy dołączyć odpowiednio dobrany rezystor polaryzujący (w podobny sposób działa polaryzacja automatyczna w układach lampowych). Spadek napięcia na tym rezystorze podnosi potencjał źródła, co prowadzi do obniżenia potencjału bramki względem tego pierwszego. Wtedy rezystor znajdujący się pomiędzy drenem a biegunem dodatnim zasilania staje się zbędny, gdyż prąd dren-źródło ograniczony jest już przez rezystor polaryzujący. Warto w tym miejscu nadmienić, że eksperymentowano z przekraczaniem maksymalnej dopuszczalnej wartości prądu

dren-źródło poprzez polaryzowanie złącza bramka-źródło niewielkim napięciem dodatnim. Oczywiście w tych warunkach badany element przypomina już bardziej tranzystor bipolarny. A jak jest z subiektywnym odczuciem ocieplenia dźwięku? Link do artykułu: <http://bit.ly/3fZHOYi>.

To jeszcze nie wszystko. Przedstawiony układ nie wykorzystuje istotnej zalety tranzystora JFET, a mianowicie znacznej impedancji wejściowej. W przypadku całego bufora o jej wartości decydują dwa ostatnie rezystory wchodzące w skład wspomnianego powyżej dzielnika. Ich wartości na schemacie są zdecydowanie za małe: powinny być przynajmniej rzędu $1 \text{ M}\Omega$ (choć nadmierne ich zwiększanie może prowadzić do wzrostu szumów).

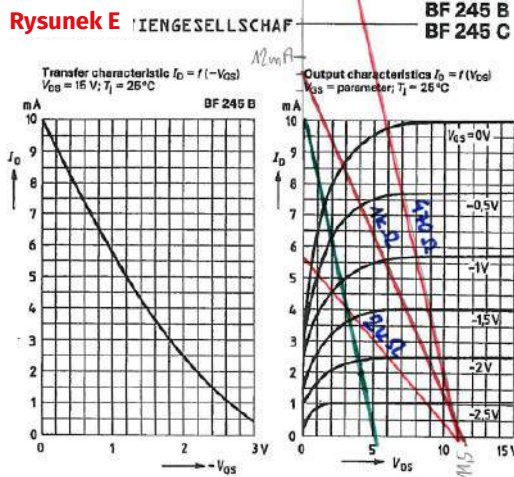
Można się zastanawiać, czy wartości rezystorów na wyjściu nie zwiększają niepotrzebnie impedancji wyjściowej układu. Ponadto polaryzacja obydwóch kondensatorów elektrolitycznych jest nieprawidłowa. Kondensator usuwający składową stałą na wejściu może mieć mniejszą pojemność i może być elementem foliowym bez polaryzacji. Należy jeszcze dodać, że tranzystor BF245 dostępny jest w trzech grupach (A, B i C) różniących się właściwościami, np. napięciem odcięcia i maksymalnym prądem dren-źródło. Na schemacie nie podano o którą grupę chodzi; dlatego nie można jednoznacznie określić wartości np. rezystora w obwodzie drenu. Na sam koniec można również wspomnieć, że w wielu JFET-ach małej mocy wyprowadzenia drenu i źródła są wzajemnie zamienne. Wtedy symbol tranzystora powinien być symetryczny (strzałka symbolizująca bramkę powinna stykać się z kanałem w połowie jego długości). Jednak uznanie tego za błąd w schemacie byłoby chyba nadużyciem. Pozdrawiam

Jeden z uczestników zadania napisał między innymi: (...) Dodatnie spolaryzowanie bramki BF245. W tym układzie na drenie odłoży się niewielkie napięcie stałe, a układ będzie praktycznie niewrażliwy na dodatnie połówki sygnału. Silnie ujemne połówki

przytają nieco tranzystor i spowodują niewielki wzrost napięcia na drenie – czyli układ zadziała jak prostownik. Żartem można uznać, że być może pomysłodawcy właśnie o to chodziło – „ocieplenie” dźwięku ma na celu wzbogacenie go w parzyste harmoniczne, a prostowanie sygnału to generacja – poza składową stałą – głównie drugiej harmonicznej ;)

(...) dzielnik wyjściowy nie jest krytycznym błędem, ale osobiście nigdy bym go tak nie zaprojektował. Takie rozwiązanie podnosi i tak już duży (rzędu kilkudziesięciu kiloomów) impedancję wyjściową – do ponad dwudziestu kiloomów. Jeżeli dzielnik ma tam być, to lepiej rezystor w obwodzie drenu podzielić na dwie części np. 5,1 kΩ (zasilanie) i 15 kΩ (dren) i z ich środka przez kondensator wyprowadzić sygnał na wyjście (rezystor 100 kΩ za kondensatorem zostaje). Przy takim rozwiązaniu tłumiąc sygnał wyjściowy zmniejszamy jednocześnie impedancję wyjściową (tu poniżej 5 kiloomów).

(...) Niejasne są dla mnie intencje zastosowania filtra dolnoprzepustowego w obwodzie ujemnego sprzężenia zwrotnego, które polaryzuje tranzystor (10 nF w połączeniu z 22 kΩ wpłynie na pasmo akustyczne). Inna sprawa że (podobnie jak cały układ!) filtr ten i tak by nie działał, bo duży kondensator wejściowy w połączeniu z (zapewne) niską impedancją wyjściową źródła skutecznie odfiltruje składową zmienną, która przejdzie przez ten filtr. Inną sprawą jest samo zastosowanie tego sprzężenia – czyżby chodziło o uproszczoną stabilizację punktu pracy tak jak robiło się to w przypadku tranzystorów bipolarnych? Tylko że w bipolarnych płynął prąd bazy, a tu prąd bramki płynąć nie powinien (...). W przypadku „ocieplacza” dźwięku dobór punktu pracy tranzystora i ewentualnych ujemnych sprzężeń jest kluczowy – zależy nam na tym, aby (przeciwnie do zwykłego wzmacniacza) był on nieliniowy

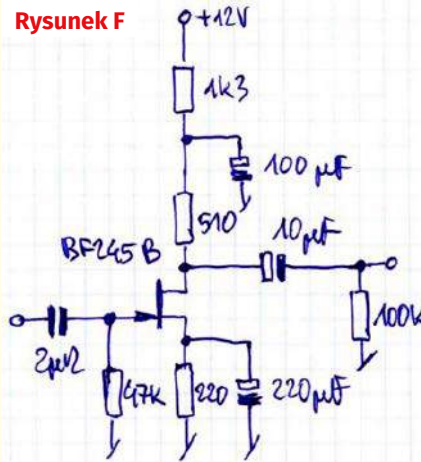


i to w pewien określony sposób – tak jak nieliniowy jest lampowy stopień wzmacniający.

Spójrzmy na charakterystyki wyjściowe BF245B – **rysunek E**. Obszar, gdzie prąd zależy od napięcia to obszar „triodowy” a na prawo od niego „pentodowy”. Wykreśliłem na nich proste dla różnych obciążeń i punktów pracy 470 Ω/–1,0 V, 1 kΩ/–0,5 V oraz 2 kΩ/–1,0 V. Widzimy że ze wzrostem rezystancji obciążenia i zmniejszaniem ujemnego przedpięcia bramki rośnie stopień „ocieplenia” i ma ono bardziej „triodowy” charakter. **Widać też że praca z obciążeniem w drenie aż 22 kΩ raczej nie ma sensu.**

Możliwości jest dużo i na coś się trzeba zdecydować biorąc pod uwagę „stopień ocieplenia”, spodziewany poziom sygnału wejściowego i pożądany wyjściowy. Nie obejdzie się zapewne bez prób odsłuchowych i wprowadzania na ich podstawie zmian.

Proponuję zacząć od układu pracującego wg. „zielonej” charakterystyki: 510 Ω/–1,0 V z obniżonym napięciem V_{DS} do 5 V – **rysunek F**. Dla tego układu dopuszczalny zakres napięć wejściowych to $\pm 1 \text{ Vpp}$, co wydaje się wystarczające (jak komuś mało, to niech zastosuje dzielnik na wejściu), pracuje on w obszarze triodowym i daje niewielkie wzmocnienie ok. 1,5 V/V (jak komuś za dużo, niech zastosuje dzielnik



wyjściowy – oczywiście taki jak opisałem wcześniej). Polaryzacja kondensatora wejściowego nie ma tu znaczenia, gdyż z obu stron jest na potencjale masy (zakładając że źródło sygnału nie wystawia składowej stałej). Ponieważ impedancja wejściowa jest stosunkowo duża, najlepiej zastosować kondensator niepolaryzowany. Można też spróbować pominąć kondensator odsprężający w źródle – wtedy sygnał na wyjściu trochę spadnie i spadnie również „stopień ocieplenia” (czytaj: zniekształcenia), co może okazać się korzystne, gdyż zniekształcenia tego układu dla sygnałów rzędu kilkuset miliwoltów będą już znaczne – może zbyt duże? (w mojej ocenie są porównywalne do wzmacniacza SE na triodzie). Jeżeli komuś bardzo zależy (...) to może wstawić tu kondensator rzędu 470 nF – jego działanie będzie zbliżone do domniemanego (czy spodziewanego/zakładanego?) działania kondensatora 10 nF woryginalie. Na koniec generalna uwaga – ponieważ nawet dla grupy BF245B prąd I_{DS} dla $V_{GS}=0 \text{ V}$ może się zmieniać od 6 do 15 mA, to dla posiadanego egzemplarza należy go najpierw zmierzyć, a następnie odpowiednio przeliczyć wartości rezystorów (np. mój projekt jest dla 10 mA). Pozdrawiam

Tyle o błędach z zadania NieGra316.

Co tu nie gra? Rozwiązanie zadania 317

Na **rysunku A** pokazany jest zamieszczony w EdW 8/2022 schemat ulepszonego przedwzmacniacza mikrofonowego najwyższej klasy. Aby uzyskać jak najlepsze właściwości, w układzie ma pracować legendarny, kosztowny ultraniskoszumny wzmacniacz AD797.

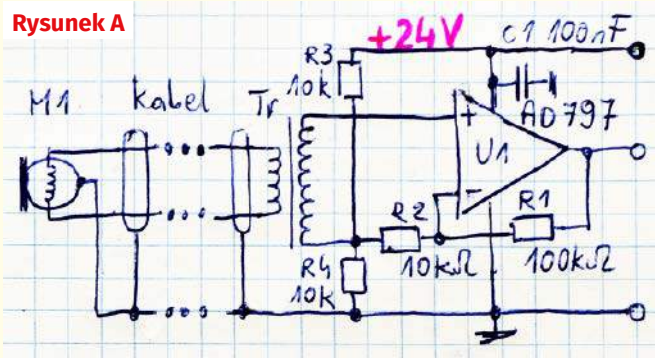
Zadanie było trudne, a wręcz bardzo trudne. Z premiedytacją narysowałem ten schemat, żeby zwrócić uwagę na podstawowy, bardzo słabo rozumiany problem. Celowo poprawiłem kardynalne błędy poprzedniego schematu z zadania NieGra312. Dlatego na pierwszy rzut oka schemat wygląda na prawidłowy

i „przyczepić” się można tylko do dwóch czy trzech szczegółów.

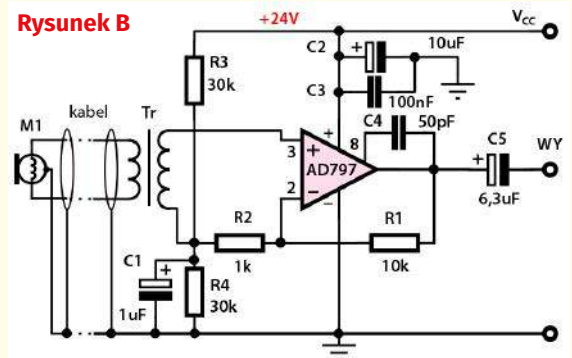
Jeden to fakt, że na wyjściu występuje duże napięcie stałe – brak bowiem obwodu RC odcinającego składową stałą. Można byłoby argumentować, że odpowiedni obwód musi występować na wejściu następnego stopnia,



Rysunek A



Rysunek B



ale rzeczywiście, jest to może nie błąd, ale ewidentna niedoróbka.

Najczęściej jednak jako błąd zgłasza się, że brak jest kondensatora lub kondensatorów równolegle do R4. Tu też można dyskutować, ponieważ w układzie z rysunku A zasadniczo dla sygnałów zmiennych rezystancja połączenia równoległego $R3||R4$ wynosząca 5 k Ω dodaje się do rezystancji R2 i trochę zmniejsza wzmocnienie. Ale sprawa jest bardziej skomplikowana, bowiem do tego dochodzi „pływające” uzwojenie wtórne transformatora. Między innymi właśnie dlatego, że jest „pływające”, niezależne od masy, teoretycznie żaden dodatkowy kondensator nie jest konieczny. Mamy tu sygnał użyteczny z „pływającego” uzwojenia wtórnego transformatora, co znacznie poprawia sytuację, a kultowy wzmacniacz AD797 ma znakomity współczynnik CMRR, więc i tu można byłoby dyskutować o szczegółach. Ale rzeczywiście, jeśli ma to być najwyższej klasy przedwzmacniacz, to takie rozwiązanie jest nie do przyjęcia, ponieważ pozwala przenikać do toru sygnałowego wszelkim śmieciom z szyny zasilania.

A przy okazji: z reguły najwyższej klasy wzmacniacze i przedwzmacniacze audio są (a przynajmniej dotąd były) zasilane symetryczne. Głównie właśnie dlatego, że nie ma wtedy obwodu sztucznej masy, która z różnych powodów bywa źródłem kłopotów, a w obwodzie prawdziwej masy nie płyną prądy zasilania, przez co obwód masy jest zdecydowanie mniej zaśmiecony. W rozwiązaniach zadania NieGra317, które do mnie dotarły, nie znalazłem wzmianki o tej kwestii. A jeżeli to ma być układ najwyższej klasy, to wręcz obowiązkowo trzeba zastosować zasilanie napięciem symetrycznym. Choć to też pole do dyskusji o różnych aspektach sprawy, brak zasilania symetrycznego można uznać za błąd. Ale nie główny błąd.

Właśnie z uwagi na fakt, że schemat wygląda na „prawie dobry”, pojawiły się też nietrafne odpowiedzi. Wróciła omówiona już w rozwiązaniu zadania 312 sprawa umieszczenia transformatora względem mikrofonu: (...) Błędem

jest włączenie transformatora dopasowującego za kablem, przez co niskie napięcie z mikrofonu są narażone na zakłócenia. Moim zdaniem, ten transformator powinien być włączony tuż za mikrofonem, jak to było pokazane w rozwiązaniu zadania NieGra312, przez co na wyjściu transformatora pojawi się wyższe napięcie, które mniej narażone na zakłócenia z kabla (...)

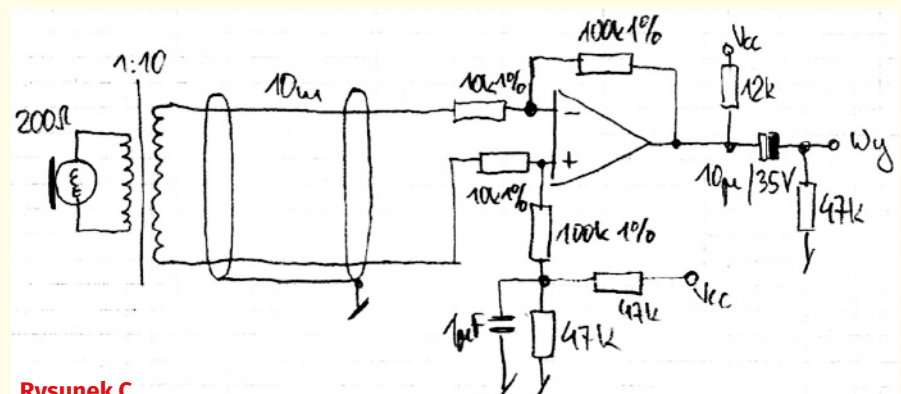
Nie będę powtarzał szczegółów, przypomnę tylko krótko: uzwojenie wtórne: wysokoohmowe i „wysokonapięciowe” nie może być obciążone dużą pojemnością długiego kabla. Transformator musi być umieszczony w pobliżu wejścia przedwzmacniacza.

Z uwagi na stopień trudności zadania pojawiły się rozwiązania częściowe, które mają zalety, ale nie do końca trafiają w sedno problemu – przykład na rysunku B.

Bardzo interesujące rozważania nadesłał jeden z uczestników: Oto moje spostrzeżenia: brak kondensatora odsprężającego dzielnik R3/R4. Tu sprawa jest bardziej złożona, bo teoretycznie jego rolę winno przejąć źródło zasilania (zawsze zakładamy że dla przebiegów zmiennych stanowi ono zwarcie). Teoretycznie więc brak tego kondensatora spowoduje tylko zmianę wzmocnienia napięciowego układu (...) Ponieważ jednak nic nie wiemy o wzajemnym rozmieszczeniu zasilacza i wzmacniacza, to zasilanie na płytce wzmacniacza powinno być „przysłabowane” kondensatorem. Co prawda mamy tu C1, ale jego wartość do tego celu

jest o wiele za mała (winien być minimum 10...22 μ F). Kondensator C1 (np. ceramiczny 100 nF) powinien zostać – producent zaleca takie „odsprężenie” łącznie z równoległym tantalowym 4,7 μ F lub większym z ewentualną niewielką szeregową rezystancją tłumiącą 1...4,7 Ω jak najbliższej pinów układu (mniej niż 5 mm). Należy też rozbudować układ wyjścia o kondensator blokujący składową stałą i rezystor nadający wyjściu potencjał masy. Pojemność tego kondensatora dobrać do spodziewanej impedancji obciążenia. Pozostawienie wyjścia na potencjale połowy zasilania z możliwością zwarcia składowej stałej jest nie do przyjęcia.

W rozwiązaniu NieGra312 napisał Pan: „...występuje kolejny poważny błąd, którego niestety nikt nie zgłosił – mianowicie występuje bardzo silne obciążenie pojemnościowe takiego wysokoimpedancyjnego źródła sygnału!”. Otóż ja zgłosiłem, tylko że moje rozwiązanie (...) pewnie nie dotarło do Pana. Odnosnie pojemności kabla napisałem wtedy: „ponieważ przewód jest dość długi oszacowałem, jaka może być jego pojemność. Na 2,5 metrowym odcinku mój pomiar wykazał sto kilkadziesiąt pikofaradów między żyłami. Na 10 metrach może to być kilkaset pikofaradów. W połączeniu z impedancją wyjściową transformatora da to filtr dolnoprzepustowy z częstotliwością graniczną w paśmie akustycznym. Raczej nie powinno być problemu



Rysunek C



(niższe szumy), uzyskamy stosując inny, znacznie tańszy wzmacniacz o małych szumach prawych – w praktyce z tranzystorami JFET na wejściach.

A jeżeli chodzi o pomysł usunięcia transformatora według rysunku D, to jak widać za przedstawionych rozważań, ma on sens! Ma, jednak jakby nie do końca. Otóż

najmniejsze możliwe szumy uzyskamy w wersji z transformatorem podwyższającym 1:10 i jakimś wzmacniaczem o małej gęstości szumów prądowych. I przy okazji dzięki obecności separującego transformatora zmniejszymy też problem brumu sieciowego.

To rzeczywiście słabo rozumiane kwestie. Przed wielu laty pisałem o tym w EP,

potem w EdW, przedstawiając kilka różnych wzmacniaczy mikrofonowych, w tym wersję z transformatorem i przedwzmacniaczem mikrofonowym LM381 (UL1321). Wygląda na to, że do tej ważnej kwestii szumów i parametrów szumowych trzeba wracać i wyjaśniać wchodzące tam w grę zagadnienia.

Policz – rozwiązanie zadania 316

W EdW 7/2022 przedstawione było zadanie **Policz316**, które brzmiało: *Chcemy przeprowadzić szereg testów różnych przedwzmacniaczy audio, żeby sprawdzić ich wpływ na dźwięk. Między innymi dla czystej ciekawości chcemy zrobić nietypowy przedwzmacniacz z jednym tylko tranzystorem wzmacniającym MOSFET N według rysunku A. Ogólnie mówi się, że tranzystory polowe mają charakterystyki inne niż tranzystory bipolarne, a podobne do charakterystyk lamp próżniowych. Dlatego chcemy się przekonać się, czy efekty będą podobne, jak w przypadku przedwzmacniacza lampowego. Możliwości konfiguracyjnych jest wiele. Układ może wzmacniać sygnał, ale równie dobrze wypadkowe wzmocnienie może być równe jedności – wtedy otrzymamy rodzaj „ocieplacza dźwięku”, podobnego do „ocieplaczy” lampowych, gdzie nie chodzi o wzmocnienie, a jedynie o modyfikację dźwięku. W układzie można zastosować rozmaite obwody pomocnicze, ale podstawowym elementem wzmacniającym, czy raczej „ocieplającym” ma być pojedynczy tranzystor MOSFET, bo to jego wpływ na dźwięk chcemy zbadać. W ramach zadania **Policz316** należy: **zapropo-***

nować schemat i wartości elementów przedwzmacniacza – ocieplacza.

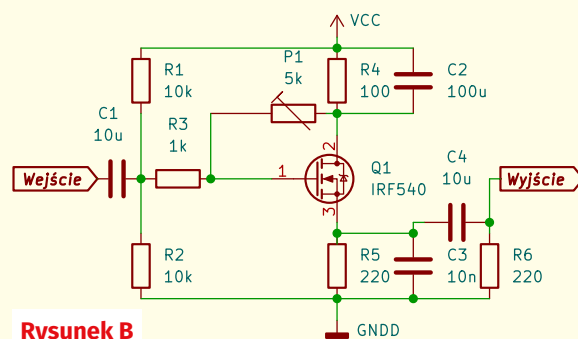
Pewien młody uczestnik przysłał schemat pokazany na **rysunku B** i napisał: *Chciałbym przedstawić moje rozwiązanie zadania Policz 316. Układ ten nie został przetestowany, jedyne testy działania tego układu odbyły się w symulatorze LTspice; niestety ten*

ocieplacz dźwięku ma wzmocnienie mniejsze od jedności.

Wykorzystanie symulacji w programie typu SPICE jest interesujące, i młodego uczestnika trzeba pochwalić za taką próbę! Pozwoli to wyeliminować ewentualne grube błędy układowe.

Jednak z kilku względów w tym przypadku symulacja ma znikome znaczenie praktyczne. Między innymi dlatego, że tranzystory polowe, nawet tego samego typu mają duży rozrzut parametrów wejściowych. Dlatego w praktyce punkt pracy „ocieplacza” trzeba będzie dobrać indywidualnie. Po drugie prosta symulacja nie pozwoli określić rzeczywistych, subiektywnie postrzeganych właściwości. Tu można byłoby jałowo dyskutować, czy ma sens symulacja polegająca na „przepuszczeniu przez schemat” przebiegu dźwiękowego w formacie WAV, a niektóre symulatory dają taką możliwość. Ogólnie biorąc, jeżeli to ma być ocieplacz, finalnie punkt pracy tranzystora trzeba dobrać eksperymentalnie.

Jeden ze stałych uczestników przysłał bardzo obszerne rozwiązanie (9 stron), zawierające mnóstwo wzorów i wyliczeń. Efektami tych wyliczeń są różne wersje ocieplaczy dźwięku na jednym tranzystorze MOSFET. Sześć zaproponowanych przez niego wersji pokazanych jest na **rysunkach C i D**. A oto drobne fragmenty opisu: (...) Na początku zrobimy kilka założeń. Napięcie zasilające przyjmę $V_{DD}=12\text{ V}$, w praktyce nie powinno być mniejsze niż 9 V. Pasma przenoszenia, nie mniej niż 20 Hz...100 kHz. Tranzystory MOSFET N pracują w obszarze nasycenia ch-ki $I_D=f(U_{DS})$. W nasyceniu prąd drenu jest w małym stopniu zależny od napięcia na drenie. Rezystancja (impedancja) źródła sygnału wejściowego w założeniu równa jest 0. W związku z tym prąd drenu możemy wyrazić zależnością



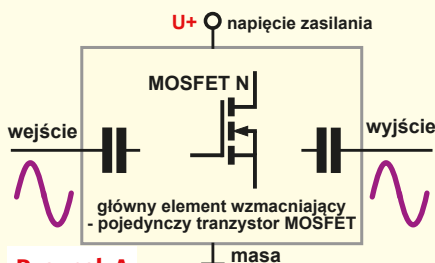
Rysunek B

$I_D = K_N (V_{GS} - V_{GS(TH)})^2$
wsp. K_N jest to parametr transkonduktancyjny tranzystora MOS zależny od wymiarów geometrycznych obszaru kanału i ruchliwości nośników w kanale. Wyznamy go z pomocą zależności

$$K_N = \frac{I_{D(ON)}}{(V_{GS(ON)} - V_{GS(TH)})^2}$$

Potrzebne dane do wyznaczenia tego wsp. znajdziemy w nocie aplikacyjnej tranzystora MOSFET N (nie wszystkie noty podają te wartości). Wybrałem tranzystor 2N7000, którego specyfikacja zawiera potrzebne dane. (...) Ze względu na rozrzut parametrów tranzystorów rozważania będą przybliżone. Ale zbliżymy nas do układu praktycznego. Od czegoś trzeba zacząć aby zbudować układ praktyczny. (...) [Dla pierwszego układu] (...) w konfiguracji wspólnego drenu (wtórnik źródłowy) (...) obliczenia potwierdzają wzmocnienie układu równe ≈ 1 . (...) [Inne układy mogą mieć] (...) wzmocnienie większe niż 1. Jakie jest to wzmocnienie, spróbuję to obliczyć. (...) Wynika z tego, że jeżeli nie potrzebujemy dużego wzmocnienia, a jedynie chodzi o modyfikację dźwięku, to możemy rozważyć, aby nie stosować kondensatora C_S . Zmieni się charakterystyka amplitudowa tj. zwiększy się dolna częstotliwość f_d pasma 3 dB. (...)

Ostateczny wybór układu będzie uzależniony od oczekiwanej wartości sygnału wyjściowego, z czym wiąże się jego wzmocnienie,



Rysunek A



impedancji wejściowej oraz parametrów szumowych, zniekształceń.

Przy okazji trzeba powiedzieć, że wybór tranzystora MOSFET w stopniach małosygnałowych to nienajlepszy wybór (sprawdzają się w stopniach końcowych). Odradza się ich stosowanie do wzmacniania najmniejszych sygnałów. Najlepiej stosować JFET-y. W zakresie niskich częstotliwości MOS-y mocno szumią. Zdecydowanie mniej FET-y, a najmniej bipolarne. [Dotyczy to także lamp i wzmacniania bardzo małych sygnałów, a problem rozwiązuje praca przy silnych sygnałach.] Na rysunkach i w obliczeniach zastosowałem 2N7000, który użyty został na potrzeby obliczeniowe.

Aby rysunki i obliczenia nie były anonimowe, a dotyczyły rzeczywistego tranzystora z uwzględnieniem noty aplikacyjnej na rysunkach i w obliczeniach podałem jego symbol. Jeżeli miałbym wybrać do praktycznego zastosowania, to wybrałbym tranzystor np. Si2312. Ma zdecydowanie mniejsze szumy niż 2N7000. (...)

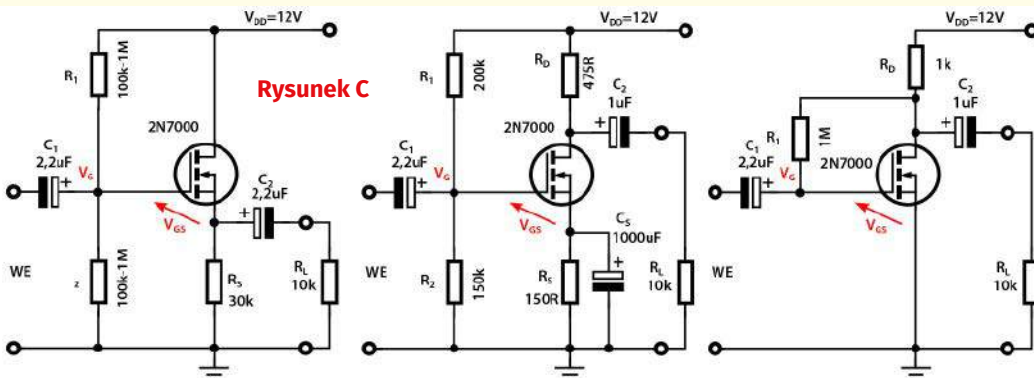
Jeden ze stałych uczestników zaczął tak: Dzień dobry. Niestety, policzenie tego zadania jest dla mnie nieco za trudne. Ale proponuję taki układ znaleziony na stronie BOSOZ PRE: (trzeba tylko poprawić błąd – C103 polipropylenowy ma 10 μ F), lista części dla wersji stereo jest tu: <http://www.pcb-audio.com/bosoz/bom-boz.txt> oraz pokaz montażu przedwzmacniacza (wersja stereo) na YT <https://www.youtube.com/watch?v=vWzcUrUbF1c>.

Sądząc z zastosowanych elementów oraz napięcia zasilania, jest to układ

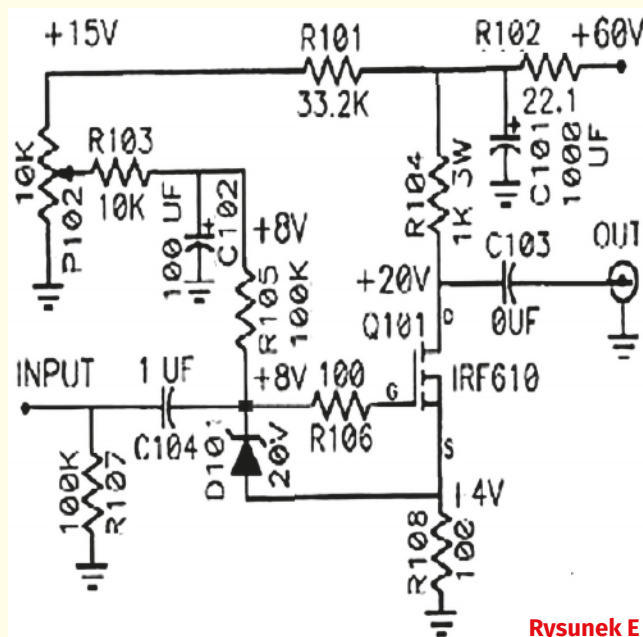
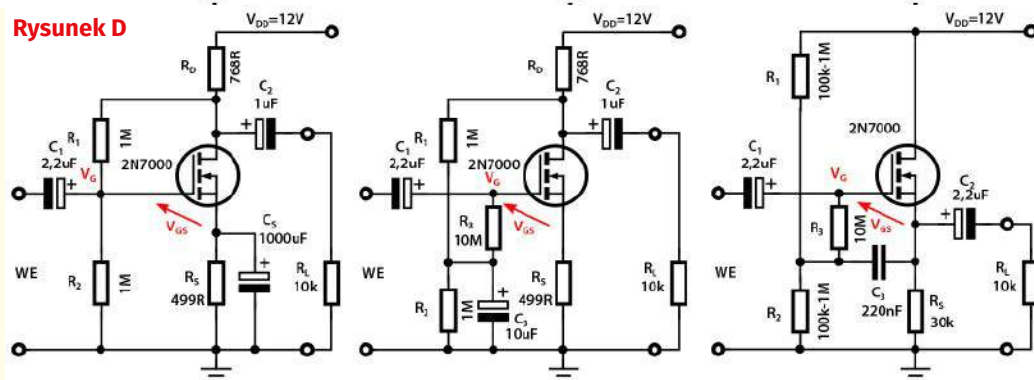
bardzo WYSOKIEJ klasy. Oczywiście, nieco majstrując przy rezystorach R102, R104 i R108, można nieco obniżyć napięcie zasilania...

Osoby zainteresowane tematem mogą porównać rysunek E – schemat ze strony http://www.pcb-audio.com/bosoz/brideofzen_1.png by dopatrzeć się (nieprzypadkowych) podobieństw i różnic z rysunkiem A z konkursu NieGra316.

Rysunek C



Rysunek D



Rysunek E

Policz – rozwiązanie zadania 317

WE dW 8/2022 przedstawione było zadanie Policz317, które brzmiało: Zachęceniem rozwiązaniem zadania głównego 312, w ramach zadania Policz317 chcemy wykorzystać kostkę nomen omen LM317 lub jej odpowiednik (LM350, LM338 lub pokrewne LDO) i z jej pomocą zrobić prosty regulowany zasilacz według rysunku A. Jego lustrzana wersja z kostką

LM337 lub podobną będzie analogicznym zasilaczem napięcia ujemnego. W ramach zadania Policz317 należy: zaproponować wartości elementów układu. Warto zacząć od transformatora. Niektórzy zechcą wykorzystać jakiś transformator już posiadany lub pochodzący z odzysku, a inni być może zaproponują zakup klasycznego transformatora. W każdym razie

najważniejsze są moc i napięcie wyjściowe transformatora. Z parametrami mostka prostowniczego i kondensatorami kłopotu nie będzie. Pewną trudnością jest dobór rezystora i potencjometru. W rozwiązaniu należałoby też podać, jaki trzeba zastosować radiator. Bez radiatora moc strat układów w obudowie TO-220 to jedynie żałosne 2 waty, które nie



pozwolą pracować przy większych prądach. Osoby bardziej zorientowane, jeśli tylko chcą, mogą ewentualnie dodać jakieś elementy lub obwody pomocnicze, np.: bezpiecznik, kontrolkę, wstępne obciążenie, ogranicznik prądu lub wskaźnik braku stabilizacji.

Zadanie wbrew pozorom było w sumie dość trudne. Trudno bowiem, a wręcz niemożliwe jest obliczenie wszystkich parametrów i wartości elementów „na sucho”, bez przeprowadzenia prób praktycznych. Szczególnie dotyczy to transformatora. Niemniej warto przeprowadzić obliczenia, by wstępnie zastosować wyliczone wartości elementów, a potem po testach modelu w razie potrzeby wprowadzić poprawki.

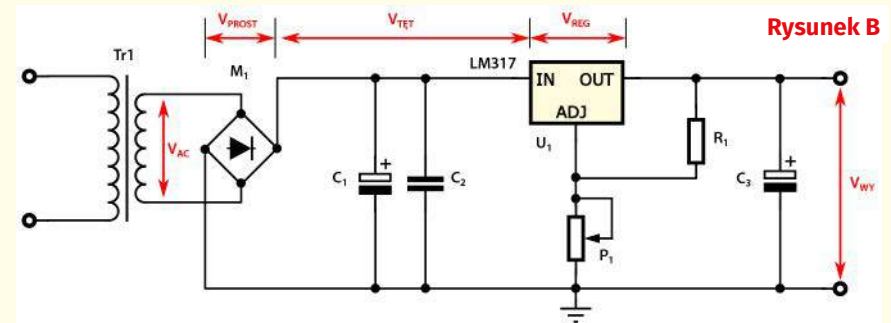
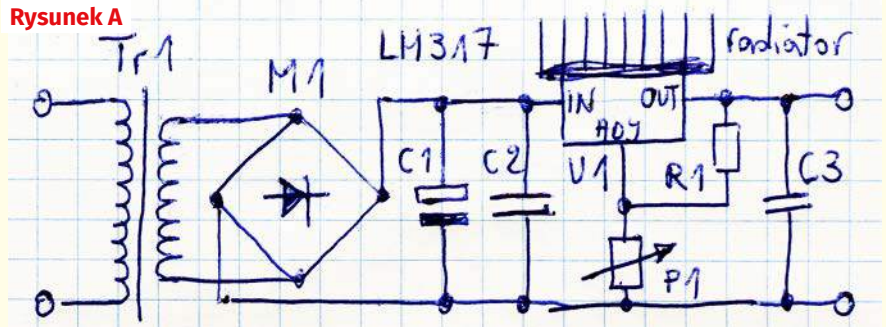
Jeden z uczestników napisał tak: (...) Rozpocznę od założeń: napięcie wyjściowe regulowane w przedziale $2 \text{ V} \dots 16 \text{ V} \pm 0,5 \text{ V}$, maksymalny prąd obciążenia $I_{\text{max}} 1,25 \text{ A}$. Tak jak Pan proponował, rozpocząłem od transformatora: **TR1** 230 V/24 V 30 VA, **M1** mostek prostowniczy 2W10M ($V_{\text{rms}} = 700 \text{ V}$, $I_{\text{max}} 2 \text{ A}$ i $V_{\text{diody}} = 1 \text{ V}$). Na **C1** po spadku napięcia na **M1** (~1 V) powinno być dostępne $V_{\text{max}} = 23 \text{ V}$, ustaliłem rozpiętość tętnienia na poziomie 4 V (Vripple) aby zagwarantować układowi LM317 $V_{\text{in min}} 19 \text{ V}$. Pojemność kondensatora policzyłem w następujący sposób $C = I_{\text{max}} / (2 * f * V_{\text{ripple}})$ co daje $1,25 / (2 * 50 \text{ Hz} * 4 \text{ V}) = 0,003125$, w związku z czym proponuję kondensator elektrolityczny 3300 μF 50 V. Do tego **C2** 0,1 μF , **C3** 1 μF .

Wiedząc, że LM317 definiuje V_{out} jako $1,25 \text{ V}(1 + P1/R1) + I_{\text{adj}} * P1$ oraz że I_{adj} to maksymalnie 100 μA (info z noty) a z V_{out} LM317 chciałbym uzyskać max 16 V. Można policzyć, że $16 \text{ V} = 1,25 \text{ V}(1 + P1/R1) \Rightarrow P1/R1 = 11,8$. Czyli jeśli wybiorę $P1 = 5 \text{ k}\Omega$, to $R1$ powinno wynosić $5 \text{ k}\Omega / 11,8 \sim 424 \Omega$, dlatego: **R1** 430R, **P1** 5 k Ω . Znając wartości rezystorów ostatecznie $V_{\text{out}} = 16 \text{ V} + I_{\text{adj}} * P1 = 16 \text{ V} + 100 \mu\text{A} * 5 \text{ k}\Omega = 16 \text{ V} + 0,5 \text{ V}$. **LM317 radiator** – tutaj przyznam się bez bicia, że nie bardzo wiem, jak zabrać się za policzenie jaki radiator jest potrzebny. Zakładając, że maksymalna moc to $23 \text{ V} * 1,25 \text{ A} = 28,75 \text{ W}$ Wzajemnie nie będzie on mały. Pozdrawiam

S ł u s z n i e !

Najtrudniejsze byłoby obliczenie radiatora! W przypadku LM317 praca przy mocy strat prawie 30 watów jest praktycznie niemożliwa. W skrajnych warunkach stabilizator nagrzej się do temperatury ponad 150 stopni.

