

ELEKTRONIKA

dla wszystkich

nr 1/2023 (324) • styczeń • www.elportal.pl

Lampowy moduł przesterowania i zniekształceń do gitary

DIY PLUS
tylko dla prenumeratorów

PROJEKTY dla elektroników

- ▶ Zdalna stacja nadzoru (monitoringu)
- ▶ Ściemniacz nowoczesnego oświetlenia zasilanego z sieci, część 2

DIY dla wszystkich

- ▶ Wykrywanie i klasyfikacja obiektów za pomocą Raspberry Pi i uczenia maszynowego wykorzystującego platformę Edge Impulse
- ▶ Prosta ładowarka bezprzewodowa do smartfonów
- ▶ Lampa-sygnalizator przelotu Międzynarodowej Stacji Kosmicznej (ISS)
- ▶ Sterowanie natężenia światła diod LED przy pomocy techniki PWM

TUTORIALE

- ▶ Szkoła Konstruktorów
- ▶ Poziomy logiczne, część 1
- ▶ Zwrotnica do mini monitora PE, część 1
- ▶ Silniki krokowe w praktyce, część 2: Wybór i identyfikacja silników krokowych
- ▶ Pokój Nauczycielski



16,90 zł (w tym 8% VAT)



EP.com.pl

Największy portal dla elektroników konstruktorów

FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

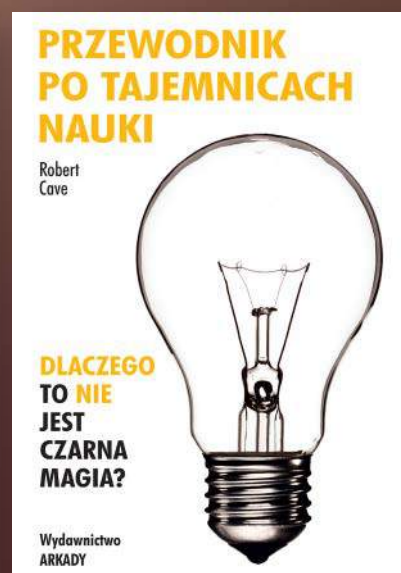
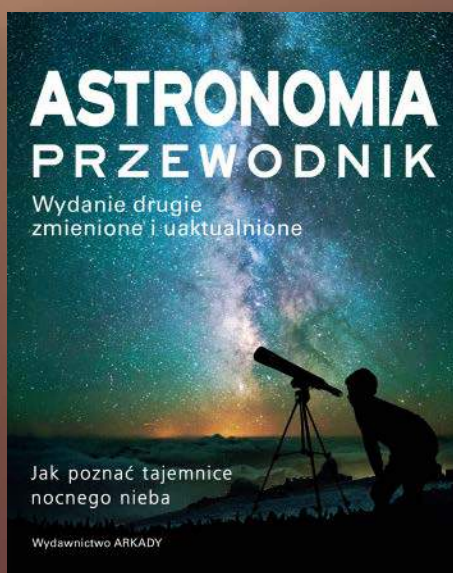
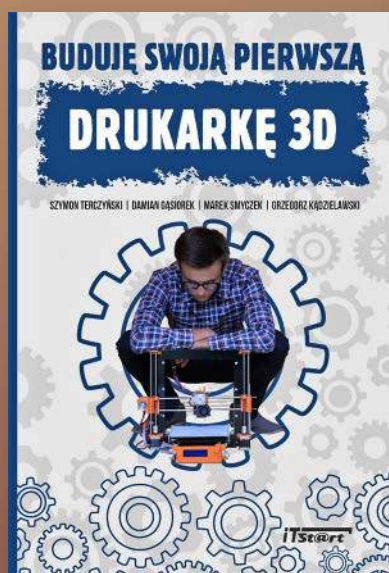
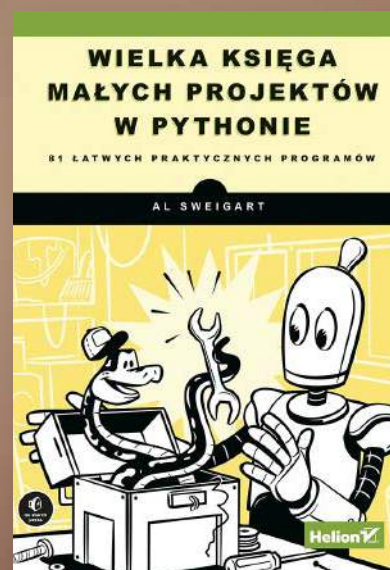
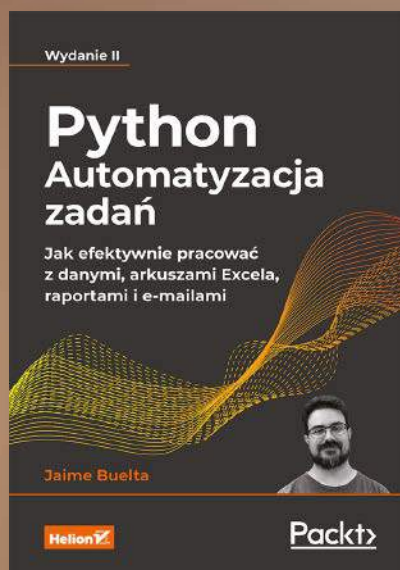
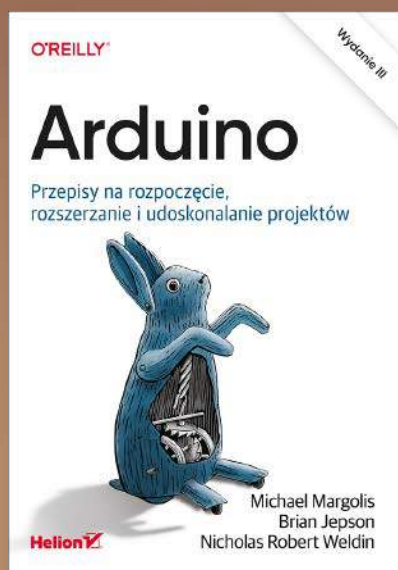
przekaźniki
półprzewodniki
złącza
przełączniki
radiatory
obudowy
i wiele więcej...

www.piekarz.pl

OLED

artronic
OPTOELEKTRONIKA
www.artronic.pl

KSIĄŻKI W ULUBIONYM KIOSKU Z RABATEM DO 30%



Zobacz pełną ofertę – ponad 500 tytułów!

Zamów wygodnie na UlubionyKiosk.pl



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

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

**na start
do 6* wydań gratis**

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

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

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

Oferta dotyczy prenumeraty drukowanej.

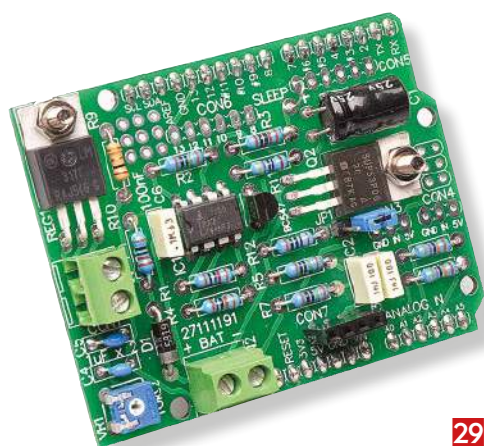
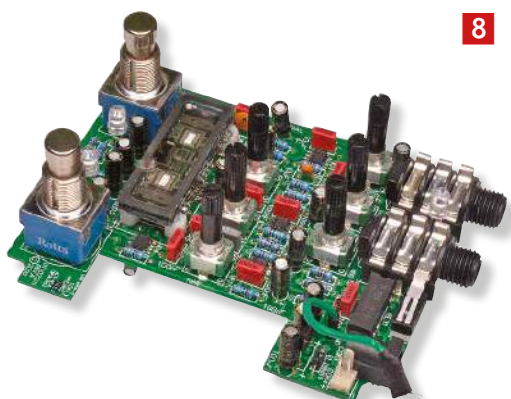
Wszystkie opcje prenumeraty i e-prenumeraty znajdziesz na stronie www.UlubionyKiosk.pl

Po opłaceniu prenumeraty przślemy 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

epresa.pl 8bfc5b24e



8 Projekty dla elektroników:

Lampowy moduł przesterowania i zniekształceń do gitary.....	8
Ściemniacz nowoczesnego oświetlenia zasilanego z sieci, część 2	21
Zdalna stacja nadzoru (monitoringu)	29

Tutoriale:

Szkoła Konstruktorów.....	34
Poziomy logiczne, część 1	51
Zwrotnica do mini monitora PE, część 1	55
Silniki krokowe w praktyce, część 2:	
Wybór i identyfikacja silników krokowych.....	61
Edukacja w EdW dla szkół i uczelni: Wykład 2 „Układy scalone”	65
Pokój Nauczycielski	76

29 DIY dla wszystkich:

Sterowanie natężenia światła diod LED przy pomocy techniki PWM.....	79
Lampa-sygnalizator przelotu Międzynarodowej Stacji Kosmicznej (ISS).....	82
Prosta ładowarka bezprzewodowa do smartfonów	86
Wykrywanie i klasyfikacja obiektów za pomocą Raspberry Pi	
i uczenia maszynowego wykorzystującego platformę Edge Impulse	88

DIY PLUS

Wysokiej klasy przedwzmacniacz mikrofonowy ze zmienną kompresją,	
regulowanym wzmocnieniem oraz funkcją redukcji szumów	91
Optycznie izolowane wejście analogowe dla Arduino	91

Rubryki stałe:

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



A za miesiąc w lutowym EdW

* Łatwe w budowie aktywne kolumny Hi-Fi z opcjonalnymi subwooferami. Nowoczesne odbiorniki telewizyjne mają doskonały obraz, ale ze względu na ich gabaryty, tj. znikomo małą grubość, nie zapewniają zadowalającego dźwięku. Praw fizyki nie da się oszukać. Dlatego warto zbudować oddzielne aktywne kolumny Hi-Fi, sterowane z wyjścia liniowego lub słuchawkowego telewizora. Konstrukcja opisana w tym artykule zapewni świetną jakość dźwięku z różnych źródeł, zarówno z odbiornika TV jak i z komputera lub odtwarzacza. Opcjonalnie zestaw kolumn można poszerzyć o subwoofery.



* **Wysokościomierz samochodowy z ekranem dotykowym**

Zwykle wysokościomierz jest stosowany w aparatach latających – szybowcach lub lotniach. Ta konstrukcja jest przewidziana do użycia w samochodzie. Zasilanie jest pobierane z gniazdka w samochodzie. Oprogramowanie dostosowano do odczytów wykonywanych w warunkach podróży samochodem. Przyrząd okaże się bardzo przydatny podczas podróży samochodem po górskich drogach, szczególnie przy określaniu wysokości punktów widokowych.

* **Programowany regulator temperatury**
Potrzeba dokładnego kontrolowania temperatury występuje nie tylko w wielu procesach przemysłowych, ale też w wielu zastosowaniach domowych, na przykład podczas wytwarzania w warunkach domowych piwa lub cydru bardzo ważne jest utrzymywanie na stałym, określonym poziomie temperatury w procesie fermentacji. Również wytwarzanie sera, przechowywanie produktów żywnościowych i ich przygotowanie do spożycia wymaga przestrzegania określonych warunków temperaturowych. We wszystkich tych zastosowaniach bardzo przydatne będzie urządzenie zbudowane na bazie elementów Peltiera.

* Plus zwykła porcja intrygujących projektów DIY.

* Plus wiele artykułów w Twoich ulubionych cyklach Tutoriali, w tym polecamy kolejny wykład w nowej rubryce „Edukacja w EdW”.

**W kioskach
od 30 stycznia**

Pamiętajmy o karpiach, które radośnie i niecierpliwie wyczekiwały świąt Bożego Narodzenia

Wszyscy znamy legendę o wynalazcy szachów, który zażądał od króla zapłaty za swój wynalazek w postaci zboża. Król miał kłaść ziarna zboża na szachownicy, zaczynając od jednego ziarenka i kładąc na kolejnych polach szachownicy dwa razy więcej ziaren niż na polu poprzednim. Ciąg 1, 2, 4, 8... na 64 polu osiąga astronomiczną wartość ponad 9 kwintylionów

9 000 000 000 000 000 000

Rozwój technologii układów scalonych przebiega również według funkcji potęgowej – co dwa lata liczba tranzystorów w układzie scalonym rośnie 2 razy (prawo Moore’a – wykład 2 w tym wydaniu EdW). Historia rozwoju układów scalonych ma już 62 lata, tj. 31 par lat, czyli jesteśmy mniej więcej w połowie układania ziarenek (tranzystorów) na polach szachownicy (w układzie scalonym). Daje to wynik rzędu 100 miliardów tranzystorów w układzie scalonym. Widać potęgę funkcji potęgowej, która zmierza do nieskończoności. Coraz wyższa skala integracji oznacza coraz większe zdolności obliczeniowe komputerów, które w roli mózgów tzw. sztucznej inteligencji (AI – Artificial Intelligence) z roku na rok wyposażają AI w coraz wyższą inteligencję. A funkcja potęgowa działa i pewnego dnia AI zrówna się z inteligencją człowieka, a potem już samodzielnie, bez udziału człowieka, będzie tworzyć „gmach wiedzy” rosnący do nieskończoności. Moment, gdy to się stanie został nazwany Osobliwością (Singularity). Takie prognozy głosi Rajmond Kurzweil, współczesny Nostradamus (przewidział m.in. powstanie Internetu i smartfona), człowiek o tak wybitnych i niepospolitych talentach, że trudno nie liczyć się z jego zdaniem. Osobliwość ma nastąpić nie później niż w 2045 roku. Kurzweil ma 74 lata i wszystko robi (łyka potworne ilości pigułek), żeby dożyć tego momentu i „przepisać” swój mózg do komputera działającego identycznie jak mózg człowieka, co ma oznaczać osiągnięcie nieśmiertelności, jeśli powłokę cielesną uznać za nieistotną. Chętnie podzielałbym entuzjazm Kurzweila wyczekującego Osobliwości, gdyby nie mąciła go myśl o losie karpi wyczekujących świąt Bożego Narodzenia. Bo do czego będzie potrzebny Człowiek, istota przeciętnie inteligentna, w świecie super inteligentnych AI. Radzę nie wypuszczać z rąk wtyczki. Chociaż AI bywają też ludzkie. Dwadzieścia lat temu podczas pobytu w USA zadzwoniłem się do informacji biletowej i gdy zorientowałem się, że rozmawiam z maszyną (nie zdała testu Turinga), trochę nerwowo (wstyd się przyznać, ale użyłem słów niecenzuralnych) zacząłem domagać się połączenia z człowiekiem, maszyna zareagowała – „OK, mister Marciniak, don't be nervous, I'm connecting you to the operator”.

Mam nadzieję, że Kurzweil się myli, przynajmniej co do daty Osobliwości. Nie widzę przestrzeni dla dalszego istotnego wzrostu skali integracji układów scalonych. Po 60 latach technologia krzemowa wyczerpała możliwości dalszego rozwoju. Tranzystor nie może być mniejszy niż atom. A pomysły na trójwymiarowe stopy chipów też mają ograniczenie w wydzielanej mocy i maksymalnej temperaturze. Potrzebna jest jakościowo nowa technologia – może komputer kwantowy, o którym mówi się już od 40 lat. Może DNA, czy szerzej bionika.

A może nastąpi okres stagnacji. Przydałby się. Mam wrażenie, jakbym siedział na krześle karuzeli, która ciągle nabiera szybszych obrotów. Mam nadzieję, że w rozwoju mikroelektroniki wchodzimy w okres dekadencji technologii krzemowej, a innej na skalę produkcyjną nie mamy.

W 1969 r., gdy Człowiek wylądował na Księżycu, było oczywiste, że najwyżej 20 lat dzieli nas od lądowania Człowieka na Marsie. Minęło 50 lat i Amerykanie mozolnie szykują się do powrotu na Księżyc. To daje do myślenia. Przykro mi, mister Kurzweil, ale pomny losu karpi, głosuję za opóźnieniem świąt Osobliwości.

Wiesław Marciniak

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



Zrozumieć tranzystory bipolarne

Jestem prenumeratorem EdW od 4 lat. Zdecydowanie odpowiada mi kierunek zmian, jakie w tym roku zaszły w EdW. Podziwiam projekty z Australii, bardzo inspirujące – wszystkie chciałoby się zrobić, ale na razie odkładam to na później. Brak czasu i z gotówką ostatnio słabo. Dwa na pewno kiedyś zrobię – Regulator obrotów silników DC i Kolumny głośnikowe. Nigdy przedtem nie czytałem artykułów tak szczegółowo opisujących projekty. Szacunek dla autorów, że nie pomijają najdrobniejszych szczegółów. Czytam też systematycznie kolejne odcinki nowych seriali tutorialowych... Ostatnio najbardziej zaciekały mnie „Silniki krokowe”. Z przyjemnością dokładnie przeczytałem wszystkie odcinki serii „Zrozumieć tranzystory”. Bardzo precyzyjny i szczegółowy opis, choć mam pewien niedosyt. Autor omawia tranzystor na poziomie modeli, a mnie brakuje podejścia fizycznego. Schemat zastępczy pozwala wszystko wyliczyć, ale nie daje poglądowego obrazu, co się dzieje w tranzystorze...

Kornel J.

Współpraca z EdW

EdW kupuję w empiku, raczej nieregularnie. W tym roku poszedłem na emeryturę. Byłem i ciągle w niewielkim wymiarze jestem nauczycielem w technikum. Mam dużo wolnego czasu. Mieszkam pod Warszawą i chętnie bym się włączył w jakąś współpracę z redakcją EdW.

A.K.

Red. Z Panem A.K. skontaktowaliśmy się i raczej nie uda nam się współpracować, ale publikujemy ten list, bo poszukujemy do współpracy redakcyjnej adiutatora przetłumaczonych tekstów. Wymagana głównie dobra znajomość terminologii stosowanej w zakresie elektroniki aplikowanej przez konstruktorów hobbistów. Oczywiście, również niezła znajomość języka polskiego i pewne obycie z fachowymi tekstami w języku angielskim. Prosimy o zgłoszenia na adres info@elportal.pl

Co z wykładami?

Z wielkim zainteresowaniem przeczytałem wykład na temat „Od mikroprocesora do mikrokontrolerów”. Umiejscowienie różnych zdarzeń na osi czasu bardzo pomaga w poukładaniu sobie tego potężnego kawałka wiedzy. Szkoda, że nie podajecie w zapowiedziach tytułu kolejnego wykładu. Najlepiej byłoby znać plan wykładów na kilka numerów w przód.

Red. Takiego konspektu wykładów nie możemy podać. Praca redakcyjna nad kolejnym numerem to wyścig z czasem i niektóre decyzje podejmowane są „w locie”. Prawdopodobnie trzeci wykład będzie o tranzystorach, ale w chwili oddawania tego numeru do druku (16.12) nie jest to na 100% pewne. A poza tym, ciągle analizujemy potrzeby zgłaszane przez Czytelników, aby dopasować do nich nasz plan tematyczny.