Rysunek A

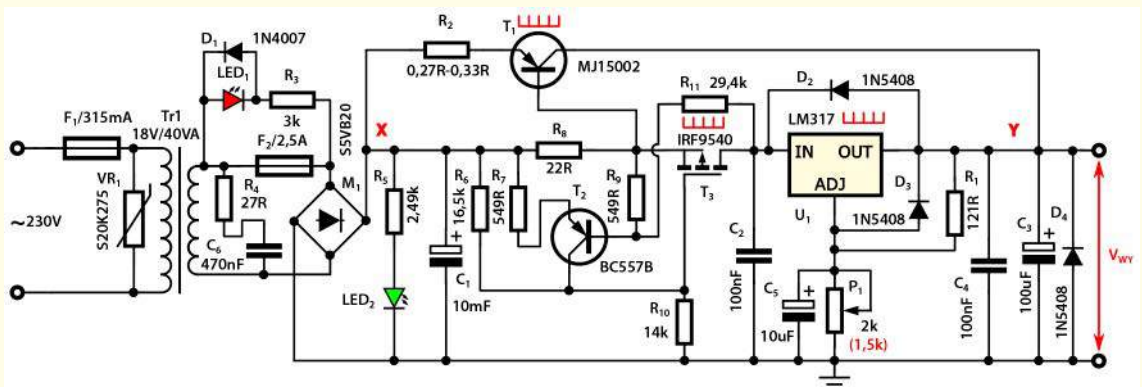


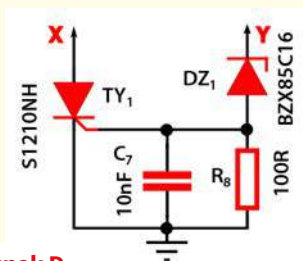
Rysunek B

Nie ulegnie uszkodzeniu, ale zabezpieczenie termiczne obniży napięcie i prąd wyjściowy, co przy rozmaitych eksperymentach może być dodatkowym powodem kłopotów.

Jeden ze stałych uczestników przysłał bardzo obszerne rozwiązanie. Oto drobne fragmenty. (...) Na początku zaczniemy od oszacowania (doboru) transformatora gdyż będzie to miało wpływ na dobór pozostałych elementów. Żle dobrany transformator zepsuje nam całą konstrukcję. Dotyczy to zarówno projektowania nowego transformatora jak również wykorzystania posiadanego, z odzysku. (...) [Pomocą może być rysunek A.] (...) dla następujących założeń: $V_{\text{WY}} = 15 \text{ V}$ – maksymalne napięcie wyjściowe, $I_{\text{WY}} = 1 \text{ A}$ – maksymalny prąd wyjściowy, $V_{\text{REG}} = 3 \text{ V}$ – minimalne napięcie na LM317, $V_{\text{PROST}} = 2,0 \text{ V}$ – spadek napięcia na mostku prostowniczym, można to określić z noty katalogowej dla danego prądu, $V_{\text{TET}} = 0,5 \text{ V}$ ($V_{\text{PP}} = 1 \text{ V}$) – napięcie tętnień za mostkiem prostowniczym. (...) Prąd obciążający transformator: $I_{\text{AC}} = 1,8 \times 1 \text{ A} = 1,8 \text{ A}$. Moc transformatora: $P = 17,5 \times 1,8 = 31,5 \text{ VA}$ (...)

[dla] niskiego napięcia sieci tj. 207 V. Należy też uwzględnić sytuację kiedy napięcie sieci może wzrosnąć do maksymalnej dopuszczalnej wartości $230 + 10\% V_N = 253 \text{ V}$. (...) LM317 może pracować w tych warunkach napięciowych. **Kondensatory** muszą być dobrane na napięcie nie niższe niż 35 V. Moc pobierana z transformatora $P = 19,25 \times 1,8 = 34,65 \text{ VA}$. **Należy wybrać transformator o mocy 40 VA, prąd uzwojenia wtórnego 1,8 A, napięcie znamionowe uzwojenia wtórnego 18 V RMS (dla napięcia sieci 230 V).** (...) **Kondensator filtru C_1** dla założonego napięcia tętnień $V_{\text{PP}} = 1 \text{ V}$ i prądu 1 A: $C = (I_{\text{DC}} \times t) / V_{\text{PP}} = (1 \times 10 \times 10^{-3}) / 1 = 10000 \mu\text{F} = 10 \text{ mF}$. (...) Dla założonej wartości napięcia wyjściowego $V_{\text{WY}} = 15 \text{ V}$ przyjmujemy wartości $R_1 = 121 \Omega$ i $P_1 = 1,5 \text{ k}\Omega$ (są takie wieloobrotowe), a gdyby był problem z uzyskaniem o tej wartości to 2 k Ω . Aby układ LM317 pracował poprawnie dla $I_{\text{WY}} = 0$ należy dobrać odpowiednio rezystor R_1 , wartość którego zagwarantuje wymagane minimalne wstępne obciążenie $I = 10 \text{ mA}$. Rezystancja R_1




Rysunek D

nie powinna być większa niż 125 Ω. (...) Moc maksymalna wydzielona wynosi: $P = (U_{INMAX} - 1,25) \times I_{DC} = (27,14 - 1,25) \times 1 = 25,89 \times 1 = 25,89 \text{ W}$ (...) ile musi wynosić rezystancja termiczna złącze – otoczenie w warunkach stosowania radiatora: $R_{thja} = \Delta T / P = (125 - 35) / 25,89 = 3,48^\circ\text{C/W}$. Wartość ta jest mniejsza od rezystancji termicznej $R_{thjc} = 5^\circ\text{C/W}$ (złącze – obudowa) podanej w nocie katalogowej (...) takiej mocy nie odprowadzimy bezpiecznie z układu. [Rozwiązaniem może być] (...) wariant z dodatkowym tranzystorem mocy (...) aby ciężar obciążenia stratami mocy przenieść na dodatkowy tranzystor T_1 . (...) to niestety układ nie posiada zabezpieczenia (...) [nadprądowego i termicznego. Dodałem] (...) układ z ogranicznikiem. (...) Inspiracją (...) było prawdopodobnie zadanie, o ile dobrze pamiętam, ze Szkoły Konstruktorów (...)

Układ z **rysunku C** możemy uzupełnić o zabezpieczenie nadnapięciowe przedstawione na **rysunku D**, które włączymy w punktach oznaczonych X-Y. Zabezpieczenie w chwili przekroczenia ustalonego napięcia wyjściowego ma spowodować przepalenie bezpiecznika. (...) Układ zabezpieczenia nadnapięciowego może być załączany jako opcja albo włączony na stałe. (...) układ zabezpieczenia nadnapięciowego można wykonać na inne napięcia albo go wyłączyć przez rozłączenie w punkcie Y. Jest to rozwiązanie inwazyjne: duży prąd płynący przez tyristor, przepalenie bezpiecznika. Poszukamy czegoś

innego. (...) Przedstawię jeszcze modyfikację (...) części dotyczącej rozwiązania zabezpieczenia nadnapięciowego. Rozwiązanie z tyristorem jest (...) zbyt inwazyjne. (...) **Rysunek E** przedstawia układ zmodyfikowany. Do takiego rozwiązania zainspirował mnie układ ze strony: <http://www.circuit-diagrams.net/page/overvoltage-protection-for-the-lm317>

Został on zmodyfikowany na potrzeby rozwiązania zadania, przedstawionego przeze mnie przykładowego układu. (...) **dobór radiatorów do tranzystorów T_1 , T_3 , T_5 , i LM317.** (...) Wyznamy wymaganą wartość rezystancji termicznej radiator-otoczenie R_{thra} (...) $R_{thra} = 4 - (0,3 + 0,875) = 2,825^\circ\text{C/W}$. Znajdę tę wartość poszukamy radiatora o takiej wartości rezystancji. Nie jest to łatwe, gdyż w większości przypadków oferenci nie podają tej wartości w powiązaniu z długością radiatora. Nasze warunki powinien spełnić radiator o symbolu **A 4129**. Podobne do niego są radiatorzy o symbolach **SK03** i **SK39** firmy **Fischer Elektronik**, które znajdziemy w katalogu tej firmy. Na podstawie ch-ki wybierzemy potrzebną długość radiatora. Wynika z tego, że wystarczą odcinki o długości 6 cm w przypadku wyboru jednego z nich. (...) tranzystor o katalogowej mocy strat równej 90 W z trudem jest w stanie odprowadzić z pomocą radiatora moc strat 41 W. Praktycznie to jego maksimum. Nie mówiąc o wielkości potrzebnego radiatora, co ma wpływ na wielkość urządzenia. W tym wypadku wybieram MJ15002. Alternatywą może być tranzystor MJ15025. Wybór tych tranzystorów gwarantuje dobór jak najmniejszego radiatora.

Tranzystor T_3 to MOSFET P – **IRF9540** (...) powinien wytrzymać bez radiatora. [Podobnie] dla **LM317** moc wydzielona wynosi ok. 1,5 W. Dla tranzystora T_5 – **BD677** moc strat wynosi ok. 5 W. Dla tego tranzystora nie znalazłem informacji

o rezystancji R_{thjc} złącze – obudowa. Podobny tranzystor do BD677, który możemy zastosować to **NTE253**. (...) W tym przypadku wybierzemy radiator **SK104 – 25STS – 14°C/W** lub **SK104 – 38STS – 11°C/W** .

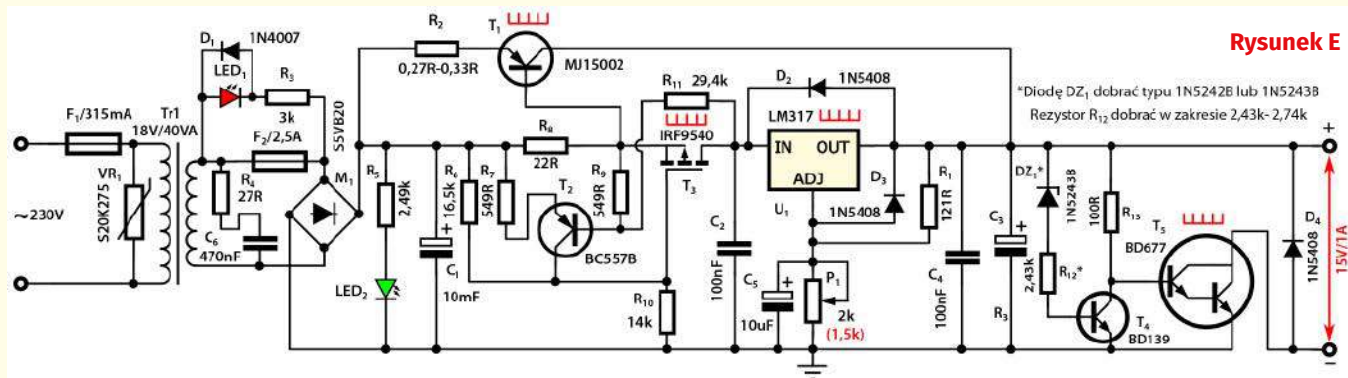
Przedstawiłem w miarę prosty przykład zasilacza o niewygórowanych parametrach. A i tak wymagało to użycia dodatkowego tranzystora aby zasilacz funkcjonował w pełnym zakresie wartości napięcia i prądu, a co jak się okazało, z samym LM317 nie jest możliwe. Według tej zasady możemy projektować zasilacz z użyciem LM317 na większe napięcia i prądy. Wybór większego napięcia wyjściowego może wymagać zastosowania układu regulatora na większe napięcie czyli LM317HV o wartości 60 V. Natomiast większe prądy wymagać będą zastosowania więcej niż jednego tranzystora. Można spotkać (np. w Internecie) całe mnóstwo układów z wykorzystaniem LM317 ale niekoniecznie są one prawidłowe. Chcąc je wykorzystać trzeba się im dobrze przyjrzeć czy aby na pewno będą działać prawidłowo.

Układ powstał na potrzeby rozwiązania zadania Szkoły Konstruktorów. Został specjalnie na tę okoliczność zaprojektowany. Nie był sprawdzony w praktyce, ale wydaje się być interesującym rozwiązaniem. (...) Zastosowałem ogranicznik nadnapięciowy, który niekoniecznie musi być zastosowany, bo może np. ograniczać możliwości zasilacza. Może być ustawiony na wyższe napięcie przez odpowiedni dobór elementów, głównie diody Zenera. (...) **Gdy wybierałem parametry zasilacza, a więc 15/1 A wydawało mi się, że są to takie warunki, które okażą się łatwe w realizacji. Okazało się inaczej.** (...)

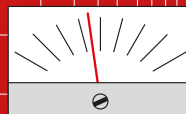
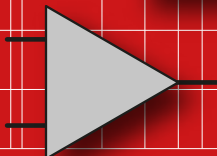
Tyle o zadaniu Polic317.

Pozdrawiam. ■

Piotr Górecki


Rysunek E

AUDIO OUT

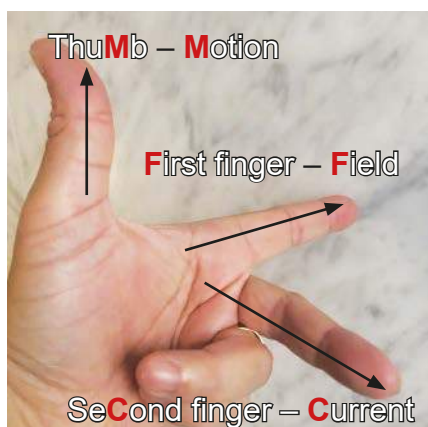


Nakrętki i śruby w głośniku, część 2

W tym miesiącu kończymy naszą krótką odskocznnię od projektu BBC LS3/5A, przedstawiając wskazówki i porady dotyczące budowania głośników.

Podłączanie głośników

System napędu głośnikowego jest praktyczną demonstracją reguły lewej ręki Fleminga, jak pokazano na rysunku 22. Ta stara zasada fizyki pokazuje, że prąd, pole magnetyczne i siła są ustawione względem siebie pod kątem prostym. Stojąc przodem do głośnika, jeśli prąd płynie przez cewkę zgodnie z ruchem wskazówek zegara (przyłożone dodatnie napięcie), a pole magnetyczne magnesu głośnika przebiega przez szczelinę od bieguna północnego (biegun wewnętrzny – „nabiegunnik”) do południowego (biegun zewnętrzny – „płyta czołowa”), wówczas siła działająca na membranę będzie skierowana na zewnątrz obudowy. Jednym z moich wykładowych trików jest włożenie niezamontowanego zespołu stożka do magnesu głośnika i powiedzenie studentom, żeby podłączyli baterię – wyskoczy (rysunek 23). Głośnik z ruchomą cewką został pierwotnie opracowany na podstawie sztuczki cyrkowej, skaczącego pierścienia



Rysunek 22. Reguła lewej ręki Fleminga odnosi się do kierunku przepływu prądu, pola magnetycznego i siły wypadkowej – wszystkie pod kątem 90° (ortogonalnie) względem siebie. Jest to zasada działania głośników z ruchomą cewką.

Olivera Lodge'a z 1898 roku. Aby cewka wyskoczyła, biegunowość magnesu, kierunek uzwojenia i przyłożone napięcie muszą być



Rysunek 23. Wyskakująca membrana głośnika. Dodatnie napięcie przyłożone do luźnej membrany powoduje jej wyskoczenie.

odpowiednie. Jeśli któryś z tych elementów jest nieodpowiednio ustawiony, cewka zostanie wessana do środka. Polaryzacja jest też szczególnie ważna w przypadku głośników stereo. Jeśli jeden głośnik pcha, a drugi ciągnie, to fale



Rysunek 24. Polaryzacja głośnika: (po lewej) dodatnie napięcie na połączeniu plusowym wypycha membranę na zewnątrz; (po prawej) odwrócenie napięcia powoduje jej wciągnięcie.



Rysunek 25. Standardowy kabel głośnikowy na wzór słynnego QED 79 firmy Pro Power. Zwróć uwagę na pasek polaryzacji. 79 żył miedzi o przekroju 0,17 mm² daje CSA 2,4 mm² i rezystancję 0,11 Ω na 10 m przewodu.

dźwiękowe, zwłaszcza długie fale basu, będą się wzajemnie znosić. Na szczęście nie ma dwuznaczności w kwestii polaryzacji głośników, jeśli do dodatniego lub czerwonego zacisku przyłożymy dodatnie napięcie baterii, wtedy membrana przesunie się na zewnątrz (z dala od magnesu), jak pokazano na **rysunku 24** (po lewej). Jeśli przyłożona polaryzacja baterii jest odwrotna, to membrana jest zasysana do środka, jak pokazano na **rysunku 24** (po prawej). Łatwo to zauważyć w przypadku głośników niskotonowych, gdzie ruch membrany jest duży. W przypadku głośników wysokotonowych jest to trudniejsze do wykrycia; ruch może wynosić tylko 0,5 mm. Najlepiej jest patrzeć na bok kopułki z prostą krawędzią tuż nad nią lub kawałkiem papieru do wykresów za nią. Nie używaj niczego ponad zużytą baterię PP3 9 V. Nie chcesz przecież rozwalić membrany po to, żeby sprawdzić czy naprawdę działa.



Rysunek 26. Kabel głośnikowy w bezbarwnej izolacji. Pro Power CPC CB08648 1,5mm², 196 x 0,1mm.



Rysunek 27. Profesjonalny kabel głośnikowy firmy Vandamme, 294 funty za rolkę o długości 100 metrów.

Rodzaje i parametry kabli głośnikowych

Głośniki to urządzenia o niskiej impedancji, często ze spadkami do kilku omów na krzywej impedancji w funkcji częstotliwości, więc wynika z tego, że napędzający je kabel powinien mieć jak najmniejszą rezystancję. Jeśli głośnik 4 Ω jest napędzany przez kabel o rezystancji 1 Ω, to 20% mocy wzmacniacza zostanie zmarnowane na podgrzewanie kabla. Ponadto, ponieważ impedancja głośnika zmienia się szalenie w zależności od częstotliwości, kabel może narzucić odchylenia od krzywej impedancji w funkcji częstotliwości. Kable głośnikowe są często dość długie, zwykle mają kilka metrów, co sprawia, że potrzeba niskiej rezystancji jest jeszcze ważniejsza. Można zastosować zwykły kabel sieciowy T&E o przekroju 2,5 mm² (CSA), który z łatwością spełnia to wymaganie, ale ma niewłaściwe cechy mechaniczne. Sztywny kabel z litym rdzeniem będzie grzechotał i brzęczał,

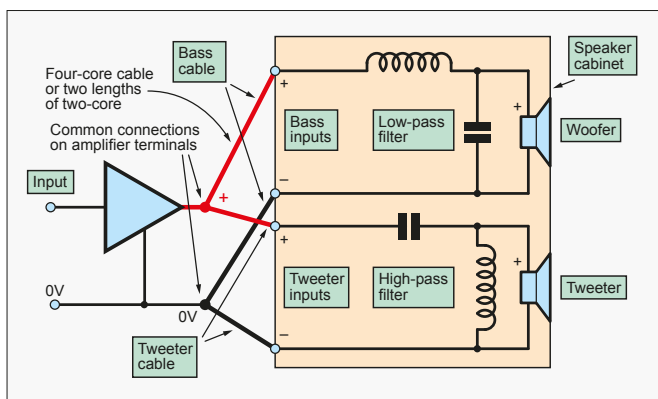
a także może fizycznie obciążać złącza, wyrывая wtyki i łamiąc terminale, jeśli będziemy poruszać głośnikami. Kabel głośnikowy musi być miękki i elastyczny, więc giętki kabel sieciowy będzie w porządku. Zapasowy przewód uziemiający jest po prostu prowadzony równolegle z przewodem neutralnym. Specjalistyczne kable głośnikowe są często płaskie, aby ułatwić przebieg pod dywanami. Konstrukcja podwójnej ósemki również zmniejsza indukcyjność, ale z drugiej strony zwiększa pojemność. Te parametry reakcyjne generalnie nie mają dużego znaczenia, ale mogą powodować niestabilność w niektórych wzmacniaczach. Kiedyś musiałem wprowadzić niewielką indukcyjność do przewodów głośnikowych niektórych modeli wzmacniaczy Naima, aby zapobiec oscylacjom wysokiej częstotliwości przy zastosowaniu kabla głośnikowego o dużej pojemności. Idealny kabel głośnikowy to klasyczny QED 79, który ma 79 żył z miedzi o przekroju 0,2 mm² w formie ósemki z paskiem wskaźnika polaryzacji po jednej stronie, jak pokazano na **rysunku 25**. Kabel głośnikowy jest często wykonany z białego PVC, który wygląda okropnie i wraz z wiekiem jeszcze się pogarsza. Kabel pokryty przezroczystym PVC (**rysunek 26**) wydaje się pasować do wiktoriańskich domów brytyjskich. Jednak w przypadku tego typu kabla często trudno jest dostrzec polaryzację. Aby obejść ten problem, niektóre kable z przezroczystego PVC wykorzystują cynowaną miedź o srebrnym wyglądzie dla ujemnej strony i różową lub niecynowaną miedź dla strony dodatniej, aby zapewnić wyraźną identyfikację polaryzacji. Popularną obsesją w środowisku Hi-Fi jest 'miedź beztlenowa' (OFC), która oznacza po prostu nieco czystsza miedź i nie robi tak naprawdę żadnej wymiernej różnicy. Metaliczna miedź nie zawiera tlenu, gdyż proces wytapiania



Rysunek 28. Takie ściągacze izolacji z marek Rolson budowlanych są doskonałe do kabli głośnikowych, jak również do okablowania domowego.



Rysunek 29. Zwróć uwagę, że okrągłe ostrze ściągacza izolacji Rolsona nie nacinają żył kabla CPC z **rysunku 26**.



Rysunek 30. Bi-wiring głośników ogranicza spadki napięcia w okablowaniu do odpowiednich obwodów głośnika niskotonowego i wysokotonowego. Dzięki temu uzyskuje się lepsze tłumienie pozapasmowe na filtrach.

usuwa go. Szczególnie lubię serię kabli OFC Vandamme Blue (patrz **rysunek 27**), ponieważ zbudowane są z przewodów 2,5 mm² CSA, są bardzo elastyczne i doskonale pasują do wtyków Speakon. Ściąganie izolacji z grubego kabla głośnikowego może być trudne; na szczęście istnieje fantastyczny żółty ściągacz izolacji, dostępny w sklepach budowlanych za około 7 funtów firmy Rolson: Automatic Wire Stripper 20857 (**rysunki 28 i 29**). Jest to udana chińska kopia amerykańskiego Ideal Stripmaster, który kosztuje pięć razy więcej. Konstrukcja nie nacina żył w porównaniu ze zwykłymi ściągaczami; patrz: <http://bit.ly/pe-dec19-rolson>.

Bi-wiring

Bardzo subtelną poprawę jakości dźwięku można uzyskać stosując bi-wiring (połączenie dwukablowe) – to znaczy odseparowując obwody głośnika wysokotonowego i niskotonowego oraz ich odpowiednie filtry pasywne i łącząc je jedynie na zaciskach wzmacniacza. Oznacza to, że skrzynka głośnikowa ma cztery zaciski z tyłu. Jest to forma połączenia w gwiazdę (**rysunek 30**) i zapewnia, że żadne

napięcie indukowane w rezystancji kabla nie będzie się nakładać na sygnał napędzający drugi przetwornik. Powiedziałbym, że jest to drobiazg dla użytkownika Hi-Fi, a inżynier dźwięku wydałby dodatkowy budżet w bardziej efektywny sposób. Przy próbie zastosowania bi-wiringu należy pamiętać, że prąd płynący do głośnika wysokotonowego jest bardzo mały, więc można do jego przesyłu użyć cienkiego, taniego drutu. To samo dotyczy głośników aktywnych. Niektórzy ekstremalni entuzjaści Hi-Fi stosują technikę zwaną „bi-amping”, w której dwa oddzielne stereofoniczne wzmacniacze mocy są używane do zasilania terminali bi-wire. Wejścia do wzmacniaczy mocy są połączone równolegle. Ten pomysł jest szalony – jeśli zamierzasz ponieść koszt użycia dwóch wzmacniaczy, czy na pewno nie lepiej jest użyć aktywnej zwrotnicy? Moim głównym źródłem kabli jest seria CPC/Farnell Pro Power, która kiedyś była tania. Niestety już nie jest, światowe ceny miedzi stale rosną. Nadal używam ich czteropinowego kabla zasilającego do małych aktywnych głośników, ponieważ mieści się on w czteropinowym XLR.



Rysunek 32. Zaciski śrubowe (po prawej) to najczęściej spotykany typ złącza głośnikowego.



Rysunek 31. Bi-wiring nie jest konieczny dla większości małych głośników 8 Ω, więc zaciski dla głośnika niskotonowego i wysokotonowego mogą być spięte.



Rysunek 33. Wtyki bananowe zostały wynalezione w latach dwudziestych ubiegłego wieku. Nadal są powszechnie stosowane w systemach głośnikowych (a także w zasilaczach i budżetowych przyrządach laboratoryjnych).

Firmy przerzucające kartony

Obecnie stało się wymogiem marketingowym, aby głośniki były przystosowane do bi-wiringu, ale zdecydowana większość użytkowników nie przejmując się tym, więc większość głośników jest dostarczana z kilkoma połączonymi stalowymi paskami łączącymi dwie pary zacisków razem, jak pokazano na **rysunku 31**. Klienci niechętnie płacą ogromne pieniądze za kable, które są często wykorzystywane do zwiększenia marży na sprzedaży kompletnych systemów. Kiedy byłem studentem, pracowałem w specjalistycznym salonie Hi-Fi w Kensington. Zdarzało się, że oferowałem „za darmo” przesadnie drogie kable, aby doprowadzić do sprzedaży systemów audio. Prawie sto LS3/5A z tanimi wzmacniaczami Technicsa i odtwarzaczami CD trafiło do posażnych pań w drogich, małych mieszkaniach. Jeśli nie zgodziły się na kosztujące (wówczas) £230 LS3/5A, sprzedawaliśmy im imitator Castle Acoustics Clyde. Jeśli to nie działało, poddawaliśmy się i sugerowaliśmy parę węgierskich Videotonów Minimax 2 ze sklepu otwartego przez niejakiego Juliana Richera



Rysunek 34. Zaciski sprężynowe są często stosowane do głośników w tanim sprzęcie; na przykład w tej skrzynce przełączającej głośniki.



Rysunek 35. 3-pinowe gniazdo XLR w starym BBC LS3/5A.



Rysunek 36. W moich małych głośnikach aktywnych stosuję 4-pinowe XLR-y.



Rysunek 37. Wtyk Cannon EP: doskonały dla wytrzymałych głośników przenośnych.

na London Bridge. On założył Richer Sounds i stał się bogaty... a ja nie.

Złącza głośnikowe

Dla osób wychowanych na nowoczesnej elektronice, złącza głośnikowe mogą wydawać się bardzo prymitywne, ale wszystko czego potrzebujemy to niskooporowe, wytrzymałe, dwukierunkowe połączenie dla grubych kabli; raczej zbliżonego do domowego okablowania niż transmisji danych. Rzeczywiście, często wystarczy po prostu kostka śrubowa. Należy jednak pamiętać, że złącza są często źródłem przecieków powietrza w obudowie, więc należy się upewnić, że są one odpowiednio uszczelnione z tyłu klejem.

Szałeństwo wyboru wtyków

Ogromna różnorodność złączy głośnikowych doskonale pokazuje, że wszelkie próby standaryzacji na przestrzeni lat zawiodły; w powszechnym użyciu jest około 7 rodzajów. Najpopularniejsze wydają się być zaciski śrubowe lub wtyki, często z gołymi przewodami i dyndającymi wtykami bananowymi

(rysunki 32 i 33), pomimo możliwości przypadkowego połączenia faz i zwarcia wzmacniaczy. (W przypadku prac rozwojowych, oferują one jednak uniwersalność). Tanie głośniki konsumenckie mają tendencję do stosowania zacisków sprężynowych, jak pokazano na rysunku 34. W oryginalnym BBC LS3/5A zastosowano męski, trzypinowy XLR, pokazany na rysunku 35. To rozwiązanie dobrze się sprawdza, pomijając ryzyko, że wyjście wzmacniacza mocy może być doprowadzone do wyjścia audio o poziomie liniowym w innym urządzeniu. Powodem zastosowania żeńskiego złącza liniowego w kablu wzmacniacza jest uniknięcie odsłoniętych pinów w złączu męskim. (W żeńskim złączu XLR gniazda są całkowicie osłonięte gumą). Nie istnieje konwencja pinów dla głośników. W BBC LS3/5A pin 1 był nieużywany, środkowy pin 3 służył jako ujemny, a pin 2 jako dodatni. Do moich aktywnych LS3/5A używam rzadko spotykanego czteropinowego XLR, pokazanego na rysunku 36. Wtyki XLR często wymagają uszczelnienia klejem z tyłu, aby zapobiec wyciekom

powietrza. W połowie lat 80. Cannon podjął próbę wprowadzenia dwupinowego XLR-a dla głośników; niestety szybko zniknął, mimo że był to doskonały projekt. Istnieje niesamowita wersja XLR zwana „Cannon EP” (pokazana na rysunku 37), która ma fantastyczny, militarny wygląd i solidną metalową konstrukcję. Używa się go w bardzo ciężkim sprzęcie nagłośnieniowym, gdzie skądinąd doskonale plastikowe złącze Speakon (poniżej) może zawieść w wyniku przejechania przez nieostrożnych kierowców ciężarówek! Do normalnego użytku, myślę, że najlepszym złączem jest zaprojektowany przez Neutrika Speakon, który jest dostępny w wersji dwu-, cztero-, a nawet ośmiopinowej. Stało się ono de facto standardem w urządzeniach profesjonalnych, takich jak systemy nagłośnieniowe. Wersja czteropinowa pokazana jest na rysunku 38. Ma je również końcówka mocy na rysunku 39. Niestety, świat Hi-Fi nie przyjął jeszcze tego rozwiązania, wciąż trzymając się ohydnie drogich złotych konektorów. W przypadku Speakona nie ma ryzyka zwarcia i można w nim zmieścić kable o grubości 15 mm. Wtyki DIN będą znane



Rysunek 38. Wtyk Speakon – złącze profesjonalne. Polaryzacja jest wyraźnie zaznaczona. Większość ma wewnątrz zaciski śrubowe.



Rysunek 39. Gniazda Speakon na wzmacniaczu mocy Yamaha. Wtyki należy obrócić po wpięciu do gniazda, aby zablokować połączenie.



Rysunek 40. Okropne złącze DIN, spotykane w niektórych europejskich głośnikach. Jest to zdecydowanie moje najmniej lubiane złącze głośnikowe.



Rysunek 41. Wysokie na dwie stopy stojaki Target z otwartą ramą; są doskonałe dla LS3/5A i innych małych głośników.

użytkownikom Quadów i Naimów; pracują zaskakująco dobrze dla małych sygnałów. Ja ich nie lubię, bo są trudne do lutowania. Istnieje złącze DIN dla głośników (rysunek 40), ale jest ono rzadko używane, ponieważ para wtyk-gniazdo ma niewystarczającą powierzchnię przewodzącą dla mocy większej niż około 10 W. Co dziwne, do wyprowadzenia głośnika na obu końcach stosuje się złącza męskie, co stwarza możliwość zwarcia wzmacniacza. Płaski pin jest ujemny. Na koniec dochodzimy do złączy 'jack', które są używane przez bractwo gitar elektrycznych, nie tylko do samych gitar, ale

także efektów i głośników. Możliwości pomyślenia okablowania są nieograniczone (i częste). Połączenie dodatnie głośnika do „końcówka” jacka, więc zwarcie z uziemionym metalem jest niemal gwarantowane. Jeśli chodzi o głośniki Hi-Fi, rada jest prosta: unikajcie wtyków – nie używajcie ich, nigdy.

Ustawienie głośników

Większość głośników Hi-Fi brzmi najlepiej na podstawkach oddalonych o pół metra od powierzchni odbijających, takich jak ściany, jak pokazano na rysunku 41. Daje to niemal swobodne warunki promieniowania dla tonów średnich i wysokich, dając gładką odpowiedź częstotliwościową. Jednakże, gdy długości fal rosną, charakterystyka promieniowania zaczyna być wzmacniana przez granice, kompensując zwijanie się basu w przypadku większości mniejszych głośników. Grube dywany z wełnianym podkładem pomagają zredukować odbicia od podłogi (patrz rysunek 42), a grube, aksamitne zasłony



Rysunek 42. Niskie standy AVF, odpowiednie do wyłożonych wykładziną dywanową małych pomieszczeń. Mają regulowany kąt, aby skierować oś głośnika do góry, co pozwala uniknąć zakłóceń w zwrotnicy.

w pokoju są zawsze dobrym pomysłem. Chociaż niektóre głośniki, takie jak LS3/5A, są często opisywane jako „półkowe”, jest to tak naprawdę odniesienie do ich rozmiaru, a montaż na półce jest kiepską pozycją akustyczną. Widziałem nawet głośniki postawione na bokach na półkach w sposób, jaki pokazano na rysunku 43. Daje to w rezultacie masywne wcięcie na zwrotnicy, brak głębi stereo i trąbiące, niskie tony średnie. Umieszczenie głośników w rogu jeszcze bardziej pogarsza sprawę. Tylko projektanci wnętrz popełniają takie akustyczne okrucieństwa.

Symetria

Głośniki powinny być umieszczone w pomieszczeniu symetrycznie, ponieważ dla dobrego obrazowania stereo istotne jest, aby lewa i prawa strona pasma przenoszenia były takie same. Jeśli masz pokój w kształcie litery L lub otwarte drzwi z jednej strony, pojawią się problemy. Aby sobie z tym poradzić, stworzyłem przedwzmacniacze z regulatorami balansu i oddzielnymi regulatorami EQ dla lewej i prawej strony. Jest to również konieczne w przypadku klientów z asymetrycznym słuchem.

Stojaki na głośniki

Stojaki na głośniki powinny być sztywne i ciężkie, aby zapobiec wibracjom głośnika spowodowanym odrzutem membrany. Daje to subtelny efekt, który jest zauważalny tylko wtedy, gdy cała podstawowa inżynieria elektroakustyczna jest zachowana. Te wtórne efekty są często przedmiotem obsesji w subiektywnym świecie Hi-Fi. Pamiętam, jak starannie unikałem poinformowania zamożnego klienta, że jego głośniki są poza fazą, zamieniając ukradkiem przewody, podczas przedstawiania jego kolumn za 1000 funtów. Często używane są kolce (patrz rysunek 44)



Rysunek 43. Nie należy zabudowywać głośników „półkowych”. Otwarte standy to jedyny sposób na uzyskanie dobrego obrazowania stereo.



Rysunek 44. Kolce takie jak te są często wkręcane w podstawy głośników i standów w celu zapewnienia ich sztywności.



Rysunek 45. Jeśli nie lubisz standów, a chcesz uzyskać lepszy bas, skuteczne są kolumny podłogowe, takie jak na zdjęciu. Głośnik centralny powinien być jednak na stojaku.



Rysunek 46. Demonstracja tłumionych (po lewej) i nietłumionych (po prawej) paneli ze sklejk (zainspirowana starym wykładem KEF-a).

do sztywnego połączenia głośnika z podstawką, a podstawki z podłogą. Kolce podłogowe sprawdzają się bardzo dobrze w przypadku podłóg litych i betonowych, ale niszczą sosnowe deski podłogowe i dywany. Ich wysokość można regulować, aby zapobiec kołysaniu

na nierównych powierzchniach. Górna część podstawki może pomóc w tłumieniu podstawy obudowy, szczególnie jeśli pomiędzy górną częścią podstawki a podstawą głośnika umieszczona jest jakaś sprężysta masa, taka jak na przykład Blu Tack. Takie podejście zmniejsza również prawdopodobieństwo strącenia głośnika z podstawy. Spawane ze sobą rury stalowe o przekroju kwadratowym, wypełnione suchym piaskiem to najbardziej opłacalna technika konstrukcyjna. Stojak powinien mieć otwartą konstrukcję, aby zapobiec odbiciom. Masywne podstawki o szerokim przekroju, choć ciężkie i bardzo sztywne, mogą zepsuć delikatną reprodukcję wokalu i przestrzeni. Rogers wyprodukował specjalne podstawki pod subwoofery LS3/5A o nazwie 'AB1', ale stwierdziłem, że zestaw brzmiał lepiej, gdy LS3/5A stały na otwartych podstawkach. Jeśli chcesz takiego układu, wybierz parę podstawek podłogowych (rysunek 45) i zaoszczędź na kosztach podstawek, a dokładnieść

przestrzenną przekuj na lepszy bas. Moje komputerowe głośniki biurkowe Wavecor są po prostu zamontowane na parze postawionych cegieł. Ostatnia sprawa: pisałem w zeszłym miesiącu obszernie o tłumieniu paneli. Oto zdjęcie (rysunek 46) z mojego wykładu demonstrującego wytłumione i niewytłumione panele drewniane. Jedna z nich „brzęczy” jak ksylofon, druga emituje tępy stukot!

Mówiąc wprost

W przyszłym miesiącu wracamy do tematu LS3 – do drewna i montażu obudowy LS3/5A. Odkryjemy też więcej tajemnic sztuki budowy głośników.. ■

Jake Rothman

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

Quiz: Sprawdź, czy coś pamiętasz z artykułów o thereminie

Aby bezbłędnie odpowiedzieć na pytania skorzystaj z wiedzy, którą znajdziesz w EdW 6, 8, 9, 10/2022

Lew Termen i Leon Theremin to ta sama osoba. Które nazwisko pojawiło się później?

- ☐ A. Lew Termen
- ☐ B. Leon Theremin
- ☐ C. trudno powiedzieć

Pierwsze koncerty publiczne na instrumencie Termenvox odbyły się:

- ☐ A. w końcu XIX wieku
- ☐ B. w roku 1920
- ☐ C. w roku 1952

Do czego służy pionowy pręt („antena”) w thereminie?