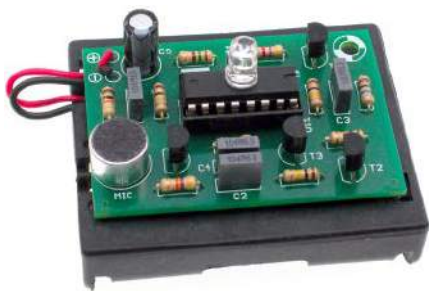
Patronat AVT

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 przesłanie listu w EdW 09/2022 (wydanie dostępne na www.ulubionykiosk.pl) i zgłoszenie akcesu do PATRONATU AVT. Zgłoszenia prosimy wysłać na adres: prenumerata@avt.pl.

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

Najbardziej popularne kity AVT

Zobacz ranking TOP-100 najbardziej popularnych kitów AVT na <https://elportal.pl/kityavt>



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



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



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



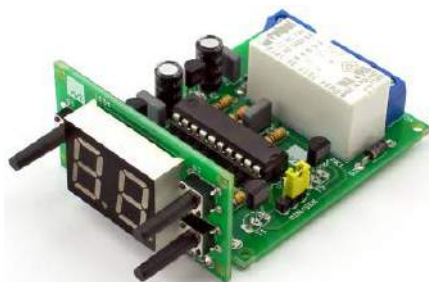
AVT5540 Radio FM z RDS
<https://sklep.avt.pl/avt5540.html>



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



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



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



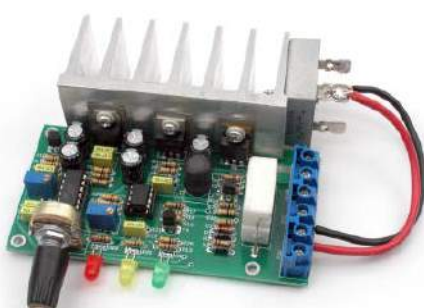
AVT990 Automatyczny włącznik świateł
<https://sklep.avt.pl/avt990.html>



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



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



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



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

Pełna oferta na: sklep.avt.pl

Lampowy moduł przesterowania i zniekształceń do gitary

Czy tęsknisz za tym prawdziwym „lampowym brzmieniem” gitary i efektami zniekształceń? Co powiesz na ten projekt – wykorzystuje on unikalną podwójną triodę niskonapięciową, więc przyznasz, że to prawdziwa okazja!

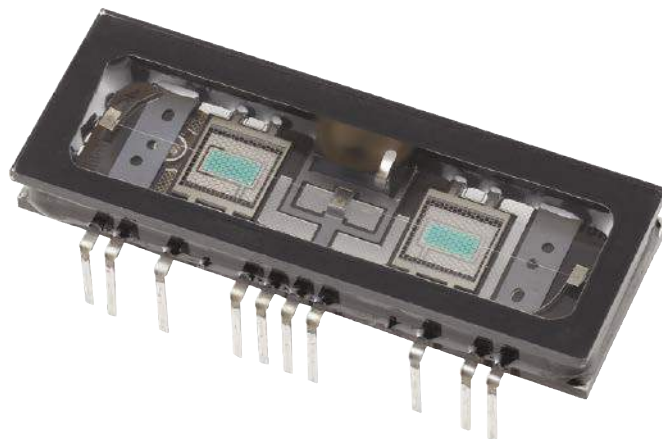
Gitary elektryczne prawie zawsze są używane (przynajmniej profesjonalnie) z jakimś rodzajem efektów dźwiękowych. Gitary akustyczne z przetwornikiem elektrycznym również mogą skorzystać z dodatkowych efektów.

Red. EdW: zdecydowaliśmy się używać terminu „przycisk (nożny)”, a nie klasycznego „pedał”, bo ten ostatni ma ściśle określone znaczenie, zwłaszcza pod względem rozmiaru i sposobu uruchamiania, jak pedały gazu czy hamulca w samochodzie. Tu w projekcie, jak widać na fotografiach, mamy pudełko z dwoma niewielkimi przyciskami, dosłownie tuż obok pokręteł. W przypadku realizacji tego projektu Red. EdW sugerowałaby jednak gorąco zastosowanie klasycznych pedałów umieszczonych obok obudowy, lub nawet rozdzielenie obudowy i pedałów i połączenie ich kablem o długości ok. 1 metra, co umożliwiłoby w wygodny sposób zmieniać nastawy, bez konieczności schylania się do podłogi. Na dodatkowej fotografii profesjonalne pedały gitarowe. Warto zauważyć, jak „delikatnie” obchodzono się z nimi na estradzie.

Wśród wielu dostępnych efektów dźwiękowych, przesterowanie i zniekształcenia są prawdopodobnie najbardziej popularne. Niektóre przystawki wytwarzają ostre zniekształcenia o jakoby gitarowym brzmieniu (jak „fuzz box”), podczas gdy inne zapewniają bardziej łagodną formę deformacji dźwięku.

Przystawki efektów powszechnie wykorzystują obwody z półprzewodnikami takimi jak JFET do zapewnienia tych zmian brzmienia, a czasami diody krzemowe do wprowadzania zniekształceń.

Jednak „Święty Graal” efektu przesterowania osiąga się w układach lampowych. Podczas gdy niektóre tranzystorowe przystawki



przesterowania próbują emulować efekt przeciągnięcia w lampach elektronowych przy przesterowaniu, nic nie zastąpi prawdziwego efektu.

Do tej pory trudno było umieścić lampy w małej obudowie efektów dźwiękowych. Ale teraz sytuacja jest inna, gdy kompaktowa niskonapięciowa podwójna trioda 6P1 jest dostępna w ofercie producenta instrumentów muzycznych, firmy Korg.

Przedstawiliśmy ją w SC w styczniu 2020 roku w naszym przedwzmacniaczu lampowym (siliconchip.com.au/Article/12217).

Ten nowy projekt SC opisany tutaj może być stosowany do wytwarzania zniekształceń lub efektu przesterowania, albo obu razem. W zestawie znajdują się dwa stopnie zniekształceń i/lub przesterowania, przy czym pierwszy stopień może być używany samodzielnie lub w połączeniu z drugim stopniem, który jest włączany przez przycisk PODBICIE (Boost albo dopalacz).

Przesterowanie vs zniekształcenia

Główna różnica pomiędzy przesterowaniem a zniekształceniami polega na rodzaju wytwarzanej deformacji dźwięku.

Z przesterowaniem mamy do czynienia, gdy wzmacniacz jest zasilany sygnałem o wysokim poziomie, co powoduje, że wyjście jest w stanie nasycenia, a sygnał wyjściowy „wygładzony” i ostatecznie ograniczony lub obcięty. Tak więc przy niskich poziomach sygnału nie ma zniekształceń lub są one niewielkie. Zniekształcenia rosną wraz ze wzrostem poziomu sygnału.

Gdy wzmacniacz staje się przesterowany, głośność pozostaje stała i nie wzrasta znacząco wraz ze wzrostem poziomu sygnału wejściowego.

Efekt ubocznym nadmiernejysterowania jest to, że ma ono tendencję do działania również jako efekt wybrzmiewania, gdzie poziom głośności pozostaje stały przez pewien czas po uderzeniu w strunę. Efekt zatkania (wybrzmiewania) trwa do momentu, gdy sygnał z gitary spadnie poniżej poziomu ograniczania.

Rodzaj zniekształceń przy przesterowaniu zależy od sposobu ograniczania poziomu sygnału we wzmacniaczu. W przypadku lamp ograniczanie sygnału jest zwykle asymetryczne, przy czym sygnał o jednej polaryzacji jest obcinany bardziej niż o przeciwnej. Red. EdW: wynika to z zasady pracy lampy elektronowej, której charakterystyka siatkowa jest w oczywisty sposób nieliniowa.

Efekt obecności zniekształceń różni się tym, że istnieje celowa metoda zniekształcenia sygnału nawet przy niskich jego poziomach, a poziom

Cechy

- Dwa stopnie zniekształceń
- Wysoka impedancja wejściowa pasuje do większości przetworników
- Regulatory wzmocnienia, poziomu wyjściowego, zniekształceń i barwy dźwięku
- Przetłączniki LINIA i PODBICIE z sygnalizacją diodami LED
- Umieszczony w wytrzymałej obudowie z aluminiowego odlew ciśnieniowego
- Brak wysokich napięć
- Wykorzystuje podwójną triodę Nutube w układzie beztransformatorowym
- Widoczna jest fluorescencja anod lampy Nutube
- Żywotność lampy Nutube 30 000 godzin
- Niski pobór mocy
- Zasilanie z baterii lub z zasilacza DC
- Zachowanie fazy sygnału od wejścia do wyjścia
- Automatyczne i ciche włączanie/wyłączanie
- Zabezpieczenie przed odwrotną polaryzacją zasilania



Materiały dodatkowe dostępne są na stronie
Silicon Chip:
<https://bit.ly/3BOEp6p>

Materiały dodatkowe są również dostępne
na stronie edw.elportal.pl:
<https://bit.ly/3jemzDo>

Lampowy moduł przesterowania i zniekształceń (dźwięku) gitary (Valve Guitar Overdrive & Distortion Pedal) jest umieszczony w solidnej obudowie z aluminiowego odlewu ciśnieniowego, nie tylko dla zminimalizowania szumów, ale także dla zapewnienia, że gitarzysta (lub gitarzystka) o ciężkiej stopie nie wyrządzi żadnych szkód podczas nawet najbardziej ekspresyjnej gry! Moduł zasilany jest napięciem 9-12 V DC (tak, naprawdę zawiera lampę!), więc możesz go używać z zasilaczem wtyczkowym lub nawet z baterią.

wyjściowy nie jest ograniczony (obcinany) tak bardzo jak w przypadku przesterowania. Innymi słowy, generalnie występują pewne zniekształcenia na wszystkich poziomach sygnału. W dalszej części artykułu zamieściliśmy kilka przebiegów z oscyloskopu, które pokazują różnice pomiędzy przesterowaniem a zniekształceniem sygnału (Oscylogram1–Oscylogram8).

Nasz moduł Guitar Overdrive & Distortion (przesterowanie i zniekształcenia dźwięku gitary) ma regulowany za pomocą pokręteł poziom przesterowania lub zniekształceń.

Jeśli pokręta zniekształceń ustawione są na minimum, a wzmocnienie zwiększone, przystawka działa w trybie przesterowania, „wyglądając” wyższe poziomy sygnału. Jeśli pokręta regulacji zniekształceń ustawione są na większe zniekształcenia, wówczas przystawka działa jak w trybie zniekształceń, przy czym poziom wzmocnienia decyduje o tym, czy generuje ona również efekt przesterowania.

Regulator poziomu zniekształceń w każdym stopniu może być ustawiony w pozycji środkowej dla uzyskania minimalnych zniekształceń, lub bliżej obu końców dla uzyskania większych zniekształceń. Przy ustawieniu w lewo, ujemna połówka sygnału akustycznego jest zniekształcona, ale dodatnia połówka już nie. I odwrotnie, w pozycji skrócenia w prawo, dodatnia połówka sygnału akustycznego jest zniekształcona, a ujemna już nie tak znacząco.

Moduł Przesterowanie i Zniekształcenia posiada dwa stopnie generujące zniekształcenia, przy czym oba razem są używane po wybraniu opcji „PODBICIE” (boost). Jeśli więc pierwszy stopień jest ustawiony na zniekształcenie dodatniej połówki sygnału, a drugi na zniekształcenie ujemnej połówki sygnału, przy włączonym podbiciu obie połówki przebiegu akustycznego będą zniekształcone. Przy wyłączonym podbiciu obecne są tylko zniekształcenia wprowadzane przez pierwszy stopień.

Różnica ta jest bardziej zauważalna, jeśli poziom sygnału podawanego na drugi stopień zostanie zredukowany tak, aby odpowiadał poziomowi sygnału podawanego na pierwszy stopień. Można to osiągnąć poprzez regulację potencjometru nastawnego wewnątrz modułu.

Moduł wyposażony jest w regulację barwy dźwięku, która umożliwia obcięcie wysokich tonów. Częstotliwość obcięcia jest regulowana w zakresie od około 2 kHz do 23 kHz. Niższa częstotliwość obcięcia redukuje zniekształcenia harmoniczne. Pozwala to uzyskać pożądane brzmienie.

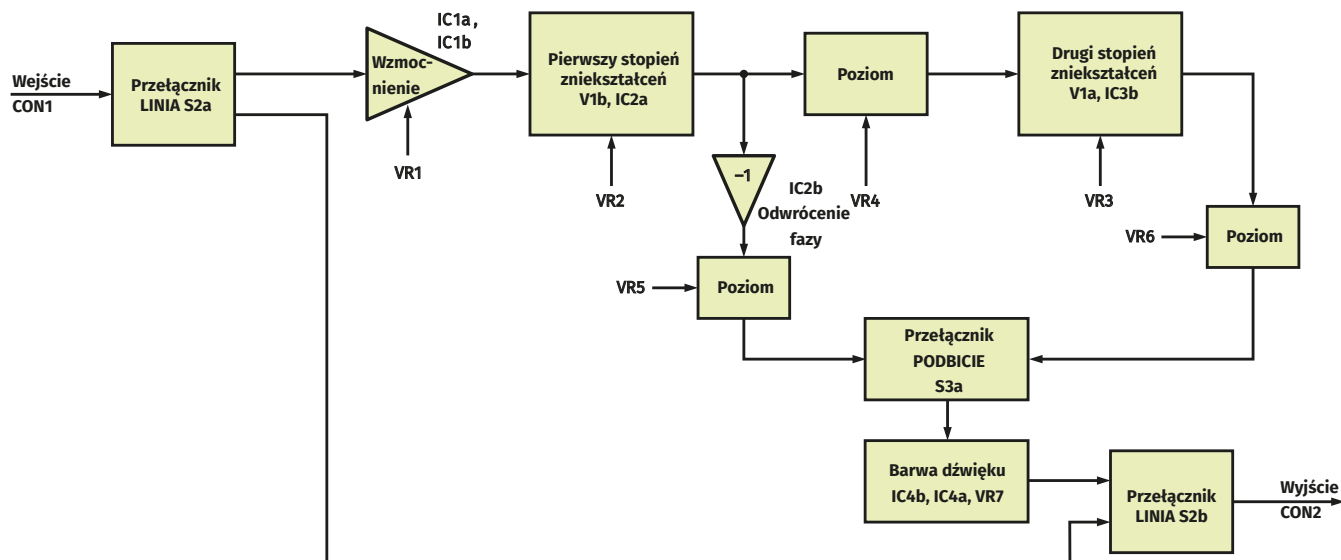
Poziomy wyjściowy, z włączonym lub wyłączonym podbiciem, są również regulowane. Sposób ich ustawienia zależy od pożądanego efektu. Poziom przy wyłączonym podbiciu jest zazwyczaj ustawiony tak, aby zapewnić taki sam poziom wyjściowy jak przy włączonym obejściu całego modułu (przycisk LINIA – Bypass).

Kiedy aktywny jest przycisk LINIA, sygnał wejściowy jest bezpośrednio podłączony do wyjścia. Gdy nie jest on włączony, sygnał przechodzi przez obwody zniekształceń i przesterowania.

Możesz ustawić poziom wyjściowy sygnału, gdy wybrane jest podbicie, na wyższy poziom, lub na taki sam jak przy wyłączonym podbiciu. Ogólnie rzecz biorąc, podbity sygnał na wyjściu i tak brzmi głośniejsze, ze względu na większe wypełnienie sygnału (przypomina on bardziej sygnał prostokątny) i dodane harmoniczne.

Prezentacja

Moduł jest umieszczony w solidnej obudowie z odlewu aluminiowego, ze względu na obsługę nogą. Posiada dwa przełączniki nożne, sześć pokręteł i trzy diody LED. Nad dwiema anodami podwójnej triody umieszczono przezroczyste oprawki, aby można było obserwować ustawienia polaryzacji siatki (więcej o tym później) i aby każdy mógł zobaczyć jak świecą Twoje lampy o magnetycznym powabie.



Rysunek 1. Schemat blokowy modułu zniekształceń Nutube. Gdy nie jest używane obejście LINIA, sygnał jest wzmacniany i buforowany przez wzmacniacze operacyjne IC1a i IC2b, a następnie dalej wzmacniany i zniekształcany przez triodę V1b. Następnie jest on podawany do triody V1a w celu dalszego wzmacnienia i deformacji, a przełącznik PODBICIE decyduje o tym, czy do sekcji regulacji barwy dźwięku i na wyjście trafi sygnał z wyjścia pierwszego czy drugiego stopnia lampowego (poprzez potencjometry regulacji poziomu VR5 i VR6).

Dwa gniazda Jack 6,35 mm (1/4") z tyłu zapewniają wejścia i wyjścia sygnału, a gniazdo DC służy do zasilania. Urządzenie może być również zasilane z wewnętrznej baterii 9 V. Zasilanie włącza się automatycznie po włożeniu wtyczki do gniazda wyjściowego.

Obsługa

Na **rysunku 1** przedstawiono uproszczony schemat blokowy modułu Guitar Overdrive & Distortion Pedal. Sygnał z gitary podłączony do gniazda CON1 może przejść bezpośrednio do wyjścia na gnieździe CON2 poprzez przełącznik LINIA (S2b). Jeśli przełącznik LINIA nie jest wybrany, sygnał przechodzi do pierwszego stopnia wzmacnienia. Składa się on ze stopnia buforującego o wysokiej impedancji wejściowej (IC1a), regulacji wzmacnienia (potencjometr VR1) oraz wzmacniacza o wzmacnieniu 21 dB (IC1b).

Pierwszy stopień zniekształcający wykorzystuje jedną z triod Nutube (V1b) do zapewnienia wzmacnienia i zniekształceń. Poziom zniekształceń wytwarzanych przez ten stopień jest regulowana za pomocą potencjometru VR2.

Wyjście triody V1b jest buforowane przez wzmacniacz operacyjny IC2a. Ponieważ lampa V1b odwraca fazę sygnału, sygnał z wyjścia IC2a jest podawany do inwertera (IC2b), przywracającego jego pierwotną fazę. Poziom wyjściowy sygnału z inwertera IC2b jest regulowany za pomocą VR5, a następnie sygnał trafia na jedną stronę przełącznika podbicia, S3a.

Sygnał wyjściowy sprzed inwertera IC2b jest również podawany na potencjometr nastawny regulacji poziomu (VR4), a następnie podawany do drugiego stopnia zniekształceń. Dzięki temu drugi stopień zniekształceń i przesterowania może mieć taki sam poziom sygnału wejściowego jak pierwszy. W takim przypadku VR4 jest tak ustawiony, aby zmniejszyć poziom sygnału z pierwszego stopnia o około 15 dB.

Alternatywnie, VR4 może być ustawiony tak, aby dostarczyć sygnał o pełnym poziomie do drugiego stopnia zniekształceń, w celu maksymalnego obciążenia i przesterowania.

Układ drugiego stopnia zniekształceń jest taki sam jak pierwszego, tyle że wykorzystuje triodę V1a i bufor IC3b. Potencjometr VR3 ustawia poziom zniekształceń, natomiast poziom wyjściowy regulowany jest potencjometrem VR6. Powstały sygnał jest podawany na drugą stronę przełącznika podbicia, S3b.

Tak więc przycisk PODBICIE pozwala wybierać pomiędzy sygnałami z pierwszego lub drugiego stopnia zniekształceń. Wybrany sygnał trafia do regulatora barwy dźwięku z nastawianym obcinaniem wysokich częstotliwości, wybranym przy pomocy potencjometru VR7.

Wyjście z regulatora barwy dźwięku trafia następnie na jedną stronę przełącznika LINIA, S2b. Przełącznik LINIA wybiera pomiędzy tym sygnałem a sygnałem wejściowym na CON1 (gdy jest w trybie liniowym).

Podwójna trioda Nutube 6P1

Jedną z rzeczy, która sprawia, że lampa Nutube jest tak wyjątkowa, jest to, że może pracować przy bardzo niskim napięciu. Tradycyjne lampy wymagają wysokiego napięcia anodowego (powyżej 100 V).

Nutube 6P1 została opracowana przez firmy Korg i Noritake Itron z Japonii. Choć jest to podwójna trioda z bezpośrednio żarzoną włóknem katody, siatką i anodą, wykonana jest w sposób, który bardziej przypomina próżniowy wyświetlacz fluorescencyjny (VFD) niż tradycyjną lampę.

Nutube ma prostokątną szklaną obudowę, a każda trioda składa się z odpowiednika jednopikselowego VFD. Wewnętrzna konstrukcja zawiera bezpośrednio żarzoną katodę w postaci drutu finezyjnie poprowadzonego z przodu. Pod katodą znajduje się metalowa siatka. Za siatką umieszczona jest anoda, pokryta luminoforem, który świeci, gdy katoda jest podgrzewana.

Drut katody jest utrzymywany w naprężeniu, wskutek czego może on drgać podobnie jak struna gitarowa. (Nutube jest zresztą sprzedawany przez producenta instrumentów muzycznych). Te wibracje są niepożądane, ponieważ mogą być źródłem mikrofonowania, kiedy zewnętrzny dźwięk może się sprzęgać z włóknem katody i zmieniać (lub modulować) wzmacniany sygnał audio. W rezultacie te wibracje są słyszalne w dźwięku.

Staranna konstrukcja modułu z Nutube może zminimalizować mikrofonowanie. Obejmuje to ochronę lampy przed wibracjami otaczającego powietrza, poprzez zastosowanie elastycznego okablowania oraz tłumiące wibracje sposobu montażu.

Podczas pracy Nutube pobiera minimalny prąd, włóknem każdej katody wymaga prądu żarzenia o natężeniu zaledwie 17 mA. Prądy siatki i anody wynoszą łącznie około 38 μ A. Nutube najlepiej pracuje z napięciem anodowym 5–30 V. Z krzywych linii obciążenia wynika,

Specyfikacja

- Zasilanie: 9–12 V DC @ 47 mA z wyłączonymi diodami LINIA i PODBICIE (+6 mA dla każdej diody)
- Wzmocnienie: maksymalnie 32 dB (PODBICIE wyłączone); do 43 dB (PODBICIE włączone)
- Pasma przenoszenia: –0,6 dB przy 20 Hz. Od góry charakterystyka częstotliwościowa jest zależna od ustawienia potencjometru barwy dźwięku
- Regulacja barwy dźwięku: filtr górnozaporowy 20 dB/dekadę, punkt –3 dB zmienia się od 2,12 kHz do 23,4 kHz
- Maksymalny poziom sygnału na wejściu i wyjściu: 2,3V RMS przy zasilaniu 9 V; 3,3 V RMS przy zasilaniu 12 V
- Minimalny poziom sygnału przy przesterowaniu: 55 mV bez podbicia, 15,5 mV z podbicciem
- Stosunek sygnału do szumu: 82 dB w odniesieniu do 55 mV na wejściu i 55 mV na wyjściu

że w tym zakresie napięć anodowych napięcie siatki musi być dodatnie w stosunku do potencjału włókna katodowego.

Jest to przeciwieństwo tradycyjnej triody, w której napięcie anody jest znacznie wyższe, a napięcie siatki jest praktycznie zawsze ujemne w stosunku do katody. Zniekształcenia wprowadzane przez Nutube można regulować zmieniając napięcie siatki.

Szczegóły układu

Schemat ideowy jest pokazany na **rysunku 2**. Widać dwie połówki lampy Nutube w pobliżu górnego środka schematu, przy czym obie pracują jako wzmacniacze ze wspólną katodą; katody są podłączone do masy poprzez styk F3. Sygnały są przykładane do siatek (G2 i G1), a wynikowy wzmocniony sygnał pojawia się na anodach, A2 i A1. Anody są obciążone rezystorami podłączonymi do dodatniego napięcia zasilającego, Vaa.

Triody Nutube mają stosunkowo niską impedancję wejściową siatki i wysoką impedancję wyjściową. Dlatego zastosowano bufor; jeden do zapewnienia niskiej impedancji dla siatki każdej triody, a drugi do utrzymania wysokiej impedancji obciążenia anod.

Zastosowane wzmacniacze operacyjne (OPA1662A) mają bardzo niski poziom szumów i zniekształceń, około 0,00006% przy 1 kHz, 3 V RMS i wzmocnieniu 0 dB. Tak więc wzmacniacze operacyjne nie wpływają w żaden sposób na brzmienie sygnału. Jakikolwiek szum lub zniekształcenia, które mogłyby one wprowadzić, są maskowane przez szum pochodzący z triod.

Ścieżka sygnału jest następująca. Gdy przełącznik LINIA (S2a) znajduje się w pozycji „praca”, sygnał przechodzi przez koralik ferrytowy FB1 i rezystor ograniczający 100 Ω. Te, w połączeniu z kondensatorem 100 pF, tłumią sygnały RF przed wejściem do modułu, co mogłoby skutkować niepożądaną detekcją i odbieraniem stacji radiowych. Kondensator 100 pF stanowi również obciążenie dla przetworników piezoelektrycznych strun gitarowych.

Poprzez sprzężenie zmiennoprądowe sygnał jest podłączony do wejścia 3 wzmacniacza operacyjnego IC1a i dodany (przesunięty napięciowo) do napięcia podkładu równego połowie napięcia zasilania ($V_{aa}/2$) przez rezystor podciągający 1 MΩ. Impedancja wejściowa modułu jest więc wysoka i wynosi 1 MΩ, co czyni go odpowiednim nawet dla przetwornika piezo gitary elektrycznej.

Szyna połówkowego napięcia zasilania ($V_{aa}/2$) jest utworzona przez dwa rezystory 10 kΩ połączone szeregowo pomiędzy zasilaniem Vaa i masą. Jest ona bocznikowana kondensatorem 100 μF w celu usunięcia szumów zasilania i buforowana przez wzmacniacz IC3a o wzmocnieniu 0 dB.

Wyjście IC1a jest sprzężone zmiennoprądowo z regulatorem poziomu, VR1, który następnie zasila IC1b. Układ IC1b zapewnia 11-krotne

wzmocnienie (21 dB). Tak więc, gdy potencjometr VR1 jest ustawiony na maksimum, sygnał wyjściowy z IC1a jest bezpośrednio podawany do wzmacniacza IC1b, w wyniku czego wzmocnienie wynosi 11 razy.

Przy pośrednich ustawieniach VR1 ogólne wzmocnienie od wejścia do wyjścia IC1b jest mniejsze.

Kondensator sprzęgający 10 μF zasila siatkę (G2) triody V1b Nutube sygnałem z wyjścia IC1b. Siatka lampy jest polaryzowana napięciem stałym poprzez rezystor 33 kΩ podłączony do obrotowego styku potencjometru VR2. VR2 pozwala ustawić punkt pracy, a tym samym zniekształcenia wytwarzane przez triodę V1b.

Zakres napięcia na styku obrotowym potencjometru VR2 jest ograniczony do 1,27–3,3 V przez dzielnik rezystancyjny 8,2 kΩ i 6,2 kΩ wraz z rezystancją VR2. Zapewnia to optymalny zakres zmian poziomu zniekształceń. Wartości rezystorów dobrano tak, aby środkowe położenie dla VR2 zapewniało najniższe zniekształcenia wprowadzane przez triodę V1b.

Wzmocniony sygnał pojawia się na anodzie V1b (A2). Obciążeniem anody jest rezystor 330 kΩ podłączony do napięcia zasilania Vaa poprzez rezystor odsprężający 150 Ω. Kondensator 100 μF bocznikuje zasilanie, aby usunąć jego tętnienia.

Sygnał anodowy o wysokiej impedancji jest ponownie sprzężony zmiennoprądowo z kolejnym buforem na wzmacniaczu operacyjnym (IC2a) poprzez kondensator 100 nF, oraz przesunięty napięciowo o potencjał połowy napięcia zasilania za pomocą rezystora 1 MΩ. Rezystor ten obciąża anodę triody V1b, a więc zmniejsza amplitudę sygnału o około 25%. Jest to nieuniknione w tak wysoko impedancyjnym układzie.

Sygnał wyjściowy z IC2a trafia do inwertera IC2b, o wzmocnieniu 0 dB, który odwraca sygnał w fazie, aby skompensować inwersję fazy sygnału przez triodę V1b. Trafia on również do siatki triody V1a poprzez potencjometr nastawny VR4. Pozwala on na tłumienie sygnału (jeśli jest to pożądane) przed podaniem go do triody. Polaryzacja siatki triody V1a jest regulowana potencjometrem VR3 w zakresie 1,96–3,48 V. Napięcia te są wyższe niż w przypadku triody V1b z powodów wyjaśnionych poniżej.

Sygnał wyjściowy z anody (A1) triody V1a jest buforowany przez IC3b, podobnie jak IC2a buforuje wyjście V1b. Sygnały z obu wzmacniaczy operacyjnych IC2b i IC3b zasilają odpowiednio potencjometry regulacji poziomu VR5 i VR6. Obrotowe styki tych potencjometrów podłącza się do obu stron przełącznika PODBICIE, S3a. S3a wybiera zatem pomiędzy wyjściami pierwszego i drugiego stopnia zniekształceń.

Zauważ, że w drugim stopniu, trioda V1a odwraca sygnał w taki sam sposób, jak robi to op-amp IC2b. Tak więc oba sygnały podawane na S3a mają tę samą fazę. Sygnał wybrany przez przełącznik PODBICIE jest podawany do bufora IC4b, zapewniając, że ani VR5 ani VR6 nie są nadmiernie obciążone. Bufor ten zapewnia również niską impedancję dla następnego stopnia regulacji barwy tonu.

Zawiera on prosty filtr górnozaporowy o częstotliwości obciążenia regulowanej potencjometrem VR7. Regulacja barwy dźwięku zapewnia obniżenie wysokich częstotliwości o 20 dB na dekadę (6 dB/oktawę). Obcięcie (punkt –3dB) rozpoczyna się przy około 23 kHz, gdy VR7 jest całkowicie obrócony w lewo, więc w tym położeniu regulacja barwy dźwięku w zasadzie nie działa.

Częstotliwość obciążenia spada do około 2 kHz, gdy VR7 jest ustawiony w pełni w prawo. Rezystancja VR7 i rezystora szeregowego 1 kΩ ustawia stałą czasową filtra RC. Punkt –3 dB można obliczyć jako $1/(2\pi RC)$, gdzie C wynosi 6,8 nF, a R zmienia się w zakresie 1–11 kΩ.

Układ IC4a buforuje wyjście regulacji barwy dźwięku typu RC. Sygnał z IC4a jest następnie sprzężony przez kondensator 10 μF, w celu usunięcia napięcia przesunięcia DC, i podany do przełącznika S2b, następnie przez przełącznik RLY1 do złącza wyjściowego CON2. Sygnał wyjściowy przechodzi przez rezystor separujący 100 Ω, aby

Gdy S2 ustawiony jest w pozycji LINIA, sygnał wejściowy na CON1 omija obwody zniekształceń/przesterowania, a wejście IC1a jest połączone z masą. Zapobiega to powstawaniu szumów przełączania w przypadku włączenia układu, utrzymując kondensator 100 nF na wejściu IC1a w stanie naładowanym.

Prąd zarzenia

Istnieją dwa sposoby zasilania włókien katod. Można doprowadzić prąd do F1 i F3 przez osobne rezystory, przy czym F2 jest wtedy połączony z masą. W tym przypadku przez każde włókno przepływa

Wskaźniki LED1–LED3 są zasilane z linii 5 V poprzez rezystory 510 Ω . LED1 jest wskaźnikiem zasilania i jest podłączona do szyny 5 V. Diody LED: LINIA (LED2) i PODBICIE (LED3) są zasilane tylko wtedy, gdy włączone są odpowiednie przełączniki.



12 Elektronika dla Wszystkich 1/2023

Zasilanie

Układ jest zasilany, gdy mikroprzełącznik S1 jest zwierany przez włożenie wtyczki Jack do CON2. Wtyczka naciska na styk masy w CON2, a to powoduje podniesienie przycisku mikroprzełącznika zasilania układu. Jest to nieco niekonwencjonalna metoda przełączania zasilania, ale działa niezawodnie.

Zdecydowaliśmy się zrobić to w ten sposób, zamiast użyć gniazda Jack montowanego na płycie drukowanej z izolowanym przełącznikiem wewnętrznym lub okablowanego gniazda montowanego na panelu, głównie dlatego, że te typy gniazd nie są powszechnie dostępne, podczas gdy typ, którego używamy, jest.

Gdy nie jest włożona wtyczka DC, gniazdo DC (CON3) łączy ujemny biegun baterii z masą, więc obwód będzie zasilany z baterii, gdy styk S1 jest zamknięty. Gdy włożona jest wtyczka zasilająca, ujemny koniec baterii jest odłączony, a urządzenie pracuje z zasilania DC doprowadzanego do CON3. W obu przypadkach dioda Schottky'ego D1 zapobiega uszkodzeniu w przypadku nieprawidłowej polaryzacji baterii lub wtyczki zasilania DC.

REG1 jest liniowym stabilizatorem 5 V o niskim spadku napięcia i niskim prądzie spoczynkowym. Jego głównym zadaniem jest utrzymanie stabilizowanego napięcia siatek dla triod Nutube oraz stabilizowanego napięcia dla żarzenia katod. Dostarcza on również

zasilanie do przekaźnika 5 V RLY1. Kondensator 100 μ F bocznkuje wejście zasilające REG1, jego napięcie wyjściowe jest filtrowane identycznym kondensatorem.