- ☐ A. do sterowania wysokością tonu
- ☐ B. do sterowania głośnością
- ☐ C. do sterowania zarówno wysokością tonu jak i głośnością

Ile oscylatorów wytwarza dźwięki w thereminie:

- ☐ A. jeden
- ☐ B. dwa
- ☐ C. trzy

Sterowanie wysokością tonu odbywa się przez zmiany:

- ☐ A. pojemności
- ☐ B. indukcyjności
- ☐ C. zarówno pojemności jak i indukcyjności

Płytką metalową służy do regulacji:

- ☐ A. niskich tonów
- ☐ B. wysokich tonów
- ☐ C. głośności

Uziemienie ręcznego theremina:

- ☐ A. nie jest potrzebne
- ☐ B. jest potrzebne i odbywa się przez pętlę metalową trzymaną w ręku
- ☐ C. jest potrzebne i odbywa się przez obudowę metalową trzymaną w ręku

Thereminy wytwarzano:

- ☐ A. wyłącznie na lampach
- ☐ B. zarówno na lampach jak i tranzystorach
- ☐ C. wyłącznie na tranzystorach

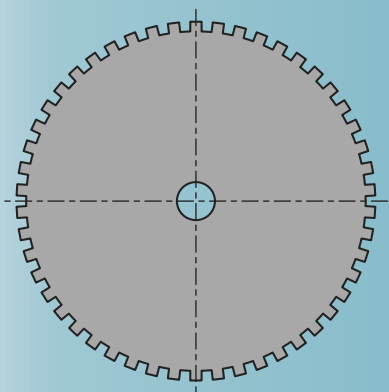
Wynalazca theremina pochodzi:

- ☐ A. z Rosji
- ☐ B. z USA
- ☐ C. z Wielkiej Brytanii

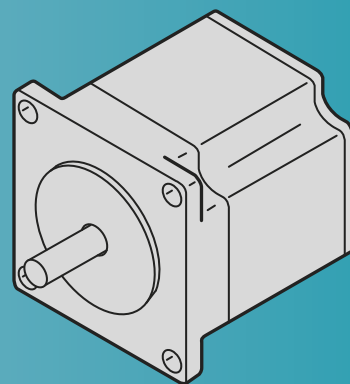
Wynalazca theremina był z wykształcenia:

- ☐ A. fizykiem
- ☐ B. muzykiem
- ☐ C. fizykiem i muzykiem

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 10.12.2022.



Silniki krokowe w praktyce



Część 1: Wprowadzenie do technologii silników krokowych

Od 18 lat odpowiadam na pytania techniczne od klientów na stronie technobotsonline.com. Technobots to firma, która dostarcza komponenty elektroniczne i mechaniczne dla hobbystów, handlu i zastosowań edukacyjnych. Powtarzającym się tematem w pytaniach od klientów są zastosowania silników krokowych i sposób ich sterowania. Pomimo pojawienia się silników krokowych około 1950 roku i powszechnego ich używania od lat 70., pozostają one elementem nieco tajemniczym dla wielu hobbystów. W tej serii artykułów opiszę zarówno teorię jak i sposoby ich użycia, a zwińczeniem będzie projekt pełnego, 3-osiowego kontrolera CNC. Po drodze zostaną omówione różne mniejsze projekty kontrolerów, więc w zależności od tego, czy chcesz po prostu uzyskać wolno obracającą się antenę radarową na swoim modelu statku, czy może „zmotoryzowany” suwak kamery, artykuły te wskażą ci drogę. Ten artykuł da Ci ogłęd na główne typy dostępnych silników krokowych.

Co to jest silnik krokowy?

Powszechnie stosowany szczotkowy silnik prądu stałego nie wymaga większego wprowadzenia; od ponad wieku napędza wszystko: od małych zabawek po potężne obrabiarki. Po prostu podłącz go do źródła prądu stałego, a będzie się on obracał z prędkością wielu tysięcy obrotów na minutę. Odwróć zaś polaryzację, a będzie się kręcić w przeciwnym kierunku. Silniki prądu stałego są zatem jednym z bardzo skutecznych sposobów zapewnienia kontroli nad prędkością obrotową. Co jednak w przypadku, gdy nie chcesz kontrolować tylko prędkości, lecz także pozycję obrotową?

Właśnie w takich przypadkach zastosowanie mają silniki krokowe. Nie są one jednak tak proste jak silniki DC; wymagają elektronicznego sterownika do zarządzania zasilaniem cewek silnika. Silnik krokowy jest bezszczotkowym silnikiem prądu stałego, który porusza się o dyskretnie kroki w każdym kierunku i może się zatrzymać i utrzymać w każdym z tych kroków. Zalety silników krokowych w porównaniu do silników szczotkowych DC będą doskonale widoczne w aplikacjach, gdzie wymagana jest kontrola pozycji lub wolna prędkość obrotowa niwelująca potrzebę czujnika pozycyjnego sprzężenia zwrotnego.

- stare akcesoria komputerowe – napędy dyskietek.

W przeciwieństwie do szczotkowych silników prądu stałego, silniki krokowe pobierają najwięcej prądu, gdy się nie obracają, ze względu na ich prąd trzymający. Niektóre sterowniki silników zmniejszają prąd trzymający silnika, aż do momentu, gdy jest on potrzebny do ponownego obrotu. Prąd trzymający i brak wentylatora wymuszającego chłodzenie oznacza, że silniki krokowe mogą mocno się nagrzewać, zwłaszcza gdy pracują w zastosowaniach stacjonarnych.

Rodzaje silników krokowych

Ta seria artykułów nie dotyczy wewnętrznej konstrukcji silników krokowych (choć temat ten jest na pewno fascynujący) ale tego, jak z nich korzystać. Jeśli potrzebujesz dobrego wprowadzenia do tego, jak działają takie silniki, to film z YouTube wymieniony poniżej w sekcji **Hybrydowe silniki krokowe** będzie doskonały. Mimo wszystko, nie możemy traktować ich wyłącznie jako czarne skrzynki, więc poniżej przedstawimy szybki przegląd trzech głównych typów silnika krokowego:

- z magnesem trwałym,
- ze zmienną reluktancją,
- hybrydowy.

Kluczowe różnice dotyczą wirnika (części silnika, która się obraca) i stojana (stacjonarne



Rysunek 1. 48-krokowy bipolarny silnik krokowy z magnesem trwałym w dwóch „puszkach”.

Zastosowania silników krokowych

Silniki krokowe mogą mieć moc od miliwatów do setek, jeśli nie tysięcy watów, i znajdują się w wielu produktach, takich jak:

- zegarki na rękę,
- drukarki 3D,
- kserokopiarki i drukarki,
- obrabiarki CNC,
- elektronika użytkowa, taka jak kontrola ostrości i powiększenia kamery,
- pompy cyfrowe do dokładnego odmierzania płynów,
- robotyka,
- automatyka,



Rysunek 2. Widok wewnętrzny silnika krokowego z magnesem trwałym 28BYJ-48 z wbudowaną przekładnią. Należy zwrócić uwagę na 'palce' wokół wewnętrznej strony stojana dające 32 kroki na obrót. Liczba kroków na wale wyjściowym wynosi w przybliżeniu $\ast 32 \times \text{przełożenie przekładni } 64 = 2048$ w trybie full-step. (*uwaga, że przełożenie przekładni nie wynosi dokładnie 64:1).

uzwojenia nawinięte wokół wirnika). Możesz spotkać się także z określeniem „miękkie żelazo”. Miękkie żelazo jest wyżarzane i w przeciwieństwie do magnetycznie „twardego” żelaza ma niską koercję, co oznacza, że nie pozostaje namagnesowane po usunięciu pola magnetycznego z cewek stojana. Jest to istotne w zastosowaniach, gdzie wymagane jest wielokrotne włączanie i wyłączanie pola magnetycznego.

Silnik krokowy z magnesem trwałym

Silniki tego typu mają osiowo namagnesowany cylindryczny wirnik z przemennymi biegunami północnymi i południowymi ułożonymi równoległe do wału wirnika. W najprostszej formie, stojan składa się zazwyczaj z dwóch uzwojeń, każde nawinięte na przeciwległym biegunie w szeregu, aby utworzyć dwufazowe uzwojenie zamknięte w metalowej puszcze. W praktyce, aby stworzyć więcej kroków na obrót, potrzeba więcej biegunów,

co jest przyczyną powstania silników typu zaprezentowanego na **rysunku 1**. Aby utworzyć więcej biegunów, ale zachować dwa uzwojenia, konieczne jest zastosowanie „palców” lub „zębów” wokół wnętrza każdej cewki. Zazwyczaj takie silniki są dostępne z 12, 20, 24, 48 i 100 kroków na obrót, czasami też więcej, ale 24 i 48 są najpopularniejsze. Niska liczba kroków może być kompensowana poprzez zintegrowanie przekładni. W silniku 28BYJ-48 (**rysunek 2**) z jego przekładni 64:1 można uzyskać za pomocą sterownika za poniżej 2 funtów (z bezpłatną dostawą na eBay) ponad 2000 kroków na obrót w tak zwanym trybie „full-step”. (Tryby krokowe zostaną opisane w późniejszym artykule.) Kluczowe zalety silników krokowych z magnesami trwałymi to wysoka przydatność do zastosowań o niskiej prędkości i umiarkowanych

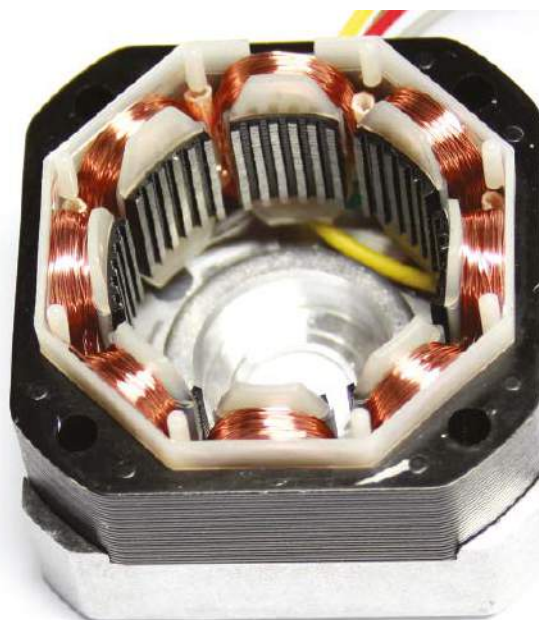
wymaganiach momentu obrotowego. Są one też zwykle tańsze niż silniki hybrydowe.

Silniki krokowe o zmiennej reluktancji

Silniki te nie mają magnesu trwałego w wirniku, lecz wykonany jest on z miękkiego żelaza/laminatów stalowych. Wirnik z miękkiego żelaza jest wielozębny i oferuje dobrą rozdzielczość kątową i prędkość obrotową, ale brakuje mu momentu obrotowego w porównaniu do innych typów silników. Jedną z istotnych różnic tego typu silnika jest to, że nie posiada on momentu postojowego. Moment postojowy to moment, który silnik wytwarza, gdy cewki stojana nie są zasilane. Można poczuć zakleszczenie spowodowane momentem postojowym



Rysunek 3. Wirnik hybrydowego silnika krokowego NEMA 17 z silnika o 200-krokach z widocznymi dwoma sekcjami wirnika, każda po 50 zębów. Przyjrzyj się uważnie, a powinieneś zobaczyć, że zęby obu sekcji wirnika są przesunięte względem siebie o pół skoku ($3,6^\circ$), tworząc efektywnie wirnik o 100 zębach.



Rysunek 4. Stojan 2-fazowego hybrydowego bipolarnego silnika krokowego NEMA 17 z usuniętym wirnikiem (patrz rysunek 3). Każdy z 8 biegunów ma uzwojenie, uzwojenia na biegunach 1, 3, 5 i 7 są połączone szeregowo tworząc fazę „A”, podobnie 2, 4, 6 i 8 są połączone szeregowo w fazę „B”. Steownik sekwencyjnie zasila fazy, które przyciągają zęby na wirniku, aż zęby wirnika zrównają się z 48 zębami stojana (dwa mniej niż wirnika). Pozostałe zęby są oddalone o $\frac{1}{4}$, $\frac{1}{2}$ lub $\frac{3}{4}$ podziałki, więc przy kolejnym kroku wirnik przesunie się o $\frac{1}{4}$ kroku $7,5^\circ$ lub o $1,8^\circ$, co odpowiada 200 krokom na obrót.



Rysunek 5. Wirnik bipolarny, 400-krokowego silnika NEMA 17, należy zwrócić uwagę na zwiększoną liczbę zębów w porównaniu do 200 krokowego wirnika z rysunku 3.

w silnikach hybrydowych i o magnesie trwałym podczas ręcznego obracania wału, gdy silnik nie jest zasilany. Ten typ silnika krokowego nie jest powszechnie stosowany i ma ograniczoną dostępność. Są one stosowane w specyficznych urządzeniach – możesz znaleźć taki silnik napędzający bęben Twojej nowej pralki. Nie będą one opisywane dalej w naszej serii.

Hybrydowe silniki krokowe

Na koniec mamy silniki hybrydowe, które wykorzystują elementy zarówno silnika z magnesem trwałym, jak i o zmiennej reluktancji. Hybrydowe silniki krokowe składają się z wirnika z magnesem trwałym z promieniowo rozłożonymi zębami z miękkiego żelaza (patrz **rysunek 3**), oraz wielobiegunowego stojana z zębami (patrz **rysunek 4**). Jest to prawdopodobnie najbardziej popularny typ silnika krokowego, ponieważ oferuje dobrą rozdzielczość kroku, moment trzymający i prędkość, ale jest też nieco droższy niż pozostałe typy silników. Najbardziej powszechne hybrydowe silniki krokowe posiadają 200 kroków na obrót, co daje $1,8^\circ$ na krok ($360^\circ/200$). Dostępne są silniki 400-krokowe, które dają długość kroku $0,9^\circ$ ($360/400$). **Rysunki 5 i 6** pokazują odpowiednio wirnik i stojan silnika 400-krokowego NEMA17.

Elektronika sterująca jest podobna zarówno dla silników krokowych z magnesem trwałym jak i hybrydowych, i może skutecznie zwiększyć liczbę kroków na obrót wiele razy za pomocą technik, które zostaną omówione w jednym z kolejnych artykułów.

Jak działa hybrydowy silnik krokowy

Większość tutoriali w Internecie wykorzystuje bardzo uproszczony widok

silnika krokowego, który nie pomaga w zrozumieniu, jak te 200-krokowe silniki faktycznie działają. Opisanie tego w niezbędnych szczegółach w naszych artykułach byłoby mniej skuteczne niż animowana grafika na doskonale przygotowanym wyjaśnieniu w filmie Nanotec, który można obejrzeć na stronie: <https://youtu.be/Ew6eVGnj7r0> (film polecany dla tych, którzy chcieliby lepiej zrozumieć, co dzieje się w hybrydowym silniku krokowym).

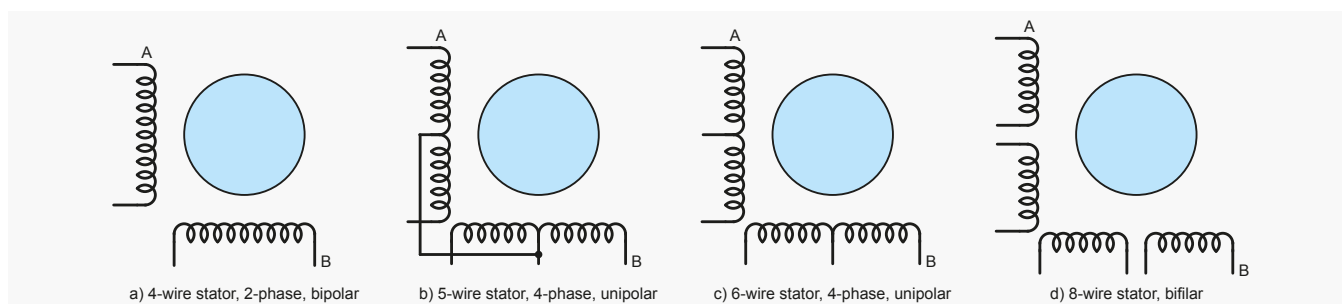
Uzwojenia unipolarne i bipolarne

Przeglądając się stałym lub hybrydowym silnikom krokowym, można zauważyć, że są one określane jako „unipolarne” lub „bipolarne”. Opisuje się w ten sposób konfigurację uzwojeń silnika, patrz **rysunek 7**. Silniki 4-fazowe unipolarne składają się z dwóch uzwojeń, z których każde posiada odczep środkowy. Sposób wyprowadzenia odczepów środkowych może być pojedynczy (6-cio przewodowy, patrz **rysunek 7c**) lub wspólny (5-cio przewodowy, patrz **rysunek 7b**). Silniki 2-fazowe bipolarne nie posiadają odczepu środkowego, wyprowadzone są tylko końce uzwojeń (4 przewodowe, **rysunek 7a**). Silniki unipolarne 6-cio przewodowe mogą być również wykorzystane jako silniki bipolarne poprzez rezygnację z dwóch



Rysunek 6. Wnętrze stojana 400-krokowego silnika bipolarnego NEMA 17, podobnie jak jego stojan, również ma zwiększoną liczbę zębów w porównaniu do 200-stopniowego stojana z rysunku 4.

odczepów środkowych. Istnieje jeszcze trzeci układ zwany „bifilarnym”, w którym zamiast nawijać każdą cewkę pojedynczym drutem, nawija się równolegle dwa druty z każdym końcem wyprowadzonym na zewnątrz (8 przewodów – **rysunek 7d**). Dzięki temu silnik może być napędzany w sposób unipolarny lub bipolarny poprzez łączenie uzwojeń szeregowo lub równolegle. Silniki unipolarne (**rysunek 8**) są łatwiejsze do napędzania gdyż układ sterujący nie musi odwracać prądu płynącego przez uzwojenie, co czyni układ znacznie prostszym do wykonania. Silnik bipolarny wymaga pełnego mostka „H”, aby umożliwić zmianę kierunku prądu, a tym samym odwrócić pole magnetyczne uzwojeń. Omówimy układy sterowników w kolejnych artykułach. Możesz się zastanawiać dlaczego nie zawsze używać prostszego podejścia do napędzania silników unipolarnych. Po prostu silniki unipolarne mają zmniejszony moment obrotowy, ponieważ tylko połowa każdego uzwojenia jest zasilana na raz, w przeciwieństwie do bipolarnych, gdzie całe uzwojenie jest zasilane. Jeśli nie budujesz swojego sterownika z elementów dyskretnych, opcja bipolarna jest właściwą drogą, ponieważ wielu producentów półprzewodników oferuje teraz dedykowane układy sterowników mostka H.



Rysunek 7. Konfiguracje połączeń silnika krokowego; najczęściej stosowany jest typ 4-przewodowy.

Tabela 1. Zależność pomiędzy numerem rozmiaru silnika NEMA a szerokością obudowy.

	NEMA 8	NEMA 11	NEMA 14	NEMA 17	NEMA 23	NEMA 34	NEMA 42
Motor face size (inches)	~0.8	~1.1	~1.4	~1.7	~2.3	~3.4	~4.2



Rysunek 8. Przykład silnika krokowego unipolarnego NEMA 17.



Rysunek 9. Silnik NEMA 17 z wydrążonym wałem – Technobots zastosował je u klienta, który potrzebował silnika do ciągłego obracania w jednym kierunku bardzo jasnego panelu świetlnego LED. Wał silnika napędzał pierścień ślizgowy, który z kolei obracał panel LED. Przewody pierścienia ślizgowego były prowadzone przez wydrążony wał do sterownika LED w taki sposób, że przewody nigdy się nie skręcały.



Rysunek 10. Bipolarny silnik krokowy NEMA 17 z wbudowaną śrubą prowadzącą T4 napędzającą mosiężną „nakrętkę”.

Rozmiary silników hybrydowych

Spośród trzech typów silników krokowych, wiemy już, że odmiana hybrydowa jest najczęściej dostępną, a odmiana bipolarna jest najczęstszą metodą sterowania. Hybrydowe silniki krokowe są dostępne w różnych rozmiarach i są identyfikowane przez numer NEMA. ‘NEMA’ to skrót od National Electrical Manufacturers Association, amerykańskiej instytucji reprezentującej ponad 300 producentów sprzętu elektrycznego z różnych sektorów rynku. Standard NEMA oparty jest na calach, ale pomimo tego jest to najbardziej rozpowszechniony sposób określający rozmiary silników. Istnieje również metryczny standard rozmiaru IEC, ale na rynku hobbystycznym większość sprzedawców używa klasyfikacji NEMA. Numer NEMA jest związany (w przybliżeniu) z szerokością kwadratowej powierzchni czołowej silnika krokowego (tabela 1). Istnieje również

rozmiar NEMA 24, który choć nie jest formalnie w standardzie, to zbliżony jest do NEMA 23, ale jest nieco szerszy, co pozwala na użycie większego wirnika i stojana, a tym samym uzyskanie większego momentu obrotowego. NEMA 17 jest szeroko stosowana przez hobbystów, szczególnie tych budujących drukarki 3D. Dla danego rozmiaru silnika NEMA, długość korpusu może być również różna, im większa długość korpusu, tym większy moment obrotowy. Inne rozmiary są używane w zastosowaniach komercyjnych, zarówno mniejsze, jak i większe niż z zakresu NEMA. Podczas gdy numer NEMA odnosi się do rozmiaru fizycznego, silniki mogą różnić się specyfikacją mechaniczną i elektryczną.

Dwa silniki o takim samym rozmiarze NEMA mogą wyglądać identycznie, ale ich parametry elektryczne mogą być bardzo różne i niekoniecznie będą wobec siebie zamienne. Dostępne są również różne rodzaje wałów dla silników krokowych: wały okrągłe, wały „D”, wały drążone (patrz rysunek 9), nawet wały śrubowe, gdzie wymagana jest konwersja ruchu obrotowego do liniowego (rysunek 10).

W następnym odcinku

W kolejnej części przyjrzymy się procesowi wyboru silnika krokowego i jak zidentyfikować wyprowadzenia silnika wydobytego z niedziałającego sprzętu. ■

Paul Cooper

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

Quiz: Silniki krokowe

Od czego zależy prędkość obrotów silnika krokowego?

- ☐ A. od napięcia impulsów wejściowych
- ☐ B. od natężenia prądu impulsów wejściowych
- ☐ C. od częstotliwości impulsów wejściowych

Kąt obrotu na jeden krok wynosi:

- ☐ A. do kilku stopni
- ☐ B. od kilku do kilkudziesięciu stopni
- ☐ C. od kilku do kilkuset stopni

Maksymalne prędkości obrotu wynoszą:

- ☐ A. kilkadziesiąt obrotów na minutę
- ☐ B. od kilku do kilkuset obrotów na minutę
- ☐ C. kilka tysięcy obrotów na minutę

Zadany kąt obrotu jest wykonywany:

- ☐ A. w pętli zamkniętej ze sprzężeniem zwrotnym
- ☐ B. w pętli otwartej, określane napięciem impulsów sterujących
- ☐ C. w pętli otwartej, określane liczbą impulsów sterujących

W stanie spoczynku przy uzwojeniach zasilanych moment jest:

- ☐ A. pełny
- ☐ B. zerowy
- ☐ C. zmienny w czasie

Czy silnik krokowy ma szczotki?

- ☐ A. ma

- ☐ B. nie ma
- ☐ C. niektóre mają, inne nie

Najbardziej popularny jest typ silnika krokowego:

- ☐ A. z magnesem trwałym
- ☐ B. o zmiennej reluktancji
- ☐ C. hybrydowy

Konfiguracja unipolarna w porównaniu z bipolarną ma moment obrotowy:

- ☐ A. większy
- ☐ B. mniejszy
- ☐ C. taki sam

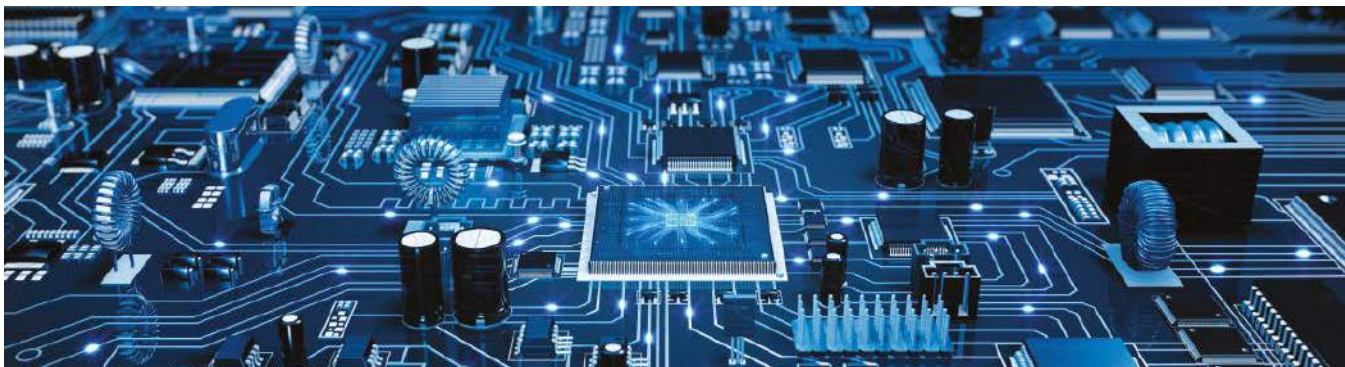
Rozmiary silników krokowych są określane według klasyfikacji NEMA. Czy dwa silniki o identycznym numerze NEMA mają całkowicie identyczne parametry elektryczne?

- ☐ A. tak
- ☐ B. mogą się różnić nieznacznie
- ☐ C. mogą być bardzo różne

Na rynku hobbystycznym silników krokowych są dwa standardy: NEMA, IEC. Który z nich dominuje:

- ☐ A. NEMA
- ☐ B. IEC
- ☐ C. żaden nie przeważa

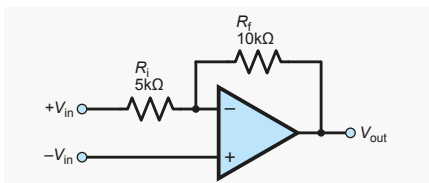
Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 17.12.2022.



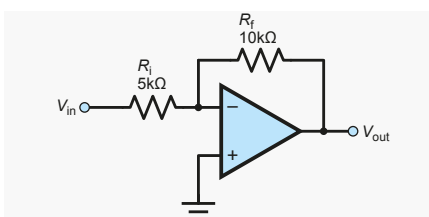
Tensometry i sygnały różnicowe

Scott Siler zamieścił jakiś czas temu na forum EEWeb pytanie dotyczące problemu, jaki pojawił się podczas tworzenia wzmacniacza dla tensometru. „Jestem inżynierem lotnictwa i pracuję nad pobocznym projektem polegającym na zbieraniu danych z tensometru za pomocą NI DAQ (urządzenie do akwizycji danych firmy National Instruments). Sygnał wyjściowy wynosi od 0 do 36 mV DC, więc potrzebuję trzykrotnego wzmocnienia, aby lepiej wykorzystać możliwości DAQ. Moje doświadczenie elektroniczne jest ograniczone do kilku zajęć, które odbyłem podczas moich studiów licencjackich. Mam wzmacniacz operacyjny AD8628 i zbudowałem podstawowy, nieodwracający układ z ujemnym sprzężeniem zwrotnym, jak na załączonym obrazku (rysunek 1). Napięcie zasilające wynosi od 0 do 3 V DC. Wydaje się, że układ nie ma żadnego wzmocnienia, a zamiast tego napięcie wyjściowe jest w rzeczywistości niższe niż wejściowe. Jestem pewien, że musiałem coś źle podłączyć, ale nie jestem w stanie powiedzieć co. Każda pomoc będzie mile widziana”.

W poście brakuje pewnych szczegółów, ale w ramach dyskusji w odpowiedzi na niego, Scott dodał trochę więcej informacji. Używa więc komercyjnego ogniwa obciążnikowego zamiast podstawowej folii tensometrycznej (wyjaśnienie poniżej) z napięciem wzbudzenia 12 V DC. Ogniwo obciążnikowe posiada sygnał wyjściowy w zakresie 0...36 mV, które Scott próbuje zwiększyć do zakresu 0...100 mV, aby lepiej wykorzystać dostępną rozdzielczość możliwości zapisu DAQ. Wykorzystuje on wejście różnicowe DAQ, które ma rozdzielczość 18 bitów. Nie dysponujemy pełnymi danymi dotyczącymi ogniwa obciążnikowego.

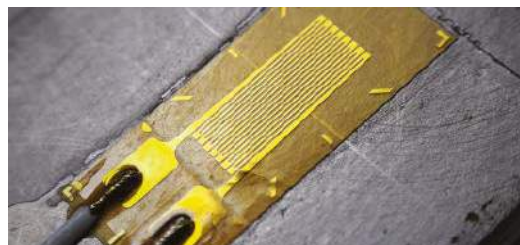


Rysunek 1. Układ Scotta

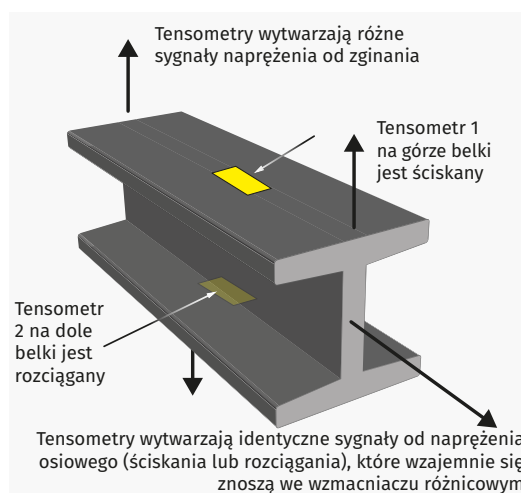


Rysunek 2. Standardowy wzmacniacz odwracający

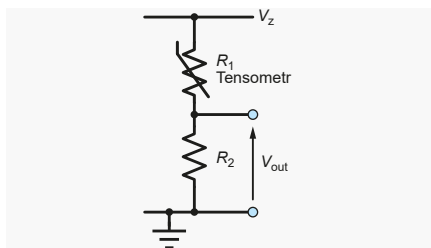
Poruszono więc w tym problemie kilka kwestii: tensometry i sposób ich wykorzystania, źródła błędów w takich układach (i ogólnie obróbka sygnałów z czujników) oraz związane z tym kompromisy projektowe. Jednak chyba najbardziej podstawową kwestią, na którą szybko zwrócono uwagę w dyskusji na forum jest to, że wyjście ogniwa obciążnikowego jest sygnałem różnicowym, a zatem wymaga wzmacniacza różnicowego – układ z rysunku 1 nie jest więc odpowiedni. Jest on bowiem oparty na standardowym wzmacniaczu odwracającym, który zwykle ma pojedyncze wejście (przez rezystor R_i) wraz z wejściem nieodwracającym, uziemionym (jak na rysunku 2) albo podłączonym do innego napięcia odniesienia. Wejścia nie są symetryczne i będą różnie oddziaływać z obwodami zewnętrznymi, zwłaszcza z układami rezystorów, które zakładamy, że są obecne w ogniwie obciążnikowym (patrz dalsza dyskusja). W tym miesiącu przyjrzymy się pokrótce tensometrom, a następnie omówimy sygnały różnicowe i sposoby ich wzmacniania.



Rysunek 3. Typowy tensometr (licencja Wikimedia Commons, Cristian V)



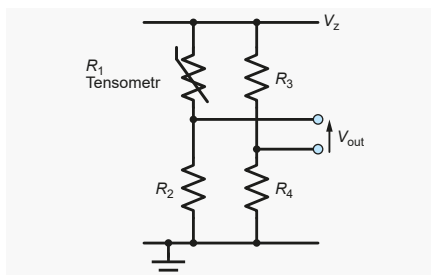
Rysunek 4. Tensometry zamontowane po przeciwnych stronach belki do pomiaru zginania pod obciążeniem. Rozciąganie lub ściskanie osiowe (wzdłuż belki) jest ignorowane, ponieważ tensometry generują identyczne sygnały, które znoszą się wzajemnie we wzmacniaczu różnicowym



Rysunek 5. Tensometr w dzielniku napięcia

Tensometry

Tensometr mierzy odkształcenie (zmianę rozmiaru lub kształtu) obiektu – zmianę położenia określonych punktów względem długości referencyjnej. Odkształcenie może wystąpić w wyniku przyłożonej siły lub zmiany temperatury obiektu. Tensometr jest czujnikiem, którego opór elektryczny zmienia się wraz z odkształceniem. Zwykle składa się on z cienkiej, elastycznej folii izolacyjnej, na której umieszczony jest długi pasek przewodzący, zazwyczaj w układzie zygzakowatym (patrz **rysunek 3**). Na dwóch końcach paska przewodzącego znajdują się styki służące do połączenia z obwodem pomiarowym. Typowa wartość rezystancji wynosi 120 Ω. Gdy tensometr jest przymocowany do obiektu, odkształca się wraz z nim, umożliwiając pomiar odkształcenia obiektu. Aby tensometr działał prawidłowo, mocowanie musi być wykonane w trwały sposób (np. za pomocą odpowiedniego kleju do danego materiału). Tensometry mogą być używane w swojej podstawowej formie do bezpośredniego pomiaru odkształcenia obiektu, na przykład do badania konstrukcji, lub mogą być wbudowane w większy element, zwane „ogniwem obciążnikowym”. W ogniwie obciążnikowym tensometr jest przymocowany do specjalnie zaprojektowanego metalowego korpusu, który odkształca się pod wpływem przyłożonej do niego siły. Ogniwia obciążnikowe są znacznie łatwiejsze w użyciu niż folie tensometryczne i mają wiele zastosowań przemysłowych w pomiarach siły i masy. Wiele tensometrów jest często używanych wspólnie w układach. Na przykład, umieszczenie dwóch tensometrów po przeciwnych stronach pręta lub belki (patrz **rysunek 4**) umożliwia pomiar odkształcenia spowodowanego zginaniem



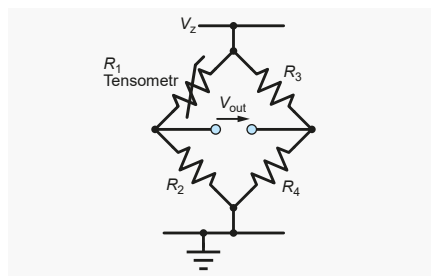
Rysunek 6. Tensometr w układzie mostka – w tym przykładzie tensometrem jest R1, ale można zastosować inne ustawienia

belki. W miarę zginania belki jeden tensometr będzie się rozciągał, a drugi ściszał. Spowoduje to przeciwne zmiany rezystancji, dając w rezultacie silniejszy sygnał niż pojedynczy tensometr. Ponadto, jeśli belka jest również poddawana ścisłaniu lub naprężaniu wzdłuż jej osi, wówczas tensometry zareagują jednakowo i efekt ten może zostać zignorowany przez obwód, który reaguje tylko na przeciwne zmiany w dwóch tensometrach – w ten sposób mierzone jest wyłącznie zginanie. Podobnie, zmiany temperatury, które są potencjalnie problematyczne w przypadku stosowania pojedynczego tensometru, zazwyczaj wpływają na każdego z nich w układzie w równym stopniu i nie mają wpływu na wyjścia z poszczególnych obwodów.

Aby zmierzyć rezystancję (uzyskać sygnał pomiarowy z tensometru) musi przez niego przepłynąć prąd. Prosty sposób na osiągnięcie tego celu jest zastosowanie dzielnika potencjału, jak pokazano na **rysunku 5** – napięcie zasilania (V_z) zapewnia przepływ prądu umożliwiający pomiar. Drugi rezystor dzielnika (R_2) jest zwykle dobierany tak, aby odpowiadał wartości tensometru w połowie zakresu pomiarowego. Napięcie wyjściowe dzielnika potencjału może być doprowadzone bezpośrednio do przetwornika analogowo-cyfrowego (ADC) lub w zależności od potrzeb buforowane przez wzmacniacz.

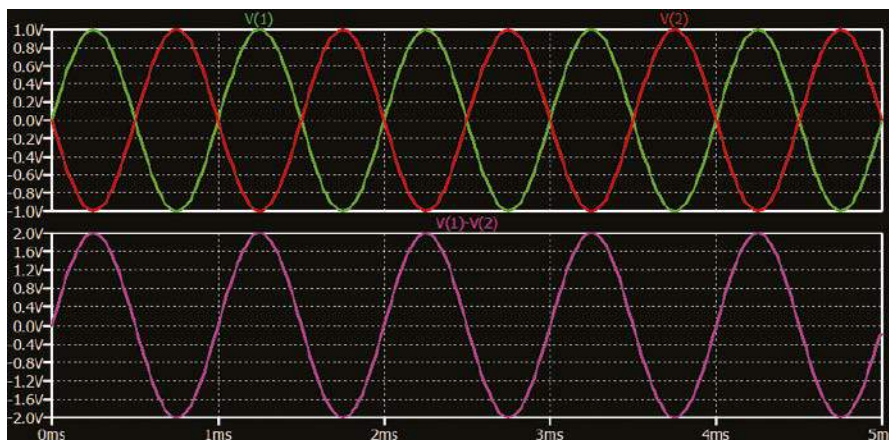
Mostek Wheatstone'a

Pojedynczy dzielnik potencjału użyty z czujnikiem, który wykazuje małe zmiany rezystancji, będzie wytwarzał napięcie wyjściowe, które zmienia się tylko w małym zakresie i to na dodatek z dużą składową stałą. Sygnał taki może być trudny w dalszym użyciu (np. do digitalizacji do użytecznej rozdzielczości). Lepszym podejściem, które umożliwia większą dokładność i jest stosowane w większości aplikacji tensometrycznych,



Rysunek 7. Ten sam obwód co na rysunku 6 narysowany w inny sposób

jest użycie obwodu mostka Wheatstone'a, jak pokazano na **rysunku 6**. Obwód mostka jest też często rysowany w postaci układu pokazanego na **rysunku 7**. Z **rysunku 6** wynika, że mostek to dwa równolegle ustawione dzielniki potencjału, a napięcie wyjściowe jest różnicą napięć obu dzielników potencjału. Napięcie różnicowe z mostka nie ma składowej stałej związanej ze zwykłym dzielnikiem potencjału, więc może być łatwo wzmacnione (przez odpowiedni wzmacniacz z wejściem różnicowym) bez składowej stałej powodującej nasycenie wzmacniacza. Obwód na **rysunku 6** ma jeden tensometr (R_1). Wszystkie pozostałe trzy rezystory w pewnym ustalonym punkcie będą odpowiadać rezystancji tensometru, w takim przypadku napięcie wyjściowe będzie równe zero w tym punkcie, ponieważ oba dzielniki potencjału będą odkładać to samo napięcie (a wyjście jest różnicą między nimi). Sygnał wyjściowy może być zarówno ujemny jak i dodatni, tak jak rezystancja tensometru, która może zmieniać się zarówno w górę, jak i w dół od punktu zerowego wyjścia wyznaczonego przez rezystancje dzielników potencjału. Wyjście ujemne nie jest napięciem poniżej masy – jest to ujemna różnica pomiędzy dwoma dzielnikami potencjału, a to, czy uznasz dane wyjście za ujemne czy dodatnie, będzie zależało od tego, jak określisz oba węzły pod względem polaryzacji. Układ

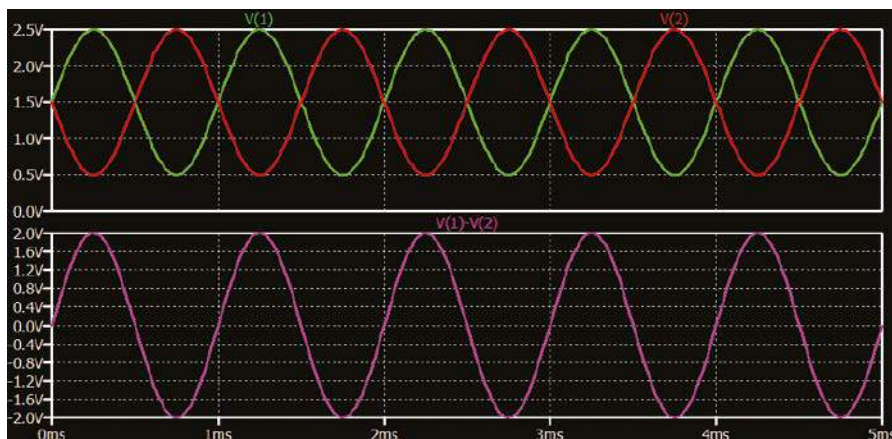


Rysunek 8. Sygnał różnicowy: napięcia na dwóch pojedynczych przewodach V1 i V2 pokazane są na górnym wykresie. Sygnał rzeczywisty jest różnicą pomiędzy V1 i V2 i jest pokazany na dolnym wykresie

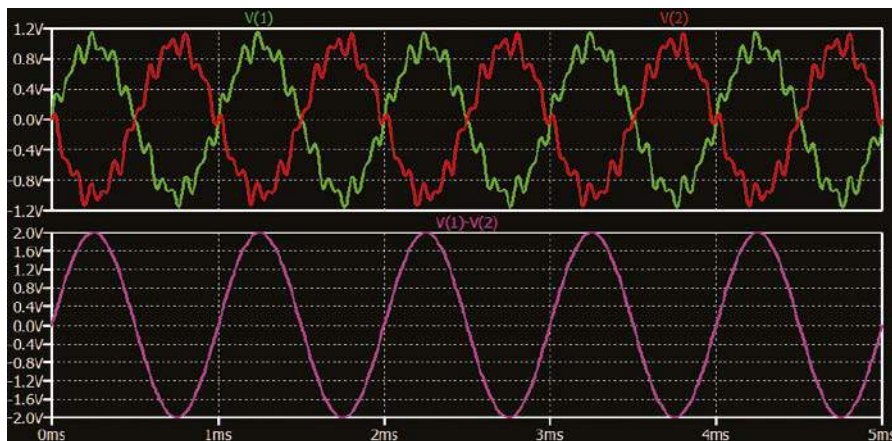
na rysunku 5 ma pojedynczy tensometr z trzema stałymi rezystorami, ale jak już zauważyliśmy, istnieje wiele możliwych scenariuszy pomiarów fizycznych (takich jak przykład z belką zginaną), które wykorzystują wiele tensometrów. Mostek może zawierać jeden, dwa, trzy lub cztery tensometry w zależności

od zastosowanej konfiguracji. W niektórych sytuacjach można zastosować dodatkowe rezystory (np. pomiędzy napięciem zasilania a mostkiem), aby dostroić zachowanie obwodu. I wreszcie, w naszym krótkim omówieniu tych obwodów należy zauważyć, że napięcie zasilania odgrywa ważną rolę w dokładności

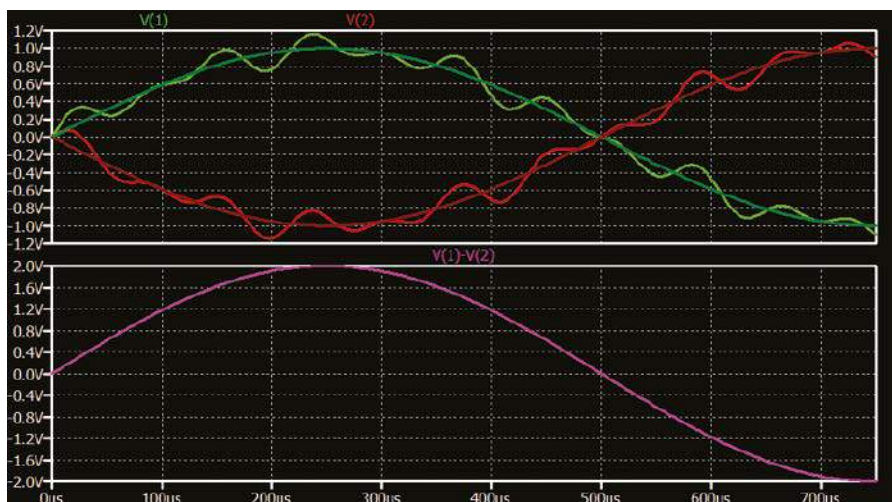
pomiaru – jeśli się zmieni, tak samo zmienia się mierzony sygnał wyjściowy. Jeśli obwód mostka jest używany z przetwornikiem ADC, napięcie zasilania może być użyte jako napięcie referencyjne ADC. Pozwoli to skompensować wszelkie zmiany tegoż napięcia. Alternatywnie, napięcie to może być również mierzone i kompensowane w oprogramowaniu.



Rysunek 9. Wykres przedstawia ten sam sygnał różnicowy co na rysunku 8, ale z innym sygnałem współbieżnym (w tym przypadku 1,5 V DC)



Rysunek 10. Ten sam sygnał różnicowy co na rysunkach 8 i 9, ale z innym sygnałem współbieżnym (w tym przypadku sygnał AC o wyższej częstotliwości)



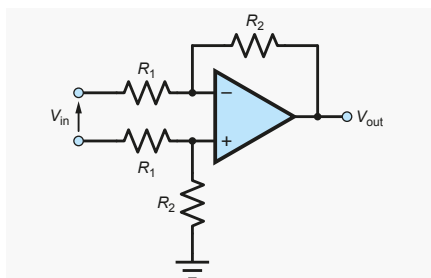
Rysunek 11. Powiększenie pierwszej części rysunku 10, aby wyraźniej zilustrować, że sygnał o wyższej częstotliwości jest sygnałem współbieżnym – tzn. jest równy i w tym samym kierunku na obu przebiegach, a zatem różnica między tymi przebiegami jest taka sama jak sygnał na rysunkach 8 i 9

Sygnały różnicowe

Jak wspomniano, wyjście z obwodu mostka Wheatstone'a jest sygnałem różnicowym. W dalszej części artykułu przyjrzymy się podstawom sygnałów różnicowych – są one bardzo szeroko stosowane w układach elektronicznych, nie tylko w obwodach tensometrycznych. Sygnał różnicowy jest przenoszony na dwóch liniach innych niż masa, więc możemy obserwować napięcie na każdym przewodzie z osobna (np. V_1 i V_2), ale rzeczywisty sygnał jest równy różnicy napięć między tymi dwoma przewodami, każdym mierzonym względem masy. Jeśli więc dwa napięcia na przewodach oznaczymy jako V_1 i V_2 , to sygnał różnicowy wynosi $(V_1 - V_2)$. Sygnały przenoszone tylko na jednym przewodzie są określane jako niesymetryczne, aby odróżnić je od różnicowych. Rysunek 8 przedstawia symetryczny sygnał różnicowy, który jest sinusoidą o częstotliwości 1 kHz i napięciu szczytowym 2 V (4 V napięcia międzyszczytowego). Dwa pojedyncze napięcia (V_1 i V_2 na dwóch przewodach przenoszących sygnał) są pokazane na górnym wykresie – są one równe i przeciwne oraz mają napięcia szczytowe o wartości 1 V. Ponieważ są one przeciwne, szczytowa różnica między nimi wynosi 2 V, co jest amplitudą sygnału różnicowego. Sam sygnał pokazany jest na dolnym wykresie – jest to $(V_1 - V_2)$.

Sygnał współbieżny

Kolejny sygnał różnicowy pokazany jest na rysunku 9 – jest to również sinusoida o częstotliwości 1 kHz i napięciu szczytowym 2 V. Dolny wykres (sygnał różnicowy) jest dokładnie taki sam jak na rysunku 8, jednak sygnały na poszczególnych liniach nie są takie same. Na rysunku 8 zarówno V_1 jak i V_2 są wyśrodkowane względem 0 V, natomiast na rysunku 9 są wyśrodkowane względem 1,5 V. Aby w pełni opisać sygnał różnicowy musimy podać dwie rzeczy – sam sygnał różnicowy, który jest różnicą napięć na dwóch przewodach, oraz napięcie, które mają one wspólne – zwane napięciem współbieżnym, lub wspólnym. Napięcie to jest średnią napięć na dwóch przewodach. Nie musi ono być napięciem stałym – może mieć dowolną formę sygnału lub kształt fali. Na rysunku 10 pokazano przykład dokładnie takiego samego sygnału różnicowego jak na rysunkach 8 i 9, ale



Rysunek 12. Podstawowy układ wzmacniacza różnicowego

z sygnałem współbieżnym o wyższej częstotliwości AC. **Rysunek 11** to powiększenie części jednego cyklu z dodanymi poszczególnymi częściami różnicowej fali sinusoidalnej, aby zilustrować, że sygnał o wyższej częstotliwości jest w tym samym kierunku na obu przebiegach i stąd jest znoszony, gdy brana jest pod uwagę różnica. Często sytuacją w układach elektronicznych jest otrzymywanie pożądanego sygnału różnicowego z niepożądaną składową współbieżną. Przykładem sytuacji, w której takie zjawisko może wystąpić jest sygnał, który jest przenoszony na długich przewodach przez elektrycznie zasumowane środowisko. Jeśli dwa przewody przenoszące sygnał różnicowy biegną blisko siebie, to efekty zewnętrzne (np. szum sieciowy, zakłócenia radiowe...) prawdopodobnie spowodują ten sam sygnał zakłóceń na każdym przewodzie – często nazywany szumem współbieżnym. Jeśli ten sygnał zakłóceń określimy jako Δ , to wtedy napięcie na przewodzie 1 wyniesie $V_1 + \Delta$, a napięcie na przewodzie 2 wyniesie $V_2 + \Delta$. Sygnał użyteczny jest różnicą między dwoma przewodami, czyli $((V_1 + \Delta) - (V_2 + \Delta)) = (V_1 - V_2)$, czyli nie ma w nim sygnału zakłóceń. Aby zakłócenia zewnętrzne wpływały na oba przewody w równym stopniu i ułatwiały ich kompensację, przewody muszą mieć taką samą impedancję – określa się to mianem połączenia symetrycznego. Z tego powodu układy symetryczne i różnicowe są często stosowane w środowiskach o dużym poziomie szumów elektrycznych lub tam, gdzie wymagana jest czystość sygnału. Szum współbieżny może być wyższej częstotliwości niż sygnał (jak na rysunku 10), ale też tej samej lub niższej.

Wzmacniacze różnicowe

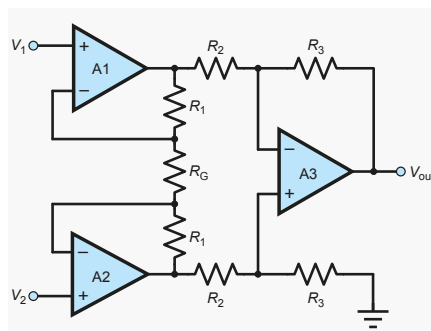
Do wzmocnienia sygnału różnicowego potrzebny jest wzmacniacz z wejściem różnicowym – nazywany, co nie jest zaskakujące, „wzmacniaczem różnicowym”. Wyjście jest często jednoprzewodowe, a wzmacniacz operacyjny (jak widać na rysunkach 1 i 2) jest nam dobrze znany. Chociaż wzmacniacz operacyjny sam w sobie jest wzmacniaczem różnicowym, to powszechnie stosowane układy zbudowane

na jego podstawie – wzmacniacz odwracający (rysunek 2) i wzmacniacz nieodwracający – nie są wzmacniaczami różnicowymi. Wzmacniacze posiadające zarówno wejścia jak i wyjścia różnicowe nazywane są wzmacniaczami w pełni różnicowymi.

Tłumienie sygnału współbieżnego

Wzmacniacz różnicowy o wejściach V_1 i V_2 wzmacnia różnicę między nimi ($V_1 - V_2$) przez swoje wzmocnienie różnicowe A_d . Tak więc dla idealnego wzmacniacza różnicowego sygnał wyjściowy wynosi po prostu $A_d(V_1 - V_2)$. W idealnej sytuacji zmiana wejściowego napięcia współbieżnego nie powinna wpływać na wyjście wzmacniacza różnicowego, ale w praktyce do pewnego stopnia tak się dzieje. Jak zauważono powyżej, składowa współbieżna jest średnią dwóch sygnałów, czyli $(V_1 + V_2)/2$. Jest ona wzmacniana przez wzmocnienie sygnału współbieżnego (A_{cm}) wzmacniacza, dając niepożądany sygnał wyjściowy $A_{cm}(V_1 + V_2)/2$. Im mniejszy wpływ sygnałów współbieżnych na wzmacniacz, tym jest on lepszy. Zdolność wzmacniacza do tłumienia sygnałów współbieżnych wyraża się jako stosunek wzmocnienia różnicowego i wzmocnienia współbieżnego; jest to współczynnik tłumienia sygnału współbieżnego (CMRR), który często wyraża się w decybelach:

$$CMRR = 20 \log_{10} \left| \frac{A_d}{A_{cm}} \right|$$



Rysunek 13. Ogólny wzmacniacz pomiarowy

Typowe wartości CMRR dla wzmacniaczy operacyjnych to 80 do 100 dB (10 000 do 1 000 000). CMRR nie jest jedyną właściwością wzmacniaczy związaną z sygnałami współbieżnymi. Ważny jest również współbieżny zakres napięcia wejściowego. Jeśli duża część sygnału stanowi sygnał współbieżny, to może on rozregulować obwody polaryzujące we wzmacniaczu, nawet jeśli sygnał różnicowy jest bardzo mały. Niektóre wzmacniacze różnicowe radzą sobie z napięciami współbieżnymi zbliżonymi do napięć zasilania, a nawet wykraczającymi poza nie, ale inne mają znacznie bardziej ograniczony zakres. Jest to coś, co zawsze powinno być sprawdzone podczas wyboru lub projektowania wzmacniacza różnicowego. Koncepcja CMRR może być również stosowana w innych sytuacjach. Przykładowo różne konfiguracje tensometrów i obwody mostków mają różne CMRR w odniesieniu do wpływu czynników wspólnych dla kilku folii

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

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

pomiarowych, takich jak zmiany temperatury lub naprężenia w belce.

Wzmacniacz różnicowy

Chociaż wymienione powyżej podstawowe układy wzmacniaczy operacyjnych nie są różnicowe, to oczywiście możliwe jest zbudowanie wzmacniaczy różnicowych z wykorzystaniem takich wzmacniaczy – na **rysunku 12** pokazano podstawowy układ wzmacniacza różnicowego. Należy jednak mieć świadomość, że użycie wzmacniacza operacyjnego o bardzo dobrym CMRR nie gwarantuje, że obwód zbudowany z jego udziałem również będzie miał dobry CMRR. Wzmocnienie różnicowe układu jest dane przez R_2/R_1 – projekt zakłada, że obie wartości R_1 , oraz obie wartości R_2 są dokładnie takie same. Niestety, CMRR układu z **rysunku 11** silnie zależy od tego dopasowania wartości rezystorów. Przy idealnym wzmacniaczu i dokładnie dopasowanych rezystorach CMRR byłby nieskończony (idealny). Jednak w przypadku rzeczywistych rezystorów otrzymujemy zmienność poszczególnych wartości, co pogarsza dopasowanie i zmniejsza CMRR. Dla przykładu rozważmy typowy wzmacniacz różnicowy o wzmocnieniu 100, w którym R_1 to 1 k Ω , a R_2 to 100 k Ω . Jeśli wartość jednego z rezystorów R_1 i jednego z rezystorów R_2 różni się o 5% od drugiego odpowiadającego mu rezystora, to otrzymamy CMRR o wartości tylko 26 dB, nawet przy idealnym wzmacniaczu operacyjnym. Dla 1% niedopasowania wartości rezystorów otrzymujemy około 40 dB CMRR, dla 0,1% około

60 dB, ale potrzebujemy rezystorów 0,01%, aby uzyskać około 80 dB (dolna granica typowych CMRR wzmacniaczy operacyjnych). Tak więc potrzebujemy precyzyjnych (a więc stosunkowo drogich) rezystorów, aby uzyskać nawet w połowie przyzwoitą wydajność tego układu. Problem ten można rozwiązać stosując scaloną wersję wzmacniacza różnicowego, gdzie układ scalony jest produkowany wraz z rezystorami o wymaganej dokładności. Mają one stałe wzmocnienie dzięki wbudowanym rezystorom, ale niektóre mają dostępnych kilka rezystorów wbudowanych, aby zapewnić różne konfiguracje wzmocnienia. Wiele takich elementów jest dostępnych od znanych producentów, takich jak Analog Devices i Texas Instruments.

Wzmacniacz pomiarowy

Innym szeroko stosowanym układem do wzmacniania sygnałów różnicowych jest wzmacniacz pomiarowy. Jest to układ z trzema wzmacniaczami, który jest szeroko stosowany w aplikacjach pomiarowych – stąd też jego nazwa. Ogólny układ wzmacniacza pomiarowego pokazano na **rysunku 13**. Wzmacniacze operacyjne A1 i A2 zapewniają bardzo dużą impedancję wejściową dzięki bezpośredniemu podłączeniu sygnału do wejścia wzmacniacza. A1 i A2 zapewniają również wzmocnienie, ustalone przez wartości rezystorów R_C i R_1 . A3 oraz rezystory R_2 i R_3 tworzą standardowy wzmacniacz różnicowy. A1 i A2 buforują wejścia do wzmacniacza różnicowego (A3) zapobiegając obciążeniu źródła wejściowego. To buforowanie mogłoby być osiągnięte przez dwa izolowane wzmacniacze, ale połączenie przez R_C powoduje, że wyjście z A1 i A2 jest zależne zarówno od V_1 jak i V_2 . Zwiększa to wzmocnienie różnicowe całego układu i poprawia współczynnik CMRR.

Można zbudować wzmacniacz pomiarowy z elementów dyskretnych, ale jego parametry mogą nie być zbyt dobre, zwłaszcza jeśli chodzi o tłumienie sygnału współbieżnego. Podobnie jak układ wzmacniacza różnicowego, wymaga on bardzo ścisłego dopasowania rezystorów (np. dwa rezystory R_3 na **rysunku 13** muszą mieć dokładnie taką samą wartość). Charakterystyki A1 i A2 również powinny być ściśle dopasowane. Podobnie jak w przypadku wzmacniacza różnicowego, dopasowanie to można osiągnąć, gdy cały obwód jest zbudowany jako układ scalony. Wiele takich układów jest dostępnych w asortymencie producentów półprzewodników.

W pełni różnicowe

Wreszcie istnieją wzmacniacze w pełni różnicowe, posiadające zarówno wyjścia jak i wejścia różnicowe. W obwodach, w których błędy muszą być utrzymane na minimalnym poziomie (takich jak systemy pomiarowe) utrzymywanie sygnału w formie różnicowej w całym procesie przetwarzania może zapewnić wiele korzyści. Eliminuje bowiem znaczną ilość niepożądanych zjawisk; jednakże ceną za to jest złożoność układów. Biorąc pod uwagę, że większość systemów pomiarowych dostarcza sygnał z czujnika do przetwornika danych, warto zauważyć, że wiele przetworników ADC ma wejścia różnicowe i są zaprojektowane do pracy z kompatybilnymi, w pełni różnicowymi wzmacniaczami – ta opcja była omawiana w odpowiedzi na forum na post Scotta o jego wzmacniaczu tensometrycznym. ■

Ian Bell

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

Quiz: Tensometry

Do czego służą tensometry?

- ☐ A. do pomiaru gęstości cieczy
- ☐ B. do pomiaru napięcia na słupach elektrycznych
- ☐ C. do pomiaru sił i naprężeń

Co to jest ogniwo obciążnikowe?

- ☐ A. podzespół elektroniczny z wbudowanym czujnikiem tensometrycznym
- ☐ B. bateria ze stałym wewnętrznym obciążeniem
- ☐ C. aktywne obciążenie wzmacniacza audio

Tensometr oporowy działa na zasadzie zmian rezystancji przy rozciąganiu:

- ☐ A. paska kauczuku metalizowanego
- ☐ B. drutu lub folii z przewodnika
- ☐ C. folii z piezoelektryka

Typowa wartość rezystancji tensometru wynosi:

- ☐ A. 12 Ω
- ☐ B. 120 Ω
- ☐ C. 12 k Ω

Typowy układ włączania tensometru to:

- ☐ A. mostek Wheatstone'a
- ☐ B. mostek Thomsona
- ☐ C. mostek Gretza

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 3.12.2022.

W układzie mostka z tensometrem sygnał użyteczny jest nazywany:

- ☐ A. sygnałem współbieżnym
- ☐ B. sygnałem niesymetrycznym
- ☐ C. sygnałem różnicowym

Podstawową zaletą układu mostkowego jest:

- ☐ A. bardzo duży sygnał użyteczny
- ☐ B. niski pobór prądu
- ☐ C. możliwość kompensacji sygnałów niepożądanych

Współczynnik CMRR wzmacniaczy operacyjnych wynosi zwykle:

- ☐ A. 40 do 80 dB
- ☐ B. 80 do 100 dB
- ☐ C. ponad 100 dB

Aby CMRR realnego układu wzmacniacza różnicowego osiągnęto wartość 80 dB, rezystory w tym układzie powinny mieć tolerancję:

- ☐ A. 0,01%
- ☐ B. 0,1%
- ☐ C. 1%

Wzmacniaczem w pełni różnicowym nazywamy:

- ☐ A. wzmacniacz z wejściem różnicowym i wyjściem niesymetrycznym
- ☐ B. wzmacniacz z wejściem niesymetrycznym i wyjściem symetrycznym
- ☐ C. wzmacniacz z wejściem i wyjściem różnicowym



MT-7071 680zł

- Akcesoria w zestawie:
- patchcord RJ45
 - patchcord RJ11
 - przewód z krokodylkami
 - słuchawki
 - etui



MT-7071K 699zł

- Akcesoria w zestawie:
- 8 dodatkowych pilotów zdalnych
 - patchcord RJ45
 - patchcord RJ11
 - przewód z krokodylkami
 - słuchawki
 - etui

Nadajnik:

- testowane typy przewodów:
 - RJ45 LAN Cat 5, 5e, 6, 7 (UTP/STP),
 - RJ11/12 tel. Cat 3 (2/4/6 pin),
 - kable koncentryczne
 - kable domofonowe
- sposób pomiaru: metoda pojemnościowa
- maksymalna odległość transmisji: 3 km (1 kHz)
- max długość testowanego przewodu LAN 300m
- test ciągłości obwodu
- identyfikacja stanu linii telefonicznej
- wymiary 138×80×35mm

Odbiornik:

- gniazdo słuchawkowe
- NCV - bezdotykowy detektor napięcia AC (AC90~1000V)
- wymiary 198×45×33 mm

Zdalny pilot:

- złącza: RJ45 (8 pin), RJ11/12 (6 pin), BNC
- wymiary 90×32×30mm



identyfikacja
przewodów



test kabli LAN,
telefonicznych
i koncentrycznych



pomiar długości
przewodu; lokalizacja
punktu przerywania



beziprzewodowa
detekcja napięcia



podświetlenie
miejsca pracy

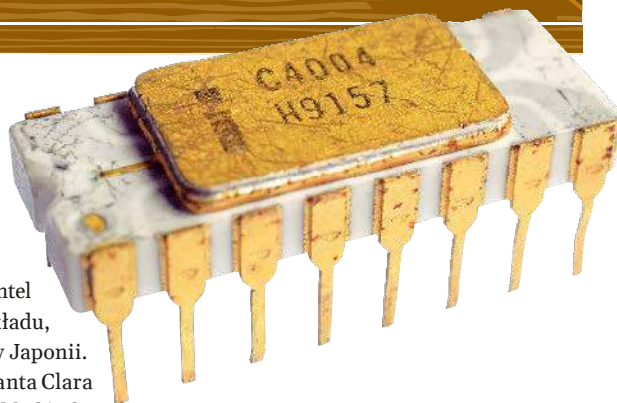
Patronat EdW nad szkołami i uczelnianymi Kołami Naukowymi rozkwita i daje redakcji EdW impulsy zachęcające do wspierania edukacji szkolnej i uczelnianej. Działa sprzężenie zwrotne. Dostajemy mnóstwo listów od uczniów, nauczycieli i studentów. Dla nich jest ta rubryka, składająca się z form edukacyjnych:

• Wykład • Konsultacje



A. Historia, czyli lektura lekka na rozgrzewkę

**Prezentujemy artykuł opublikowany w 2019 roku w „Elektro-
nice Praktycznej” z okazji jubileuszu 50-lecia mikroprocesora.**



Nie do wiary, to już pół wieku?! Wprawdzie pierwszy mikroprocesor, słynny 4004, Intel wprowadził na rynek wiosną 1971 roku, to jednak koncepcja tego układu, wbrew amerykańsko-centrycznej wizji świata, zrodziła się wcześniej w Japonii. W czerwcu 1969 r. Dr Masatoshi Shima z firmy Busicom zawiązał w Santa Clara i przedstawił w firmie Intel zamówienie na wyprodukowanie zestawu układów logicznych do kalkulatora planowanego do produkcji w Busicom. Schematy, które przywiózł ze sobą Japończyk, były rewolucyjne – zrywały z programowaniem wszytym w hardware, proponując architekturę komputerową z jednostką arytmetyczną o dostępie swobodnym, realizującą zadania logiczne zarządzane przez software. Koncepcja architektury logicznej, rozdzielającej dane od adresów i zawierającej blok CPU (Central Processor Unit), blok pamięci ROM, blok pamięci RAM i blok interfejsów peryferyjnych, zakładała de facto budowę małego komputera na układach scalonych. Właśnie 4-bitową jednostką centralną z tego projektu stał się pierwszy na świecie mikroprocesor oznaczony jako 4004. Japończycy wiedzieli, co chcą osiągnąć, ale nie mieli do tego technologii. Mieli ją Amerykanie w firmie Intel, która wówczas istniała dopiero od niespełna roku, czyli według obecnych pojęć była start-upem. Skąd w startupie najnowocześniejsza na świecie technologia? Otóż przywędrowała w głowach założycieli Intela, a wcześniej pracowników firmy Fairchild – Gordona Moore’a i Roberta Noyce’a, za którymi przyszli do Intela czołowi technologowie Fairchilda – A. Grove i L.L. Vadasz. Była to technologia SGT (Silicon Gate Technology), opracowana w 1967 roku w firmie Fairchild przez Federico Faggina, który również przeszedł do Intela w roku 1970.

Intel zaczął się od zdrady Fairchilda przez jego najlepszych technologów, której przyczyną bezpośrednią były niezaspokojone ambicje Roberta Noyce’a – chciał zostać prezesem Fairchilda, a właściciele firmy mieli inne plany. Nie mogąc się powstrzymać, by nie napisać w tym miejscu paru zdań o historii kilku zdrad, które stały się motorem postępu technologicznego uwieńczonego powstaniem mikroprocesora.

Zaczął się od „zdradzieckiej ósemki”, tj. przejścia w 1957 roku ośmiu najzdolniejszych uczniów odkrywcy tranzystora Williama Shockleya z jego firmy Shockley Semiconductor Laboratory do firmy Fairchild. Tej ósemce przewodził Robert Noyce, który 11 lat później wraz z Gordonem Moore’em odszedł „na swoje”, zakładając firmę Intel. Obie „zdrady” miały swoje przyczyny. W 1957 r. R. Noyce z siedmioma kolegami mieli dość apodyktycznego szefa (W. Shockley był geniuszem – laureatem Nagrody Nobla za odkrycie



Gordon Moore



Robert Noyce



Andrew S. Grove (András István Gróf)



Leslie László Vadász

Fot. Steve Jurvetson from Menlo Park, USA
– Permission to Speak Freely, CC BY 2.0

tranzystora – ale człowiekiem bardzo trudnym we współpracy) i chcieli rozwijać prace nad technologią krzemową, z czym nie zgadzał się W. Shockley. Przejęcie „zdradzieckiej ósemki” do Fairchilda uczyniło z tej firmy, produkującej wcześniej przyrządy optyczne, światową potęgę w technologii półprzewodnikowej. Na początku lat sześćdziesiątych dwie firmy przewodziły światowej mikroelektronice: Fairchild w segmencie układów scalonych analogowych i Texas Instruments w segmencie układów scalonych cyfrowych, przy czym w innowacjach technologicznych zdecydowanym liderem światowym był Fairchild. Wynalazkiem fundamentalnym, bez którego nie byłaby możliwa produkcja układów scalonych, okazała się technologia planarna, użyta przez R. Noyce’a w roku 1959 do wyprodukowania pierwszego układu scalonego. Układy scalone produkowane w pierwszej połowie lat sześćdziesiątych bazowały na tranzystorach bipolarnych, chociaż już w 1963 r. były dobrze znane i opisane właściwości tranzystora MOS, w szczególności inwertera CMOS i doskonale zdawano sobie sprawę z kolosalnej przewagi tych tranzystorów nad bipolarnymi w zastosowaniu do układów scalonych. Jednak rozwój technologii układów scalonych MOS napotkał na przykrą przeszkodę w postaci ich niestabilności czasowej. W rozwiązaniu tego problemu największe zasługi ma dr Bruce Deal pracujący w firmie Fairchild. Pierwsze układy scalone w technologii PMOS z bramką aluminiową nie dość, że mogły po jakimś czasie przestać pracować (niestabilność), to wymagały zasilania napięciem 24 V i były dość powolne. Dopiero opracowanie technologii NMOS z tak zwaną samocentrującą bramką polikrzemową (SGT) zwiększyło szybkość działania układów MOS około 5-krotnie i pozwoliło zmniejszyć napięcie zasilania do poziomu kompatybilnego ze standardem zasilania (+5 V) bardzo wówczas popularnych układów cyfrowych TTL.

Autorzy „drugiej zdrady” – R. Noyce i G. Moore – mieli gotową technologię do ekspansywnego rozwoju produkcji układów scalonych LSI (zawierających ponad 1000 tranzystorów) w ich własnej firmie Intel. Zaczęli od pamięci RAM 1 kbit, a wizyta dr. Shimy z zamówieniem na mikroprocesor otworzyła przed ich firmą fantastyczne perspektywy. Trzeba było tylko przeciągnąć z Fairchilda do Intela twórcę technologii NMOS Silicon-gate Federico Faggina, który okazał się również wybitnie uzdolnionym projektantem na poziomie elementarnych układów logicznych. Japoński pomysł wyartykułowany na poziomie schematów blokowych architektury został doprowadzony do postaci projektów układu 4004 na poziomach logicznym, elektrycznym i technologicznym. Dokonał tego Federico Faggin i w topologii ścieżek czipa tego układu wpisał swoje inicjały FF, uważając siebie, nie bez podstaw, za twórcę pierwszego mikroprocesora.

Można też mówić o „trzeciej zdradzie”, a może po prostu o amerykańskiej mobilności, gdyż po kilku latach F. Faggin odszedł z Intela i założył firmę Zilog produkującą kultowy mikroprocesor Z80, stosowany w legendarnych minikomputerach ZX Spectrum, Amstrad, Commodore 128 (jako drugi procesor).

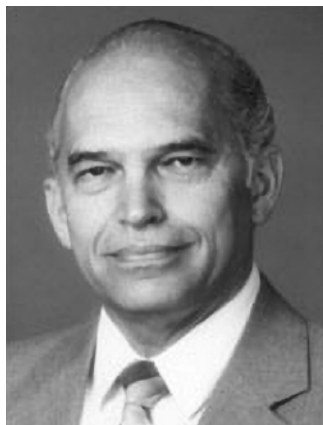
Bohater tej opowieści – mikroprocesor 4004 – nie odegrał wielkiej roli. Japończycy zastosowali go w swoim kalkulatorze i użyto go jeszcze w kilku innych aplikacjach, ale dopiero jego następca 8-bitowy Intel 8008 i kolejny 8080 na dobre otworzyły nieograniczone możliwości aplikacji i rozwoju elektroniki.

Układ 4004 zawierał 2300 tranzystorów i miał ścieżki o szerokości 10 μm . Po 50 latach w technologii CMOS wymiar charakterystyczny jest 1000-krotnie mniejszy, czyli wynosi 10 nm, a gęstość upakowania jest milion razy większa, czyli układy scalone zawierają miliardy tranzystorów. A jednak podstawowe idee nie zmieniły się – nadal rzeźbimy ścieżki w krzemie i budujemy wielkie systemy z klocków znanych ponad 50 lat temu, tj. z inwerterów CMOS.

Jak już tyle napisałem, to może jeszcze dodam wątek polski i mój osobisty. W 1973 roku na IV Konferencji Mikroelektroniki w Toruniu gościliśmy czterech przedstawicieli firmy Fairchild, wśród nich był wspomniany już wcześniej dr Bruce Deal i dr Harry Sello (uczeń i ulubieniec W. Shockleya). Nikt tak pięknie nie potrafił opowiadać o pierwszym mikroprocesorze jak Harry Sello. Nie była to wizyta przypadkowa. Od wielu miesięcy uzgadniane były 2 licencje CEMI (Centrum Mikroelektroniki na Służewcu w Warszawie, gdzie obecnie jest galeria handlowa) z Fairchildem, w tym jedna na układy scalone MOS LSI. Tak się złożyło, że organizatorzy konferencji powierzyli mi eskortowanie dwóch Amerykanów w podróż pociągiem z Torunia do Warszawy. Jednym z nich był dr B. Deal – mój guru, gdyż mój doktorat i habilitacja, którą wówczas robiłem, dotyczyły niestabilności struktur MOS, tj. tematyki, w której dr Deal był światowym autorytetem. Rozmowa z nim była dla mnie wielką gratką. Okazało się szybko, że ja znam każdy jego artykuł, a on nie wie nic o moich pracach. Żartuję, rzecz jasna, choć przecież miałem publikacje w pismach światowych, a w artykułach Philipsa powoływano się na moją metodę diagnostyczną. W połowie drogi do Warszawy, gdzieś w okolicach Kutna, stało się jednak jasne, że dr Deal nie tylko nie słyszał o moich pracach, ale w ogóle nie zna żadnych prac poza swoimi. Wtedy zrozumiałem, że poza Silicon Valley jest pustynia, którą nie warto się interesować.



William Shockley



Bruce E. Deal



Federico Faggin



Harry Sello

A z dr. Harrym Sello, który w Toruniu miał porywający wykład na temat mikroprocesora, choć ciągle reprezentował firmę Fairchild, spotkałem się po latach. Około 1990 r. dr H. Sello wizytował Instytut Technologii Materiałów Elektronicznych, którym wówczas kierowałem. Poza rozmową biznesową oddawaliśmy się nostalgicznym wspomnieniom konferencji w Toruniu i opłakiwaliśmy okrutny los firmy Fairchild, która z potęgi światowej stoczyła się do niebytu. My friend Harry (z Amerykanami można się zaprzyjaźnić błyskawicznie) wspominał m.in. swoją misję we Włoszech. Otóż wśród wielu nietrafionych działań zarządu Fairchilda była nieudana próba zbudowania joint venture z Włochami. Na paroletnią „zsyłkę” do Włoch wysłano Harry’ego, o czym chętnie wspominał, opowiadając różne anegdoty. Na przykład taką: „zwolniłem z pracy jegomościa, który uprawiał seks w miejscu pracy. Następnego dnia złożyła mi wizytę jego żona i zwymyślała mnie najgorszymi wyzwiskami. A ja wszystko zrozumiałem. Byłem zachwycony, że po 2 latach pobytu tak dobrze opanowałem język włoski”. Taki był Harry. Zmarł 2 lata temu i już nikt na świecie nie opowie Wam tak jak on o pierwszych mikroprocesorach.

Dodam, że z licencji Fairchilda przygotowanej w 1973 roku, ostatecznie nic nie wyszło. Gdy w Pentagonie dowiedzieli się, że kraj komunistyczny dostanie najnowszą technologię (2–3 lata po USA), to złapali się za głowy i prezes Fairchilda Lester Hogan zapłacił za to stołkiem prezesowski, a gotowy już kontrakt został unieważniony. Młodym Czytelnikom wyjaśniam, że w czasach zimnej wojny w USA ustalano normy sygnowane przez COCOM (Coordinating Committee for Multilateral Export Control), które precyzyjnie określały, jakie technologie są dostępne dla krajów komunistycznych, tak aby utrzymywać stałe opóźnienie technologiczne tych krajów wobec USA od 10 do 15 lat. Normy te obowiązywały cały Zachód i nie były to żarty. Osobiście znałem biznesmena z zachodniej Europy, który z nami niejawnie współpracował i w USA czekał na niego wyrok 150 lat. Dodam dla jeszcze większej jasności, że nasz Wielki Brat (ZSRR) też stosował wobec nas embargo technologiczne, choć czynił to obłudnie pod płaszczykiem koordynacji prac rozwojowych z planami militarnymi w ramach Układu Warszawskiego. Może kiedyś napiszę w jakich warunkach moje pokolenie Polaków podejmowało wysiłki, by nadążyć w technologii, łamiąc embargo amerykańskie w tajemnicy przed ZSRR. Kluczową rolę w tych działaniach motywowanych patriotyzmem odgrywały służby, którym szczególnie satysfakcję sprawiało wyprowadzanie w pole Radzianów, jak pogardliwie nazywali „bratnie” służby ZSRR. Ale to temat historyczny na inną okazję. ■

B. Meritum, czyli sedno tematu

To nasz pierwszy wykład w nowej rubryce „Edukacja w EdW”, więc nie od rzeczy będzie skomentować przyjęty na stałe podział wykładu na dwie części:

- A. – Historia, czyli lektura lekka na rozgrzewkę,
- B. – Meritum, czyli sedno tematu.

Część A mamy już za sobą.

Zaczynając część B mamy ambitny cel trafić w samo sedno. Im większe sedno, tym łatwiej weń trafić. Genialnie ujął to Wojciech Młynarski w zacytowanym obok wierszu „Sedno” – „Czy to sedno się zrobiło takie duże?” Żeby nie ułatwiać sobie nadmiernie zadania spróbujmy jednak ograniczyć sedno.

O czym będzie mowa na tym wykładzie?

Mikroprocesory (μP) i mikrokontrolery (μC , lub inaczej MCU) to temat na opasłą książkę.

Na pewno nie będziemy omawiać zagadnień fizyko-technologicznych, związanych z wytwarzaniem μP i MCU. Zajmiemy się tym tematem w przyszłości, w ramach wykładu o układach scalonych. Nie będziemy też omawiać zagadnień układowych, ani na poziomie układów elementarnych, ani na poziomie architektury. Nie zamierzamy też dotykać tematu programowania μP i μC . To coś nam pozostaje?