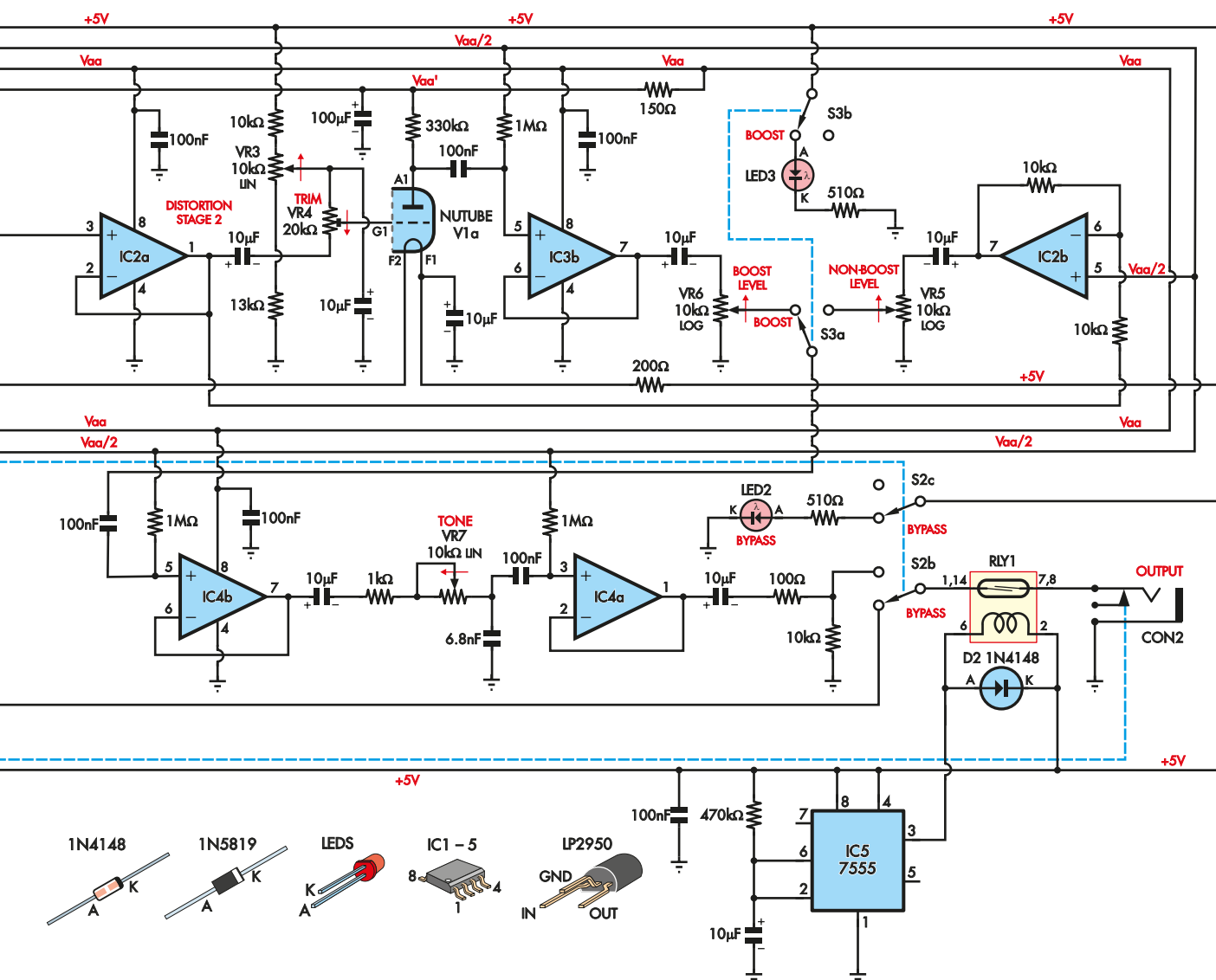
Opóźnienie przekaźnika

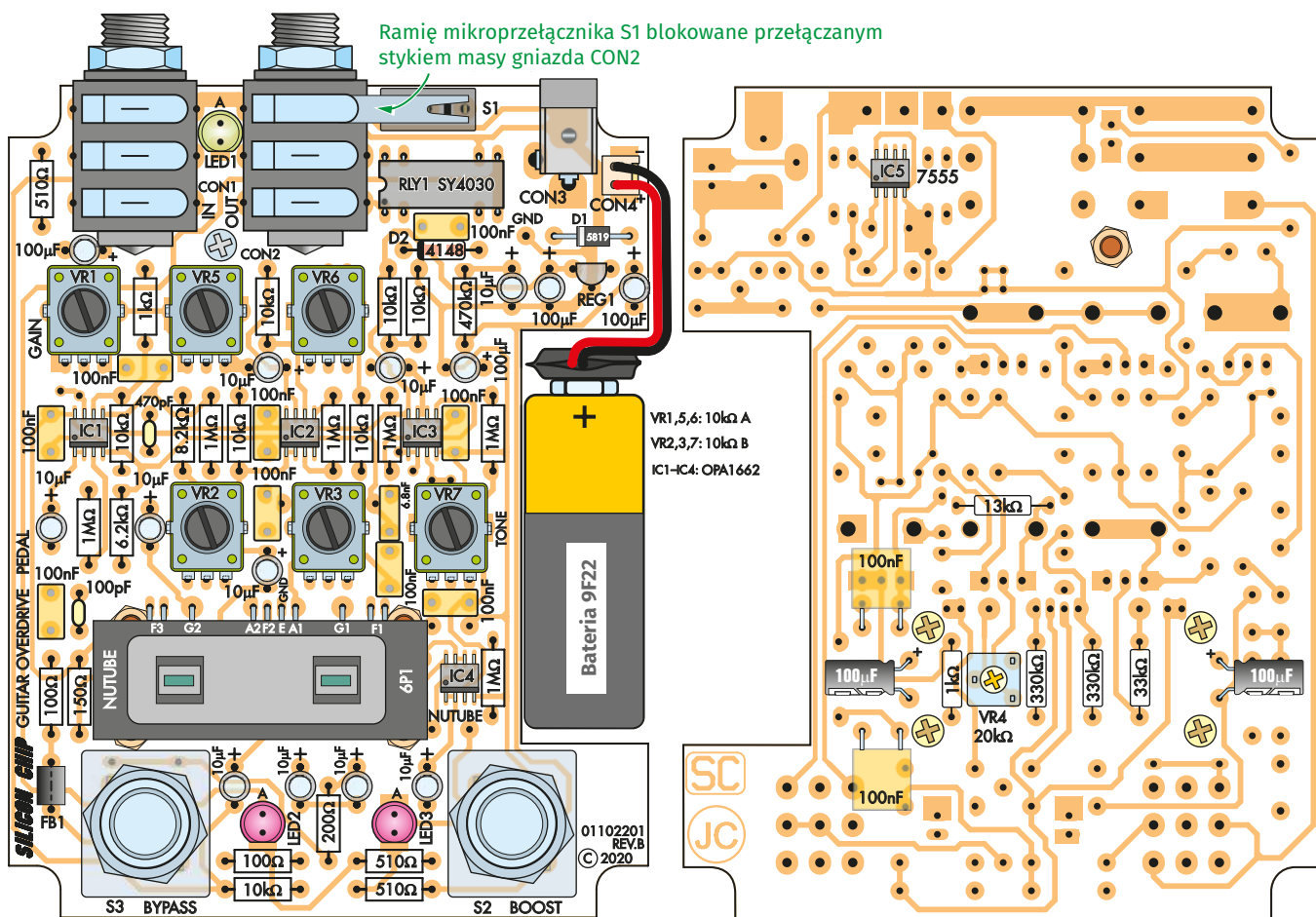
Jak wspomniano, RLY1 włącza się z opóźnieniem przy pierwszym włączeniu zasilania. To opóźnienie zapewnia IC5, wersja CMOS timera 555. Przy pierwszym włączeniu zasilania kondensator 10 μ F na wejściu wyzwalającym (końcówka 1) i progowym (końcówka 6) jest rozładowany. Na wyjściu 3 jest napięcie 5 V, które zasilą końcówkę 6 przekaźnika. Na cewce przekaźnika nie występuje napięcie, więc jest on wyłączony.

Gdy kondensator 10 μ F naładowuje się do 67% napięcia zasilania 5 V (czyli do 3,33 V), osiągnięte zostaje napięcie progowe i wyjście (końcówka 3) przechodzi w stan niski, załączając cewkę przekaźnika.

RLY1 jest przekaźnikiem kontaktronowym o niewielkim prądzie cewki równym 10 mA, więc IC5 możeysterować cewkę bezpośrednio. Dioda D2 tłumi impuls samoindukcji cewki, gdy RLY1 jest wyłączany.

Należy pamiętać, że RLY1 zapobiega doprowadzeniu sygnału LINIA do wyjścia, gdy moduł jest wyłączony. Ponieważ jednak zasilanie włącza się automatycznie po włożeniu wtyczki do gniazda wyjściowego CON2, a nie da się uzyskać sygnału z urządzenia bez podłączenia do CON2, nie stanowi to żadnego problemu.

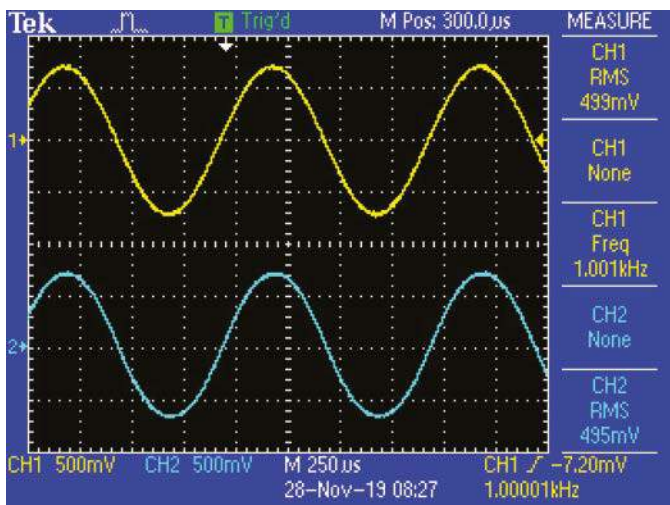




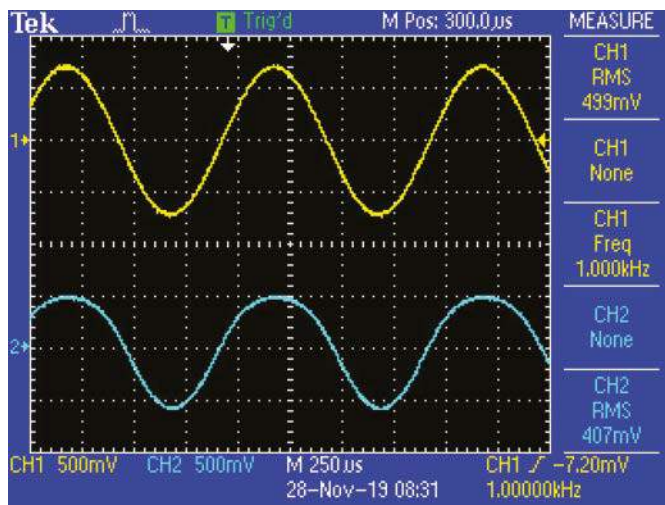
Góra płytki drukowanej

Spód płytki drukowanej

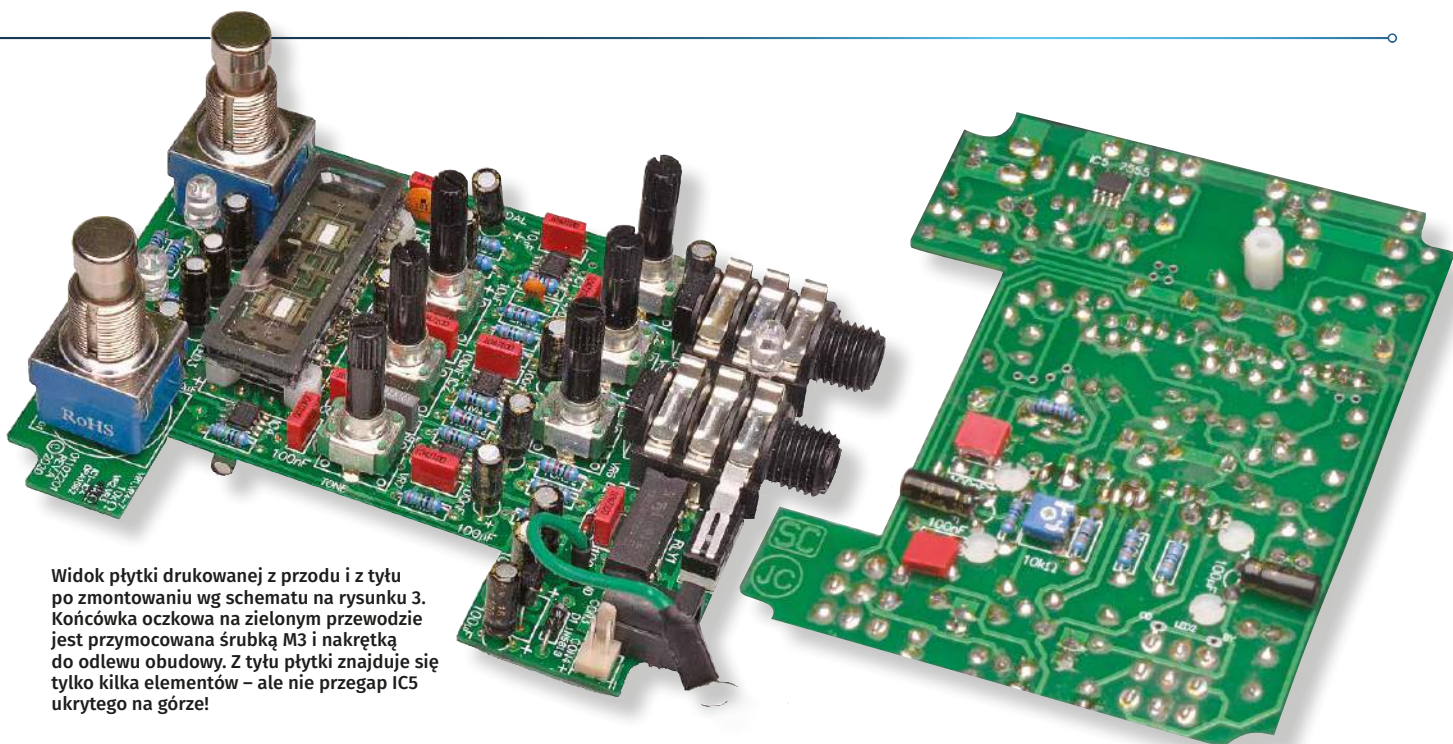
Rysunek 3. Te schematy montażowe podspółotów na PCB pokazują miejsce lutowania części po obu stronach płytki. Zwróć uwagę jak dzwignia mikroprzełącznika S1 jest podpięta pod styk masy gniazda CON2 (zobacz też zdjęcia). Chociaż potencjometry VR1-VR3 i VR5-VR7 wyglądają identycznie i wszystkie są potencjometrami 10 kΩ, to niektóre z nich są liniowe, a niektóre logarytmiczne, zgodnie z opisem znajdującym się obok płytki. Pamiętaj, aby odpowiednio zorientować układy scalone, diody, diody LED, kondensatory elektrolityczne i RLY1 jak pokazano powyżej. Pomiędzy CON1 i CON2 widoczne mocowanie poliamidowej tulejki dystansowej 9 mm oraz po prawej 4 tły poliamidowych śrubek mocujących poliamidowe kołki dystansowe pod lampę Nutube. Patrz również fotografie na stronie 15.



Oscylogram 1: sygnał wejściowy jest pokazany na górze, a sygnał wyjściowy na dole. Tutaj pierwszy stopień regulacji zniekształceń jest ustawiony na minimalne zniekształcenia (pozycja środkowa), przy czym regulacja wzmocnienia jest ustawiona tak, że nie ma przesterowania. Dlatego przebieg wyjściowy jest podobny do wejściowego, praktycznie nie widać deformacji sygnału. Zwróć uwagę, że fazy obu przebiegów są identyczne.



Oscylogram 2: przy użyciu tych samych ustawień co w przypadku Oscylogramu 1, z tym że potencjometr regulacji pierwszego stopnia zniekształceń jest obrócony całkowicie w prawo. Dolny przebieg pokazuje obcięcie i spłaszczenie (flat-topping) dodatniej połówki sygnału sinusoidalnego, dające znaczne zniekształcenia.



Widok płytki drukowanej z przodu i z tyłu po zmontowaniu wg schematu na rysunku 3. Końcówka oczkowa na zielonym przewodzie jest przymocowana śrubką M3 i nakrętką do odlewu obudowy. Z tyłu płytki znajduje się tylko kilka elementów – ale nie przegap IC5 ukrytego na górze!

Budowa

Moduł Guitar Overdrive & Distortion zmontowany jest na dwustronnej płytce drukowanej o kodzie 01102201 i wymiarach 86×112 mm. Całość umieszczona jest w odlanej obudowie aluminiowej o wymiarach 119×94×34 mm. **Rysunek 3** pokazuje schemat montażowy płytki drukowanej.

Zacznij od zamontowania elementów SMD na górnej stronie płytki, czyli IC1–IC4, a następnie IC5 na spodniej stronie. Nie są one trudne do przylutowania przy użyciu lutownicy z precyzyjnym grotem.

Przyda się dobre widzenie z bliska; być może będziesz musiał użyć lupy montażowej z oświetleniem lub okularów, aby widzieć wystarczająco dobrze szczegóły lutowania.

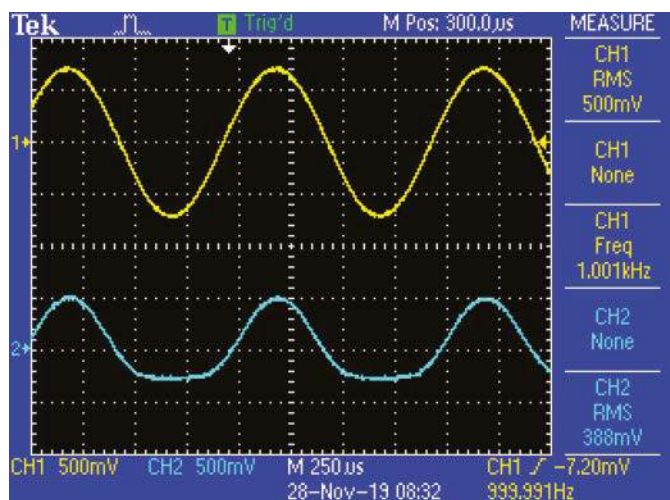
Upewnij się, że te elementy są prawidłowo zorientowane przed wlutowaniem na miejsce. Sprawdź również, czy IC5 to timer 555. Dla każdego podzespołu przylutuj najpierw jedno wyprowadzenie i sprawdź ustawienie obudowy.

W razie potrzeby skoryguj położenie elementu poprzez ponowne podgrzanie spiny lutowniczej przed przylutowaniem pozostałych końcówek. Jeśli któreś z końcówek są zwarte przez lut, użyj plecienki lutowniczej, aby usunąć jego nadmiar.

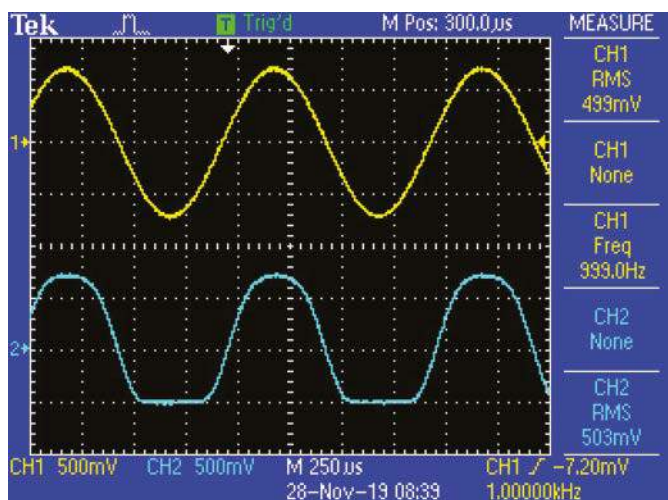
Zauważ, że sąsiadujące wyprowadzenia 1 i 2 układów IC1, IC2 i IC4 oraz wyprowadzenia 6 i 7 układów IC3 i IC4 łączą się ze sobą na płytce drukowanej, więc zwarcie lutem między tymi końcówkami jest dopuszczalne. Nie wygląda to jednak elegancko ani profesjonalnie.

Kontynuuj budowę montując rezystory na górnej stronie PCB (użyj DMM do sprawdzenia wartości), a następnie wlutuj koralik ferrytowy (FB1). Przeprowadź obcięty drut rezystora przez koralik i wygnij go tak, aby pasował do otworów na płytce PCB. Przed przylutowaniem wyprowadzeń koralika dociśnij je tak, aby przylegał do płytki PCB.

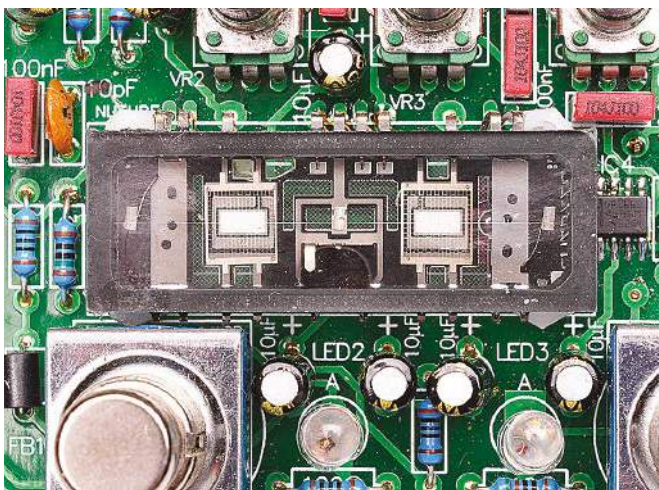
Teraz można zainstalować rezystory, które znajdują się na spodzie płytki drukowanej. Przylutuj je od górnej strony płytki i przytnij wyprowadzenia jak najbliżej płytki. Mogą zostać teraz zamontowane



Oscylogram 3: potencjometr regulacji zniekształceń pierwszego stopnia jest teraz skręcony w pełni w lewo. Górny przebieg to sygnał wyjściowy, podczas gdy dolny przebieg pokazuje spłaszczony szczyt sygnału sinusoidalnego przy jego ujemnych poziomach.



Oscylogram 4: wzmacnienie jest zwiększane aż do uzyskania przesterowania z potencjometrem regulacji pierwszego stopnia zniekształceń ustawionym na minimum (pośrodku). Poziom wyjściowy jest ustawiony tak, aby zmniejszyć poziom sygnału na wyjściu, kompensując wysokie wzmacnienie na wejściu. Należy zwrócić uwagę na to, jak płaska jest ujemna część sygnału; zwiększenie poziomu sygnału zwiększyłoby jego amplitudę i zaczęło spłaszczać również część dodatnią.



diody D1 i D2 – zauważ, że są to różne typy i zwróć uwagę na ich prawidłową orientację.

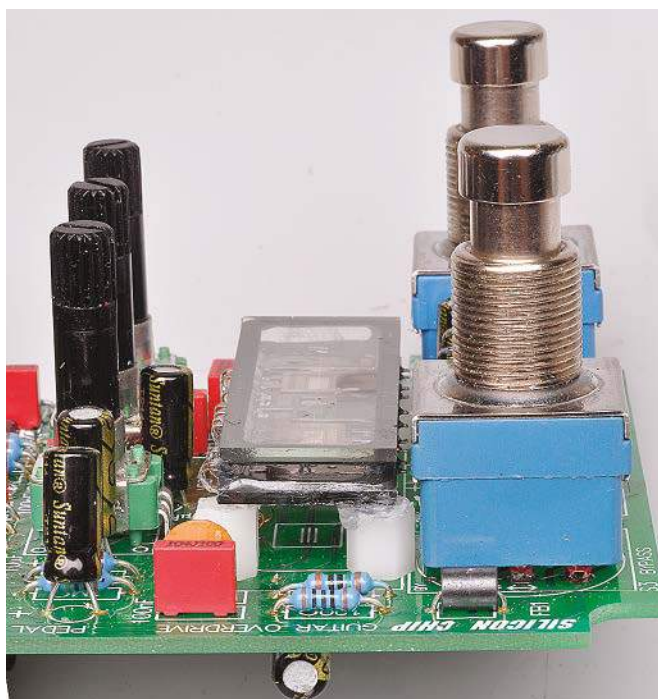
Teraz zamontuj kondensatory MKT i dwa kondensatory ceramiczne, a następnie kondensatory elektrolityczne, które są spolaryzowane. Ich dłuższe wyprowadzenia trafiają w otwory oznaczone „+” na płytce. Dwa kondensatory 100 nF i dwa 100 μ F, które montuje się na spodzie płytki, muszą leżeć podłużnie na boku.

Następnie zamontuj potencjometr nastawny VR4 na spodzie, lutując jego końcówki na górnej stronie. VR4 może być oznaczony raczej jako 203 a nie 20 k Ω .

Następnie zamontuj potencjometry VR1–VR3 i VR5–VR7, zwracając uwagę, że VR1, VR5 i VR6 są typu logarytmicznego (oznaczone jako A), a VR2, VR3 i VR7 są typu liniowego (oznaczone jako B). Mogą być oznaczone jako 103 zamiast 10 k Ω .

Kolejnym krokiem jest dopasowanie REG1 poprzez lekkie rozchylenie jego wyprowadzeń, aby pasowały do otworów na płytce drukowanej. Należy również zamontować kołek PC w punkcie pomiarowym GND. Teraz można zamontować złącze wtykowe przewodów baterii CON4 (wtyk 403-2 dwuszpilkowy prosty do druku), następnie przekaźnik kontaktowy RLY1, dwa gniazda sygnałowe Jack i gniazdo zasilania DC.

Mikroprzełącznik NC S1 jest zamontowany tak, że jego dźwignia jest wsunięta pod przełączany styk masy tulei gniazda CON2. Zapewniliśmy



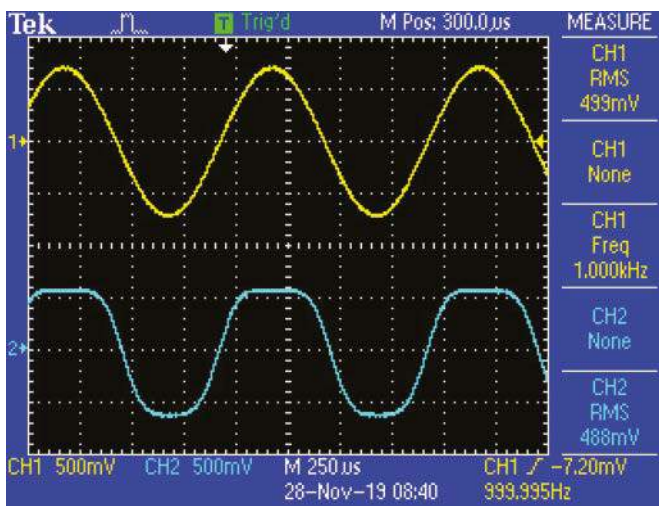
Lampa 6P1 jest montowana na czterech poliamidowych kołkach dystansowych o długości 6,3 mm, jak widać na zdjęciach. Pozwala to zminimalizować mikrofonowanie, które w przeciwnym razie mogłoby być problemem.

podłużne otwory, aby można było włożyć i wsunąć przełącznik w takiej pozycji, aby jego dźwignia trafiła pod styk.

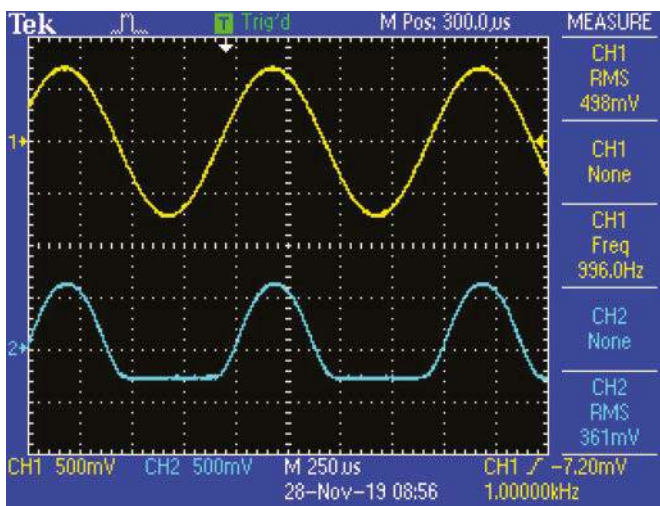
Sprawdź, czy przełącznik ma otwarty obwód, pomiędzy dwoma zewnętrznymi wyprowadzeniami gdy nie ma włożonego wtyku Jack. Po włożeniu wtyku Jack dwa zewnętrzne wyprowadzenia muszą być zwarte.

Być może trzeba będzie nieco wygiąć dźwignię S1, aby przełącznik działał niezawodnie.

Teraz należy zamontować przełączniki nożne S2 i S3. Upewnij się, że są one ustawione idealnie pionowo przed przyłutowaniem ich końcówek. Diody LED zostaną zamontowane później, po zamontowaniu płytki w obudowie.



Oscylogram 5: ustawienia takie same jak w przypadku Oscylogramu 4, ale z potencjometrem regulacji zniekształceń pierwszego stopnia ustawionym do końca w prawo. Powoduje to powstanie przesterowania w formie zbliżonej do przełączania sygnału; wchodząca fala sinusoidalna jest zamieniana na coś w rodzaju zaokrąglonej fali prostokątnej.



Oscylogram 6: ustawienia takie same jak w przypadku Oscylogramu 4 i Oscylogramu 5, ale z potencjometrem regulacji zniekształceń pierwszego stopnia skróconym do końca w lewo. Przebieg wyjściowy jest teraz bardzo mocno wypłaszczony, praktycznie odcięty, przy ujemnych poziomach, ale prawie nie zniekształcony przy dodatnich poziomach sygnału dźwiękowego.



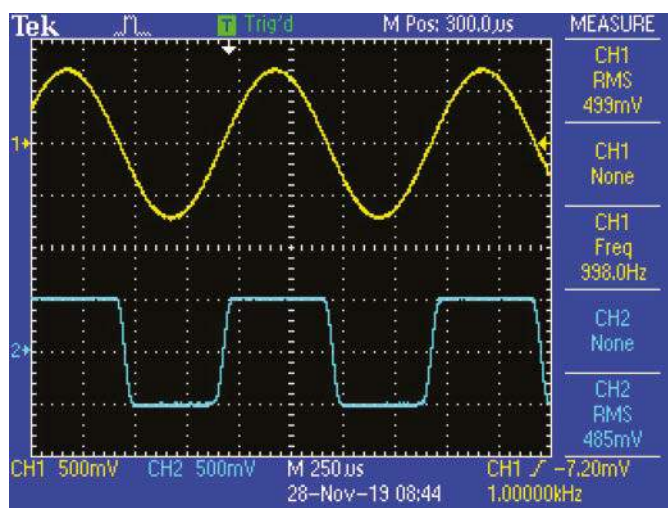
W odlewanej obudowie wywiercono 14 otworów i dwa wycięcia. Należy zwrócić uwagę, że otwory te znajdują się w dolnej i krótszej bocznej części obudowy. (Patrz zwymiarowany schemat wiercenia na stronie 20).

Okablowanie

Lampa Nutube jest montowana z obudową równoległą do PCB. Jej wyprowadzenia są przylutowane do pól lutowniczych na górze PCB za pomocą krótkich odcinków emaliowanego drutu miedzianego. Drut ten pomaga zapobiegać mikrofonowaniu Nutube, zapewniając elastyczne połączenie.

Zagnij wyprowadzenia Nutube pod korpus i przylutuj 20 mm odcinki emaliowanego drutu miedzianego o średnicy 0,25 mm do każdego wyprowadzenia Nutube. Stopiony lut trzymany nad końcem drutu wypali emalię tak, że drut będzie można przylutować.

Na końcach Nutube znajdują się dwa wyprowadzenia dla F1 i dwa wyprowadzenia dla F3. Są one połączone razem, więc do podłączenia każdej pary do PCB potrzebny jest tylko jeden przewód.



Oscylogram 7: jak poprzednio, z włączonym podbiciem, sygnał jest teraz tak przesterowany i obciążony, że przebieg wyjściowy jest prawie prostokątny.



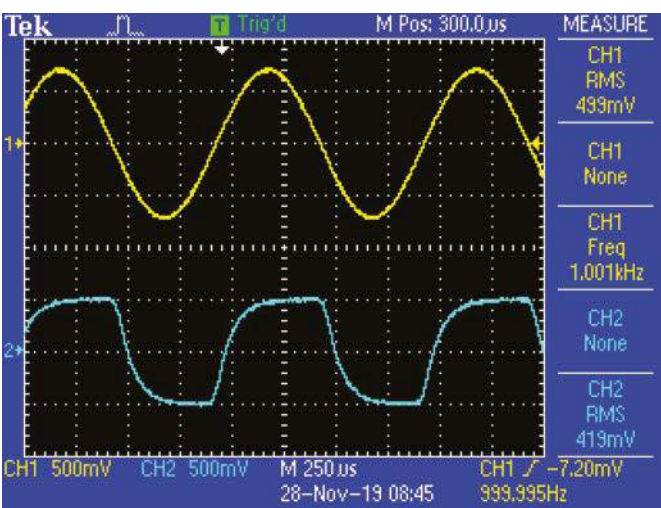
Płytką drukowaną jest montowana w obudowie do góry nogami, jak widać tutaj, przy czym pokrywa obudowy staje się podstawą. Wszystkie elementy sterujące wychodzą przez ściankę, która była dnem obudowy – teraz jest to panel przedni! Pięć przezroczystych opravek w panelu pokazuje stan diod LED i anod podwójnej triody 6P1.

Przymocuj cztery poliamidowe tulejki dystansowe 6,3 mm do PCB pod miejscem mocowania lampy Nutube, używając śrub poliamidowych lub poliwęglanowych.

Umieść małą ilość neutralnego samoutwardzalnego kitu silikonowego na górze każdej tulejki, a następnie ułóż na nich lampę Nutube.

Pomiędzy każdą tulejką a spodnią stroną obudowy Nutube powinna być 1 mm warstwa silikonu. Upewnij się, że lampa jest prawidłowo umieszczona i poczekaj aż silikon się utwardzi (ok. 24 godzin).

Następnym krokiem jest przycięcie przewodów baterii na długość 60 mm, a następnie ich zaciśnięcie lub przylutowanie do dedykowanych styków gniazda.



Oscylogram 8: pokazuje efekt działania regulacji barwy dźwięku przy ustawieniu na maksymalne obcięcie wysokich tonów. Ustawienia są takie same jak na Oscylogramie 7, z wyjątkiem tej zmiany. Zwróć uwagę na różnicę pomiędzy prostokątnym przebiegiem na Oscylogramie 7 a efektem przypominającym tutaj falę morską, wskutek działania regulacji barwy dźwięku.

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

- 1 szt. dwustronna płytka drukowana o kodzie 01102201 i wymiarach 86×112 mm
- 1 szt. etykieta na panel
- 1 szt. obudowa z aluminiowego odlewu ciśnieniowego o wymiarach 119×94×34 mm [Jaycar HB5067]
- 1 szt. podwójna trioda Korg Nutube 6P1 [RS Components 144-9016] (V1)
- 2 szt. gniazda Jack 6,35 mm do wlotowania w PCB [Jaycar PS0195] (CON1, CON2)
- 1 szt. 2-szpilkowy wtyk męski do montażu na PCB raster 2,54 mm typ 403-02 [Jaycar HM3412, Altronics P5492] (CON4)
- 1 szt. gniazdo zasilania DC montowane na PCB [Jaycar PS0520, Altronics P0621A] (CON3)
- 1 szt. obudowa gniazda 402-2 dwustykowa na przewód [Jaycar HM3402, Altronics P5472] plus
- 2 szt. metalowe styki do gniazda 402 (z listwy 10-stykowej, np. 2×P5470A)
- 1 szt. mikroprzełącznik ZMA03A150L30PC lub równoważny [np. Jaycar SM1036] (S1)
- 2 szt. przełączniki nożne 3PDT (potrójny dwupozycyjny) [Jaycar SP0766, Altronics S1155] (S2, S3)
- 1 szt. przełącznik kontaktowy 5V DIL [Jaycar SY4030, Altronics S4100] (RLY1)
- 6 szt. pokręteł o średnicy 11,5 mm i wysokości 6 mm z 18 zębami typu „bezkońcowego” [RS Components 299-4783] (patrz tekst)
- 1 szt. koralik ferrytowy o średnicy zewn. 4 mm, długości 5 mm [Altronics L5250A, Jaycar LF1250] (FB1)
- 5 szt. przezroczystych opravek LED 5 mm [RS Components 171-1931]
- 1 wtyczka mono Jack 6,3 mm lub przewód z wtyczką Jack (do testowania przełączania zasilania)
- 1 bateria 6F22 9 V
- 1 złaczce zatraskowe do baterii 9 V
- 1 kawałek blachy aluminiowej o grubości 1-1,5 mm i wymiarach 9x45 mm
- 1 szt. katek PC (GND)
- 1 szt. końcówka oczkowa na przewód (do uziemienia obudowy)
- 4 szt. przyklejane gumowe nóżki LUB
- 4 szt. śruby nylonowe M4×10 – patrz tekst
- 4 szt. poliamidowe katek dystansowe z gwintem M3 o długości 6,3 mm (do umieszczenia pod Nutube)
- 4 szt. poliamidowe lub poliwęglanowe śruby M3x6 (do podkładek Nutube)
- 1 szt. poliamidowy katek dystansowy z gwintem M3 i długości 9 mm (podpora dla PCB)
- 2 szt. śruby M3×6 (do końcówki oczkowej i katek dystansowych 9 mm)
- 1 szt. nakrętkę M3 i podkładka gwiazdkowa (do końcówki GND)
- 1 odcinek 160 mm długości drutu miedzianego emaliowanego (DNE) o średnicy 0,25 mm
- 1 odcinek 50 mm zielonego (lub żółto-zielonego) przewodu przyłączeniowego (do GND)
- 2 szt. opaski kablowe 100 mm

Półprzewodniki:

- 4 szt. podwójnie wzmacniacze operacyjne OPA1662AID, SMD SOIC-8 [RS Components 825-8424] (IC1-IC4)
- 1 szt. timer CMOS ICM7555CBA, SMD SOIC-8 (IC5)
- 1 szt. dioda Schottky’ego 1A 1N5819 (D1)
- 1 szt. dioda 1N4148 (D2)
- 1 szt. regulator LDO 5 V LP2950CT-5.0 (REG1)
- 3 szt. 5 mm diody LED o wysokiej jasności (zalecana jedna zielona i dwie czerwone)

Kondensatory:

- 6 szt. kondensatorów elektrolitycznych 100 µF/16 V PC MB
- 10 szt. kondensatorów elektrolitycznych 10 µF/16 V PC MB
- 11 szt. kondensatorów poliestrowych 100 nF MKT
- 1 szt. kondensator poliestrowy 6,8 nF MKT
- 1 szt. kondensator ceramiczny (lub KSF) 470 pF
- 1 szt. kondensator ceramiczny (lub KSF) 100 pF

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

- | | | | |
|--------------|---------------|---------------|---------------|
| 6 szt. 1 MΩ | 1 szt. 470 kΩ | 2 szt. 330 kΩ | 1 szt. 33 kΩ |
| 1 szt. 13 kΩ | 7 szt. 10 kΩ | 1 szt. 8,2 kΩ | 1 szt. 6,2 kΩ |
| 1 szt. 1 kΩ | 3 szt. 510 Ω | 1 szt. 200 Ω | 1 szt. 150 Ω |
- 2 szt. 100 Ω
 - 1 szt. 20 kΩ miniaturyowy potencjometr nastawny poziomy [Altronics R2481B, Jaycar RT4362] (VR4)
 - 3 szt. 10 kΩ pionowe potencjometry 9 mm logarytm. (A) [Altronics R1958] (VR1, VR5 i VR6)
 - 3 szt. 10 kΩ pionowe potencjometry 9 mm liniowe (B) [Altronics R1946] (VR2, VR3 i VR7)

Różne:

Lut, plecionka lutownicza Cu, bezbarwny neutralny samoutwardzalny kit silikonowy (np. silikon łożeniowy)



Wykonany przez nas element maskujący zakrywający wycięcia (jak widać obok). Rysunek 4 (poniżej) pokazuje wymiary.

Włóż te styki w odpowiednie ustawieniu do obudowy gniazda (obudowa gniazda 402-02 dwustykowa). Zatrzaśnij styki w obudowie upewniając się, że masz czerwony i czarny przewód w prawidłowej pozycji dla zasilania: „+” do czerwonego i „-” do czarnego.

Niezbędny jest przewód uziemiający do połączenia obudowy z zaciskiem GND na płycie drukowanej. Zapobiega to szumom przenikającym do układu przez obudowę. Przylutuj przewód do końcówki oczkowej na jednym końcu i do metalowego kołka GND płytki drukowanej na drugim.

Można użyć rurki termokurczliwej do osłony końcówką oczkowej i kołka GND na PCB. Po zmontowaniu, końcówka oczkowa jest przymocowana do obudowy za pomocą śruby M3×6, podkładki gwiazdkowej i nakrętki M3.

Włączanie zasilania i testowanie

Jeśli planujesz używać baterii, podłącz ją teraz. Alternatywnie, podłącz zasilanie 9–12 V DC do CON3. Włóż wtyczkę Jack do CON2, aby włączyć zasilanie.

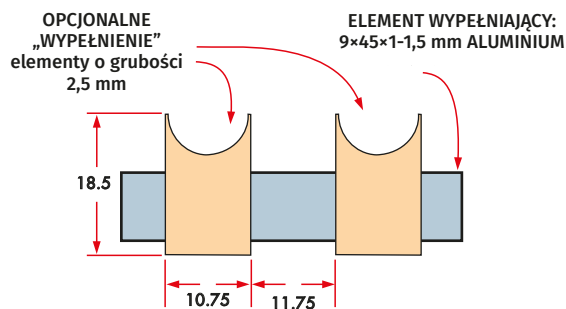
Ustaw swój multimetr na odczyt napięcia stałego, podłącz ujemną sondę do zacisku GND i zmierz napięcia wejściowe i wyjściowe stabilizatora napięcia REG1. Na wejściu powinno być około 0,3 V poniżej napięcia zasilania DC. Napięcie na wyjściu stabilizatora powinno zawierać się w przedziale od 4,95 V do 5,05 V.

Sprawdź również, czy RLY1 włącza się po około pięciu sekundach. Powinieneś usłyszeć ciche kliknięcie.

Ustaw VR2 pośrodku tak, aby lewa anoda Nutube świeciła najjaśniej. Podobnie wyreguluj VR3 tak, aby prawa anoda Nutube świeciła najjaśniej. Zauważ, że kiedy sygnał przechodzi przez przystawkę, blask anod nieco przygasa. Na razie ustaw VR4 w pełni w prawo.

Obudowa

Jako podstawę wykorzystujemy pokrywę obudowy aluminiowego odlewu ciśnieniowego, a główny korpus staje się górą. Schemat wiercenia (rysunek 5) pokazuje miejsca wykonania otworów w podstawie i boku obudowy, może być również użyty jako szablon. Otwory są wymagane dla osi potencjometrów, opravek LED, otworów pod lampę Nutube oraz przycisków nożnych na obszarze panelu głównego.