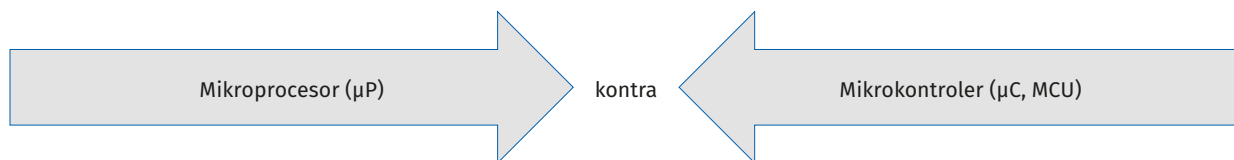
Wiersz Wojciecha Młynarskiego „Sedno”

Miłe Panie i Panowie bardzo mili,
czy to w upał, czy w czas mroźnych gołedzi
coraz częściej słyszę da się
w telewizji, radio, prasie
cieple, szczerze i otwarte wypowiedzi.
A w tych wypowiedziach, wyznam wam, panowie,
nieodmiennie wciąż uderza mnie to jedno,
to uderza mnie, panowie, że ktokolwiek się wypowie,
trafia swoją wypowiedzią w samo sedno.
Już od dawna nie słyszałem wypowiedzi
co by błędu zawierała choć zarodek,
czasem zda mi się na oko,
że ktoś myli się głęboko,
a on sunie bez wysiłku w sedno srodek!
Celnym owych wypowiedzi słucham przeto,
rozmyślając sobie nie raz przy ich wtórze,
czy ci ludzie w jednej chwili
tacy mądrzy się zrobili,
czy to sedno się zrobiło takie duże?

(tekst wg piosenki:
<https://www.youtube.com/watch?v=Y0N46aDPf2M>)

Tylko to, co najbardziej jest potrzebne konstruktorowi, czyli Czytelnikowi EdW. Zajmiemy się przeglądem μP i μC , o których warto wiedzieć, jeśli się kreatywnie projektuje urządzenia elektroniczne z zastosowaniem tych układów. Przed tym jednak wyjaśnijmy sobie parę spraw terminologicznych.

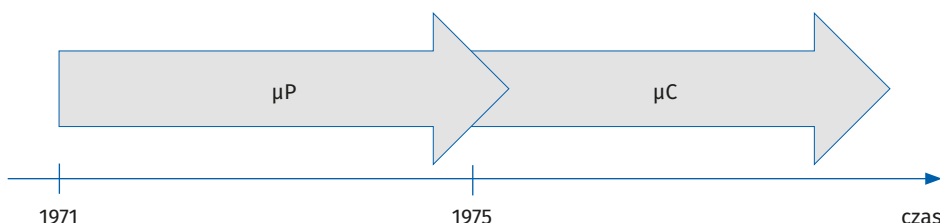
Mikroprocesor kontra mikrokontroler



Rysunek 1.

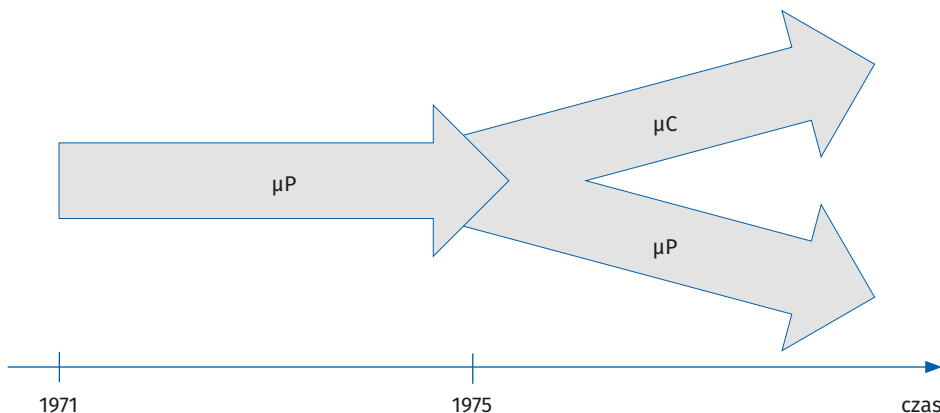
Czy właściwe jest słowo „kontra”, czyli przeciwstawianie μP i μC ? Chyba nie, przecież to jedna rodzina układów, a μC zrodził się z μP w miarę rozwoju technologii i niejako μC zawiera w sobie μP .

To może właściwa byłaby grafika jak na **rysunku 2**.



Rysunek 2.

Mikroprocesor zawiera w sobie tylko procesor (CPU – Central Processor Unit) i trochę rejestrów pamięci, a w miarę rozwoju technologii do jednego krzemowego chipa udało się „zapakować” zarówno procesor jak też pamięć RAM, ROM oraz interfejsy Wejścia/Wyjścia (I/O). w ten sposób powstał jednoukładowy komputer (single chip computer), któremu dano nazwę mikrokomputer, później zmieniony na mikrokontroler. Można by sądzić, że dalej rozwijały się już tylko μC , bo po co produkować μP , które wymagają dołączenia zewnętrznych układów pamięci i układów I/O, jeśli mamy już układy zawierające cały system komputerowy w jednym chipie. Otóż to nie jest tak. Ewolucja przebiegała w taki sposób jak na **rysunku 3**.



Rysunek 3.

Mikrokontrolery wyłoniły się z rozwoju technologicznego mikroprocesorów i ciągle rozwijają się, ale mikroprocesory też ciągle się rozwijają. Jak to jest możliwe? Otóż μP i μC rozwijają się niezależnie w dwóch całkowicie różnych obszarach zastosowań. Mikroprocesory są stosowane w komputerach ogólnego przeznaczenia, natomiast μC w komputerkach przeznaczonych do zagnieżdżenia (embedded) w określonych urządzeniach (pralka, samochód, itp.), w których spełniają funkcje pomiarowo-sterownicze, stąd przyjęła się nazwa mikrokontroler, a nie pierwotnie stosowana mikrokomputer. A jeśli tak różne są pola zastosowań μP i μC , to zapewne ich parametry i funkcje różnią się zasadniczo, zatem pierwsza grafika ze strzałkami w kontrze (**rysunek 1**) całkiem nieźle oddaje rzeczywistość. Podsumujmy więc te różnice:

- μP jest sercem komputera ogólnego przeznaczenia, głównie w komputerach osobistych, natomiast μC jest wyspecjalizowanym jednocipowym komputerkiem zagnieżdżonym w określonym sprzęcie (urządzeniu, gadżecie). Uwaga – zawsze są jakieś odstępstwa od reguły i tak jest w tym przypadku. Niektóre typy mikroprocesorów są wytwarzane do zastosowań zagnieżdżonych.
- wobec znacznej ceny komputera osobistego μP – kluczowy podzespół komputera – nie musi być bardzo tani, natomiast μC to na ogół element groszowy, stosowany często w bardzo tanich gadżetach (oczywiście cena zależy od parametrów układu i bez trudu można znaleźć μP tańsze od μC).

Spytacie, dlaczego tak jest? Przecież μP zawiera tylko procesor, a μC ma w sobie zarówno procesor jak i pamięci oraz interfejsy I/O, zatem μC powinien być droższy niż μP . Tylko procesor procesorowi nie jest równy. Procesor w μP to mózg, który musi szybko rachować (przetwarzać wielką liczbę instrukcji), by podołać uniwersalnym zadaniom peceta biurkowego czy laptopa. Natomiast procesor w μC to mózdzek realizujący ograniczony zakres zadań urządzenia, w którym jest zagnieżdżony. Stąd wynika szereg różnic, o których warto wiedzieć:

- μP działa z zegarem o większych częstotliwościach (obecnie do kilku GHz, a w μC ok. kilkuset MHz)
- μP czerpie zasilanie w pececie i jego moc nie jest tak krytycznym parametrem jak moc μC zasilanego często z małątkiej baterijki. W rozwoju μC cały czas idzie walka o jak najmniejszy pobór mocy, natomiast dla μP liczy się coraz większa szybkość działania, a moc ograniczają tylko warunki odprowadzania ciepła.

Działanie μP opiera się na architekturze von Neumanna, a μC stosuje architekturę harwardzką, choć nie zawsze tak musi być.

No tak, mamy dwa nowe terminy do wyjaśnienia.

Różnice między architekturą von Neumanna i harwardzką pokazują **rysunki 4, 5**. W architekturze von Neumanna dane i instrukcje są przekazywane między pamięcią i procesorem tą samą magistralą. Może to powodować ograniczenia transferu wspólną magistralą. To ograniczenie usuwa architektura harwardzka, w której pamięć danych jest oddzielona od pamięci instrukcji i transfer odbywa się odrębnymi magistralami danych i instrukcji.

Jeszcze słów parę o mikroprocesorach multiprocesorowych, które poprawniej nazywane są multirrdzeniowymi (multi-core). Siłą napędową rozwoju mikroprocesorów jest ciągłe dążenie do zwiększania ich mocy obliczeniowej. Do pewnego czasu postęp osiągnano przez zwiększanie częstotliwości taktowania procesora. Ta ścieżka postępu natrafiła na barierę – ograniczenie możliwości odprowadzania ciepła. Każde przełączenie z 0 na 1 lub odwrotnie dokonuje się fizycznie w podstawowej cegielce układów logicznych, tj. w inwerterze. Dokonuje się to przez przełączanie kluczy tranzystorowych, a każde przejście tranzystora ze stanu przewodzenia do nieprzewodzenia, lub odwrotnie łączy się z poborem prądu w procesie przejściowym. To główna przyczyna strat energii w układach cyfrowych i wydzielania ciepła. Jeśli zwiększamy częstotliwość taktowania μP , to wzrasta liczba przełączeń inwerterów w czasie sekundy, czyli wydziela się coraz więcej ciepła. Układy scalone zawierające setki milionów lub miliardy tranzystorów podczas pracy z ich pełną szybkością stają się małymi grzejnikami i nie pomagają żadne metody chłodzenia. Dlatego na poziomie częstotliwości taktowania ok. 3 GHz zatrzymał się marsz ku coraz większym mocom obliczeniowym przez zwiększanie częstotliwości taktowania. Wymyślono inną drogę postępu – równoległą pracę wielu procesorów, a właściwie wielu rdzeni, to jest tych części procesora, w których realizuje się operacje obliczeniowe.

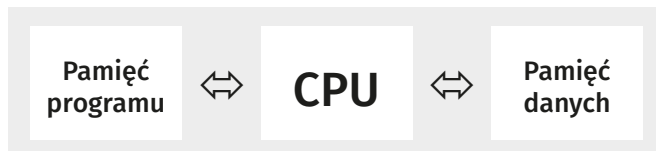
Na przykład μP dwurdzeniowy może wykonywać operacje dwukrotnie szybciej bez zwiększania częstotliwości taktowania. Można powiedzieć, że w przeszłości mikroprocesory miały tylko jeden rdzeń, choć określenie rdzeń pojawiło się wtedy, gdy zaczęto multiplikować zdolności obliczeniowe μP przez zrównoleżenie operacji. Na tym nie koniec pomysłów w dziedzinie multiplikacji zdolności obliczeniowych. Wymyślono też wirtualny rdzeń, kojarzony z pojęciem informatycznym nazywanym wątkiem (ang. thread). Wątek to potok (strumień) instrukcji. Otóż do jednego rdzenia można skierować jednocześnie dwa wątki, tj. dwa potoki instrukcji, które będą realizowane w reżimie współdzielonym. Nie daje to podwojenia zdolności obliczeniowych, ale wzrost mocy obliczeniowej o 10–30% też jest nie do pogardzenia.

Wróćmy jeszcze do porównania ceny, która przecież łączy się również z kosztami wytwarzania chipów. Super szybki i wydajny μP , szczególnie multiprocesorowy, zawiera o rzędy wielkości więcej tranzystorów niż prosty μC , a więc powierzchnia jego struktury krzemowej jest znacznie większa. To oznacza, że z jednej płytki krzemowej można wyprodukować mniej takich złożonych mikroprocesorów, niż prostych mikrokontrolerów. W związku z tym koszt wytworzenia samej struktury krzemowej μP jest wyższy niż dla struktury μC . Dodajmy, że obudowa μP jest też bardziej kosztowna, gdyż połączenia magistrali danych i programu z zewnętrznymi pamięciami wymagają wielu końcówek obudowy. To też czynnik kosztowy składający się na wyższą cenę μP .

Popularne μC można kupić już w cenie kilku złotych (**rysunek 6**), podczas gdy mikroprocesor kosztuje od kilkudziesięciu do kilkuset złotych (**rysunek 7**). Mikroprocesory wieloprocessorowe, a właściwie moduły procesorów o najbardziej wysrubowanych parametrach do zastosowań militarnych i kosmicznych (odporne na promieniowanie) mogą kosztować nawet kilka do kilkudziesięciu tysięcy dolarów (**rysunek 8**).



Rysunek 4. Architektura von Neumanna



Rysunek 5. Architektura harwardzka



Rysunek 6. Przykładowy mikrokontroler 8-bitowy MCU 8051 (50 MHz, 64 kB Flash, 4,25 kB RAM) dostępny w cenie ok. 8 zł



Rysunek 7. Przykładowy mikroprocesor do PC-tów, Intel Core i5 11400(F), dostępny w cenie ok. 800 zł

Mikrokontrolery – jak wybierać i co wybrać

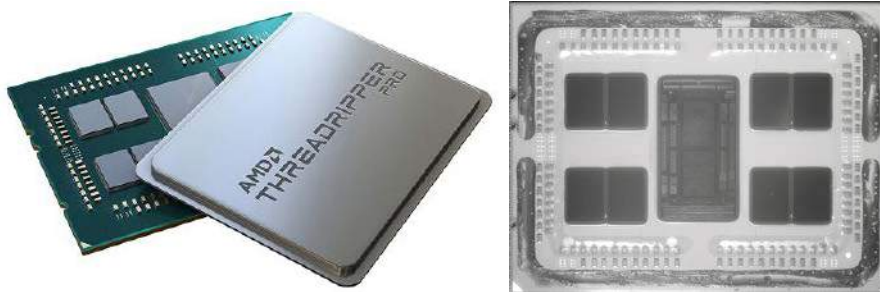
Zostawiamy mikroprocesory i dalej będziemy zajmować się wyłącznie mikrokontrolerami, jeśli nie trafi nam się jakiś mikroprocesor do zastosowań zagnieżdżonych. Wiedza o mikroprocesorach jest potrzebna konstruktorom komputerów, pecetów i laptopów, a mało prawdopodobne jest, by czytelnicy EdW podejmowali takie zadania. Temat dobrego wyboru

mikrokontrolera jest też mało istotny dla hobbistów software'owych, którzy bazują na płytkach Arduino, Raspberry Pi, itp. Zakładamy jednak, że Czytelnicy EdW są zainteresowani projektowaniem układów optymalnych dla wielu różnorodnych zastosowań, nie tylko w pojedynczych egzemplarzach, ale też produktów seryjnych. Jest taka pokusa, żeby pójść drogą „na skróty”, tj. wybrać mikrokontroler bez starannej analizy wymagań, jakie stawia system/urządzenie, w którym mikrokontroler ma realizować swoje zadania pomiarowe i kontrolne. Nie idźmy tą drogą. Wszak różnorodność zastosowań mikrokontrolerów jest przeogromna. Są stosowane w pojazdach, robotach, sprzęcie biurowym, medycznym, komunikacyjnym, domowym, zabawkach itd. Praktycznie w każdym urządzeniu/sprzęcie mającym „inteligencję” znajdziemy MCU. Dlatego zanim zaczniemy wybierać μC najpierw trzeba mieć projekt systemu/urządzenia i wynikającą stąd specyfikację wymagań dla mikrokontrolera, które porządkujemy według określonych kryteriów. By uniknąć nieporozumień wróćmy jeszcze na chwilę do zdania o płytkach Arduino i Raspberry Pi. To zupełnie różne platformy sprzętowe, przeznaczone do realizacji różnych zadań projektowych. Arduino to płytka uruchomieniowa z prostym mikrokontrolerem, natomiast Raspberry Pi to w istocie całkowicie funkcjonalny komputer z systemem operacyjnym i procesorem o potężnej (w porównaniu do MCU w Arduino) mocy obliczeniowej. Na przykład Raspberry Pi 3 ma 64-bitowy, 4-rdzeniowy procesor o częstotliwości taktowania 1200 MHz, podczas gdy wiele wersji Arduino zadowala się mikrokontrolerem 8-bitowym o częstotliwości taktowania 16 MHz.

Kryteria wyboru mikrokontrolera w 7 krokach

1 Sporządzamy listę interfejsów ze światem zewnętrznym.

Chodzi o interfejsy (porty) fizyczne (inaczej sprzętowe) z urządzeniami peryferyjnymi, służące do wymiany informacji między tymi urządzeniami i mikrokontrolerem. Jeśli na przykład μC ma zbierać dane analogowe z kilku sensorów, to musi mieć odpowiednią liczbę przetworników ADC, a do sterowania szybkością obrotów silnika DC musi mieć interfejs PWM. Wiedza o interfejsach jest dość rozległa i dostępna z wielu źródeł. W tym miejscu ograniczymy się do rozróżnienia dwóch rodzajów interfejsów. Pierwszy to interfejsy komunikujące mikrokontroler z urządzeniami zewnętrznymi, tj. USB, I²C, UART, SPI (patrz słowniczek akronimów). Te interfejsy mają decydujący wpływ na wymaganą objętość pamięci programów w mikrokontrolerze. Drugi rodzaj interfejsów to wejścia i wyjścia cyfrowe, wejścia ADC, PWM itp. Z listy interfejsów wynika liczba pinów (nóżek) mikrokontrolera, chociaż w zależności od przyjętych rozwiązań MCU mają od kilku do kilkudziesięciu pinów. Żeby zademonstrować, jak ważne jest sporządzenie listy interfejsów pokazujemy przykładową tabelkę interfejsów 8-pinowego mikrokontrolera STM8S001 (tabela 1) i przypisanie peryferiów/portów do pinów tego mikrokontrolera (rysunek 9).



Rysunek 8. Przykładowy procesor o wielkiej mocy obliczeniowej do zastosowań militarnych, AMD Ryzen Threadripper PRO 64 2,7 GHz, dostępny w cenie ok. 40000 zł

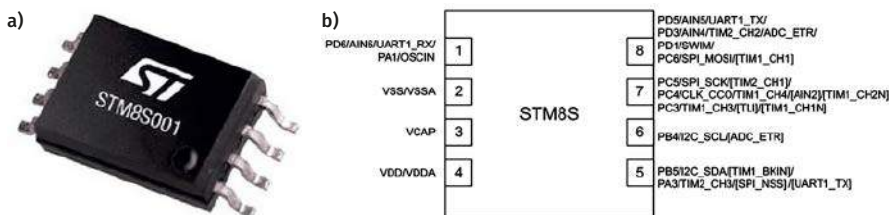
Porównanie μP z μC w skrócie encyklopedycznym

Mikroprocesor

- duża moc obliczeniowa (do 10000 MIPS)
MIPS – milion instrukcji na sekundę
- wielordzeniowość
- przetwarzanie wielopotokowe
- dostępne instrukcje zmiennoprzecinkowe
- wymaga dodatkowych peryferii (kontroler DMA, pamięć programu i danych, kontroler przerwań)
- praca w trybie rzeczywistym i wirtualnym
- zastosowania główne: komputery (PC, stacje robocze, serwery)
- główni producenci mikroprocesorów: IBM, Intel, AMD, Motorola

Mikrokontroler

- mała moc obliczeniowa (do 30 MIPS)
- jeden procesor (rdzeń)
- przetwarzanie jednopotokowe
- na ogół niedostępne instrukcje zmiennoprzecinkowe
- zawiera pamięci i peryferia (liczniki i układy czasowe, przetworniki A/C i C/A, interfejsy)
- praca tylko w trybie rzeczywistym
- zastosowania główne: zagnieżdżone w różnorodnych urządzeniach służące do pomiarów i sterowania
- producenci mikrokontrolerów – patrz ramka „Najbardziej popularne rodziny mikrokontrolerów”



Rysunek 9. a) 8-pinowy mikrokontroler STM8S001, b) przypisanie peryferiów/portów do wyprowadzeń w układzie STM8S001J3

Tabela 1. Lista portów μ C STM 8S001J3

Numer pinu na obudowie S08	Nazwa pinu	Domyślna funkcja	Funkcja alternatywna	Funkcja alternatywna „remap”
1	PD6/ AIN6/ UART1_RX	Port PD6	Analog input 6/UART1 data receive	–
	PA1/ OSCIN	Port PA1	External clock input (HSE)	–
5	PA3/ TIM2_CH3 SPI_NSS/ UART1_TX	Port PA3	Timer 2 channel 3	SPI master/slave select/ UART1 data transmit
	PB5/ I2C_SDA/ TIM1_BKIN	Port PB5	I ² C data	Timer 1 break input
6	PB4/ I2C_SCL/ ADC_ETR	Port PB4	I ² C clock	ADC external trigger
7	PC3/ TIM1_CH3/ TLI/ TIM1_CH1N	Port PC3	Timer 1 channel 3	Top Level Interrupt/ Timer 1 channel 1 inverted
	PC4/ CLK_CCO/ TIM1_CH4/ AIN2/ TIM1_CH2N	Port PC4	Configurable clock output/ Timer 1 channel 4	Analog input 2/ Timer 1 channel 2 inverted
	PC5/ SPI_SCK/ TIM2_CH1	Port PC5	SPI clock	Timer 2 channel 1
8	PC6/ SPI_MOSI/ TIM1_CH1	Port PC6	SPI master/slave in	Timer 1 channel 1
	PD1/ SWIM	Port PD1	SWIM debug interface	–
	PD3/ AIN4/ TIM2_CH2/ ADC_ETR	Port PD3	Analog input 4/ Timer 2 channel 2/ ADC external trigger	–
	PD5/ AIN5/ UART1_TX	Port PD5	Analog input 5/ UART data transmit	–

2 Szacujemy, jak dużą **moc obliczeniową** powinien mieć mikrokontroler.

Punktem wyjścia są wymagania stawiane przez software. Należy określić niezbędną szybkość przetwarzania. Zdecydować, czy wystarczy jeden rdzeń, czy powinny być dwa, co jednak będzie oznaczać większy pobór mocy (parametr krytyczny dla zastosowań w IoT z zasilaniem baterijnym). Wymagania na dużą moc obliczeniową zdecydowanie rosną, gdy na przykład potrzebne jest przetwarzanie grafiki (GPU – Graphics Processing Unit). Zagadnienie mocy obliczeniowej łączy się z takimi parametrami jak **częstotliwość taktowania** i **długość słowa**. Zegary mikrokontrolerów osiągają częstotliwości taktowania do kilkuset MHz, jednak na ogół taktowanie może być znacznie powolniejsze. Dla ustalenia częstotliwości taktowania istotniejsza niż moc obliczeniowa jest wymagana szybkość transmisji danych. **Długość słowa** może wynosić 4, 8, 16, 32 lub 64 bity. Najbardziej popularne w zastosowaniach hobbystycznych są mikrokontrolery 8 i 16-bitowe. Nadal stosowane są 4-bitowe, choć to architektura prehistoryczna sprzed pół wieku. W szczególnie wymagających projektach amatorskich stosuje się też mikrokontrolery 32-bitowe. Niekiedy wybór architektury jest procesem iteracyjnym. W pierwszym przybliżeniu może wystarczać 16 bitów, ale po dokładniejszym przeanalizowaniu wszystkich potrzeb dla danej aplikacji może się okazać, że mikrokontroler 32-bitowy jest lepszym wyborem.

3 Szacujemy **potrzeby pamięci**.

W mikrokontrolerze mamy dwa rodzaje pamięci – **Flash** i **RAM**. Pamięć Flash to rodzaj pamięci nieulotnej, przechowującej program, czyli rozkazy (dawniej były

Podręczny słowniczek akronimów

Rodzaje pamięci stałych w μ C

- **ROM** (Read Only Memory) – użytkownik nie ma możliwości zmiany zawartości pamięci ROM, może ją tylko odczytywać. Jest zaprogramowana przez producenta kodem firmowym oprogramowania dostarczanego wraz z μ C. Są to najczęściej: program diagnostyczny, program ładowania, program uruchomienia i nadzorowania programów użytkownika (program monitora).
- **EPROM** (Erasable Programmable ROM) – służą do przechowywania programów i danych użytkownika. Zawartość można skasować naświetlaniem (przez okienko) UV i zapisać nową zawartość specjalnym programatorem.
- **OTP** (One Time Programmable) – programowana jednokrotnie elektrycznie przez zmianę jedynek na zera.
- **EEPROM** (Electrically Erasable PROM) – pamięć nieulotna, którą można przeprogramować wielokrotnie elektrycznie (bez naświetlania). Przechowuje dane i programy użytkownika.
- **FLASH** (bulk erasable non-volatile memory) – dominująca obecnie technologia pamięci nieulotnej o cechach podobnych do EEPROM, ale jest znacznie szybsza i zapewnia większą liczbę cykli zapisu i kasowania.

Interfejsy μ C

- **GPIO** – General-Purpose Input/Output – wejście/wyjście ogólnego przeznaczenia
- **SPI** – Serial Peripheral Interface – szeregowy interfejs urządzeń peryferyjnych
- **UART** – Universal Asynchronous Receiver-Transmitter – uniwersalny asynchroniczny nadajnik-odbiornik
- **USB** – Universal Serial Bus – uniwersalna magistrala szeregową
- **PWM** – Puls – Width Modulation – modulacja szerokości impulsów
- **I²C** – Inter – Integrated Circuit – pośrednik między układami scalonymi, jest to szeregową, dwukierunkową magistralą służącą do przesyłania danych w urządzeniach elektronicznych
- **ADC** – Analog Digital Converter – przetwornik analogowo-cyfrowy
- **DAC** – Digital Analog Converter – przetwornik cyfrowo-analogowy

w tym celu stosowane różne rodzaje pamięci ROM – patrz słowniczek akronimów). Określenie nieulotna oznacza, że informacje przechowywane w tej pamięci nie znikają po wyłączeniu zasilania układu. Wymagana pojemność tej pamięci zależy od potrzeb (wielkości) zapisywanego w niej programu. Dodajmy, że w pamięci Flash oprócz programu przechowywane są również tablice stałych i w segmencie bootloadera wgrywane są nowe programy bez użycia programatora. Pamięć RAM (Random Access Memory) to pamięć operacyjna o dostępie swobodnym, w której zapisywane są dane chwilowe zmiennych, które są traczone po wyłączeniu zasilania układu.

Określając niezbędne pojemności pamięci Flash i RAM warto mieć na uwadze pewien zapas na ew. korekty czy modyfikacje programowe w przyszłości.

4 Wybieramy architekturę mikrokontrolera.

Architektura mikrokontrolera odnosi się do jego struktury wewnętrznej. Dwa podstawowe rodzaje architektury to:

- architektura von Neumanna (rysunek 4)
- architektura harwardzka (rysunek 5)

W **architekturze von Neumanna** wspólną magistralą przesyłane są instrukcje programu (z pamięci Flash do CPU), jak też dane (między pamięcią RAM i CPU w obu kierunkach).

W **architekturze harwardzkiej** są dwie oddzielne magistrale do przesyłania instrukcji programu z pamięci Flash i danych z i do pamięci RAM. Architektura harwardzka ma istotną zaletę w porównaniu ze znacznie starszą (z roku 1945) ideą von Neumana. Otóż architektura harwardzka pozwala uzyskać tę samą szybkość obliczeń przy mniejszej częstotliwości taktowania (operacje na instrukcjach z pamięci i danych przebiegają równocześnie, podczas gdy w architekturze von Neumanna niezbędny jest podział czasu między tymi operacjami). Mniejsze częstotliwości taktowania to mniejszy pobór mocy, co jest szczególnie ważne w projektach z zasilaniem bateryjnym.

Warto dodać, że architektura von Neumanna też ma swoją istotną zaletę – jest nią mniejsza powierzchnia struktury krzemowej (chipa), co przekłada się na niższą cenę mikrokontrolera. Z punktu widzenia programisty istotny jest podział mikrokontrolerów na dwie architektury różniące się typem listy instrukcji:

- RISC (Reduced Instruction Set Computer)
- CISC (Complex Instruction Set Computer)

Architektura RISC oznacza mniejszą liczbę instrukcji (kilkadziesiąt) i przetwarzanie potokowe, co przyczynia się do zwiększenia szybkości wykonywania programu. Jedna instrukcja jest wykonywana w jednym cyklu zegara. **Architektura CISC** operuje dużo większą liczbą instrukcji (kilkaset). Jeden rozkaz jest wykonywany w kilka do kilkunastu cyklach zegara, zatem działanie μC jest powolniejsze niż w architekturze RISC. Natomiast zaletą architektury CISC jest wspieranie języków wysokiego poziomu, co ułatwia programowanie. Dlatego w niektórych popularnych procesorach (np. firm Intel i AMD) stosuje się takie rozwiązanie, że z punktu widzenia programisty są widziane jako CISC, ale ich rdzeń jest typu RISC. A legendarny mikrokontroler 8051 przeszedł ewolucję przez 40 lat z architektury CISC do RISC w niektórych wersjach w ostatnich latach.

5 Koszt mikrokontrolera staje się istotny, gdy projektujemy układ do produkcji seryjnej. Na koszt wpływa wiele parametrów, poza już omówionymi, jak liczba portów, moc obliczeniowa, szybkość taktowania, pojemności pamięci, również parametry techniczne – rodzaj



www.elportal.pl – filmy Mouser Electronics o procesorach RISC-V

RISC-V: o co chodzi?

W ostatnich latach wielkim rozgłosem cieszy się nowa architektura listy instrukcji procesora (ISA – Instruction Set Architecture), opracowana w laboratoriach Uniwersytetu Kalifornijskiego w Berkeley w 2010 roku i rozwijana w trybie open – source pod nazwą RISC-V. Chodzi o przyciągnięcie maksymalnej liczby projektantów hardware'u i software'u do rozwoju tej architektury, aby wyciągnąć wnioski z błędów popełnionych przy tworzeniu innych ISA procesorów. Magnesem przyciągającym tak duże zainteresowanie jest koncepcja umożliwiająca każdemu projektowanie, produkcję i sprzedaż chipów i oprogramowania RISC-V. Wywołało prawdziwe pospolite ruszenie, powstała Fundacja RISC-V, której członkami są uniwersytety i firmy technologiczne, między innymi Google, IBM, Microsoft, NVIDIA i Oracle). Celem tych działań jest stworzenie stabilnego ekosystemu, który uczyni procesory o dużej mocy obliczeniowej bardziej dostępnymi i użytecznymi, również w zastosowaniach zagnieżdżonych. Ważne jest osiągnięcie kompatybilności przestrzeni adresowych 32-bitowych, 64-bitowych i 128-bitowych. Już we wstępnej fazie eksploatacji rdzenia o architekturze RISC-V w Uniwersytecie Princeton wykryto ponad 100 błędów, wpływających na nieprawidłowe komunikowanie się procesora z pamięcią. W kolejnych wersjach produkcyjnych wykryte błędy są usuwane. Dodajmy, że V w akronimie to nie jest litera, tylko rzymskie 5, bo była to piąta generacja architektury RISC, opracowana w Uniwersytecie Kalifornijskim.

obudowy, moc rozpraszana, zakres temperatury pracy, napięcie zasilania (5 V lub 3,3 V). stąd rozrzut cen różnych rodzajów mikrokontrolerów jest potężny. W literaturze amerykańskiej mają na to zgrabną frazę: „można kupić sto chipów za jednego dolara lub jeden chip za sto dolarów”.

6 Wsparcie społeczności i producenta.

Nie należy się śpieszyć za bardzo z aplikowaniem najnowszych typów mikrokontrolerów. Roztropnie jest odczekać kilka miesięcy, aż testy użytkowników na wielu różnorodnych aplikacjach ujawnią ewentualne niedoróbki producenta i rozwiną się źródła znacznych zasobów oprogramowania. Bardzo ważne znaczenie ma też dostępność opracowanego przez producenta kitu ewaluacyjnego. Dysponując płytką rozwojową można sprawnie i szybko wykonać prototyp projektu. Wreszcie warto przy wyborze mikrokontrolera brać pod uwagę renomę firmy producenta.

7 Bezpieczeństwo.

W niektórych zastosowaniach (na przykład w samochodzie) mikrokontrolery mogą być narażone na ataki hakerskie. Producenci oferują specjalne mikrokontrolery mające zabezpieczenia kryptograficzne lub fizyczne (hardware'owe) przed atakiem hakerskim. Są to μ C certyfikowane zgodnie z najnowszymi standardami bezpieczeństwa. Takie mikrokontrolery można znaleźć w ofercie firm: ST Microelectronics, Microchip, NXP, Texas Instruments, Renesas, Cypress, Infineon.

Co wybrać – przegląd mikrokontrolerów popularnych w społeczności hobbistów

Wybór mikrokontrolera do projektu opiera się głównie na:

- dostępności zakupu pojedynczych egzemplarzy lub niewielkich serii,
- dostępności tanich narzędzi programistycznych i programatorów,
- popularności określonych typów mikrokontrolerów w społeczności pasjonatów elektroniki.

Ten ostatni czynnik (popularność) daje dość ograniczony obraz rynku mikrokontrolerów na świecie. Jest ich znacznie więcej niż tylko układy AVR, ARM, ESP, PIC, które najczęściej są stosowane w projektach hobbistycznych. Spróbujmy rozejrzeć się nieco szerzej, a okaże się, że dorobek trwającej już prawie pół wieku ery mikrokontrolerów jest znacznie bogatszy i, o dziwo, ciągle są dostępne i chętnie stosowane niektóre „prehistoryczne” rodziny mikrokontrolerów. Oczywiście, nie ma już technologii sprzed 50 lat, więc te stare, zasłużone typy mikrokontrolerów są wytwarzane współczesną technologią i znacznie górują swoimi parametrami nad ich protoplastami. Technologia przez 50 lat rozwijała się według słynnego prawa Moore'a. Od lat siedemdziesiątych do chwili obecnej wymiary liniowe w układzie scalonym (szerokości ścieżek i wyodrębnionych warstw) zmniejszyły się ok. 1000 razy, od kilku mikrometrów do kilku nanometrów. Zatem powierzchnia struktury krzemowej tranzystora zmniejszyła się 1000^2 , czyli milion razy. Stąd skala integracji wzrosła ok. milion razy, tj. w pojedynczym chipie są obecnie miliardy tranzystorów, a 50 lat temu były to tysiące. Pierwszy komercyjnie dostępny mikrokontroler firmy Texas Instruments TMS 1000 (w patencie USA nazwany microcomputer), który pojawił się w 1974 roku, zawierał 8000 tranzystorów (4-bit, 28-pin). W dostępnym już w 2017 roku 8-rdzeniowym mikroprocesorze Ryzen 5 (rysunek 10) CPU na powierzchni struktury krzemowej ok. 2 cm² zawiera 4,8 miliarda tranzystorów. Jest to oczywiście mikroprocesor 64-bitowy. A w module Apple M1 Ultra, na dwóch chipach mieści się 114 miliardów tranzystorów. Wydawałoby się, że już dawno powinniśmy zapomnieć o architekturach 4-bitowych, czy 8-bitowych z lat siedemdziesiątych.

A jednak...

4 bity wiecznie żywe

Ciągle na rynku są oferowane mikrokontrolery 4-bitowe. Dlaczego? Można by przypuszczać, że decydują o tym koszty. Im mniejsza powierzchnia struktury krzemowej tym więcej chipów uzyskuje się z jednej płytki krzemowej, a więc powinna spadać cena jednostkowa chipa. Jednak ceny nie są szczególnie niskie – kilka do kilkadziesiąt USD/szt. przy minimalnych zamówieniach 10 szt. Prawdopodobnie producenci (małe firmy tajwańskiej i chińskiej – Nyquest, Cicotex, Tens Technology, CR Micro) liczą głównie na nabywców poszukujących części zamiennych dla starych urządzeń (na przykład w sprzęcie medycznym).

Kultowy 8051

Mikrokontroler 8051 został opracowany w firmie Intel i wprowadzony na rynek w roku 1980. Zapoczątkował rodzinę mikrokontrolerów MCS-51 (Micro Computer System) – udoskonalanych do różnorodnych zastosowań, ale zawsze z tym samym rdzeniem o ośmiobitowej architekturze z kompatybilną w stosunku do pierwowzoru listą rozkazów. Od 2007 roku Intel nie produkuje już chipów z serii MCS-51, ale kilka firm nabyło od Intela licencję i pochodne 8051 są ciągle produkowane w różnych wersjach, a wielu konstruktorów z lubością projektuje nowe urządzenia z tym mikrokontrolerem. Można to zrozumieć. Przecież przez 40 lat użytkownicy 8051 opracowali bogate oprogramowanie i know-how, a tego nie wyrzuca się tylko dlatego, że pojawił się lepszy mikrokontroler. Zresztą obecnie produkowane 8051 biją na głowę szybkością działania te z lat osiemdziesiątych. Niektóre obecnie produkowane wersje 8051 pracują z częstotliwością zegara 450 MHz, podczas gdy 40 lat temu była to częstotliwość 12 MHz.

8051 ma strukturę zmodyfikowanej architektury harwardzkiej z listą rozkazów CISC, chociaż mikroarchitektura MCS 8051 jest chronioną własnością intelektualną Intela. Liczba 12 cykli na instrukcję wskazuje na listę rozkazów CISC, ale są też najnowocześniejsze wersje 8051, które w opisie danych mają „1 T cykl instrukcji”, co oznacza, że wymagają tylko jednego cyklu na instrukcję, a to cecha listy rozkazów RISC. W starych wersjach 8051 przy częstotliwości zegara 12 MHz i 12 cyklach na instrukcję mieliśmy szybkość działania 1 milion instrukcji na sekundę. Natomiast w najnowszych wersjach o częstotliwości zegara 450 MHz i 1 cyklu



Rysunek 10. Mikroprocesor Ryzen 5 zawiera 4,7 miliarda tranzystorów

na instrukcję, mamy szybkość działania 450 milionów instrukcji na sekundę. Różnica oszałamiająca. Obecnie mikrokontrolery 8051 są bardzo tanie, więc znajdują zastosowanie w gadżetach produkowanych masowo, takich jak szczoteczki do zębów, klawiatury, piloty, zabawki, myszy komputerowe, itp.

Wiele firm produkuje 8051 jako samodzielny chip lub jako rdzeń na licencji IP (Intellectual Property – własność intelektualna). Można tu wymienić zarówno pierwszoligowe firmy światowe – Atmel, Infineon, Philips, Siemens, jak i wiele firm azjatyckich – np. Sonix, WCH, STC, Silan.

Z80

Mikroprocesor Z80 odegrał wielką rolę w minikomputerach lat siedemdziesiątych i osiemdziesiątych, ale też ciągle jest stosowany jako rdzeń mikrokontrolerów do zastosowań w systemach zagnieżdżonych. Opracowany przez zespół pod kierownictwem Federico Faggina (uznanego też za twórcę pierwszych mikroprocesorów 4004, 4040 oraz mikroprocesora 8080) pojawił się na rynku w 1976 jako pierwszy produkt nowej firmy Zilog, założonej przez grupę byłych pracowników Intelu. Sukces był oszałamiający do tego stopnia, że start-up Zilog już po roku mógł utworzyć własną fabrykę chipów, zatrudniającą 1000 osób. Anegdotyczną ciekawostką jest umieszczenie w layoutie topologii chipa sześciu zbędnych tranzystorów „dla zmylenia przeciwnika”. Była to pułapka dla firm uprawiających „reverse engineering”, tj. kopiujących cudze projekty.

Później się dowiedziano, że japońska firma NEC, słynna z klonowania projektów amerykańskich bez licencji, straciła pół roku na rozgryzienie zagadki 6-ciu tranzystorów. Zresztą Z80 też w dużym stopniu „czerpał” z projektu mikroprocesora 8080, opracowanego wcześniej przez Faggina w firmie Intel. Kompatybilność programowa Z80 z 8080 była wielkim atutem, bo użytkownicy Z80 od razu mogli korzystać z bogatego środowiska programowego, jakie wcześniej powstało dla 8080. Z80 był licencjonowany, kopiowany i klonowany przez wielu producentów na całym świecie, między innymi powstały klony w krajach bloku komunistycznego (NRD, Rumunia, ZSRR). Zaowocowało to ogromną bazą użytkowników. Z czasem powstały różne wersje pochodne, jak np. rodziny eZ80 (zegar do 50 MHz), Z8, czy eZ8 Encore. Rdzeń oparty na architekturze Z80 zastosowano w wielu mikrokontrolerach zagnieżdżonych w sprzęcie telekomunikacyjnym, konsolach do gier, różnych gadżetach elektronicznych. Zarówno mikroprocesory Z80 jak i mikrokontrolery z rdzeniem Z80 są ciągle produkowane i bardzo popularne.

Trzeci muszkieter – 6502

Do legendarnej trójki mikroprocesorów, które zrewolucjonizowały elektronikę lat siedemdziesiątych i osiemdziesiątych, poza 8051 i Z80 trzeba zaliczyć 6502 – produkt Motoroli. Wszyscy trzej muszkieterowie żyją i mają swoich fanów nie tylko „20 lat później”, ale nawet 40 lat później, czyli tu i teraz. Historia 6502 jest podobna jak Z80. W przypadku Z80 grupa projektantów mikroprocesora 8080 odeszła z Intelu mając zaawansowany projekt nowego mini mikroprocesora, który zrealizowała pod nazwą Z80 w nowej firmie Zilog. Podobna historia, w tym samym roku 1975, zdarzyła się w firmie Motorola, gdzie grupa projektantów (ośmiu z siedemnastu-osobowego zespołu, który opracował świetny mikroprocesor 6800) odeszła z prawie gotowym projektem nowego mikroprocesora, który zrealizowała pod nazwą 6502 w firmie MOS Technology. Warto dodać, że był to okres ciężkiego kryzysu w amerykańskim przemyśle półprzewodnikowym. W kryzysie zwykle powstają napięcia między „rozwojowcami”, którzy chcą środków na realizację ich wspaniałych pomysłów, a zarządami firm ledwo wiążącymi koniec z końcem. Tak też było w tym przypadku. Zarząd Motoroli podjął błędną decyzję o wstrzymaniu prac nad nowym mikroprocesorem i wielki sukces 6502 stał się udziałem innej firmy (MOS Technology), która przyciągnęła zespół konstruktorów. 6502 jest mikroprocesorem 8-bitowym z 16-bitową szyną adresową. Celem konstruktorów było opracowanie mikroprocesora znacznie tańszego od obecnych na rynku Intel 8080 i Motorola 6800. Aby to osiągnąć zawalczono o jak najmniejszą powierzchnię struktury chipa. Udało się zmieścić układ na powierzchni ok. 16 mm², prawie dwukrotnie mniejszej niż dla 6800. Układ 6502 wszedł na rynek w cenie ok. jednej szóstej ceny konkurentów. Nie dziwne, że zdobył rynek. Wkrótce został użyty w słynnych komputerach tamtych lat – Commodore 64, Atari, BBC Mikro, Apple II. Jego udoskonalona wersja w technologii CMOS, o niskiej mocy, wciąż ma się dobrze, mimo że jest dość kosztowna. Wiele firm stosuje w swoich produktach licencyjny rdzeń 6502. Jest on wykorzystywany zarówno w mikrokontrolerach jak też w układach niestandardowych FPGA (Field-Programmable Gate Array) i ASIC (Application-Specific Integrated Circuit).

Dużą popularnością cieszą się mikrokontrolery firmy WDC ze zmodyfikowanym rdzeniem 6502, np. mikrokontroler W65C2655 z 16-bitowym procesorem W65C816S, który jest w pełni kompatybilny z 8-bitowym W65C02S, pracującym już przy zasilaniu 1,8 V.

Fascynujące nazwy AVR, ARM, PIC

AVR to rodzina mikrokontrolerów wprowadzona w 1996 roku do produkcji w firmie Atmel, przejętej w 2016 roku przez firmę Microchip. Atmel twierdzi, że AVR to nie jest akronim i nie oznacza nic konkretnego, jednak ogólnie uznaje się, że AVR pochodzi od imion dwóch studentów norweskich, twórców architektury AVR (Alf-Egil Bogen, Vegard Wollan). Zatem AVR można rozwinąć jako Alf, Vegard, RISC. Inną cechą architektury AVR było zastosowanie zmodyfikowanej architektury harwardzkiej, w której dane i programy są przechowywane w odrębnych pamięciach z różnymi przestrzeniami adresowymi, ale za pomocą specjalnych instrukcji możliwy jest odczyt danych z pamięci programu. Poza tym w mikrokontrolerach 8-bitowych AVR zastosowano innowacyjne w połowie lat dziewięćdziesiątych rozwiązania – architektura RISC i pamięci Flash, zamiast stosowanych wówczas powszechnie pamięci ROM, EPROM i EEPROM. Mikrokontrolery AVR stały się sławne w świecie hobbistów po ich zastosowaniu w płytkach Arduino.

A co oznaczają równie popularne akronimy ARM i PIC? Akronim ARM oznacza Advanced RISC Machines, a pierwotnie Acorn RISC Machine, od nazwy brytyjskiej firmy projektowej Acorn Computers, która zaprojektowała nową architekturę procesora i w roku 1985, we współpracy z firmą technologiczną VLSI Technology, wyprodukowała procesory ARM1. Ogromne zainteresowanie innych producentów tym rozwiązaniem, przede wszystkim firmy Apple, spowodowało, że w roku 1990 zespół projektowy firmy Acorn został wyodrębniony i utworzył nową firmę ARM Ltd. Ta nowa firma we współpracy z Apple rozwijała koncepcję procesorów z architekturą

ARM i udzielała wielu firmom licencje na IP (własność intelektualna). Do chwili obecnej powstało ponad 20 modeli architektury ARM, w ostatnich latach oferowanych jako rdzeń Cortex. Na licencjach ARM ponad 20 firm produkuje mikroprocesory i mikrokontrolery. Ocenia się, że ponad 75% rynku 32-bitowych CPU w systemach wbudowanych stanowią procesory z architekturą ARM (telefony komórkowe, laptopy, routery, itd.).

Akronim PIC pierwotnie oznaczał Peripheral Interface Controller. Powstał w firmie General Instrument w 1976 roku jako dodatkowy mikrokomputer (jak wtedy nazywano układy dziś nazywane mikrokontrolerami) współpracujący z produkowanym przez General Instrument pierwszym na świecie 16-bitowym mikroprocesorem CP 1600. Ostatecznie CP1600 nie utrzymał się na rynku, za to PIC pozostał jako samodzielny produkt. W roku 1985 od firmy General Instrument odpęczkował dział technologiczny tworząc nową firmę Microchip. W tej nowej firmie PIC rozkwitał, ale akronim zmienił znaczenie na Programmable Intelligent Computer. Microchip sprzedaje obecnie rocznie ok. miliarda układów PIC w bardzo szerokiej gamie wersji, 8-bit, 16 bit i 32-bity, w obudowach od 6-pin SMD i 8-pin DIP do 144-pin SMD. W świecie hobbistów PIC zyskał ogromną popularność, na równi z legendarnym mikrokontrolerem Intela 8051. Zagadką, której nie rozwiłałem, był podział zwolenników PIC i 8051 według klucza geograficzno-politycznego. Na Zachodzie bardziej popularny był PIC, a w Europie Wschodniej 8051.

Najbardziej popularne rodziny mikrokontrolerów

- Atmel
 - AT tiny
 - AT mega
 - AVR Dx
 - AVR 32
- Microchip
 - 8-bitowe PIC 16, PIC 18
 - 16-bitowe ds PIC 33/PIC 24
 - AVR DD
- ST Microelectronics
 - STM 8 (8b), STM 32 (32b)
- Texas Instruments
 - MSP 430 (16b)
 - AM 24 x, TMS 320 F28 x
- NXP Semiconductors
 - MCX
 - rodzina LPC
- Motorola 68 HC 11 (8b) (ostatnio produkowany w NXP)
- TSMC
 - ESP
- ROHM Semiconductors
 - 16-bitowy ML 62 Q1000
- Renesas Electronics
 - RL 78 Low Power 8 & 16-bit
 - RA ARM Cortex – M
 - RX 671

Nazwy firm w tej ramce mają charakter orientacyjny, gdyż na ogół ten sam μC jest produkowany na licencji przez różne firmy.

Polecam artykuł „10 najbardziej popularnych mikrokontrolerów wśród hobbistów” – www.elportal.pl/czytelnia/mikrokontrolery-mcv-mcc/3575-10-najbardziej-popularnych-mikrokontrolerow-wsrod-hobbistow

Mikrokontrolery o niskim i ultra niskim poborze mocy

Konstruktorzy zawsze troszczą się o niski pobór mocy mikrokontrolerów. Ale jak największa szybkość działania też jest celem projektantów, a te dwa parametry stoją w sprzeczności. Im większa częstotliwość taktowania tym większy jest pobór mocy dynamiczny, związany z prądem przesunięcia płynącym przez pojemności pasożytnicze tranzystorów MOS. Dlatego od dawien dawna stosowano w mikrokontrolerach dwie metody redukcji poboru mocy:

- użycie dwóch zegarów taktujących o różnych częstotliwościach, aby w przedziałach czasu, gdy nie jest potrzebna maksymalna szybkość działania, przełączać pracę mikrokontrolera na zegar o mniejszej częstotliwości taktowania;
- hibernacja, inaczej uśpienie mikrokontrolera, gdy nie musi być aktywny. W stanie uśpienia układ pobiera tylko prąd statyczny, który dla inwerterów CMOS jest ekstremalnie niski.

Jednak obecnie, wobec szybkiego rozwoju zastosowań mikrokontrolerów w segmentach IoT i wearables (Internet Rzeczy i elektronika noszona) niski pobór mocy staje się parametrem krytycznym, ze względu na zasilanie bateryjne lub nawet zasilanie energią pozyskiwaną z otoczenia (energy harvesting). Tam, gdzie można sobie na to pozwolić, zmniejsza się częstotliwość taktowania i np. mikrokontroler zdolny do pracy przy częstotliwości 1 GHz, w wersji ultra niskiej mocy pracuje z częstotliwością zegara 50 MHz. Kolosalne znaczenie dla poboru mocy ma napięcie zasilania, wszak moc zależy od napięcia w kwadracie. Na przykład zmniejszenie napięcia zasilania z 5 V do 1 V daje prawie 25-krotną redukcję poboru mocy. Pół wieku rozwoju technologii MOS to nie tylko walka o coraz mniejsze rozmiary tranzystora, ale też ciągłe modyfikacje technologiczne, które by pozwoliły zmniejszyć napięcie progowe U_T tranzystorów MOS. W pierwszej technologii PMOS napięcia progowe tranzystorów wynosiły kilka V i potrzebne było napięcie zasilania 24 V. Po zmniejszeniu U_T do ok. 1 V (technologia NMOS i CMOS) można było produkować układy scalone MOS zasilane napięciem 5 V, kompatybilnym z TTL. Teraz możliwe jest już stosowanie napięcia zasilania ok. 1 V przy U_T rzędu 0,2 V. Przy tak niskim napięciu progowym problemem staje się pełne wyłączenie tranzystora dla napięcia 0 V na bramce. Tranzystor pracuje wtedy w zakresie nazywanym przedprogowym, tj. przy niepełnym wyłączeniu, co oczywiście będzie powodować wzrost poboru prądu statycznego. Kołdra jest krótka.

Zakończenie

Jeśli ktoś myślał, że po przeczytaniu tego wykładu nauczy się czegoś konkretnego o mikroprocesorach i mikrokontrolerach, to z góry przepraszam za doznane rozczarowanie. Niestety, pozyskanie wiedzy użytecznej dla praktycznych działań projektowych wymaga znacznie większej inwestycji w naukę. Trzeba solidnie przestudiować przynajmniej jeden dobry podręcznik na poziomie podstawowym, a potem selektywnie uczyć się wybranej rodziny mikrokontrolerów. Na przykład podręcznik firmy Microchip poświęcony tylko jednej rodzinie AVR DD ma 614 stron. Najłatwiej mają hobbisci sprawnie aplikujący Arduino lub Raspberry Pi. Hardware mają gotowy, a programy tworzą z gotowych klocków. Niewiele muszą się uczyć. Ale sprawny „DIY maker” i konstruktor elektronik to całkiem różne plemiona. A po co był ten wykład? Ot tak, do poczytania, żeby zachęcić do nauki.

Przegląd tak szerokiej tematyki siłą rzeczy prowadzi do potraktowania większości zagadnień „po łebkach” i uproszczeń nie zawsze akceptowalnych. Będę wdzięczny za wszelkie uwagi, komentarze i pytania. Na pytania ukierunkowane w głąb poszczególnych zagadnień chętnie odpowiem w rubryce „Konsultacje”.

Dla przykładu w tym wydaniu odpowiadamy na pytanie dotyczące MCU o ultra niskim poborze mocy. Używam liczby mnogiej „odpowiadamy”, bo sam korzystam ze wsparcia konsultacyjnego wybitnych specjalistów współpracujących z naszymi redakcjami „Elektroniki dla Wszystkich”, „Elektroniki Praktycznej” i „Elektronika”. ■

Wiesław Marciniak

Konsultacje

Student. Zainteresowałem się mikrokontrolerami (MCU) o ultra niskim poborze mocy. Jednak wybór nie jest prosty. Prawie każdy producent twierdzi, że ma produkty tej kategorii. Czym się kierować, by ocenić, który mikrokontroler rzeczywiście spełnia kryteria układu o ultra niskim poborze mocy? I jak porównywać różne MCU pod względem poboru mocy?

Konsultant. Rzeczywiście prawie każdy producent jest skłonny twierdzić, że produkuje MCU „o najniższej mocy na świecie”. Pobór mocy MCU jest dość złożoną funkcją kilku jego parametrów i warunków pracy. Dlatego trzeba uważać, żeby nie porównywać gruszek z jabłkami, czyli nie wystarczy tylko rzut oka na kartę katalogową z danymi o parametrach. Trzeba te dane odnieść do rzeczywistych warunków pracy MCU. Dalej wyjaśnię szczegółowo te dość ogólnikowe stwierdzenia. Najpierw ustalmy pojęcia podstawowe.

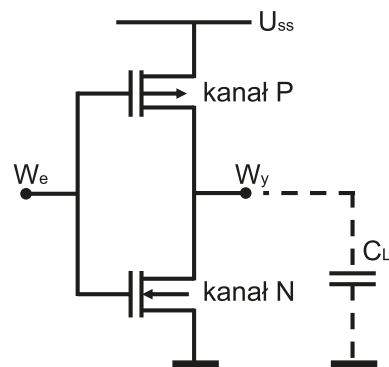
Otóż całkowita moc pobierana przez MCU jest sumą dwóch składowych:

$$P_{\text{całkowita}} = P_{\text{spoczynkowa}} + P_{\text{aktywna}}$$

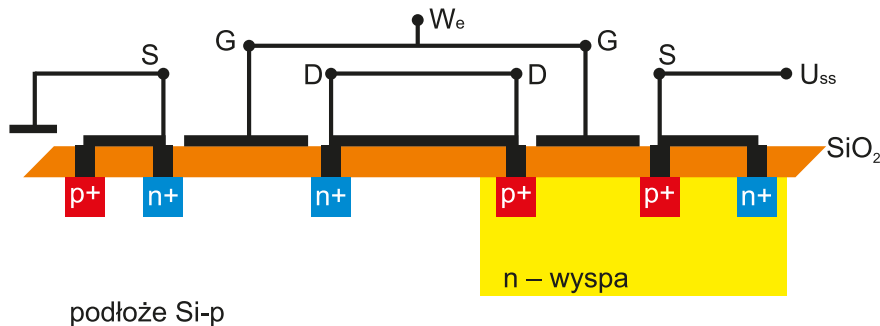
Moc spoczynkowa wynika z prądów statycznych płynących przez inwertery. Wiadomo, że w stanie nie statycznym zawsze jeden z dwóch tranzystorów inwertera CMOS jest w stanie nieprzewodzenia (**rysunek 1**), zatem przez inwerter płynie tylko znikomy prąd wsteczny złącza p – n, nazywany prądem upływu inwertera CMOS.

Student. O jakim złączu p – n teraz mówimy?

Konsultant. Proszę zwrócić uwagę na strukturę fizyczną inwertera CMOS. Narysujmy ją w dużym uproszczeniu (**rysunek 2**). Od masy do USS mamy po kolei strukturę n+ – p – n+ tranzystora z kanałem N, połączoną szeregowo ze strukturą p+ – n – p+ tranzystora z kanałem p. Niech USS wynosi +5 V. Wówczas przy poziomie logicznym 1, tj. +5 V na wejściu W_e tranzystor „dolny” jest w stanie przewodzenia, tj. ma strukturę n+ – n (kanał) – n+, czyli nie ma tu żadnego złącza p – n, natomiast tranzystor „górny” jest zatłaczany i jego złącze p+ (źródło) – n (wyspa) jest spolaryzowane w kierunku zaporowym. Natomiast przy poziomie logicznym 0, tj. 0 V na wejściu W_e



Rysunek 1.



Rysunek 2.

zycznych struktury CMOS w danej technologii) i oczywiście od liczby inwerterów. Natomiast analiza mocy aktywnej jest bardziej złożonym zadaniem. Pobór mocy aktywnej następuje w momencie przełączania inwertera, gdy zmienia się stan logiczny na wejściu i tranzystory zmieniają swój stan przewodzenia. Na przykład przy przejściu na wejściu z „jedynek” na „zero” tranzystor „górny” przechodzi ze stanu nieprzewodzenia do przewodzenia, a „dolny” odwrotnie, ze stanu przewodzenia do nieprzewodzenia. W momencie przełączania pojawia się impuls poboru prądu ze źródła zasilającego, gdyż po pierwsze przez chwilę oba tranzystory przewodzą, po drugie następuje przeładowanie pojemności obciążającej inwerter (C_L na rysunku 1). Jest to pojemność bramek tranzystorów kolejnego inwertera, podłączonego do wyjścia rozpatrywanego inwertera. Te pojemności są bardzo małe, ale po przemnożeniu przez miliony inwerterów w układzie wychodzą znaczne wartości prądów przełączania. Ponieważ pobór mocy MCU zależy od liczby inwerterów, to jest oczywiste, że 32-bitowe MCU pobierają większą moc niż 8-bitowe MCU. Pobór mocy aktywnej zależy w oczywisty sposób od częstotliwości zegara taktującego. Ponieważ prąd aktywny jest pobierany w chwilach przełączania inwerterów, to rośnie liniowo w miarę wzrostu częstotliwości przełączeń.

Kolejnym parametrem wpływającym na pobór mocy aktywnej jest napięcie zasilania. Jest to zależność bardzo silna – pobór mocy aktywnej rośnie proporcjonalnie do kwadratu napięcia. Zależność mocy spoczynkowej od napięcia jest dużo słabsza, gdyż prąd wsteczny złącza p – n w niewielkim stopniu zależy od napięcia. Jeśli mamy świadomość, że sprzęt, w którym jest zagnieżdżony MCU, będzie pracował w trudnych warunkach środowiskowych, np. na ulicy w Bagdadzie, to musimy uwzględnić zmiany poboru mocy w funkcji temperatury. W skrajnych temperaturach pracy urządzenia pobór prądu, a więc również pobierana energia, może się różnić o rząd wielkości.

Pobór mocy aktywnej w decydującym stopniu zależy od trybu pracy MCU. Porównując MCU o małym poborze mocy trzeba uwzględnić w jakich odcinkach czasu MCU będzie się znajdować w określonych trybach pracy. Na przykład MCU zastosowany w czujniku ruchu PIR większość czasu pracuje w trybie uśpienia, budząc się jedynie od czasu do czasu po aktywacji impulsem wyzwalamym. W tym przypadku głównym wyznacznikiem poboru mocy jest niewielki prąd pobierany w stanie uśpienia. Inna jest sytuacja dla zastosowania MCU w czujnikach procesów przemysłowych, gdzie MCU „pracuje” bez przerwy, zatem liczy się przede wszystkim pobór

mocy aktywnej. Poszczególne MCU mogą się znacznie różnić funkcjonalnością, tj. rodzajami trybów aktywnych, co zasadniczo wpływa na ocenę porównawczą MCU w odniesieniu do poboru mocy w konkretnej aplikacji. Pobór mocy aktywnej jest standardowo podawany w jednostkach pobieranego prądu na 1 MHz, dla określonego napięcia zasilania i temperatury pracy. To logiczne, skoro pobór mocy aktywnej zależy od częstotliwości, napięcia zasilania i temperatury. Wiele MCU ma zmienne tryby taktowania, pozwalające użytkownikowi wybrać częstotliwość taktowania, źródło zegara taktującego (oscylator wewnętrzny lub zewnętrzny) oraz sposób zasilania układów peryferyjnych. Przykładem elastyczności oferowanej projektantowi przez producenta MCU jest seria niskoenergetycznych układów MSP 430 firmy Texas Instruments. W kilku wybieranych przez użytkownika trybach pracy układy te pobierają od 375 do 100 mA/MHz przy 3 V zasilania, w zależności od użycia pamięci FRAM lub SRAM (FRAM – Ferroelectric RAM, SRAM – Static RAM). Porównując pobór mocy różnych MCU trzeba zwrócić uwagę na porównywalność różnie definiowanych stanów uśpienia. Może być kilka trybów, np. „deep sleep”, „sleep”, „standby”, różniących się określeniem, które bloki MCU są w stanie wyłączenia. Na przykład seria nanoWatt XLP (eXtreme Low Power) firmy Microchip szczyci się bardzo niskim poborem prądu 9 nA w najgłębszym trybie uśpienia, ale jeśli jest wymagany reset typu brown-out, to ten MCU pobiera 45 nA. Jeśli jest wymagany timer watchdog, to pobór wzrasta do 200 nA, natomiast tryb uruchamiający zegar/kalendarz czasu rzeczywistego daje wzrost poboru prądu do 400 nA.

Inny przykład – seria 8-bitowych, energooszczędnych MCU firmy ST Microelectronics o oznaczeniu STM 8 L 101. W trybie najniższego poboru mocy pobiera prąd 350 nA.

Student. Czy to znaczy, że ten układ wielokrotnie ustępuje poprzednio omawianemu, który miał prąd zaledwie 9 nA?

Konsultant. Otóż nie. Naprawdę trudno jest porównywać gruszki z jabłkami. Układ firmy ST Microelectronics ma wewnętrzny stabilizator. Ma kilka trybów do wyboru: wait, active-halt i halt. W trybie wait procesor jest zatrzymywany, ale urządzenia peryferyjne nadal pracują, podczas gdy MCU czeka na zewnętrzne zdarzenia. W trybie active-halt procesor jest zatrzymany, ale auto-wake up i niezależny watchdog są nadal aktywne (jeśli są włączone). W trybie halt zegary procesora i układów peryferyjnych są całkowicie zatrzymane.

Student. Myślę, że dla uzyskania najniższego poboru energii właściwe jest użycie trybu najgłębszego uśpienia.

Konsultant. To nie jest takie pewne. Może się okazać, że wybudzanie MCU z głębokiego snu będzie się wiązało z czekaniem na ustabilizowanie się oscylatora krystalicznego i będzie to czas stracony, w którym pobór prądu był bezużyteczny.

Energia zużyta podczas budzenia się MCU jest marnowana. Należy zatem zwrócić uwagę na czas potrzebny na przejście między trybami; pożądane są urządzenia o szybkim czasie wybudzania, które spędzają minimalny czas na przejściu. Przykładowo, dla wspomnianej serii ST8L101, prąd zasilania podczas wybudzania z trybu active-halt wynosi 2 mA, a czas wybudzania z trybu active-halt do trybu run wynosi 4 ms. Czas ten nie uwzględnia jednak czasu rozruchu oscylatora krystalicznego, a ten może być znacznie większy od czasu rozpoczęcia wykonywania programu przez procesor.

Podsumowując, należy zbadać pobór mocy przez urządzenie dla wszystkich trybów pracy, z których będziemy korzystać, a także ile energii jest zużywane podczas przechodzenia między tymi trybami. Aby porównać MCU pod kątem konkretnego zastosowania, należy rozważyć, który z trybów niskiego poboru mocy spełnia wymaganą funkcjonalność i jak długo układ będzie spędzać w każdym trybie. ■

Quiz: Mikroprocesory (μP), Mikrokontrolery (μC)

Architektura RISC w porównaniu do architektury CISC charakteryzuje się:

- ☐ A. mniej skomplikowanymi instrukcjami i powolniejszym działaniem μC
- ☐ B. mniejszą liczbą instrukcji i szybszym działaniem μC
- ☐ C. mniejszą liczbą instrukcji i powolniejszym działaniem μC

Pierwsze mikroprocesory 4004 i 4040 zostały zaprojektowane przez Federico Faggin. Jakie inne słynne mikroprocesory/mikrokontrolery również on zaprojektował:

- ☐ A. 6800, 6502
- ☐ B. TMS 1000, PIC
- ☐ C. 8080, Z8

Liczba tranzystorów we współczesnych mikroprocesorach o najwyższej skali integracji osiąga wielkość:

- ☐ A. do miliona
- ☐ B. do miliarda
- ☐ C. ponad miliard

Zegary mikrokontrolerów osiągają częstotliwość taktowania rzędu:

- ☐ A. kilku MHz
- ☐ B. kilkudziesięciu MHz
- ☐ C. kilkuset MHz

Programy w mikrokontrolerze są przechowywane w pamięci:

- ☐ A. Flash
- ☐ B. RAM
- ☐ C. Flash i RAM

Jeśli w mikrokontrolerze dane i programy są przesyłane wspólną magistralą, to mówimy o architekturze:

- ☐ A. von Neumanna
- ☐ B. oxfordzkiej
- ☐ C. harwardzkiej

Pierwszy mikrokontroler TMS 1000, który pojawił się w 1974 roku zawierał:

- ☐ A. 8 000 tranzystorów
- ☐ B. 25 000 tranzystorów
- ☐ C. 80 000 tranzystorów

Kultowy mikrokontroler 8051 został opracowany w firmie:

- ☐ A. Microchip
- ☐ B. Texas Instruments
- ☐ C. Intel

Innowacją mikrokontrolerów AVR było zastosowanie:

- ☐ A. Zmodyfikowanej architektury harwardzkiej
- ☐ B. Zmodyfikowanej architektury von Neumanna
- ☐ C. Zmodyfikowanej architektury oksfordzkiej

W akronimie ARM litera R pochodzi:

- ☐ A. od akronimu RISC
- ☐ B. od słowa Reverse
- ☐ C. od imienia współzałożyciela Intela Roberta Noyce'a

Mikrokontroler PIC powstał w firmie:

- ☐ A. Intel
- ☐ B. Microchip
- ☐ C. General Instrument

Moc pobierana przez mikrokontroler zależy od częstotliwości taktowania i napięcia zasilającego, przy czym jest to zależność:

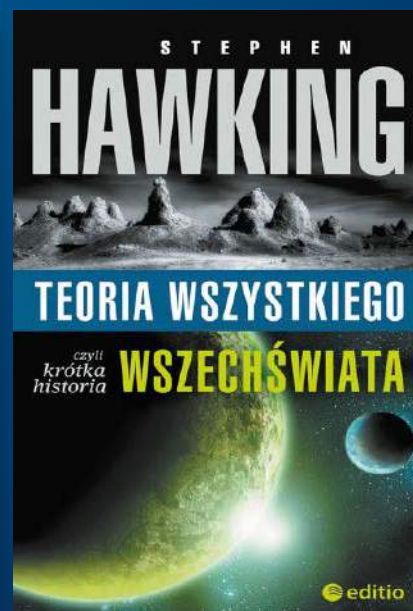
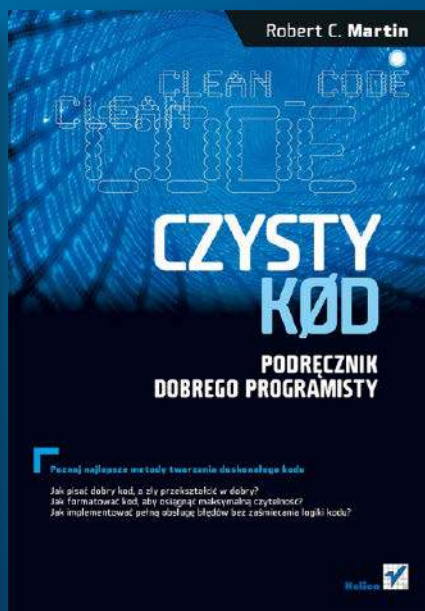
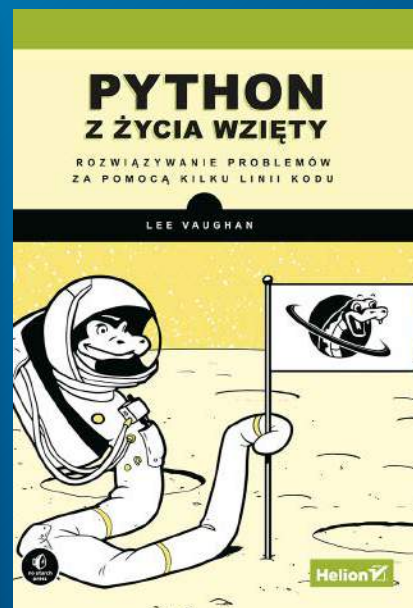
- ☐ A. liniowa od częstotliwości i napięcia
- ☐ B. liniowa od częstotliwości i kwadratowa od napięcia
- ☐ C. kwadratowa od częstotliwości i liniowa od napięcia

Rozwiązanie znajdziesz na www.elportal.pl/quiz od dnia 26.11.2022.

REKLAMA

ELPORTAL.pl – portal z informacjami dla każdego elektronika

KSIĄŻKI W ULUBIONYM KIOSKU Z RABATEM DO 30%



Zobacz pełną ofertę – ponad 500 tytułów!

Zamów wygodnie na UlubionyKiosk.pl

System ostrzegania o przekroczeniu temperatury z sygnalizacją na smartfonie



Projekt przedstawia aplikację na Androida alarmującą na smartfonie, gdy temperatura rozważanego procesu przekroczy zdefiniowaną przez użytkownika temperaturę maksymalną.

Systemy oprzyrządowania alarmowego i sygnalizacyjnego są szeroko stosowane w sterownikach zakładów przetwórczych. Jeżeli zmienne procesowe monitorowane przez różne czujniki polowe przekraczają wartości maksymalne (krytyczne), to na komputerze PC lub sygnalizatorze w sterowni generowany jest alarm. Pomaga to poinformować operatora o konieczności podjęcia niezbędnych działań w celu uniknięcia katastrofy przemysłowej. Prezentowany projekt to system ostrzegania o przekroczonej temperaturze oparty na specjalnie stworzonej

aplikacji pracującej pod systemem Android oraz modułach Arduino Uno i Bluetooth. Autorski prototyp projektu pokazano na **rysunku 1**.

Układ i jego działanie

Schemat układu ostrzegającego o nadmiernej temperaturze pokazano na **rysunku 2**. Jest on zbudowany z użyciem modułu Arduino Uno (Board1), modułu Bluetooth HC-05 oraz scalonego czujnika temperatury LM35.

System składa się z nadajnika i odbiornika. Układ pomiaru temperatury jest nadajnikiem informacji, podczas gdy smartfon z systemem Android i aplikacją jest odbiorcą informacji. Dane o temperaturze, zebrane przez czujnik LM35 (IC1) z obszaru lub otoczenia, są przesyłane do smartfona przez Bluetooth. LM35 daje na wyjściu sygnał analogowy proporcjonalny do temperatury, a Arduino dokonuje niezbędnej konwersji napięcia na odpowiadającą mu temperaturę, przed wysłaniem wzorca bitowego szeregowo do modułu Bluetooth.

Komponenty użyte w tym projekcie są opisane poniżej.

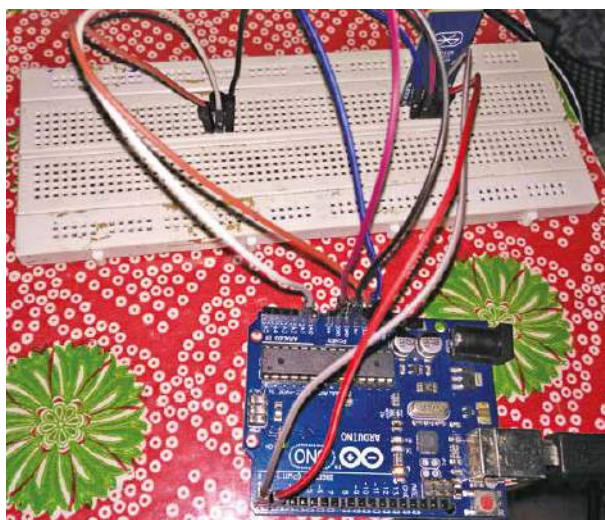
LM35. Jest to precyzyjny scalony czujnik temperatury, którego napięcie wyjściowe jest proporcjonalne do temperatury (w °C). LM35 wykazuje niskie samonagrzewanie, a jego pełny zakres pomiarowy wynosi od -55°C do 150°C (Red. EdW: przy sposobie podłączenia LM35 użytym w projekcie, jego zakres pomiarowy wynosi od 2° do 150°C). W tym projekcie, końcówka wyjściowa układu LM35 jest podłączona do wejścia analogowego A1 modułu Arduino. Zasilanie LM35

jest doprowadzone z końcówek +5 V oraz GND tego modułu.

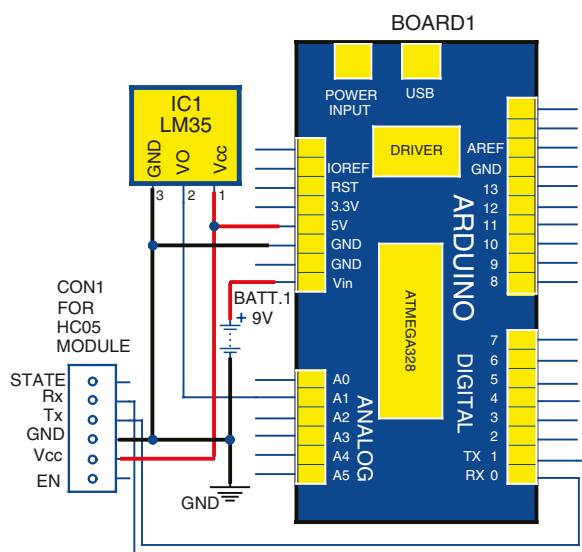
Arduino Uno. Arduino Uno to płytką rozwojową opartą na mikrokontrolerze AVR ATmega328P (MCU) z sześcioma wejściami analogowymi i czternastoma dwustanowymi końcówkami wejścia/wyjścia. MCU posiada 32 kB pamięci flash programowalnej w systemie (ISP), 2 kB RAM i 1 kB EEPROM.

W tym projekcie Arduino służy do konwersji napięcia wyjściowego czujnika LM35 na odpowiednią temperaturę (w °C) i przesyłania uzyskanej wartości do modułu Bluetooth poprzez końcówkę Tx (wyjście szeregowe).

HC-05 moduł Bluetooth. Moduł HC-05, to łatwy w użyciu moduł Bluetooth SPP (Serial Port Protocol) przeznaczony do zapewnienia bezprzewodowego połączenia szeregowego. Jest to w pełni kwalifikowana modulacja Bluetooth V2.0+EDR (Enhanced Data Rate) o przepływności 3 Mb/s, z kompletnym nadajnikiem-odbiornikiem radiowym pracującym na częstotliwości 2,4 GHz i obsługą warstwy baseband. Wykorzystuje jednokładowy



Rysunek 1. Prototyp autora połączony na płytce stykowej



Rysunek 2. Schemat połączeń części nadawczej układu



Rysunek 3. Wygląd ekranu smartfona w czasie działania aplikacji

```
float pin=A1; // Wejście z czujnika LM35
float value=0; // Wartość odczytana z wejścia A1

void setup()
{
  Serial.begin(9600); // Prędkość transmisji szeregowej
}
void loop()
{
  value=analogRead(pin); // Pomiar napięcia z LM35
  Serial.println((value*500)/1023); // Wysyłanie przeliczonej wartości
  delay(1000); // Odstęp czasowy pomiędzy pomiarami
}
```

Listing 1. Kod Arduino

system Bluetooth BlueCore4-Ext firmy CSR z technologią CMOS i adaptacyjną funkcją zmian częstotliwości nośnej (AFH).

W tym projekcie moduł Bluetooth służy do bezprzewodowego przesyłania zmierzonej temperatury do smartfona.

Smartfon z systemem Android. Aplikacja do generowania alarmów o nadmiernej

temperaturze jest oparta na systemie Android i została zaprojektowana przy użyciu platformy App Inventor opracowanej przez MIT, dostępnej pod adresem: <https://appinventor.mit.edu/> (z możliwością obsługi także w języku polskim). App Inventor to intuicyjne, wizualne środowisko programistyczne, które pozwala użytkownikowi

**Kod źródłowy
tego projektu jest dostępny
do pobrania ze strony
<http://bit.ly/3WX8tpv>**

tworzyć w pełni funkcjonalne aplikacje na smartfony i tablety metodą składania gotowych bloczków realizujących różne funkcje bez wnikania w szczegóły realizacyjne. Gdy ta specjalnie zaprojektowana aplikacja działa na smartfonie, to dane wysyłane przez moduł Bluetooth HC-05 są odbierane i wyświetlane na jego ekranie. Gdy zmierzona temperatura przekroczy ustawioną wartość krytyczną, to na smartfonie włączy się alarm dźwiękowy. Opisana aplikacja została przetestowana na smartfonie Redmi 3S.