Rysunek 4. Wytnij kawałek aluminium jak na rysunku, aby częściowo zakryć wycięcia, a następnie przykładaj do niego dwa opcjonalne kawałki tw. sztucznego, aby całkowicie zakryć wycięcia.

Są również wymagane, na krótszym boku obudowy, wycięcia dla dwóch gniazd Jack oraz otwory gniazda zasilania DC i uziemienia obudowy. Wycięcia, a nie otwory, są potrzebne, aby można było manewrować gniazdami Jack przy montażu.

Aby zapobiec przedostawaniu się brudu i innych zanieczyszczeń do wnętrza obudowy, wykonaliśmy zaślepkę o wymiarach 45 mm × 9 mm z arkusza aluminium o grubości 1–1,5 mm. Zakrywa ona wycięcia od wewnątrz, po włożeniu gniazd Jack. Dodaliśmy również kilka kształtek z tworzywa sztucznego, aby wypełnić wycięcia do tego samego poziomu co zewnętrzna część obudowy.

Jest to opcjonalne; kawałki wypełnienia mogą być przyklejone do elementu nośnego, jak pokazano na rysunku i zdjęciu.

Dobrym pomysłem jest dodanie gumowych nóżek, aby obudowa nie przesuwiała się podczas użytkowania. Chociaż można przykleić gumowe nóżki do pokryw, nie byliśmy jednak przekonani, że pozostaną one dalej przyklejone podczas ostrego użytkowania.

Wymieniliśmy więc oryginalne śruby mocujące pokrywę na poliamidowe śruby M4 z łbem walcowym. Łby śrub odstają od powierzchni o kilka milimetrów i dlatego pełnią rolę nóżek. Jednakże, aby to umożliwić, otwory w obudowie dla oryginalnych śrub mocujących musiały zostać rozwiercone do średnicy 3,5 mm, a następnie nagwintowane gwintownikiem M4.

Rysunek 6 pokazuje projekt panelu pokryw, który przygotowaliśmy dla obudowy. Można go skopiować z tego schematu, lub pobrać ze strony internetowej Silicon Chip i wydrukować (plik do pobrania zawiera również szablony do wiercenia).

W celu zabezpieczenia etykiety można ją wydrukować na folii do rzutnika jako lustrzane odbicie, tak aby po naklejeniu toner czy atrament znalazł się pomiędzy obudową a folią.

Użyj folii do rzutnika, która jest odpowiednia dla Twojej drukarki (atramentowej lub laserowej) i przymocuj ją za pomocą bezbarwnego, neutralnego samoutwardzalnego kitu silikonowego. Wyciśnij grudki i pęcherzyki powietrza, zanim silikon się utwardzi. Po utwardzeniu, wytnij otwory w folii za pomocą noża modelarskiego lub introligatorskiego.

Więcej szczegółów na temat tworzenia etykiet znajdziesz na stronie www.siliconchip.com.au/Help/FrontPanels.

Montaż płytki drukowanej

Przymocuj 9 mm tulejkę dystansową z gwintem M3 do tylnej części płytki PCB za pomocą śrubki M3 przechodzącej przez jej górną część. Otwór znajduje się pomiędzy CON1 i CON2. Ten element dystansowy utrzymuje płytkę drukowaną na miejscu, opierając się o pokrywę, gdy obudowa jest złożona.

Jeśli jeszcze tego nie zrobiłeś, przylutuj przewód masy do kołka GND PC na górze PCB i obkurcz krótki odcinek rurki termokurczliwej nad kołkiem. Miejsce montażu oczkowej końcówki przewodu masy znajduje się obok gniazda DC. Przed włożeniem płytki PCB do obudowy należy końcówkę przykręcić



Gniazda wejścia/wyjścia 6,35 mm muszą być wsunięte na miejsce, co wymaga raczej wycięć niż otworów (jak widać na zdjęciu z wierceniem na stronie 20). Ze złomu aluminiowego (patrz obok) uformowaliśmy element wypełniający tej samej wielkości co wycięcia, utrzymywany w miejscu przez same gniazda, ich podkładki/nakrętki i pokrywę – płytę dolną.

śrubką M3, używając podkładki gwiazdkowej i nakrętki.

Przytnij końcówkę oczkową tak, aby przewód znajdował się jak najbliżej obudowy i nie dotykał żadnych elementów na płytce PCB.

Włóż oprawki LED od zewnątrz obudowy. Wizjery Nutube również wymagają oprawek, aby zapobiec przedostawaniu się brudu i kurzu. Oprawki można unieruchomić za pomocą małych opasek kablowych, zaciskając je od wnętrza obudowy, a następnie przykleić kitem silikonowym.

Przed włożeniem płytki do obudowy włóż diody LED do otworów w płytce. Dłuższe przewody anodowe muszą wejść w otwory oznaczone „A” na PCB. Nałóż poliamidowe podkładki przycisków nożnych na trzpień każdego z nich, a następnie zamontuj płytkę w obudowie.

Wciśnij diody LED w ich oprawki, aby je unieruchomić, a następnie przylutuj przewody LED od tyłu PCB.

Komora baterii ma postać prostokątnego wycięcia w płytce PCB. Baterię można zabezpieczyć przed przemieszczaniem się, umieszczając wokół niej trochę piankowego opakowania lampy Nutube.

Włóż piankę pomiędzy koniec baterii a krawędź płytki drukowanej. Jeśli nie używasz baterii, odłącz gniazdo 402 baterii od CON3 i wyjmij ją, aby zapobiec zwarceniu styków z płytką.

Gałki

Ponieważ ośki potencjometrów wystają tylko nieco powyżej 9 mm ponad górę modułu, nie możesz użyć standardowych gałek z kołnierzem. Kołnierze mają za zadanie zakryć nakrętkę zabezpieczającą potencjometr, ale w tym przypadku nie ma żadnych nakrętek, co powoduje, że wewnętrzne rowki są za krótkie do zamocowania gałek na oškach.

Można to obejść na dwa sposoby: albo użyć gałek bez kołnierza, albo odciąć kołnierze. Gałki wymienione na liście części nie mają kołnierzy.

Jeśli z jakiegoś powodu nie możesz ich dostać, możesz kupić gałki Jaycar z serii HK7730-7734 (polecamy niebieskie Cat HK7733) i odciąć dolny kołnierz za pomocą pilki do metalu.

Na koniec zamocuj pokrywę na miejscu za pomocą oryginalnych śrub lub wspomnianych wcześniej poliamidowych śrub M4. Przyklej gumowe nóżki do podstawy, jeśli nie używasz śrub nylonowych jako nóżek.

Zdejmowanie pokręteł

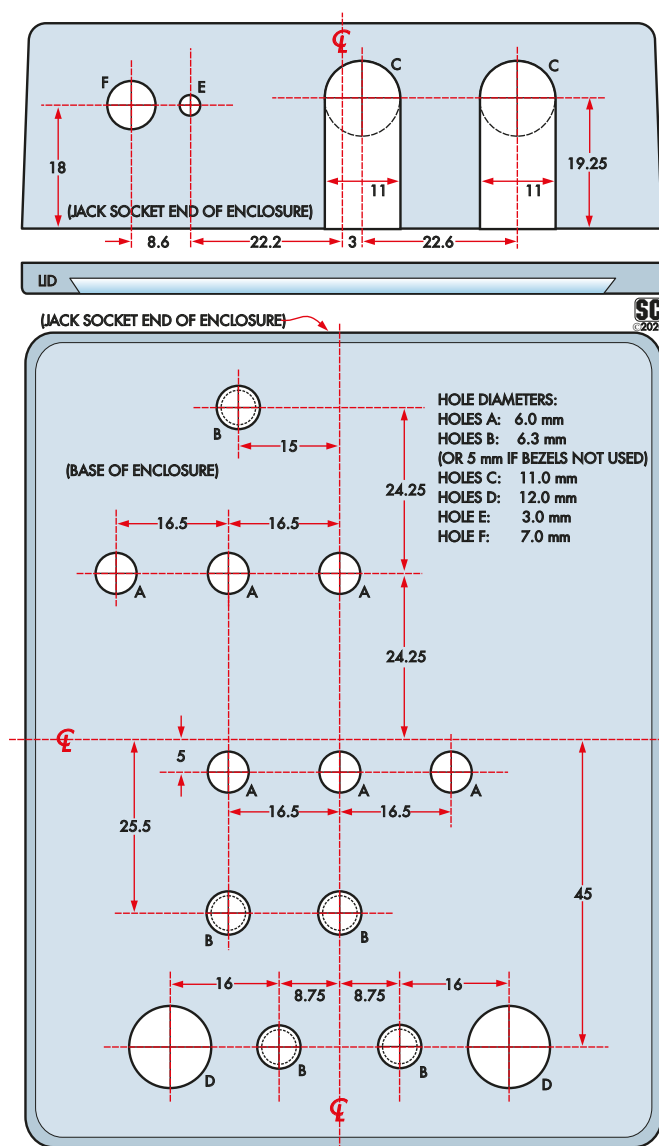
Gałki mogą być trudne do zdjęcia przez pociągnięcie; może być konieczne ich odebranie za pomocą dźwigni. Włóż arkusz cienkiego tworzywa sztucznego pomiędzy dźwignię (np. płaski śrubokręt) a obudowę, aby zapobiec uszkodzeniu panelu.

Używanie

Jest to w zasadzie tylko kwestia pokręcania gałkami aż do uzyskania pożądanego brzmienia. Jedynym regulatorem, który nie jest dostępny z zewnątrz, jest potencjometr nastawny VR4, więc warto zastanowić się co chcemy z nim zrobić, przed zamknięciem obudowy. Należy jednak pamiętać, że moduł jest tak



Zamiast przyklejać nóżki do pokryw obudowy (która staje się podstawą!) użyliśmy czterech poliamidowych śrub z łbem walcowym M4, które działają jak całkiem solidne nóżki, a ich łby lekko wystają z powierzchni. Uznaliśmy, że przyklejone nóżki prawdopodobnie nie wytrzymają długo, ale śruby powinny wytrzymać.



Rysunek 5. Wywierć otwory w dnie i krótszym boku obudowy zgodnie z rysunkiem. Dwa z otworów w boku muszą mieć kształt podłużnych wycięć, aby gniazda CON1 i CON2 dały się wsunąć na miejsce. Jedyny otwór wymagany w pokrywie jest opcjonalny, aby uzyskać dostęp do potencjometru nastawnego VR4; użyj płytki drukowanej jak szablonu, aby zlokalizować ten otwór, jeśli zdecydowałeś się go wywiercić.

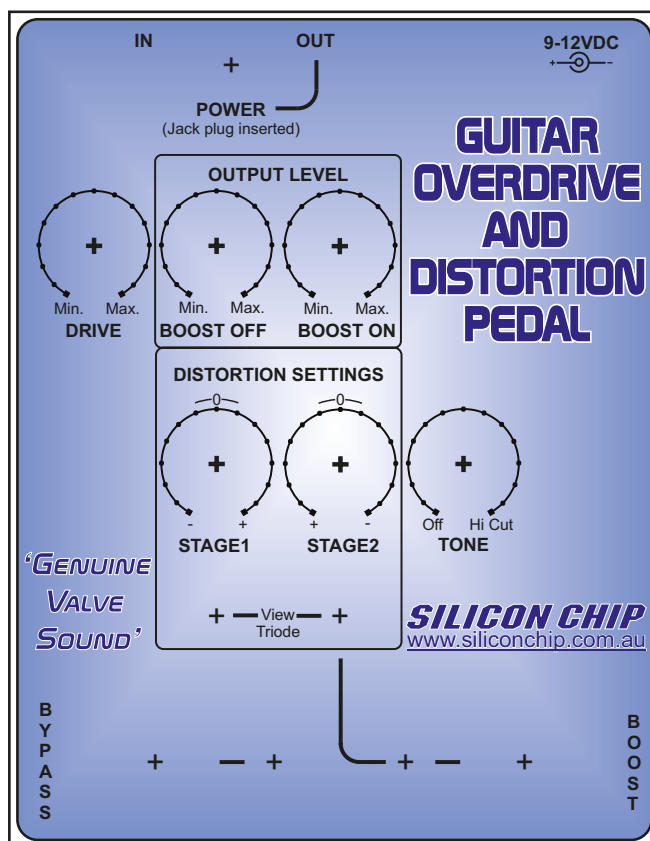
zaprojektowany, że można wywiercić otwór w podstawie, aby z zewnątrz regulować VR4 za pomocą śrubokręta.

My preferujemy pozostawienie VR4 w pełni skróconego w prawo, aby zapewnić znaczne obcinanie sygnału podczas podbicia. Możesz jednak chcieć wyregulować VR4 tak, aby drugi stopień zniekształceń miał podobny poziom wejściowy do pierwszego, i aby łączył się bardziej równomiernie z regulacją stopnia zniekształceń. Jest to kwestia osobistych preferencji.

REKLAMA



Elementy do Twojego projektu możesz kupić na:
sklep.avt.pl



Rysunek 6. Tej samej wielkości etykiety na panel przedni, która pasuje na spód odlewanej obudowy (który oczywiście staje się wierzchem!) Najłatwiej jest wyciąć otwory po przyklejeniu etykiety. Zwróć uwagę na nasze komentarze dotyczące trwałości tego opisu – prawdopodobnie będzie on poddawany dość ostremu traktowaniu!

Wiele wzmacniaczy do instrumentów muzycznych posiada przełącznik pętli uziemienia, który pozwala na uziemienie lub pozostawienie bez podłączenia wspólnego ekranu wyprowadzenia Jack. W przypadku użycia z gitarą, która ma przetworniki piezoelektryczne, powinieneś uzyskać mniejszy szum, gdy jest ekran jest podłączony do uziemienia (masy obudowy).

Fragmenty zrzutów ekranu oscyloskopu Oscylogram1–Oscylogram8 pokazują jak zmienia się przebieg sygnału wyjściowego przy różnych ustawieniach regulacji. Zobacz te zrzuty ekranowe, aby uzyskać więcej informacji. ■

John Clarke

Adaptacja do wydania polskiego – Andrzej Nowicki

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

Ściemniacz nowoczesnego oświetlenia zasilanego z sieci, część 2

W zeszłym miesiącu opisaliśmy, w jaki sposób nasz nowy ściemniacz, obcinający zboczne opadające („trailing edge”), może regulować jasność ściemnianych diod LED i świetlówek kompaktowych, a także tradycyjnych żarówek i zasilanych przez odpowiednie transformatory halogenów, podczas gdy starszy ściemniacz obcinający zboczne narastające („leading edge”) jest nieprzydatny.

Nasze rozwiązanie to elegancka i nowoczesna wyglądająca konstrukcja, która może być sterowana za pomocą jednego lub kilku paneli dotykowych, lub cienkiego pilota na podczerwień. Teraz przejdziemy do budowy i okablowania.

W pierwszym artykule (EdW, grudzień 2022 r.) wyjaśniliśmy, dlaczego do sterowania nowoczesnym oświetleniem LED potrzebny jest ściemniacz obcinający zboczne opadające sinusoidalnego przebiegu napięcia sieciowego.

Starsze ściemniacze wykorzystywały triaki, a to powodowało konieczność włączania zasilania lampy (lamp) w trakcie trwania połówki sinusoidy sieci, z samoczynnym wyłączeniem przy przejściu sinusoidy przez zero.

Takie rozwiązanie nie nadaje się dla urządzeń wykorzystujących zasilacze impulsowe, takich jak diody LED, CFL i lampy halogenowe z transformatorami elektronicznymi. Generuje ono bardzo duże skoki prądu, które szybko zniszczą ich elektronikę. Ściemniacz z obcinaniem zbocza opadającego nie ma tej wady i wiele nowoczesnych lamp jest przystosowanych obecnie do ściemniania przez urządzenia tego typu.

Ponieważ w zeszłym miesiącu omówiliśmy obszernie zasady pracy ściemniaczy, nie mieliśmy miejsca, aby pokazać rzeczywiste przebiegi napięć w trakcie ich pracy. Robimy to teraz, na przykładzie oscylogramów 1–5.

Oscylogram 1 pokazuje starszy ściemniacz pracujący z obciążeniem w postaci zwykłej żarówki. Oscylogram 2 pokazuje ten sam ściemniacz przy próbie regulowania ściemnianych diod LED. Widać, że nie działa on zbyt dobrze!

Z kolei na oscylogramach 3–5 widać przebiegi podawane na ściemnianą lampę LED przez ściemniacz obcinający zboczne opadające. Można zauważyć, że te przebiegi są regularne, zgodne z oczekiwaniami, i jasność lampy zmienia się tak, jak sobie tego życzymy,

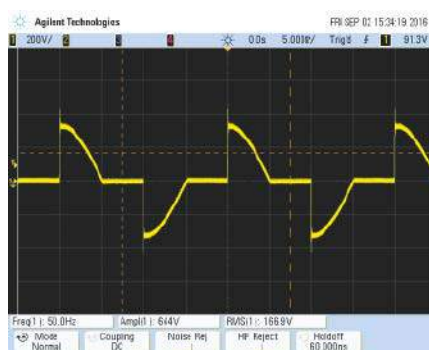
od niskiego poziomu przy oscylogramie 3 do umiarkowanie wysokiej jasności przy oscylogramie 5.

Teraz, gdy wiesz jak działa ten ściemniacz i przeczytałeś o wszystkich jego wspaniałych zaletach, naturalnie chcesz zbudować jeden (lub kilka).

Możesz kupić płytki PCB i trudno dostępne części w sklepie internetowym Silicon Chip (patrz lista części w zeszłym miesiącu), a pozostałe części w sklepie AVT. Następnie możesz zacząć składać płytki, korzystając z poniższych instrukcji.

Czy budowa jest legalna?

Zanim zaczniemy, zwróć uwagę, że choć z pewnością możesz zbudować ten ściemniacz samodzielnie (i włożyliśmy sporo wysiłku, aby uczynić budowę tak prostą, jak to tylko możliwe), w Australii samodzielne podłączenie go jest nielegalne. Będziesz potrzebował licencjonowanego elektryka, aby to zrobić.

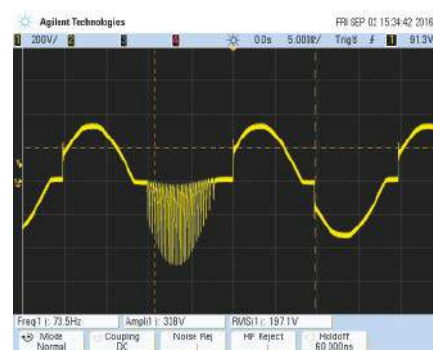


Oscylogram 1. Zwykła żarówka ściemniona do pełnej jasności za pomocą staromodnego ściemniacza z obcinaniem zbocza narastającego. Widać jak napięcie na lampie skacze nagle od bliskiego zera do pełnego napięcia szczytowego sieci (ok. 325 V DC), kiedy triak się włącza. Spowodowałoby to ogromny prąd rozruchowy w typowej lampie LED, która ma kondensatorowe obciążenie zasilacza impulsowego. Prawdopodobnie zniszczyłoby to lampę w krótkim czasie; a jeśli nawet nie, to prawdopodobnie migałaby ona jak stroboskop.



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://bit.ly/3YCNzws>
Materiały dodatkowe są również dostępne na stronie edw.elportal.pl: <https://bit.ly/3WwL6Ce>

Jeśli jednak masz szczęście mieszkać w Nowej Zelandii (albo w Polsce), możesz legalnie wykonać podłączenie ściemniacza do przewodów w puszcze montażowej, dlatego pokazaliśmy odpowiednie schematy elektryczne. Jednak zawsze będziemy nalegać, aby skorzystać z pomocy uprawnionego elektryka, zwłaszcza w przypadku starszych instalacji, mogących mieć przewody



Oscylogram 2. Oto lampka LED zasilana ze ściemniacza jak po lewej, ustawionego prawie na pełną jasność. To jest przykład tego, czego nie należy robić! Widać, że choć w tym przykładzie skoki napięcia nie są tak duże jak na oscylogramie 1, to jednak w całym drugim półcyklu sieci lampka zachowuje się nienormalnie, szybko włącza się i wyłącza, pobierając duże impulsy prądowe. Jej elektronika nie wytrzyma długo w takich warunkach.



Trzy zmontowane płytki PCB użyte w tym projekcie. Po lewej stronie znajduje się płytka główna, która zawiera PIC, mikrokontroler sterujący wszystkim, wraz z transformatorem (który sam nawiniąłeś) i MOSFET-ami mocy, plus, oczywiście, blok zacisków śrubowych (tylna strona tej płytki nie jest pokazana – SZKODA!). W środku znajduje się płytka montażowa, już z przylutowanymi nakrętkami M3, która zapewnia izolację elektryczną, natomiast po prawej stronie znajduje się opcjonalny moduł dodatkowej płytki dotykowej ściemniacza (która jest potrzebna tylko wtedy, gdy stosowane są dodatkowe płytki dotykowe).

o innych oznaczeniach barwnych, lub nawet bez nich. Bezwzględnie skorzystaj z pomocy elektryka, gdy trzeba wykonać lub zmodyfikować ścienne połączenia kabli elektrycznych czy zainstalować nowe puszkę. Wszelkie prace podłączeniowe można wykonywać TYLKO po wyłączeniu napięcia w sieci.

Budowa

Ściemniacz zmontowany jest na płytce drukowanej oznaczonej kodem 10111191, o wymiarach 66×104 mm. Zespół PCB jest skręcany „na kanapkę” z oddzielną płytką nośną o kodzie 10111192 i wymiarach 58,5×104 mm, a cały moduł instalowany jest wewnątrz pustej obudowy Clipsal

Classic, z dopasowaną gładką aluminiową płytką dotykową.

Red. EdW: Rozmiary gotowego ściemniacza są znacznie większe, niż rozmiary typowych polskich puszek ściennych o średnicy 70 mm, stosowanych do montażu wyłączników światła. Nie stanowi to problemu w przypadku ścian gipsowo-kartonowych; w przypadku ścian z cegły znalezienie odpowiednich puszek może być problemem.

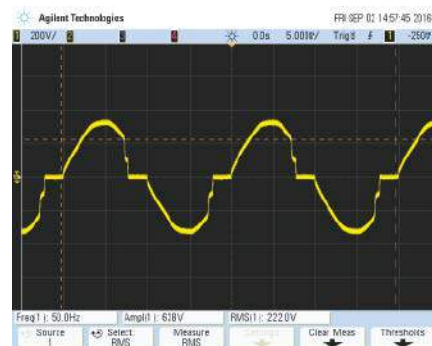
Gotowy ściemniacz może być zamontowany w plastikowej puszcze w ścianie z cegły. W przypadku puszek metalowej niezbędna będzie izolacyjna wkładka dystansowa, utrzymująca odstęp co najmniej 30 mm pomiędzy elementami ściemniacza, a metalowymi ścianami

puszki, w przeciwnym razie obwód elektryczny może zetknąć się z metalową puszką, co stanowiłoby zagrożenie.

Ściemniacz można zamontować bezpośrednio do kołków rozporowych w ścianie z płyt gipsowo-kartonowych przy użyciu standardowego sprzętu montażowego. Alternatywnie może być umieszczony w cienkiej lub standardowej głębokości puszcze z tworzywa sztucznego do montażu powierzchniowego.

Na początek należy zapoznać się ze schematem montażowym elementów na PCB, rysunek 3. Rysunek 4 pokazuje jak wygląda mniej więcej pusta płytka nośna.

Montaż głównej płytki PCB należy rozpocząć od wlutowania MOSFET-ów Q1 i Q2.



Są to elementy montowane powierzchniowo, które są przylutowane do górnej strony płytki. Wystające metalowe płytki obudowy należy przylutować za pomocą gorącej lutownicy.

Dobrze jest rozprowadzić trochę topnika na płytce przed wlutowaniem dwóch mniejszych wyprowadzeń, a następnie zakończyć lutowanie płytki. Upewnij się, że podgrzewasz płytkę na tyle długo, aby lut prawidłowo spłynął zarówno na radiator MOSFET-a jak i na pole lutownicze PCB, tworząc ładne, gładkie wypełnienie.

Następnie można wutować elementy z wyprowadzeniami osiowymi – tj. rezystory, diody Zenera i diody. Tabela kodów barwnych rezystorów pokazuje kolory pasków, ale dobrze jest, tak dla pewności, użyć multimetru cyfrowego do zmierzenia każdej wartości przed lutowaniem (często kolory pasków rezystorów można pomylić, zwłaszcza przy słabym oświetleniu).

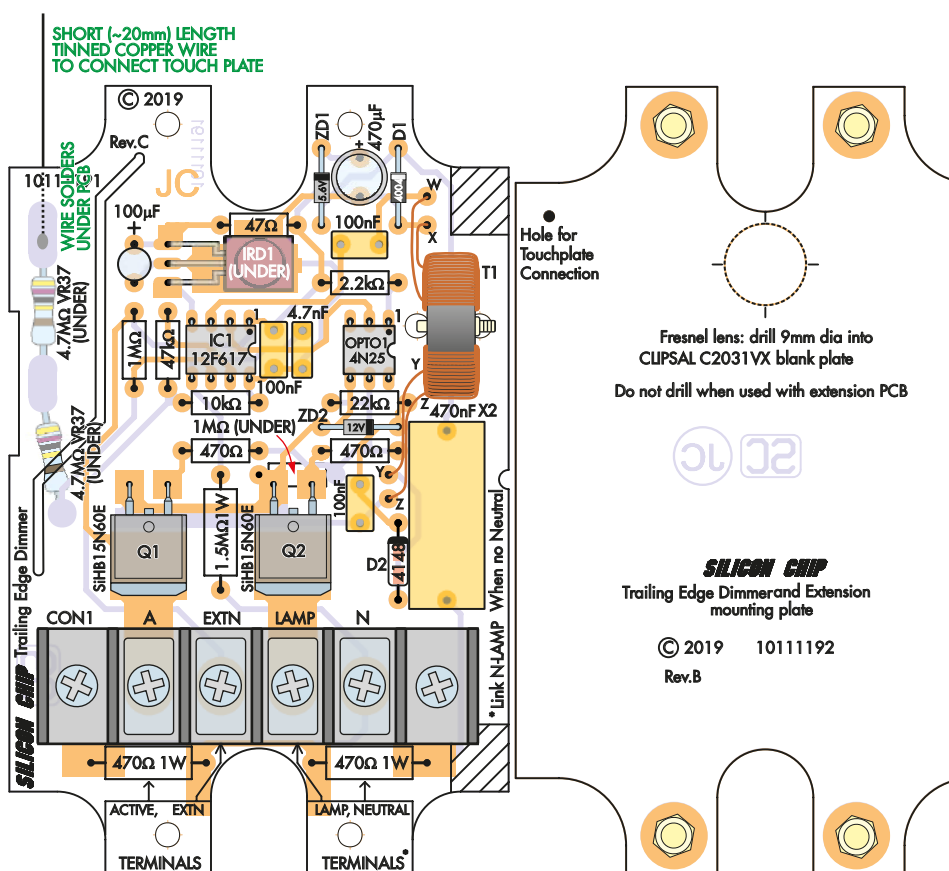
Zwróć uwagę na rezystory 4,7 MΩ szczególnie wyróżnione w tabeli – muszą być one podanego typu, a dla rezystorów 1 W brak jest 5-paskowego kodu, ponieważ są one bardzo rzadkie.

Rezystorów 4,7 MΩ i 1 MΩ na razie nie montuj, zrobisz to później na spodzie płytki drukowanej.

Diody D1 i D2 można łatwo odróżnić, ponieważ D2 jest znacznie mniejsza od D1, ale ZD1 i ZD2 mogą wyglądać podobnie, więc należy uważać, aby ich nie pomylić i zamontować diody Zenera 5,6 V i 12 V w miejscach pokazanych na rysunku 3.

W następnej kolejności zamontuj mikrokontroler i transoptor. Najlepiej byłoby, gdyby IC1 był zamontowany w podstawie, aby w razie potrzeby ułatwić ponowne programowanie, natomiast OPTO1 powinien być przylutowany bezpośrednio do płytki. Następnie, zakładając, że chcesz mieć zdalne sterowanie na podczerwień, spiłuj górne ostre rogi plastikowej obudowy odbiornika podczerwieni, tak aby pasował do wnętrza soczewki Fresnela. Odbiornik podczerwieni jest zamontowany płasko na spodzie płytki drukowanej, a soczewka znajduje się wzdłuż środkowej pionowej linii płytki. Na schemacie płytki drukowanej pokazano prawidłową pozycję montażyową.

Zegnij wyprowadzenia odbiornika pod kątem prostym i przeprowadź je przez pola lutownicze PCB, a następnie przylutuj je od góry. Jeśli nie chcesz korzystać z opcji zdalnego sterowania na podczerwień, powinieneś



Rysunek 3. Schemat montażowy PCB głównej płytki ściemniacza, który pomoże Ci podczas budowy. Odbiornik podczerwieni IRD1 i trzy rezystory (jeden 1 MΩ i dwa 4,7 MΩ) są zamontowane na spodzie płytki (nie pokazane osobno). Rezystory te są przylutowane na powierzchni płytki w stylu SMD, pomimo tego, że są elementami z wyprowadzeniami osiowymi do montażu przewlekanego.

MOSFET-y Q1 i Q2, typu SMD, są przylutowane do górnej części płytki. Zwróć uwagę na krótki odcinek pocynowanego drutu miedzianego przylutowanego do spodu płytki – wygina się on o 90° (tj. pionowo do płytki), aby przejść przez otwór w drugiej płycie, nośnej, a następnie przez otwór wywiercony w obudowie Clipsal, aby zapewnić kontakt z płytką dotykową.

Rysunek 4. Druga płytką, nośna (pokazana po prawej), nie ma żadnych elementów, ale ma cztery podwójne nakretki przylutowane do górnej części płytki, do których można przykręcić główną (lub dodatkową) płytkę PCB.

zamiast tego zamontować rezystor 1 k Ω pomiędzy dwoma skrajnymi polami montażowymi dla IRD1.

Pamiętaj, aby przed przylutowaniem obu układów poprawnie zorientować je tak, aby wcięcie lub kropka na końcówce 1 znajdowały się tak, jak pokazano na schemacie montażowym.

Teraz można wltować kondensatory, zaczynając od mniejszych MKT, następnie większy kondensator MKP-X2 i na końcu kondensatory elektrolityczne.

Tylko kondensatory elektrolityczne (100 μF i 470 μF) są spolaryzowane; ich dłuższe wyprowadzenia wchodzą w otwory oznaczone symbolem „+” na rysunku 3 i na opisie płytki.

Następnie zamontuj dużą, czterodrożną listwę zaciskową CON1. Przymocuj ją do płytki drukowanej za pomocą dwóch śrubek M3 o długości około 20 mm przeprowadzonych przez skrajne otwory montażowe w listwie,

w razie potrzeby wywierć w PCB 2 otwory o średnicy 3,2 mm w miejscach skrajnych pól; i dwóch nakrętek sześciokątnych M3 (nie były wymienione na liście części w zeszłym miesiącu). Gdy listwa jest już pewnie przymocowana do płytki, przyłutuj cztery środkowe zaciski, używając dużej ilości lutu, aby zapewnić niezawodne połączenia. Teraz możesz usunąć niepotrzebne już mocowanie M3 – patrz fotografie trzech płytek na początku tego artykułu. Widać wyraźnie, szczególnie po prawej, otwory w PCB o średnicy >3 mm.

Uzwojenie transformatora T1

T1 składa się z toroidalnego rdzenia ferrytowego i uzwojeń wykonanych z emaliowanego drutu miedzianego (DNE) o średnicy 0,25 mm. Uzwojenie pierwotne ma 12 zwojów, natomiast wtórne 48 zwojów, jak pokazano na **rysunku 5**.

Uzwojenia: pierwotne i wtórne są nawinięte po przeciwnych stronach toroidu, w celu

Czy ten ściemniacz może być używany ze zwykłą lampą itp.

Już nas o to pytano (!)... a co jeśli masz lampę, która jest zwyczajnie podłączona do gniazdka w ścianie (np. standardowa lampa stojąca, lampka na biurko, itp. – taka, która nie jest na stałe włączona w okablowanie domu)? Czy ten ściemniacz nadaje się do tego typu lamp?

Piękno zaprojektowanego ściemniacza polega na tym, że pasuje on do tak wielu typów lamp (żarówek, ściemnialnych LED, ściemnialnych CFL i tak dalej), zatem w większości przypadków będzie idealny.

Oczywiście trzeba się upewnić, że użyta obudowa lampy jest w 100% izolowana i, jeśli jest metalowa, uziemiona. A przede wszystkim – jest wystarczająco duża, aby ściemniacz się w niej zmieścił.

Sugerujemy kierowanie się opisem „Testowanie przed instalacją” i schematem z rysunku 9d, ponieważ zapewnia on regulację od zera do 100% normalnej jasności. Fragment „Do reszty instalacji” oznacza wtedy standardową wtyczkę do gniazdka z bolcem ochronnym. Jeśli w tym momencie nie rozumiesz, jak zainstalować ściemniacz wewnątrz lampy, to znaczy, że powinieneś powierzyć tę pracę fachowcowi.

Alternatywnie, wg tego samego rysunku 9d, choć nie będzie to tak eleganckie, można zbudować przedłużacz sieciowy zakończony obudową z tworzywa sztucznego, wystarczająco dużą, aby pomieścić zarówno ściemniacz, jak i gniazdko elektryczne z bolcem ochronnym do podłączenia ściemnianej lampy. W takim rozwiązaniu złącze opisane jako „Zaciski lampy” jest wtedy zastąpione przez gniazdko z bolcem ochronnym.

Trzeba zauważyć, że podłączenie za pomocą wtyczki do gniazdka z bolcem ochronnym możliwe jest tylko w jednym położeniu, co powinno zapewnić połączenie przewodów: fazowego i neutralnego z identycznymi przewodami w gniazdku ściennym. Niestety, olbrzymim problemem mogą być bliźniacze gniazdko podwójne, z ich budowy wynika, że każde z dwóch gniazdek ma inaczej rozmieszczone przewody, w stosunku do drugiego. Tak naprawdę, to w przypadku zarówno lampy z wbudowanym ściemniaczem, jak i przedłużacza ze ściemniaczem, najbezpieczniej jest używać ich tylko z jednym, specjalnie oznaczonym, i uprzednio dokładnie sprawdzonym, gniazdkiem sieciowym z bolcem ochronnym.

Wykaz elementów, kupuj w sklep.avt.pl

(W-wa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl);
Elementy, których brakowało na liście w zeszłym miesiącu:

2 śruby M3×20 plus 2 nakrętki M3 (zarówno dla głównej jak i dodatkowej płytki PCB)

wzajemnej izolacji. Red. EdW: radzimy zacząć od uzwojenia wtórnego, zabezpieczając jego początek na zewnątrz toroidu odrobiną kleju momentalnego. Drut nawijaj tak, aby na wewnętrznej powierzchni toroidu każdy zwój przylegał ściśle do sąsiedniego. Również koniec uzwojenia zalecamy zabezpieczyć klejem. Następnie, na środku wolnej części toroidu, nawiń podobnie uzwojenie pierwotne.

Skręć dwa końce uzwojenia pierwotnego razem kilka razy i zrób to samo z końcami drutu uzwojenia wtórnego (dla przejrzystości nie jest to pokazane na rysunku 5). Odetnij nadmiar długości drutu, upewniając się, że pozostało go wystarczająco dużo, aby osiągnąć dedykowanych pól PCB, a następnie użyj papieru ściernego lub noża, aby usunąć emalię z końców każdego drutu, w celu ich pocynowania. Upewnij się, że lutowie prawidłowo pokryło drut.

Aby zapewnić, że oba uzwojenia pozostaną oddzielone, przełóż dwie opaski kablowe przez rdzeń ferrytowy i zaciśnij je mocno po obu stronach uzwojenia pierwotnego, a następnie odetnij ich nadmiar.

Następnie przytnij kawałek rurki termokurczliwej o średnicy 16 mm i długości większej niż szerokość rdzenia ferrytowego, nasuń ją na rdzeń tak, aby końce uzwojeń przechodziły przez otwory rurki po przeciwnych stronach, i obkurcz ją, aby się nie ruszała.

Gdy już to zrobisz, wytnij lub wybij kilka otworów na dole, aby umożliwić przesunięcie opaski kablowej. Możesz to zrobić za pomocą śrubokręta, ale uważaj, aby nie uszkodzić rdzenia lub któregoś z uzwojeń.

Montaż T1

Przełóż opaskę zaciskową przez otwory wykonane w rurce termokurczliwej, a następnie przeciągnij ją przez 3 mm otwory w płytce PCB, które mają za zadanie utrzymać transformator na miejscu. Kwadratowy koniec opaski powinien znaleźć się na górze płytki, po jednej stronie rdzenia toroidalnego. Płytki drukowanej nie da się prawidłowo zamontować, jeśli zaciskowy koniec opaski znajdzie się na spodzie płytki.

Przylutuj oba końce uzwojenia pierwotnego do pól oznaczonych jako W i X; orientacja drutów nie ma znaczenia. Podobnie przylutuj końcówki uzwojenia wtórnego do pól oznaczonych jako Y i Z.

Teraz możesz zamontować trzy rezystory, które znajdują się na spodzie płytki. Rezystor 1 MΩ ma otwór na jedno z wyprowadzeń i pole lutownicze SMD na drugie, ale oba wyprowadzenia są przylutowane po spodniej stronie płytki. Upewnij się, że przycięłeś na krótko wyprowadzenie przechodzące na górę PCB zaraz po przylutowaniu.

Połączenie z płytką dotykową

Lutowanie dwóch rezystorów Vishay 4,7 MΩ serii VR37 jest krytyczne. Wygnij i przytnij ich wyprowadzenia tak, aby leżały płasko na okrągłych polach lutowniczych SMD na spodzie PCB, a następnie przylutuj je na miejscu, jak przy montażu powierzchniowym.