Procedura testowa

Etapy testowania są następujące:

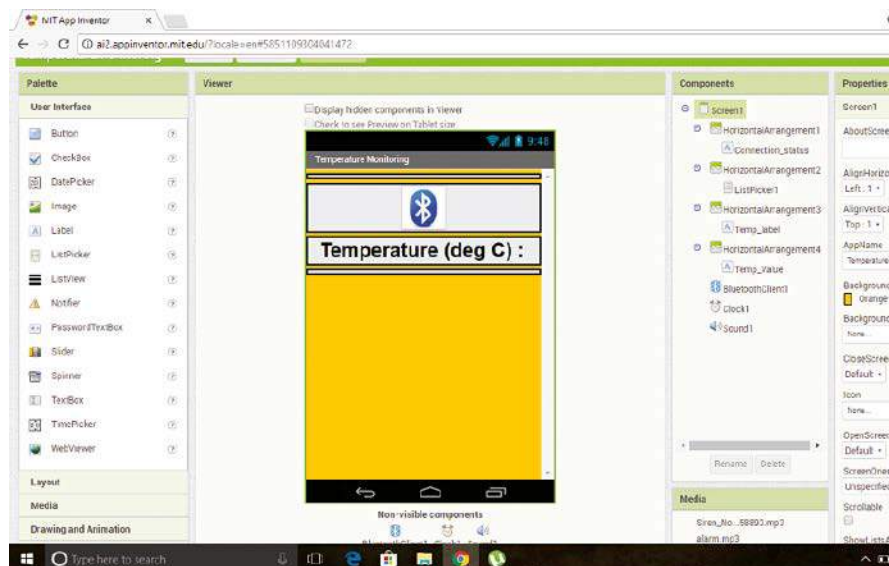
1. Prześlij plik z kodem źródłowym programu, plik bt.ino, do Arduino używając Arduino IDE. W programie tym funkcja Serial.println() pozwala przesyłać wartość temperatury, w postaci szeregowej, do modułu Bluetooth.
2. Zainstaluj aplikację Temp_Monit.apk na smartfonie. Po udanej instalacji włącz połączenie Bluetooth smartfona. Po pomyślnym sparowaniu modułu Bluetooth HC-05 z Bluetooth smartfona, otwórz aplikację w smartfonie i naciśnij przycisk z logo Bluetooth.

Zlisty dostępnych urządzeń wybierz HC-05. Po pomyślnym nawiązaniu połączenia w aplikacji zostanie wyświetlony komunikat „Connected”. Aplikacja wyświetli na ekranie telefonu informację o temperaturze. Obraz ekranu smartfona pokazującego aktualną temperaturę pokazano na **rysunku 3**.

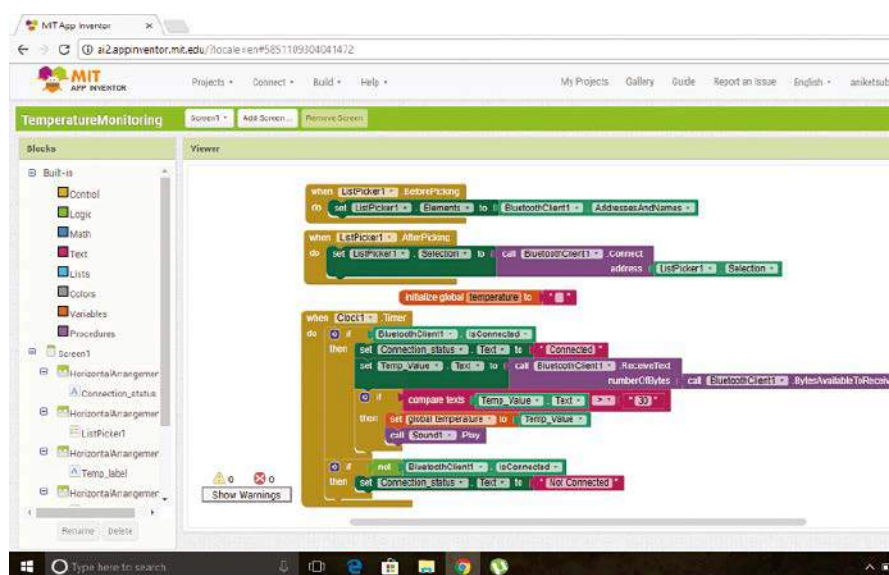
Zrzuty ekranu z sekcji Designer (wygląd aplikacji) i Blocks (program), z fazy projektowania przy użyciu aplikacji MIT App Inventor, pokazano odpowiednio na **rysunekach 4 i 5**. Jeśli temperatura przekroczy wartość krytyczną (tutaj do testów ustawiono ją na 30°C), smartfon wygeneruje dźwięk syreny.

Chociaż kod źródłowy Androida używa temperatury 30°C jako wartości przegrzania do generowania alertu, możesz użyć dowolnej innej wartości, zmieniając liczbę 30 w wierszu „Compare texts Temp_Value.Text > „30” w sekcji Blocks. ■

Aniket Subham
Shibendu Mahata



Rysunek 4. Obraz z sekcji projektowania wyglądu aplikacji w App Inventorze



Rysunek 5. Obraz z sekcji projektowania kodu aplikacji (Blocks)

Rejestrator danych pomiarowych z pamięcią USB



Prezentowany tutaj układ jest rejestratorem danych pracującym w czasie rzeczywistym. Układ jest oparty na: mikrokontrolerze PIC18F452 (MCU), układzie kontrolera hosta USB VNC1L-1A firmy FTDI, zegarze czasu rzeczywistego DS1302 (Maxim) i wyświetlaczu LCD. System pozwala elastycznie ustawiać czas rozpoczęcia rejestracji i okres jej trwania.

Rejestratory tego typu są niezbędne dla zdalnych systemów gromadzenia danych i ogólnie wymagają dużych nośników pamięci do przechowywania pozyskanych informacji. Nośniki, takie jak pendrive'y, wymagają do współpracy urządzenia z funkcją hosta USB lub komputera PC do transferu zebranych danych. Mikroprocesory ogólnego przeznaczenia nie obsługują funkcji hosta USB. Oto nowy system akwizycji danych z układem z rodziny Vinculum, który zapewnia funkcjonalność hosta USB. Ponadto z pomocą układu zegara czasu rzeczywistego DS1302 i klawiatury matrycowej 4×4 można ustawić czas rozpoczęcia pomiarów i czas ich wykonywania.

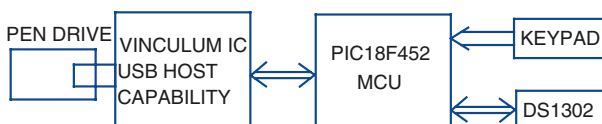
Schemat blokowy rejestratora pokazano na **rysunku 1**. Prototyp wykonany przez autora pokazano na **rysunku 2** (Uwaga prototyp wykonano na innej płytce niż ta pokazana na **rysunkach 4 i 5**).

Układ rejestratora i jego działanie

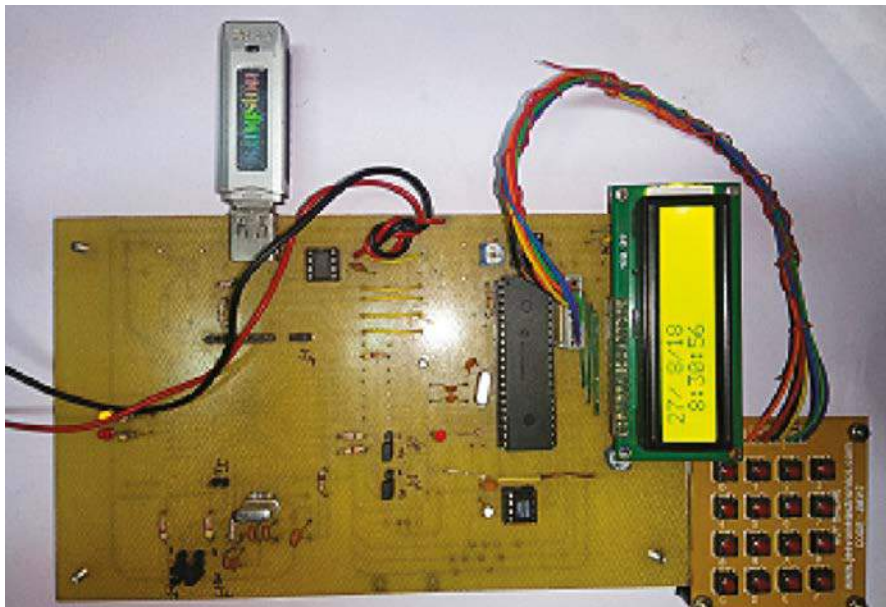
Schemat ideowy układu rejestratora danych pokazany jest na **rysunku 3**. Układ składa się: stabilizatora napięcia 3,3 V typu MAX882 (IC1), czujnika temperatury LM35 (IC2), kontrolera hosta USB – VNC1L-1A (IC3), mikrokontrolera PIC18F452 (IC4), układu zegara czasu rzeczywistego DS1302 (IC5) z funkcją podtrzymania baterijnego, wyświetlacza ciekłokrystalicznego 16×2 (LCD1) oraz kilku innych komponentów.

Mikrokontroler PIC18F452 zajmuje się akwizją danych. Układ VNCIL-1A jest połączony z PIC18F452 poprzez interfejs szeregowy UART.

Układ VNC1L-1A. Posiada dwa porty USB oraz uniwersalny blok interfejsowy mogący pracować jako: UART, SPI lub FIFO. Końcówki USB1DP i USB1DM obsługują odpowiednio sygnały danych USB D+ i D– portu USB1.



Rysunek 1. Schemat blokowy rejestratora danych w pamięci USB



Rysunek 2. Prototyp autora

Końcówki USB2DP i USB2DM obsługują sygnały danych USB D+ i D– portu USB2.

Port-1 obsługuje peryferyjne urządzenia podrzędne USB (nieużywane w tym projekcie). Port-2 obsługuje urządzenia klasy BOMS (pamięci masowe), takie jak dyski flash USB. Rodzaj szeregowego interfejsu UART, SPI lub FIFO do komunikacji pomiędzy VNC1L-1A i PIC18F452 można wybrać za pomocą pary końcówek ACBUS5 Ic3(46) i ACBUS6 Ic3(47), podłączając je do masy lub 3.3 V przez rezystory 47 kΩ.

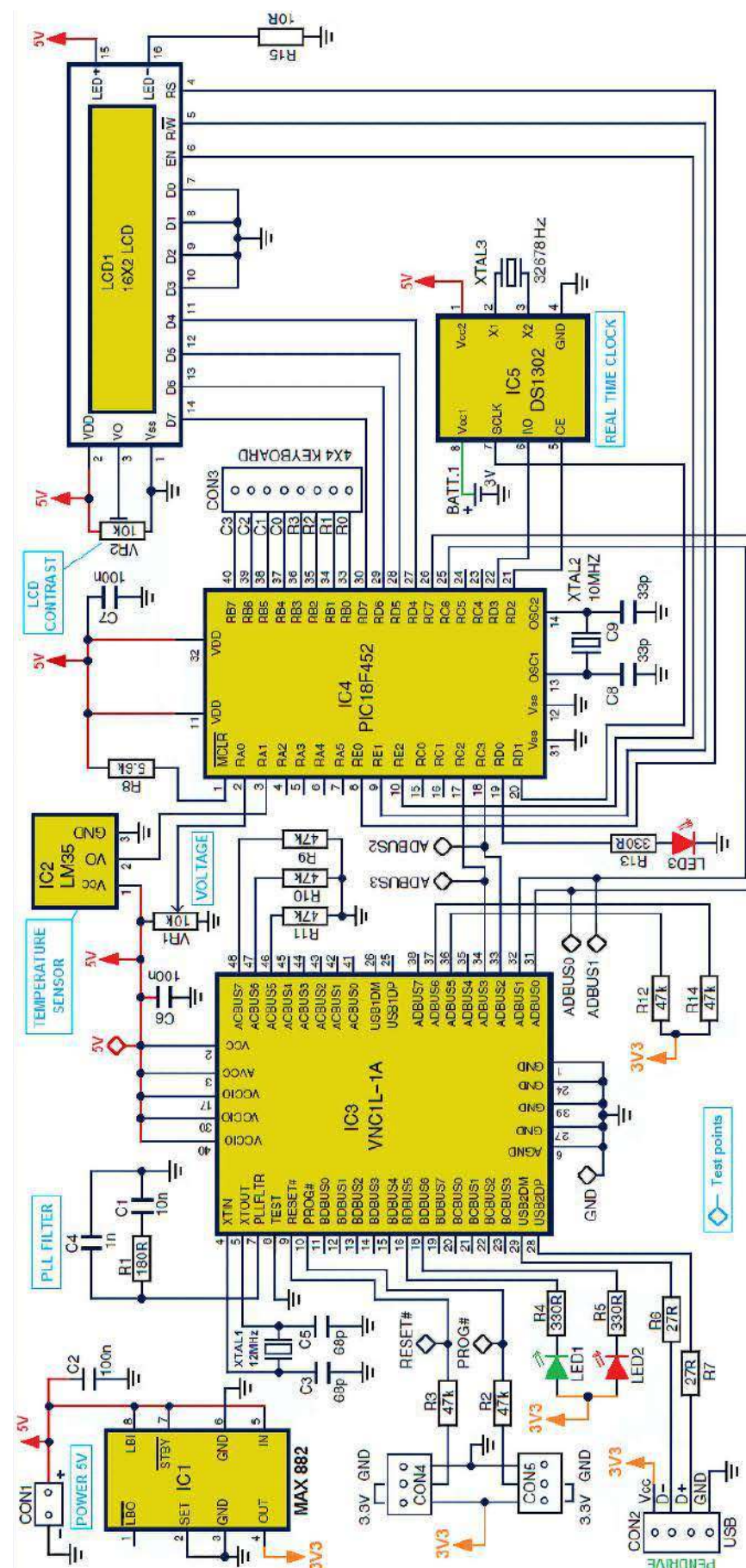
Interfejs UART wybiera się poprzez podłączenie końcówek 46 i 47 układu IC3 do masy przez rezystory 47 kΩ. ADBUS0 Ic3(31) i ADBUS1 Ic3(32) przenoszą sygnały TxD i RxD, a ADBUS2 Ic3(33) i ADBUS3 Ic3(34) przenoszą odpowiednio sygnały RTS# i CTS# interfejsu UART. ADBUS7 wybiera, przez zwarcie go rezystorem 47 kΩ do masy, taktowanie układu zewnętrznym rezonatorem kwarcowym 12 MHz.

Vinculum oferuje, dla układu VNC1L-1A, kilka wstępnie skompilowanych

standardowych wersji oprogramowania układowego (firmware). W obecnej aplikacji używana jest wersja oprogramowania VDAP (dysk i urządzenia peryferyjne). Oprogramowanie układowe umieszcza się w pamięci flash układu VNC1A-1A (o pojemności 64 kB) i zawiera ono interpreter poleceń, który umożliwia zewnętrznemu mikrokontrolerowi PIC wysyłanie i odbieranie poleceń/danych za jego pośrednictwem.

Interpreter obsługuje polecenia zarządzania dyskami podobne do systemu DOS, które umożliwiają: wykonywanie operacji katalogowych (tworzenie, usuwanie lub wyświetlanie katalogów), operacje na plikach (tworzenie, otwieranie, odczytywanie, zapis, zamykanie, usuwanie lub zmianę nazwy plików) oraz debugowanie (identyfikowanie lub pobieranie wersji oprogramowania układowego). Zawiera również funkcje umożliwiające wykrywanie podłączenia i odłączenia pamięci flash do VNC1L-1A.

Kiedy oprogramowanie układowe wykrywa dysk flash podłączony do portu USB, to wysyła przez interfejs UART, podobnie jak DOS, znak zachęty D:\>. Steruje też miganiem diody



Rysunek 3. Schemat ideowy rejestratora

LED2 podłączonej do końcówki 18 układu IC3 (BDBUS6), gdy dane są zapisywane do pamięci USB.

Interpreter poleceń działa w trybie rozkazów lub danych. W trybie rozkazów, VNC1L-1A interpretuje kody z portu UART jako polecenia i wykonuje je. W trybie danych przekazuje dane z portu monitora do pamięci USB. Początkowo, po uruchomieniu, VNC1L-1A pozostaje w trybie rozkazów. Następnie sygnalizuje końcówkami VNC1L-1A ADBUS5 Ic3(36) i ADBUS6 Ic3(37), które przenoszą sygnały DTR# i DSR# interfejsu szeregowego UART i przełącza tryb pracy. W opisywanym zastosowaniu VNC1L-1A jest ustawiany do pracy w trybie poleceń poprzez podłączenie ww pinów do napięcia 3,3 V przez rezystory R12 i R14.

Końcówki: ADBUS0 do ADBUS3, PROG#, RESET#, Vcc i GND układu VNC1L-1A służą do programowania go przez programator FTDI.

Usuń zworki na złączach CON4 i CON5 podczas programowania, w pozostałych przypadkach piny 9 i 10 powinny być podłączone do 3,3 V.

Więcej szczegółów na temat kontrolera hosta USB układów rodziny Vinculum, w tym programowania FTDI, można znaleźć wraz z kodem źródłowym na stronie source.efymag.com

Opis oprogramowania

W tym miejscu opiszemy oprogramowanie obsługujące próbkowanie w regularnych odstępach czasu, dwóch napięć analogowych przyłożonych do końcówek 2 i 3 mikrokontrolera PIC18F452 oraz zapis wyników pomiarów do pliku w pamięci flash USB dołączonej do VNC1L-1A. Oprogramowanie to zostało napisane przy użyciu kompilatora języka C firmy CCS i jest wgrywane do pamięci flash mikrokontrolera PIC18F452.

Zawiera ono moduły realizujące następujące zadania:

Synchronizacja komunikacji UART. Komunikacja UART pomiędzy VNC1L-1A i PIC18F452 jest zapewniona przez odbiór echa znaków „E” oraz „e” wysyłanych z PIC18F452.

Obsługa przerwań odbioru danych dostępnych (INT_RDA). Za każdym razem, gdy odebrany bajt z VNC1L-1A, znajdzie się w buforze odbiorczym interfejsu UART układu PIC18F452, generowane jest przerwanie. Program odczytuje wtedy znak z bufora i wysyła go do pamięci USB (pendrive).

Próbkowanie kanałów analogowych. Oprogramowanie obsługuje próbkowanie kanałów analogowych 0 i 1. Ponadto kompilator C ma wbudowane funkcje do konfigurowania wbudowanego w PIC przetwornika ADC, do wybierania kanałów wejściowych i odczytywania danych po konwersji analogowo-cyfrowej.

Funkcje te mogą być wykorzystywane przez program do próbkowania kanałów analogowych.

Obsługa zapisu danych na dysk Flash USB. Oprogramowanie zawiera funkcje służące do wysyłania poleceń utworzenia pliku Results.dat na dysku flash, zapisania w nim dwudziestu wartości próbkowanych danych oraz zamknięcia pliku.

Zegar czasu rzeczywistego i interfejs użytkownika

W niniejszym projekcie interfejs użytkownika jest wyposażony w klawiaturę do wprowadzania danych oraz wyświetlacz LCD do monitorowania czasu trwania pomiarów. Układ DS1302 dostarcza do systemu informację o aktualnym czasie. Połączenie DS1302 z MCU jest uproszczone dzięki zastosowaniu synchronicznej komunikacji szeregowej. Do jej realizacji potrzebne są tylko trzy przewody dla sygnałów: wyboru układu (CE), wejścia/wyjścia (linia danych) i zegara szeregowego (SCLK).

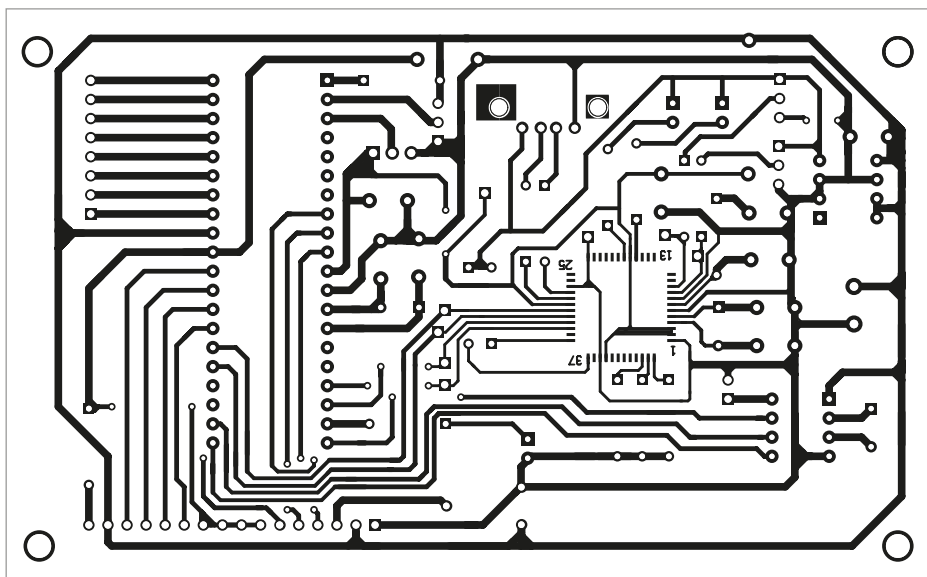
Podczas operacji odczytu lub zapisu sygnał CE musi być ustawiony w stan wysoki, ponieważ końcówka obsługująca ten sygnał, ma wewnętrzny rezystor 40 kΩ, ściągający to wejście do masy. SCLK służy do synchronizacji ruchu danych w interfejsie szeregowym. Co więcej, DS1302 został zaprojektowany do pracy przy bardzo niskim poborze mocy, dzięki czemu zachowuje dane i informacje o czasie, przy zużyciu energii mniejszym niż 1 μW.

Kompilator CCS C posiada gotowe funkcje obsługi dla układu DS1302 oraz klawiatury 4×4. Ponadto MCU ma wbudowane rezystory podciągające dla pinów portu B, które można włączyć za pomocą poleceń programowych. Stąd połączenie między pinami portu B MCU a klawiaturą jest zrealizowane bezpośrednio. DS1302 ma możliwość ustawienia czasu, w postaci: minuta, godzina, data i dzień, aby dane mogły być rejestrowane w sposób ciągły przez cały dzień, jeśli zapewnione jest wystarczające zasilanie. Ustawienia czasu można śledzić na wyświetlaczu LCD1.

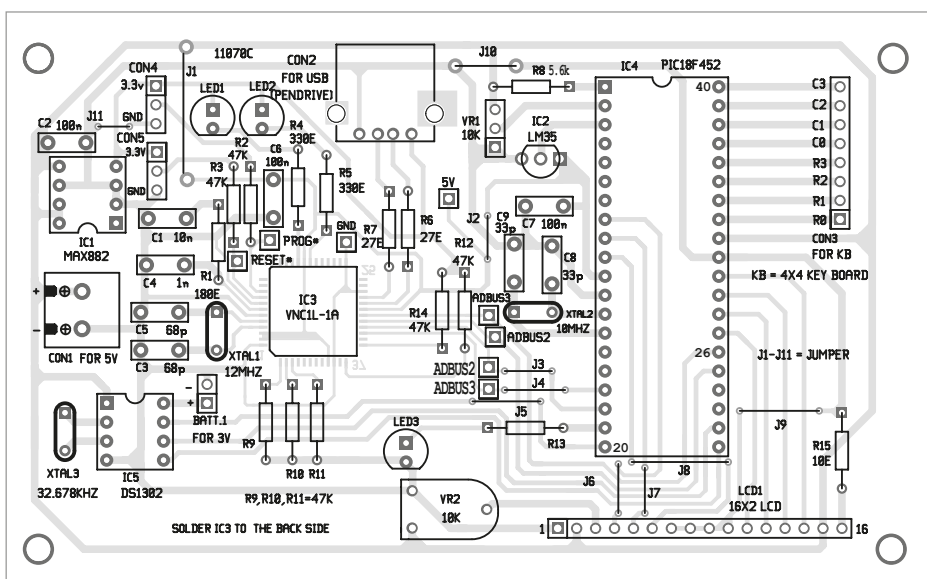
Budowa i testowanie

Układ ścieżek na płytce drukowanej rejestratora danych pokazano na rysunku 4, a układ elementów na rysunku 5. Zmontuj obwód na PCB.

Po zakończeniu montażu PCB, użyj kodu z pliku DataLogger.hex do zaprogramowania mikrokontrolera PIC18F452, za pomocą



Rysunek 4. Projekt płytki drukowanej rejestratora – ścieżki



Rysunek 5. Rozłożenie elementów na płytce drukowanej

odpowiedniego programatora. Po zaprogramowaniu układu wyjmij go z programatora i włóż go do podstawki na płytce drukowanej. Twój obwód jest teraz gotowy do użycia.

Włóż pendrive do złącza CON2. Podłącz napięcie 5 V DC do CON1 (uwaga na biegunowość, układ nie ma zabezpieczenia przed odwrotnym podłączeniem zasilania), aby

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

Inne:

CON1: 2-pinowe złącze zaciskowe
CON2: złącze USB typu A
CON3: złącze 8-pinowe
CON4, CON5: złącza 3-pinowe
XTAL1: rezonator kwarcowy 12 MHz
XTAL2: rezonator kwarcowy 10 MHz
XTAL3: rezonator kwarcowy (zegarkowy) 32678 Hz
podstawa precyzyjna DIP40
przewody potężeniowe
pendrive
płytką drukowaną

Półprzewodniki:

IC1: MAX882 – stabilizator o niskim spadku napięcia
IC2: LM35 – czujnik temperatury
IC3: VNC1L-1A MCU – kontroler hosta USB
IC4: PIC18F452 – mikrokontroler
IC5: DS1302 – zegar czasu rzeczywistego

LCD1: wyświetlacz ciekłokrystaliczny 16×2
LED1...LED3: diody świecące, średnica 5 mm

Rezystory: (wszystkie 0,25 W, ±5 % węglowe)

R1: 180 Ω
R2, R3, R9...R12, R14: 47 kΩ
R4, R5, R13: 330 Ω
R6, R7: 27 Ω
R8: 5,6 kΩ
R15: 10 Ω
VR1, VR2: potencjometry montażowe 10 kΩ

Kondensatory:

C1: 10 nF – ceramiczny
C2, C6, C7: 100 nF – ceramiczne
C3, C5: 68 pF – ceramiczne
C4: 1 nF ceramiczny
C8, C9: 33 pF ceramiczny

**Kod źródłowy
tego projektu jest
dostępny do pobrania
ze strony
<https://bit.ly/3soblh2>**

zasiłnić obwód. LED1 i LED2 będą migać na przemian przez dwie sekundy, aż do podłączenia pendrive'a. Czujnik temperatury LM35 automatycznie mierzy temperaturę otoczenia. Świecenie diody LED3 podłączonej do końcówki 19 PIC18F452 wskazuje, że pomiar temperatury jest w toku.

Potencjometr VR1 (10 kΩ) podłączony do pinu 2 mikrokontrolera PIC18F452 ustawia wstępnie zdefiniowany odczyt (powiedzmy 2 V). Oznacza to, że pobierane są dwa odczyty i wysyłane na dysk USB – jeden z LM35, a drugi z suwaka potencjometru VR1.

Ustawienia i testowanie

Podłącz klawiaturę 4×4 do CON3 w celu ustawienia czasu. Poniżej podano przykład wprowadzania danych wejściowych z klawiatury. Wykonaj poniższe czynności, aby przetestować obwód.

1. Sprawdź wszystkie połączenia w układzie
2. Podłącz płytkę drukowaną do zasilacza 5 V i wydajności 500 mA DC (CON1)
3. Na wyświetlaczu LCD zostanie wyświetlony komunikat „Starting”. Poczekać chwilę, aż na wyświetlaczu LCD1 pojawi się komunikat „Press 1 to Start”
4. Naciśnij 1 na klawiaturze

Clipboard		Font		Alignment		Number						
A1		2.33										
	A	B	C	D	E	F	G	H	I	J	K	L
1	2.33	31.76										
2	2.33	31.76										
3	2.33	30.79										
4	2.32	32.25										

Rysunek 6. Wynik pracy rejestratora w pliku Excel

5. Wprowadź aktualny rok: np. 22 (tj. dwie ostatnie cyfry roku)
6. Wprowadź aktualny miesiąc: np. 10
7. Wprowadź aktualny dzień miesiąca, np. 20
8. Wprowadź dzień tygodnia: 05 (01 – oznacza niedzielę)
9. Wprowadź aktualną godzinę, np. 08
10. Wprowadź aktualne minuty, np. 30
11. Wprowadź godzinę startu zapisu, np. 08
12. Wprowadź minuty startu, np. 32
13. Wprowadź godzinę końca, np. 08
14. Wprowadź godzinę końca, np. 35
15. Wprowadź odstęp między pomiarami: godziny, tu: 00
16. Wprowadź odstęp między pomiarami: minuty, tu: 01

W powyższym przykładzie zostanie ustawiony aktualny czas i data: czwartek 20 października 2022 r., godzina 8.30. Rzeczywiste rejestrowanie danych rozpoczyna się o 8.32 i trwa do 8.35. Odstępy pomiędzy pomiarami wynoszą 0h 01m.

Po zakończeniu zapisu pomiarów, wyjmij pendrive'a z gniazda CON1 i włóż go do gniazda

USB w komputerze PC. Teraz możesz wczytać zarejestrowane dane za pomocą arkusza kalkulacyjnego Microsoft Excel lub innego kompatybilnego z nim programu. Plik z wynikami otwierany za pomocą programu Excel pokazano na **rysunku 6**. Pierwsza kolumna pokazuje napięcie z potencjometru VR1, a druga wartości temperatury z czujnika LM35 (IC2). Płyta prototypowa została przetestowana z pamięciami USB: Scandisk 16 GB, Sony 32 GB i HP 16 GB/32 GB, żeby wymienić tylko kilka.

W ten sposób system akwizycji danych, wykorzystujący pamięć USB i system z kontrolerem hosta Vinculum, próbkuję sygnał w określonych odstępach czasu i przechowuje dane na dysku flash.

Dane wyjściowe zapisywane są w formacie Microsoft Excel, co ułatwia szybki dostęp do danych. Ponadto użyty format Excela, ułatwia dalszą analizę danych. ■

A. Robson Benjamin

Ten artykuł był wcześniej opublikowany na łamach „EFY”, listopad 2019 (efymag.com)

REKLAMA

m.technik
Ciekawi świata są zawsze młodzi

przejrzyj i kupisz na
www.ulubionykiosk.pl

www.elportal.pl

epresa.pl d96e5f8e8c

93

Zegar analogowy i termometr z wyświetlaczem TFT LCD



Prezentowany układ wyświetla zegar analogowy synchronizowany sygnałem GPS i temperaturę pokojową na wyświetlaczu TFT. Czas GPS to precyzyjny standard czasu związany z uniwersalnym czasem koordynowanym (UTC). W większości projektów hobbystycznych jako wyświetlacze wykorzystywane są ekrany LCD, GLCD, OLED lub TFT. Projekt obecny bazuje na wyświetlaczu TFT o przekątnej 8,99 cm (3,5 cala) i rozdzielczości 480×320 pikseli, ze sterownikiem ILI9486 (rysunek 1).

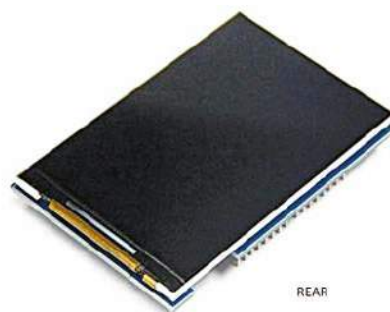
Najnowsze chińskie wyświetlacze TFT są dość tanie, ale doskonale współpracują z płytami Arduino i Raspberry Pi. Dostępne są dwa różne typy nakładek TFT: jeden z 26 końcówkami (1×2 DIL) dla Raspberry Pi i drugi jako nakładka Arduino TFT dla płytki Arduino Uno.

Wyświetlacz TFT nakładka na Arduino

Ta wersja wyświetlacza pasuje idealnie do złącza na płycie Arduino Uno. Ale główną wadą tego rozwiązania jest to, że po zamontowaniu nakładki na płycie Arduino trudno jest wykorzystać jej piny GPIO do jakiegokolwiek innej aplikacji. Nakładka TFT ma dołączoną kartę micro SD, która łączy się z pinami szeregowego interfejsu peryferyjnego (SPI) w celu komunikacji z mikrokontrolerem. W tym projekcie karta ta nie jest wykorzystywana. Pewnym problemem w przypadku użycia tego wyświetlacza może być to, że nie jest on powszechny. Należy go poszukać przez Google („3.5 tft lcd shield arduino”). Na szczęście plik nagłówkowy `mcufriend_kbd.h` jest dostępny bezpłatnie i gotowy do zainstalowania dla tego i wielu innych podobnych wyświetlaczy. Drugim plikiem nagłówkowym wymaganym do tego projektu jest `Adafruit_GFX.h`. Oba te pliki nagłówkowe, wraz z głównym kodem źródłowym, są dostępne na stronie EFY. Program wyświetla tarczę zegara analogowego wraz z cyfrowym wyświetlaniem daty, czasu oraz temperatury zmierzonej za pomocą czujnika temperatury LM35 lub TMP36. Sygnał czasu pochodzi z modułu odbiornika GPS U-Blox NEO-6M. Poznaawszy zasadę działania całości można ją wdrożyć w wielu innych aplikacjach.



FRONT



REAR

Rysunek 1.

Budowa zegara

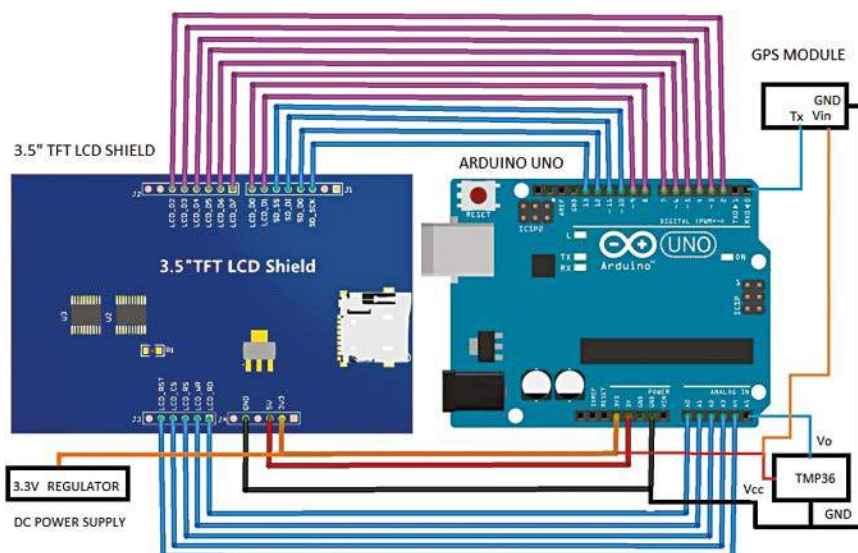
Projekt zegara analogowego GPS z wyświetlaczem temperatury wymaga zgromadzenia następujących elementów:

1. ekranu TFT 8,99 cm (3,5 cala) 480×320, ze sterownikiem ILI9486,
2. płytki Arduino Uno,
3. odbiornika GPS U-Blox NEO-6M,
4. czujnika temperatury LM35 lub TMP36,
5. stabilizatora lub zasilacza prądu stałego 3.3 V.

Połączenia są łatwe do wykonania i pokazano je na rysunku 2. Bezpośrednie połączenie

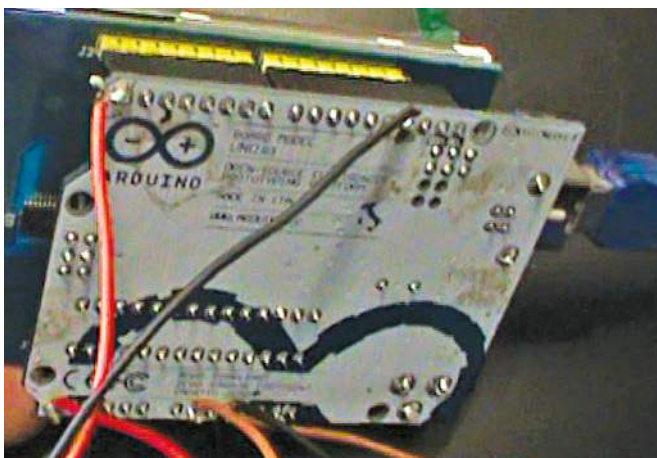
nakładki Arduino TFT jest proste; wystarczy zamontować ją na górze płytki Arduino Uno. Szczegóły połączeń między ekranem TFT a Arduino Uno podano w tabeli.

Ponieważ górna część Arduino Uno jest przykryta ekranem TFT, to połączenia dla czujnika temperatury i odbiornika GPS są wyprowadzone (przylutowane bezpośrednio do punktów lutowanych złączy) od spodu płytki Arduino Uno (patrz rysunek 3). Jeśli chcesz zwolnić niektóre piny Arduino, powinieneś zajrzeć do pliku `mcufriend_shield.h` i odpowiednio go zmodyfikować.



Rysunek 2.

**Kod źródłowy
tego projektu jest
dostępny do pobrania
ze strony
<https://bit.ly/3UEvjJK>**



Rysunek 3.



Rysunek 4.

Połączenia pomiędzy nakładką TFT	
TFT	Arduino
R & D	A0
WR	A1
RS	A2
CS	A3
RST	A4
D0	8
D1	9
D2	2
D3	3
D4	4
D5	5
D6	6
D7	7
3,3 V	3,3 V
GND	GND
5	5

Ekran jest w zasadzie przeznaczony do pracy przy napięciu 3,3 V, ale może pracować też do 5 V. Jednak dłuższa praca na 5 V nie jest zalecana, ponieważ przegrzewa się on przy tym napięciu. Zasilanie 3,3 V pochodzi ze stabilizatora napięcia LD1117V33.

Oprogramowanie

Pisanie kodu-szkicu Arduino (GPS_analog_clock.ino) do projektu to prawdziwa frajda! Poprzez zmiany w kodzie możesz stworzyć obraz na wyświetlaczu TFT na wiele różnych sposobów. Stworzenie grubej linii ograniczającej czy uczynienie aby wskazówki godzinowe i minutowe poruszały się płynnie, było dość trudne, ponieważ biblioteka Adafruit_GFX nie jest zbyt dopracowana. Trygonometria ze szkoły średniej to wszystko, czego potrzebujesz, aby to się udało. Otwórz kod GPS_analog_clock.ino w Arduino IDE i dołącz pliki nagłówkowe. Skompiluj i prześlij kod na płytkę Arduino Uno.

Dołączony jest również inny szkic: GPS_analog_clock2.ino, zmieniający kolor wskazówki minutowej co minutę.

Eksplotacja i testowanie

Po wgraniu kodu przylutuj piny TMP36 i modułu GPS do płytki Arduino. Następnie zamontuj ekran TFT na górze płytki Arduino. Po wykonaniu wszystkich połączeń zgodnie z **rysunkiem 2**, podłącz układ do źródła zasilania 3,3 V DC. Moduł GPS potrzebuje kilku minut na odebranie sygnału z satelitów.

Większość odbiorników GPS ma wbudowaną antenę, która może bardzo ułatwić odbiór

sygnału satelitów GPS na niskiej orbicie okołoziemskiej (LEO), nawet jeśli okna w pokoju są zamknięte. Po poprawnym odebraniu sygnału z dwóch takich satelitów na tarczy analogowej pojawia się aktualny czas. Jednocześnie też wyświetlane zostaną cyfrowo: data, godzina i temperatura, po prawej stronie wyświetlacza TFT.

Czujnik temperatury TMP36 działa na 3,3 V. Zamiast tego można również użyć LM35, ale do niego potrzebujemy zasilania 5 V DC. ■

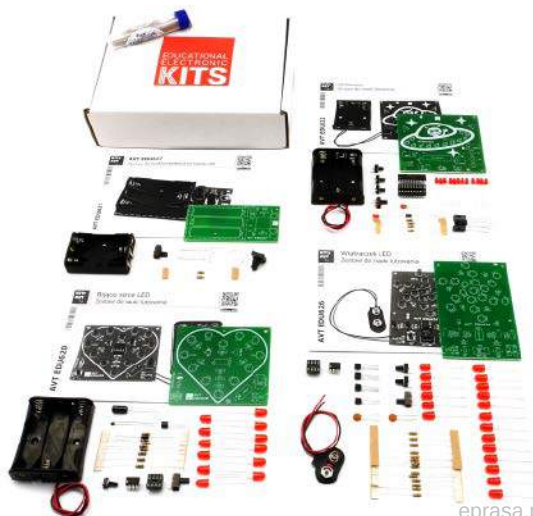
Somnath Bera

Red. EdW: Ten sam projekt był opublikowany w magazynie Elektor i jest tam też spora dyskusja z autorem na temat różnych problemów związanych z projektem <https://www.elektormagazine.com/labs/4-tft-analog-gps-clock-on-arduino>

Ze względu na częste zmiany układów sterujących wyświetlaczami, tu można znaleźć biblioteki sterujące: https://github.com/Bodmer/TFT_eSPI.

Ten artykuł był wcześniej opublikowany na łamach „EFY”, kwiecień 2020 (efymag.com)

REKLAMA



AVTEDU4PAKIET

AVTEDU4PAKIET – to zestaw 4 kitów DIY do nauki lutowania:

- AVTEDU620 – Bijące serce LED
- AVTEDU626 – Wiatraczek LED
- AVTEDU627 – Zestaw do budowy podręcznej latarki LED
- AVTEDU632 – UFOledek



sklep.avt.pl / Allegro Sklep-AVT
lub 03-197 Warszawa,
ul. Leszczynowa 11

Inteligentny sterownik oświetlenia



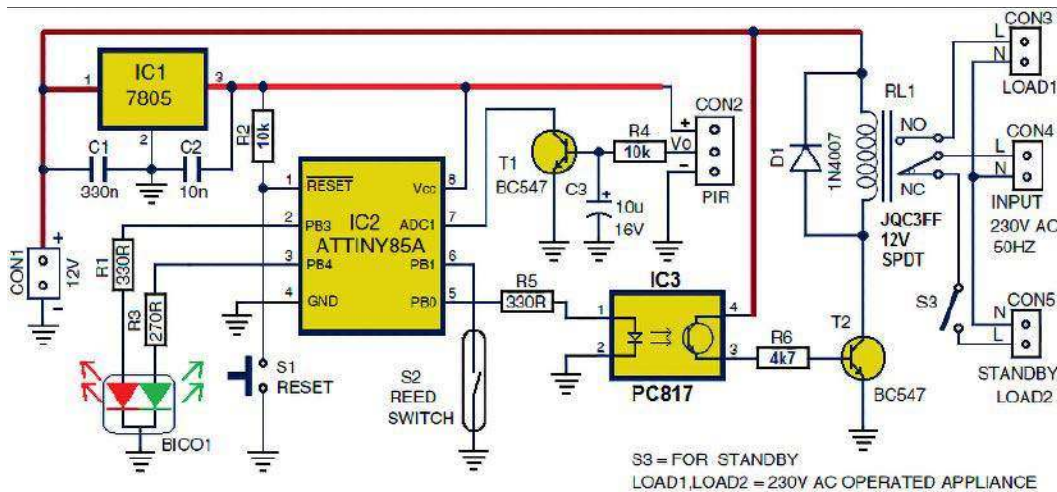
Większość sklepów z tekstyliami ma pomieszczenia przebieralni, w których klienci mogą przymierzyć nowe ubrania przed ich zakupem. Pomieszczenia te są oświetlone przez cały dzień, niezależnie od tego czy są używane, czy nie. Prowadzi to do marnowania cennej energii elektrycznej. Prezentowany tutaj układ, włącza światło tylko wtedy, gdy ktoś wejdzie do przebieralni. Obwód sygnalizuje też czy pomieszczenie jest wolne czy zajęte. Sterownik może być również stosowany do sterowania oświetleniem łazienek czy toalet publicznych lub innych podobnych miejsc.

Opis układu i jego działania

Schemat obwodu inteligentnego sterownika oświetlenia pokazano na rysunku 1. Składa się on z regulatora napięcia 5 V typu 7805 (IC1), mikrokontrolera AVR (MCU) ATtiny85A (IC2), transoptora PC817 (IC3), czujnika ruchu PIR, dwukolorowej diody LED ze wspólną katodą (BICO1), kontakttronu normalnie otwartego (S2), przekaźnika na 12 V, z pojedynczym stykiem przełącznym (RL1), dwóch tranzystorów npn typu BC547 (T1 i T2) oraz kilku innych elementów.

Mikrokontroler IC2 to mózg sterownika, który, na podstawie stanu dwóch wejść informujących o: ruchu w pomieszczeniu i położeniu drzwi, steruje przekaźnikiem włączającym światło oraz informacyjną diodą LED. Obecność osoby w pomieszczeniu wykrywana jest przez czujnik ruchu typu PIR, natomiast położenie drzwi wykrywa czujnik kontakttronowy. Jeśli czujnik wykryje jakikolwiek ruch, jego wyjście przyjmuje stan wysoki. W przeciwnym razie wyjście pozostaje w stanie niskim.

Styki kontakttronu pozostają rozwarne, gdy drzwi są otwarte. Gdy drzwi są zamknięte, styki się zamykają. Stany czujnika PIR, kontakttronu i wynikające z tego stany przekaźnika oraz diody LED podane są w tabeli nr 1. Podaje ona stany wejściowe i odpowiadające im stany wyjściowe dla każdej sytuacji. Jeśli pomieszczenie jest puste, a drzwi są zamknięte, to dioda LED świeci na zielono. Przekaźnik włącza się, a światło jest gaszone aby oszczędzać energię. Jeśli ktoś otworzy drzwi, dioda LED miga naprzemiennie na czerwono i zielono, do momentu aż osoba wejdzie do pomieszczenia przebieralni.



Rysunek 1. Schemat ideowy układu sterownika oświetlenia

Przekaźnik włącza się na około dziesięć sekund, a następnie wyłącza się. W tym czasie zapala się żarówka (Load1) załączana przez przekaźnik. Jeśli osoba wejdzie do pomieszczenia w ciągu dziesięciu sekund po otwarciu drzwi, to przekaźnik pozostaje aktywowany, a dioda LED emituje kolor „pomarańczowy” (świecą obie sekcje czerwona i zielona) aż do momentu zamknięcia drzwi. Następnie LED świeci na czerwono, sygnalizując, że pomieszczenie jest zajęte, a drzwi są zamknięte.

Wyjście mikrokontrolera jest izolowane od zasilania przekaźnika za pomocą transoptora IC3. Przekaźnik jest sterowany przez transoptor. Chroni to MCU przed wszelkimi zakłóceniami w obwodzie zasilania 12 V przekaźnika. Dioda gasząca D1 tłumi wysokie napięcie indukowane na cewce przekaźnika, gdy jest on wyłączany. W obwodzie sterowanym: Load1 jest światłem głównym, a Load2 jest opcjonalnym światłem czuwania.

Oprogramowanie

Kod źródłowy (trial_room.ino) został stworzony przy użyciu Arduino IDE. Arduino IDE normalnie nie obsługuje mikrokontrolerów ATtiny85A i dlatego musimy dodać bibliotekę go obsługującą.

Dodanie wsparcia dla ATtiny

Aby dodać obsługę ATtiny do Arduino IDE, wykonaj poniższe czynności.

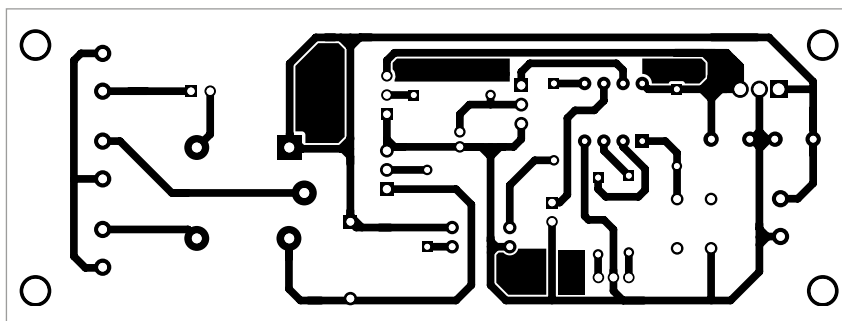
Wybierz z menu:

1. Plik → Preferencje i w Adresach URL Menedżera dodatkowych płytek wpisz adres: <https://bit.ly/3Tlfopn>.
2. Wybierz: Narzędzia → Płytki → Menedżer płytek. Przewiń listę i znajdź ATtiny autorstwa Davisa A. Mellisa. Kliknij na tę pozycję i zainstaluj ją. Teraz ATtiny pojawi się na liście w menu Płytki.
3. Wybierz ATtiny25/45/85 w menu Narzędzia → Płytki. Wybierz ATtiny85 w menu Narzędzia → Procesor. Teraz wybierz częstotliwość 8 MHz (oscylator wewnętrzny) w menu Narzędzia → Zegar.

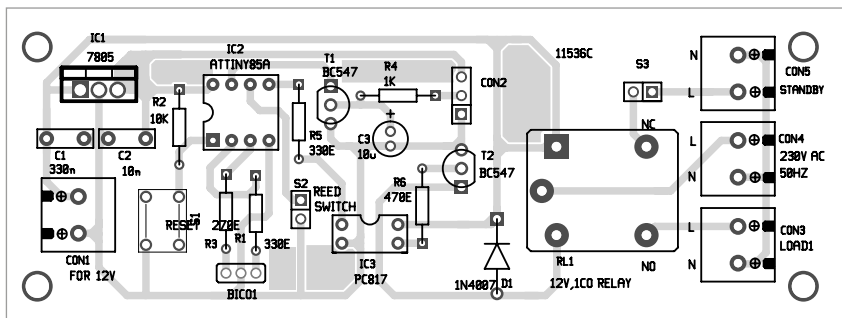
Arduino jako programator (ISP)

Aby Arduino Uno działało jako programator ISP, podłącz je do komputera (laptopa) i otwórz Arduino IDE. Upewnij się, że wybrałeś Arduino/Genuino Uno jako swoją płytkę.

**Kod źródłowy
tego projektu jest
dostępny do pobrania
ze strony
<https://bit.ly/3zXkVeY>**



Rysunek 2. Układ ścieżek na płycie drukowanej



Rysunek 3. Rozmieszczenie elementów na płycie drukowanej

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

Półprzewodniki:

IC1: stabilizator napięcia 7805, obudowa TO220
IC2: mikrokontroler Attiny85A, DIP8
IC3: tranzystor BC177, DIP4
T1, T2: tranzystory BC547 npn, TO92
D1: dioda prostownicza 1N4007
BICO1: LED dwukolorowa, wspólna katoda
PIR: czujnik PIR (na 5 V)

Rezystory: (wszystkie 0,25 W, ±5% węglowe)

R1, R5: 330 Ω
R2, R4: 10 kΩ
R3: 270 Ω
R6: 4,7 kΩ

Kondensatory:

C1: 330 nF, ceramiczny
C2: 10 nF, ceramiczny
C3: 10 μF elektrolityczny

Inne:

CON1, CON3–CON5: 2-pinowe złącze zaciskowe, raster 5,08 mm
CON2: 3-pinowe żeriśkie złącze Berg podstawka DIP8 (precyzyjna)
S1: mikroprzycisk do druku (uwaga, rozmiar 6 mm × 4 mm)
S2: czujnik kontaktronowy drzwiowy, np. AVT HO-03 BIAŁY
S3: włącznik/wyłącznik na 230 V
RL1: przekaźnik JQC3FF/012/1Z – 12 V, styk przełączny
LOAD1, LOAD2 – lampa 230 V AC
– zasilacz 12V DC (lub akumulator 12 V)
– płytka drukowana (wymiar 111,76 mm × 41,402 mm)

Tabela 1. Stany wejść i wyjść sterownika

Sytuacja	Wejścia		Wyjścia	
	Czujnik PIR	Kontaktron	Dioda LED	Przełącznik
Pomieszczenie	czyha	zwarty	zielona	wyłączony
Drzwi	czyha	rozarty	naprzemienie zielona + czerwona	włączony na 10s
Osoba wchodzi lub wychodzi	wykrywa ruch	rozarty	zielona + czerwona	włączony
Pomieszczenie	wykrywa ruch	zwarty	czerwona	włączony

Otwórz przykładowy szkic Arduino ISP i prześlij szkic Arduino ISP do Arduino Uno. Arduino Uno jest teraz programatorem ISP.

Red. EdW: Obszerny opis programowania ATtiny85 można znaleźć tutaj: <https://elportal.pl/projekty/arduino/105-budujemy-programator-attiny85-do-arduino-uno>.

Programowanie ATtiny85A

Domyślnie, układ ATtiny85A działa z częstotliwością 1 MHz. Aby działał z prędkością 8 MHz, wybierz z menu: Narzędzia → Nagraj Bootloader. Upewnij się, że płytka

Arduino Uno jest zaprogramowana do pracy jako ISP i jest połączona z ATtiny85A poprzez piny ISP. Teraz prześlij kod źródłowy (trial_room.ino) do ATtiny85A za pomocą płytki Arduino Uno.

Montaż i testowanie

Wygląd płytki drukowanej (rozmiar 1:1) pokazano na rysunku 3, a rozmieszczenie elementów na niej pokazano na rysunku 4. Po zmontowaniu obwodu na płycie drukowanej, należy zamknąć go w odpowiedniej obudowie.

Red. EdW: W miejscu S3 na płycie należy włutować zworę, a wyłącznik S3 należy zainstalować w obwodzie poza płytka.

W miejsce S2 najlepiej zastosować jest gotowy kontaktronowy czujnik drzwiowy co rozwiąże problem jego montażu i problemy estetyki. Część z magnesem mocujemy do skrzydła drzwi, a część z kontaktronem do ich framugi. Kontaktron podłączamy do sterownika przewodem dwużyłowym w izolacji. Kontaktron i magnes muszą być odpowiednio ustawione względem siebie, tak aby po zamknięciu drzwi S2 był zamknięty. Load1 to lampa główna używana jako inteligentne oświetlenie, a Load2 to lampka sygnalizująca gotowość. Możesz włączyć żarówkę małej mocy w trybie gotowości, aby była włączona, dopóki wyłącznik S3 nie będzie włączony, a przekaźnik nie zostanie zasilony (włączony przekaźnik RL1 oznacza, że Load1 jest włączony). ■

A. Samiuddhin

Ten artykuł był wcześniej opublikowany na łamach „EFY”, grudzień 2019 (efymag.com)

REKLAMA

Świat projektantów i programistów dla elektroniki w nowej odświeżonej. Odwiedź wiecznie młody

ELPORTAL.pl

Inteligentny sterownik urządzeń domowych z wykorzystaniem platformy Blynk



Artykuł opisuje sterownik do inteligentnego domu oparty na Internecie Rzeczy (IoT). Układ pozwala na sterowanie aplikacją w Smartfonie, czterech dowolnych domowych urządzeniach elektrycznych, takich jak: lampy, wentylatory, klimatyzatory czy pompy wodne. W tym projekcie sterowaliśmy dwoma żarówkami oraz elektryczną pompą ogrodową. Użyliśmy jednego czujnika (DHT-11) do monitorowania temperatury i wilgotności w pomieszczeniu, jednego czujnika wilgotności gleby oraz jednego czujnika ultradźwiękowego (HC-SR04) do pomiaru poziomu wody w wysoko zawieszonym zbiorniku. Stan wymienionych urządzeń i czujników, w tym: włączenie/wyłączenie światła, temperatura w pomieszczeniu, poziom wilgotności gleby oraz poziom wody, mogą być graficznie przedstawione na ekranie Smartfona.

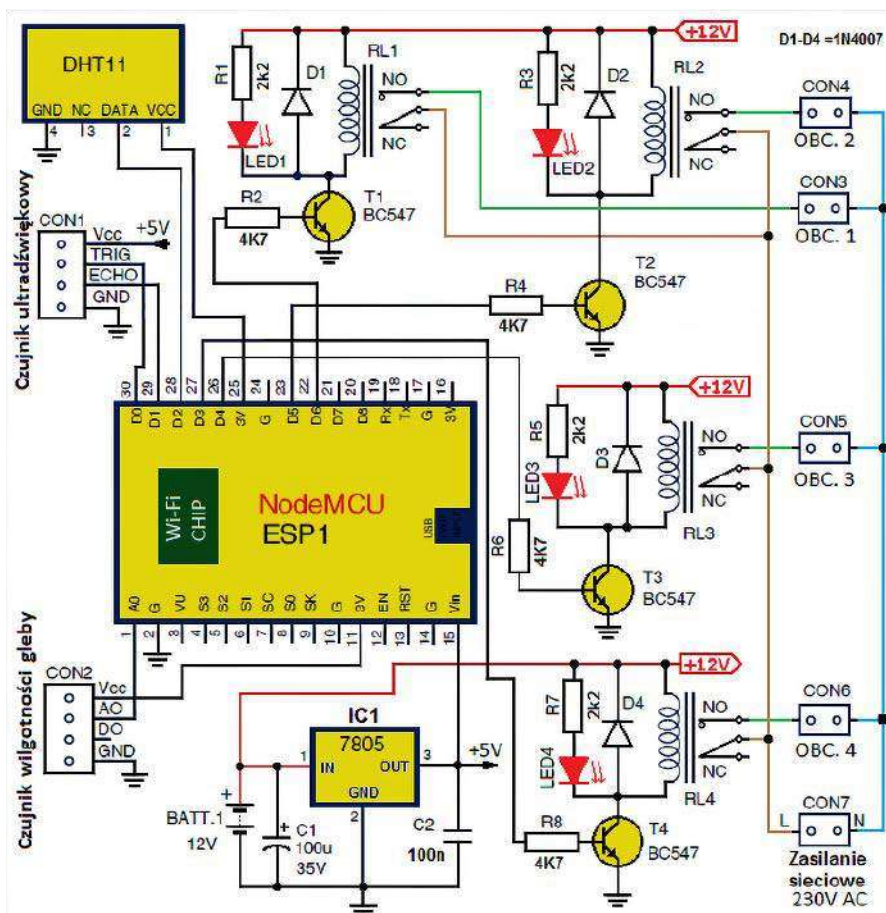
Układ i jego działanie

Schemat ideowy obwodu sterownika urządzenia inteligentnego domu opartego na IoT pokazano na rysunku 1. Jest on zbudowany z wykorzystaniem modułu NodeMCU (ESP1), czterech przełączników (RL1 do RL4) z cewkami na 12 V, sterowanych kluczami tranzystorowymi, stabilizatora 7805 (IC1), czujnika temperatury i wilgotności (DHT11), czujnika

wilgotności gleby na napięcie 5 V oraz czujnika ultradźwiękowego (HC-SR04) i kilku innych elementów. W tym projekcie Obciążenie 3 (Lampa 1) i Obciążenie 2 (Lampa 2) są złączane odpowiednio, przez przełączniki RL3 i RL2. Czujnik DHT11 służy do monitorowania temperatury i wilgotności w pomieszczeniu. W aplikacji Blynk wirtualne wskaźniki V5 i V6 służą do wyświetlania

zmierzonych: odpowiednio wilgotności i temperatury. W tym celu użyliśmy w kodzie poleceń: **Blynk.virtualWrite(V5, h)** i **Blynk.virtualWrite(V6, t)**, odpowiednio.

Czujnik wilgotności gleby służy do określenia stopnia wilgotności gleby w ogrodzie. Jeśli gleba jest sucha, to automatycznie włącza się silnik pompy ogrodowej i następuje nawadnianie roślin.



Rysunek 1. Schemat ideowy sterownika urządzeń domowych

Różne funkcje aplikacji Blynk i końcówki NodeMCU

Funkcja	Końcówka
Tab1: Pin Load3	D4 (NodeMCU)
Tab1: Pin Load2	D5 (NodeMCU)
Silnik pompy	D3 (NodeMCU)
Tab2: V7	Wirtualny poziom wilgotności
Tab2: V8	Wirtualny poziom wody

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

Półprzewodniki:

ESP1: Moduł Wi-Fi NodeMCU
IC1: 7805, stabilizator 5 V, obudowa TO-220
T1...T4: tranzystory npn, BC547B
D1...D4: diody prostownicze, 1N4007
LED1...LED4: diody LED, średnica 5 mm

Rezystory:

(wszystkie 0,25 W, $\pm 5\%$ węglowe)
R1, R3, R5, R2: 2,2 k Ω
R2, R4, R6, R8: 4,7 k Ω

Kondensatory:

C1: 100 μ F, 25 V elektrolityczny
C2: 0,1 μ F ceramiczny

Inne:

CON1, CON2: 4-pinowe żeńskie złącze Berg
CON3...CON7: 2-pinowe złącze zaciskowe
BATT1: 2-pinowe złącze zaciskowe akumulatora 12 V
DHT11: 4-pinowe żeńskie złącze Berg dla DHT11
RL...RL4: przełączniki 12 V, JQC3FF/012/1Z Hongfa
bateria lub akumulator 12 V
4-stykowy czujnik temperatury i wilgotności DHT11
4-stykowy czujnik wilgotności gleby na 5 V
4-stykowy moduł ultradźwiękowy HC-SR04



Rysunek 2. Tworzenie nowego projektu na platformie Blynk



Rysunek 3. Wybór NodeMCU jako modułu wykonawczego



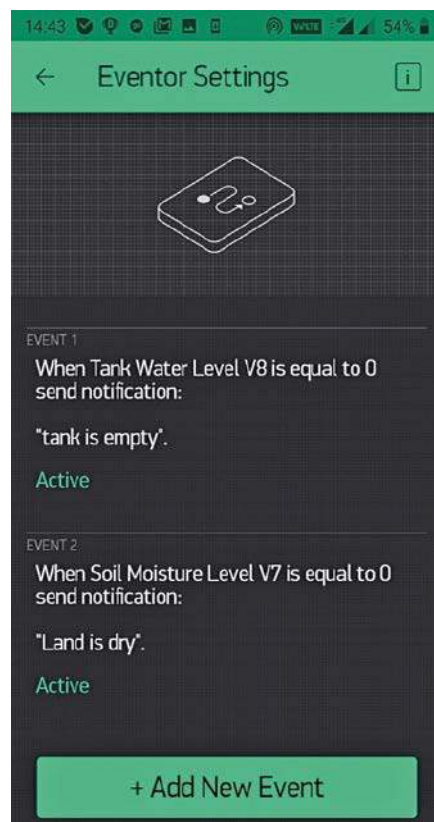
Rysunek 4. Pole widżetów do dodawania przycisków sterujących



Rysunek 5. Dodawanie zakładek sterujących



Rysunek 6. Ustawianie Tab2 dla poziomu wilgotności gleby



Rysunek 7. Ustawienia obsługi zdarzeń (eventor)

Czujnik ultradźwiękowy (HC-SR04) służy do pomiaru poziomu wody w zawieszonym wysoko zbiorniku wody (patrz rysunek 11). Wysokość zbiornika użytego w prototypie wynosiła 20 cm. Jeśli odległość do lustra wody, zmierzona przez czujnik, wynosi 0 cm, to aplikacja pokaże zbiornik jako pełny; a jeśli ta odległość wynosi 20 cm, to jako pusty. Jak pokazano na schemacie, końcówki D0 i D1 modułu NodeMCU są podłączone odpowiednio do końcówek Trig i Echo czujnika HC-SR04 mierzącego poziom wody, a pompa wody może być sterowana przez przełącznik RL1, który jest podłączony do końcówki D6 modułu NodeMCU przez klucz tranzystorowy (T1).

Czujnik DHT11 jest podłączony do końcówki D2 modułu NodeMCU. Silnik pompy ogrodowej jest załączany przez przełącznik RL4, który jest podłączony do końcówki D3 modułu NodeMCU przez tranzystor T2. Po włączeniu zasilania (poprzez podłączenie akumulatora 12 V), NodeMCU łączy się z siecią Wi-Fi.

Otwórz aplikację Blynk na telefonie komórkowym, aby rozpocząć sterowanie urządzeniami elektrycznymi. Wcześniej jednak potrzebujesz zainstalować oprogramowanie w NodeMCU. Do wpisania kodu dla Arduino i przesłania go na płytkę NodeMCU należy użyć Arduino IDE. Aplikacja Blynk jest wymagana na Smartfonie z systemem Android.

Pamiętaj, że kod autoryzacyjny wygenerowany w aplikacji Blynk będzie wymagany w kodzie Arduino. W tym celu utwórz najpierw projekt w aplikacji Blynk.

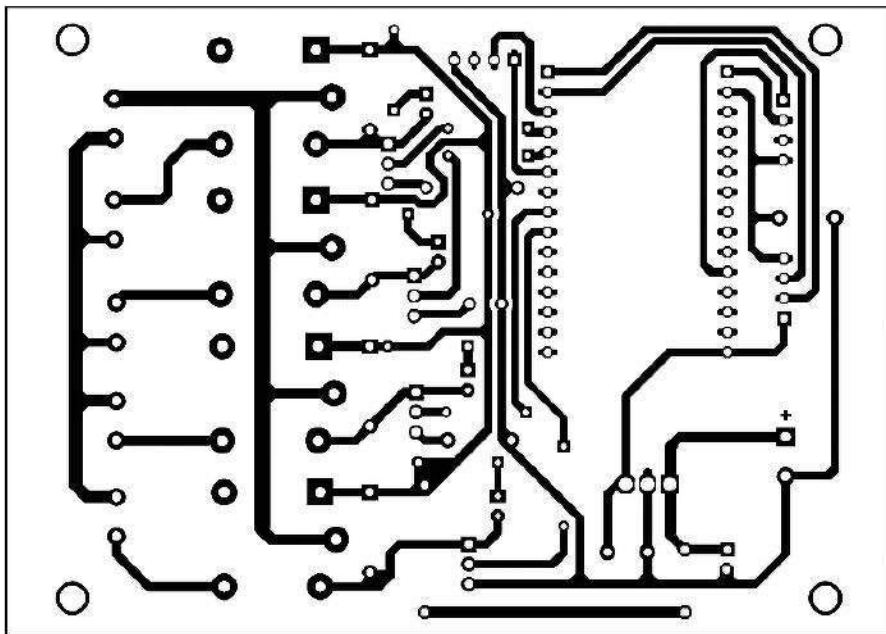
Tworzenie projektu na platformie Blynk IoT

Kroki tworzenia projektu w Blynk podano poniżej. Zainstaluj aplikację Blynk ze Sklepu Play, na Smartfonie z systemem Android. Zarejestruj się w aplikacji. Otrzymasz token uwierzytelniający na skrzynkę e-mail użytą podczas procesu rejestracji. Zanim użyjesz tokena, ponieważ będziesz go później potrzebował do kodu Arduino. Po rejestracji zaloguj się do aplikacji Blynk i utwórz swój projekt, wybierając polecenie **Nowy projekt** (rysunek 2). Nadaj projektowi odpowiedni tytuł i wybierz jako sprzęt moduł NodeMCU (rysunek 3). Teraz otwórz widżet i wybierz **Buttons** (przyciski) jak pokazano na rysunku 4. Zobaczysz teraz kolejny ekran, na którym możesz dodać etykiety wartości, aby móc wyświetlać temperaturę, wilgotność, super wykres dla reprezentacji graficznej zmian parametrów itp.

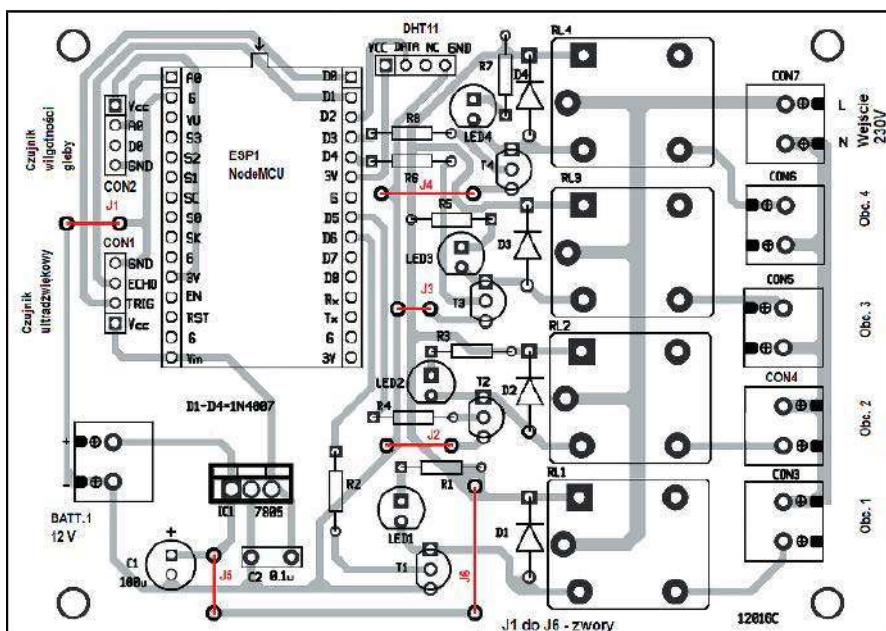
Dodaj dwie zakładki kontrolne: Tab1 i Tab2 (rysunek 5). W Tab1 dodaj Light1 (Lampa 1) i Light2 (Lampa 2) do sterowania obwodów, odpowiednio: Obciążenie 3 i Obciążenie 2. W Tab2 dodaj dwa wskaźniki poziomu: V7 i V8, aby wyświetlić odpowiednio rezystancję gleby



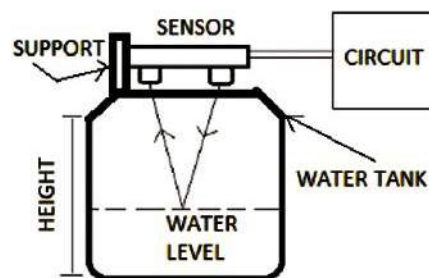
Rysunek 8. Prototyp autorów



Rysunek 9. Projekt płytki drukowanej – widok od strony ścieżek



Rysunek 10. Rozmieszczenie elementów na płytce drukowanej



Rysunek 11. Montaż czujnika ultradźwiękowego na zbiorniku wody

(proporcjonalną do jej wilgotności) i poziom wody w zbiorniku. Możesz także dodać obsługę zdarzeń (eventor) i komunikaty, aby otrzymywać powiadomienia o zmianach parametrów.

Teraz przyporządkuj kontrolki w aplikacji Blynk oraz obsługiwane nimi końcówki NodeMCU, jak podano w tabeli powyżej.

Prototyp autorów pokazano na rysunku 8.

Następnie otwórz kod `smart_home.ino` w Arduino IDE. Skonfiguruj NodeMCU z poświadczeniami Wi-Fi i wklej kod autoryzacyjny wygenerowany wcześniej w aplikacji Blynk. Skompiluj kod i prześlij go na płytkę NodeMCU.

Montaż układu i jego testowanie

Układ można zmontować na płytce uniwersalnej lub na specjalnie zaprojektowanej płytce drukowanej pokazanej na rysunkach 9 i 10. Układ ścieżek pokazano na rysunku 9, a rozmieszczenie elementów na rysunku 10. Po zmontowaniu układu zamknij go w odpowiedniej plastikowej obudowie (jest to konieczne ze względu na to, że na płytce część ścieżek jest pod napięciem 230 V). Włącz zasilanie układu (12 V). Otwórz projekt z aplikacji Blynk, aby sprawdzić, czy wszystkie urządzenia są połączone z Internetem. Jeśli są połączone, to czerwony znak w aplikacji zniknie.

Teraz możesz rozpocząć sterowanie urządzeniami, dotykając odpowiednich przycisków w aplikacji. Silnik pompy ogrodowej jest sterowany przez końcówkę D3 NodeMCU i tranzystor T4. Możesz go włączyć lub wyłączyć za pomocą aplikacji lub ręcznie. Możesz monitorować poziom wody w zbiorniku i jeśli zbiornik jest pełny, możesz ręcznie wyłączyć silnik pompy. Jeśli chcesz, aby silnik działał automatycznie, dodaj odpowiednie zdarzenie w Menedżerze zdarzeń, obsługiwane gdy zbiornik jest pusty lub pełny. ■

K. Muni Sekhar Reddy
Palla Siva Kumari

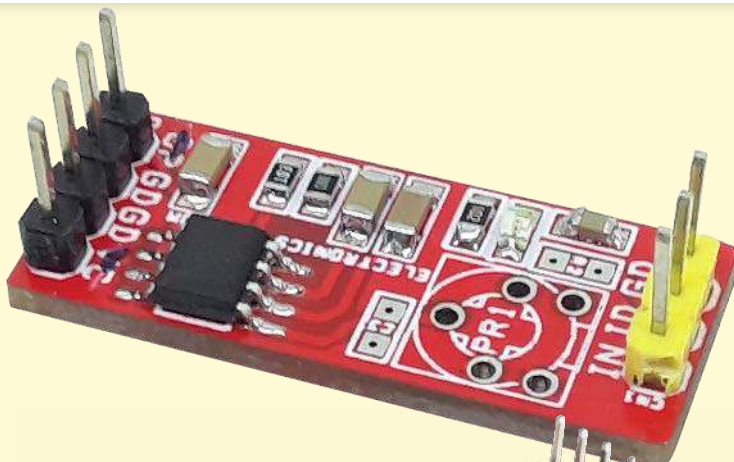
Ten artykuł był wcześniej opublikowany na łamach „EFY”, marzec 2020 (efymag.com)

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**.

Przetwornik częstotliwość/napięcie – obrotomierz z czujnikiem magnetycznym o zmiennej reluktancji

Obwód niżej pokazany jest przetwornikiem częstotliwości na napięcie, który może mieć zastosowanie w wielu aplikacjach. Zasadniczo jest to układ kondycjonowania sygnału współpracujący z czujnikiem pola magnetycznego, który wykorzystuje zjawisko zmiennej reluktancji. Predysponowany jest do detekcji i pomiaru prędkości obrotowej wału silników i podobnych maszyn.

Dokończenie artykułu na stronie:
<http://bit.ly/3WTP5hX>



Sterownik termoelektrycznej chłodziarki

Pokazany tu moduł jest kompletnym driverem mocy przeznaczonym dla 3-ampereowej TEC (Thermoelectric Cooler) chłodziarki (działającej w oparciu o efekt Peltiera). Sygnałem sterującym jest napięcie DC, a wyjście jest prądowe dla wykonawczego elementu, złącza Peltiera.

Dokończenie artykułu na stronie:
<http://bit.ly/3GhZl9G>



Niektóre projekty aktualnie dostępne tylko dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl:

1. RPi – stacja pogodowa IoT
2. Niskobudżetowy monitor jakości powietrza IoT oparty o RaspberryPi 4
3. Automatyczny system ogrodnictwa z NodeMCU i Blynk, ArduFarmBot 2
4. TinyML – Rozpoznawanie ruchu przy pomocy Raspberry Pi Pico
5. Wzmacniacz piezoelektryczny do gitary i skrzypiec
6. Wysokowydajny i niezawodny sterownik bipolarnego silnika krokowego
7. Sterownik silnika prądu stałego z wykorzystaniem przełącznika i mosfetu – interfejs Arduino
8. Przedwzmacniacz do mikrofonu MEMS
9. Super prosty czuły wykrywacz metali
10. Stymulator czaszkowy Arduino (Bio-BrainTuner)
11. Izolowany obwód wykrywania napięcia 250 V AC z pojedynczym wyjściem (wejście 250 V prądu przemiennego, wyjście 5 V)
12. Generator sygnałów AD9833
13. Obserwacja charakterystyk tranzystora
14. Wyświetlacz EKG z użyciem Arduino
15. Łatwy do zbudowania robot kroczący
16. Sonarowy theremin MIDI
17. Zamek elektroniczny na kod
18. Prosty tester tranzystorów
19. Zegar binarny z użyciem Microbit
20. Przetwornik częstotliwości na napięcie (tachometr) – przetwornik częstotliwości na napięcie z czujnikiem magnetycznym o zmiennej reluktancji

Miesięcznik „Elektronika dla Wszystkich” (12 numerów w roku) jest wydawany we współpracy z kilkoma redakcjami zagranicznymi



Wydawnictwo:
AVT-Korporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Adres redakcji:
03-197 Warszawa, ul. Leszczyńska 11
e-mail: edw@elportal.pl, www.elportal.pl

Redaktor merytoryczny
Paweł Sujko

Dział Reklamy:
Katarzyna Gugala
katarzyna.gugala@elportal.pl, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobański
jakub.sobanski@elportal.pl

Copyright AVT-Korporacja Sp. z o.o., Warszawa, ul. Leszczyńska 11. Projekty publikowane w „Elektronice dla Wszystkich” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki dla Wszystkich”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice dla Wszystkich” jest dozwolone wyłącznie po uzyskaniu pisemnej zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice dla Wszystkich”.

DTP, okładka, redakcja strony internetowej www.elportal.pl:
MAD Sp. z o.o.

Prenumerata:
W Wydawnictwie AVT, e-mail: 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

Sięgnij po archiwalne wydania
ELEKTRONIKI DLA WSZYSTKICH

Każde
wydanie
przejrzyj
on-line



Szybka
wysyłka
GRATIS

Zamów wygodnie na www.UlubionyKiosk.pl