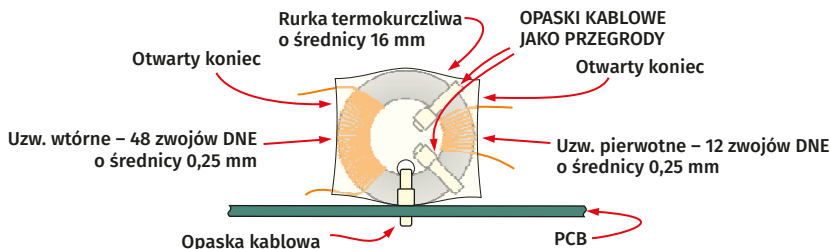
Upewnij się, że są one umieszczone we właściwej pozycji i nie zastępuj tych elementów żadnymi innymi. Rezystory te, wybrane specjalnie dla bezpieczeństwa, są koloru jasnoniebieskiego i mają napięcie znamionowe 2,5 kVRMS.

Są one zamontowane w taki sposób, że żadne połączenia na górnej części płytki drukowanej nie są prowadzone po ich przeciwną stronę, a dodatkowo ich początek i koniec oddzielone są szczeliną izolacyjną wyfrezowaną w PCB. Dzięki temu przewody rezystorów są w pełni odizolowane od elementów znajdujących się na górze płytki.

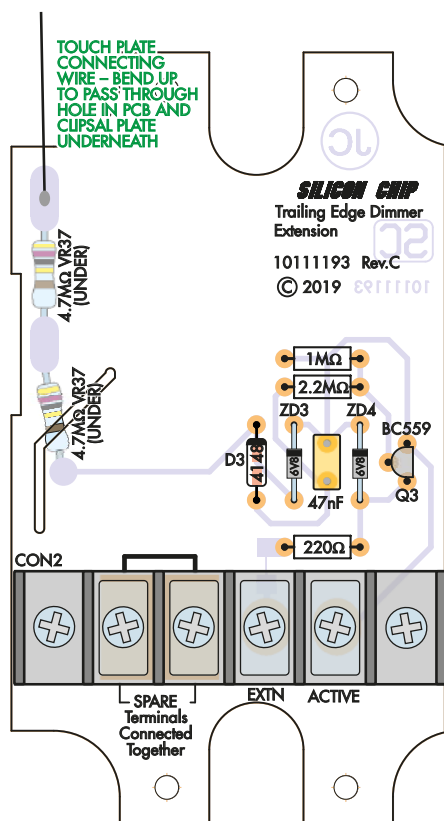
Zapewnia to również wysoki stopień izolacji prądowej pomiędzy płytką dotykową a obwodami napięcia sieci.

Szeregowe rezystory bezpieczeństwa utrzymują kontakt z płytką dotykową poprzez krótki odcinek (powiedzmy około 20 mm lub coś koło tego) pocynowanego drutu miedzianego. Jest on przylutowany do najwyższego pola SMD na spodzie po lewej stronie płytki.

Drut ten jest wygięty pod kątem 90°, aby przejść przez otwór w drugiej płytce (nośnej), a następnie przez odpowiedni maleńki otwór wywiercony w płytce Clipsal.

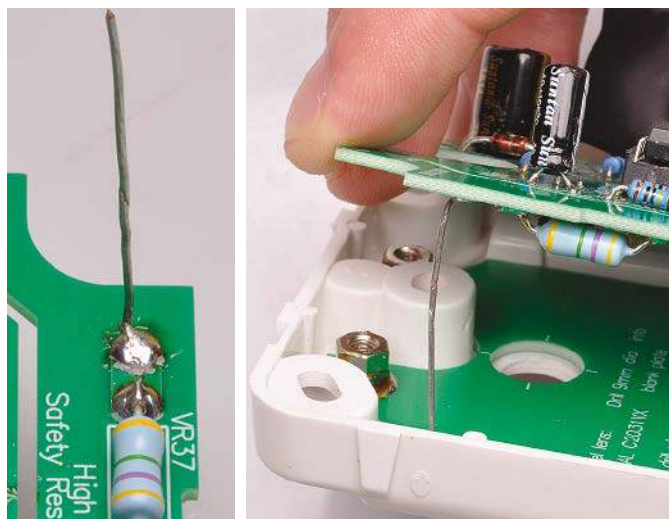


Rysunek 5. Pokazany jest sposób nawinięcia transformatora T1 na toroidalnym rdzeniu ferrytowym emaliowanym drutem miedzianym (DNE) o średnicy 0,25 mm. Po wykonaniu obu uzwojeń, zaciśnij pomiędzy nimi dwie opaski kablowe jako przegrody i odetnij ich końce, następnie nasuń na transformator rurkę termokurczliwą i obkurcz ją. W rurce, przez środek toroidu, zrób otwór śrubokrętem (przebijakiem) i przymocuj ją do PCB jak na rysunku, przed przylutowaniem skręconych wyprowadzeń T1 do płytki.



Rysunek 6. Schemat montażowy dodatkowej płytki PCB ściemniacza, stosowanej, gdy chcesz mieć dwa lub więcej ściemniaczy dotykowych sterujących tym samym światłem lub zestawem świateł. Znajduje się na niej tylko kilka elementów, więc powinna być łatwa i szybka w montażu, o ile będziesz uważał, aby umieścić elementy w miejscach i w połączeniach pokazanych na rysunku. Identyfikacja jak na rysunku 3, krótki odcinek pocynowanego drutu miedzianego zagina się pod kątem 90°, prostopadłe do PCB, aby przeprowadzić go przez drugą płytkę (nośną), a następnie przez obudowę Clipsal, aby umożliwić kontakt z aluminium, zewnętrzną płytką dotykową.

Najłatwiejszym sposobem wywiercenia tego ostatniego jest umieszczenie płytki nośnej w obudowie Clipsal i wywiercenie otworu o średnicy 0,9 mm na wylot, tzn. użycie płytki nośnej jako szablonu.



Podczas późniejszego montażu, przewód jest zaginany tak, aby znajdował się na równi z powierzchnią płytki Clipsal, dzięki czemu, gdy aluminiowa płytka dotykowa jest zatrzaśnięta na miejscu, styka się z przewodem.

(Przewód ten nie jest przylutowany ani w żaden inny sposób przymocowany do aluminiowej płytki dotykowej).

Programowanie IC1

Jeśli zakupiłeś wstępnie zaprogramowany mikrokontroler PIC w sklepie internetowym SILICON Chip, możesz go teraz włożyć do podstawki. Upewnij się, że jego kropka przy końcówce 1 jest zorientowana jak na rysunku 3.

Jeśli masz czysty układ scalony PIC12F617, musisz pobrać firmware (plik 1011119A.hex lub 1011119B.hex) ze strony internetowej SC lub <https://edw.elportal.pl/materiały-dodatkowe>, a następnie załadować go do układu za pomocą programatora uniwersalnego lub programatora szeregowego PICkit 3 albo PICkit 4 (lub podobnego) z płytką interfejsu. Wersja pliku *.hex zależy od wersji stosowanego pilota IR, patrz errata w SC z sierpnia 2019 r., str. 112.

Odpowiednia jest np. płytka z adapterem do programowania PIC/AVR z wydania SC z maja i czerwca 2012 r. (siliconchip.com.au/Series/24).

Jeśli używasz uniwersalnego programatora, użyj dostarczonego z nim oprogramowania. Dla PICkit 3 i PICkit 4 można użyć MPLAB IPE (integrated programming environment), części MPLAB IDE (integrated development environment), które można bezpłatnie pobrać ze strony Microchip i są dostępne dla Windows, MacOS i Linux.

Montaż płytki

Płytkę nośną PCB jest zwymiarowana tak, aby dokładnie pasowała do obudowy Clipsal C2031VX.

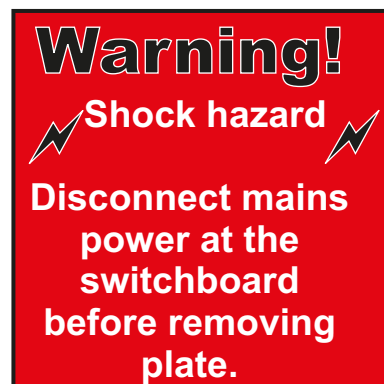
BRĄZOWY i NIEBIESKI... czy.... CZERWONY i CZARNY... a może inny?

Na naszych schematach elektrycznych przewód fazowy sieci jest oznaczony kolorem BRĄZOWYM, a neutralny NIEBIESKIM. Odpowiada to kodom barwnym przewodów w Polsce wg aktualnych przepisów – patrz wklejony rysunek.

Niestety, mogą być problemy w starszych instalacjach, więc w ich przypadku konieczna jest identyfikacja.

Pozwala to na zamontowanie głównej płytki drukowanej ściemniacza. Dopasuj płytkę nośną do obudowy Clipsal, zwracając uwagę na to, że strona z nadrukiem powinna być widoczna po zakończeniu pracy; płytka PCB będzie pasować tylko w jednej orientacji.

Zaznacz środek otworu wymaganego dla soczewki, która ma pasować do obudowy



Rysunek 7. Ta nalepka ostrzegawcza powinna być skopiowana lub wydrukowana i przyklejona na powierzchni czotowej obudowy Clipsal, przed założeniem aluminiowej ozdobnej płytki dotykowej (upewnij się, że nie zakrywa ona drutu, który dotyka tej aluminiowej płytki). Jest to przypomnienie dla każdego, kto demontuje ściemniacz, że wewnątrz znajdują się przewody i obwody pod napięciem. (Ostrzeżenie! Niebezpieczeństwo porażenia prądem. Przed zdjęciem płyty należy odłączyć zasilanie sieciowe w rozdzielni.)



Te trzy zdjęcia pokazują lokalizację i montaż drutu łączącego obwód PCB z zewnętrzną aluminiową płytką dotykową. Przechodzi on przez otwórki w płytce nośnej i w obudowie Clipsal oraz styka się z aluminiową płytką dotykową po jej zatrzaśnięciu w obudowie. Uwaga – nie pokazano soczewki Fresnela!

Testowanie przed instalacją

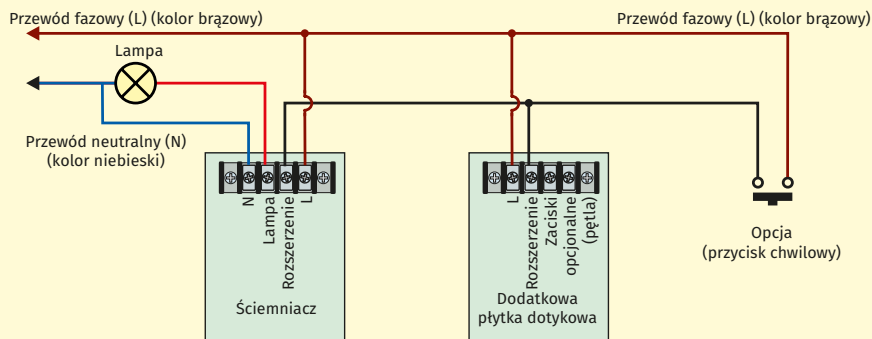
Ponieważ prawdopodobnie będziesz musiał zapłacić elektrykowi, aby przeszedł do Twojego domu i zainstalował ściemniacz(e), będziesz chciał najpierw upewnić się, że działają poprawnie. Najłatwiejszym bezpiecznym sposobem jest użycie montowanej powierzchniowo ramki czy puszek z tworzywa sztucznego, w której umieścisz ściemniacz z płytką dotykową, przykręconą do kawałka materiału izolacyjnego (np. płyty MDF) wystarczająco dużego, aby zastąpić zaciski napięciowe ściemniacza.

Potrzebny będzie również przedłużacz sieciowy przecięty na pół, aby zapewnić zasilanie obwodu (z gniazdka sieciowego) oraz lampa (typu, którego używasz) z wtyczką do podłączenia do gniazdka.

Zdejmij zewnętrzną i wewnętrzną izolację z końcówek przeciętego przedłużacza sieciowego i wywierć dwa otwory w ramce, na tyle duże, aby przeszedł przez nie kabel sieciowy. Prześledź procedurę instalacji w tekście głównym niniejszego artykułu, upewniając się, że przeprowadziłeś kontrole bezpieczeństwa zgodnie z opisem. Za pomocą izolowanego złącza śrubowego lub zatrzaskowego połącz ze sobą żyły przewodów ochronnych (żółto-zielone) obu połówek kabla.

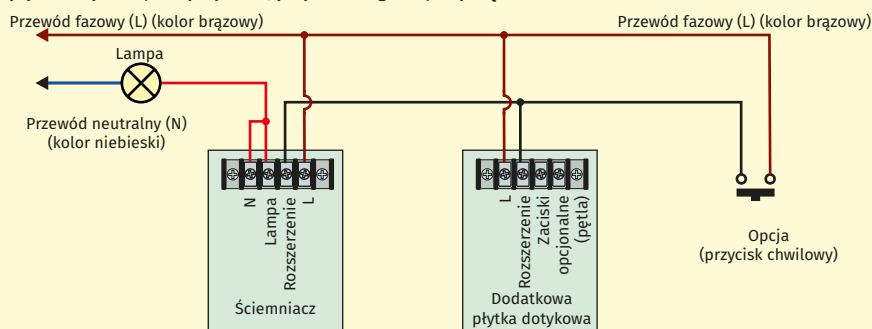
Końcówka przedłużacza sieciowego z gniazdkiem będzie użyta do podłączenia lampy. Większość ściemnialnych lamp LED ma dotaczoną wtyczkę sieciową, więc można ją po prostu włączyć. Jeśli używasz innego typu lampy, będziesz potrzebował odpowiedniej oprawy i bezpiecznego sposobu podłączenia jej do wtyczki sieciowej.

Podsumowując, jeżeli Twoja instalacja końcowa ma dostępny przewód neutralny sieci, możesz podłączyć przewód fazowy (brązowy) przewodu wtykowego do zacisku „L” ściemniacza, przewody neutralne przewodu wtykowego i gniazdkowego (niebieskie) do zacisku „N” (zastosowana listwa zaciskowa z łatwością pomieści dwa przewody w jednym zacisku), a przewód fazowy (brązowy) kabla gniazdkowego do zacisku „LAMP”. Jeśli w ostatecznej instalacji nie będzie dostępny przewód neutralny sieci, połącz ze sobą przewody neutralne (niebieskie)



GDY DOSTĘPNY JEST PRZEWÓD NEUTRALNY

Rysunek 8. Sposób tymczasowego podłączenia ściemniacza do testów (lub alternatywnie do użytku z lampą wtykową). U góry pokazany jest sposób sterowania pojedynczą lampą, gdy mamy do dyspozycji zarówno przewód fazowy (brązowy) jak i neutralny (niebieski), natomiast dolny schemat pokazuje połączenia, gdy nie jest dostępny przewód neutralny. Jeśli nie zamierzasz używać dodatkowej płytki dotykowej lub przycisku, po prostu zignoruj te połączenia.



GDY NIE MA DOSTĘPU DO PRZEWODU NEUTRALNEGO

wtyczki i gniazdka (ponownie użyj izolowanego złącza śrubowego lub zatrzaskowego), przewód fazowy wtyczki (brązowy) podłącz do zacisku „L” ściemniacza, a przewód fazowy (brązowy) gniazdka do zacisków „N” i „LAMP”, stosując krótką zworę z odizolowanego na końcach drutu, jak pokazano poniżej. Upewnij się, że urządzenie nie jest podłączone do prądu podczas powyższego montażu!

Przymocuj ramkę ze ściemniaczem do płyty MDF (itp.) tak, aby żadne z przewodów sieciowych nie były odstłonięte. Następnie można podłączyć lampę do gniazdka (kierunek podłączenia nie ma znaczenia), a wtyczkę do gniazdka (**UWAGA!** Musisz mieć

pewność, że przewód fazowy we wtyczce trafi na przewód fazowy w gniazdku ściemniacz, tak samo przewód neutralny) i odczekać co najmniej dziewięć sekund (aby pominąć etap kalibracji, jak wyjaśniono w tekście). Następnie możesz przetestować, czy działa płytka dotykowa, i (jeśli jest zamontowany), pilot na podczerwień.

Jeśli lampa, której używasz do testowania, jest tą samą, która zostanie użyta w ostatecznej instalacji, możesz również przeprowadzić procedurę kalibracji – patrz opis na końcu artykułu. Łatwiej będzie, jeśli zrobisz to teraz, ponieważ na tym etapie dużo łatwiej jest włączyć i odłączyć ściemniacz od zasilania.

Clipsal i zwróć uwagę, że ten otwór nie jest wiercony, gdy budujesz płytkę dodatkową lub gdy zdecydowałeś się zrezygnować z funkcji pilota na podczerwień. Nanieśliśmy sitodruk krzyża celowniczego, aby pokazać wymaganą pozycję środka otworu na płytce obudowy Clipsal. Wywierć otwór o średnicy 9 mm.

Takiej samej wielkości otwór należy wywiercić w ozdobnej, wierzchniej płytce aluminiowej. Ostrożnie wywierć go na drewnianej podkładce; zaczynając od wiertła o mniejszej średnicy i rozwiercając otwór do 9 mm; uzyskasz lepsze wykończenie otworu niż przy użyciu od razu wiertła 9 mm.

Wywierć również otwór o średnicy 0,9 mm na drut łączący płytkę dotykową z elektroniką. Ten otwór jest wykonywany tylko w plastikowej płytce Clipsal, a nie w płytce aluminiowej!

Zdejmij płytkę nośną i włóż cztery śrubki M3×10 od spodu płytki w każdym narożu; na każdą śrubkę nakręć po dwie nakrętki M3, od strony opisu płytki. Dokręć mocno pierwszą nakrętkę, ale drugą nakrętkę dokręć tylko na styk. Zlutuj ze sobą razem każde dwie nakrętki i przyłutuj dolne nakrętki do PCB. Gdy lutowane połączenia ostygną, usuń śrubki M3. Na pierwszej fotografii w tej części artykułu, w środku, widać przyłutowane nakrętki.

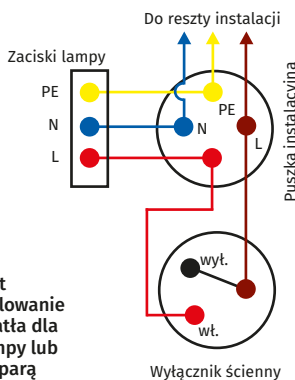
Przyłutuj, w stylu SMD, 15–20 mm odcinek pocynowanego drutu miedzianego od spodu głównej płytki ściemniacza, na końcu pola lutowicznego górnego rezystora zabezpieczającego 4,7 MΩ. Drut powinien znajdować się bezpośrednio naprzeciwko otworu w płytce nośnej (i otworu w obudowie Clipsal) służącego do podłączenia aluminiowej płytki dotykowej.

Montaż końcowy

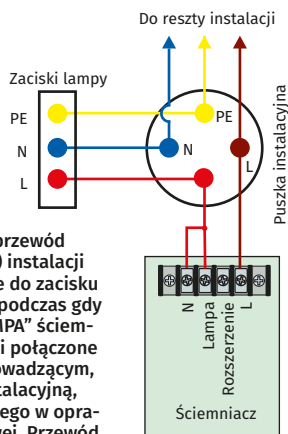
Dopasuj płytkę nośną do obudowy Clipsal i wcisnij ją tak, aby płytka ściśle przylegała do jej wewnętrznej płaszczyzny.

Możesz zabezpieczyć ją silikonem lub uszczelniającym poliuretanowym, aby upewnić

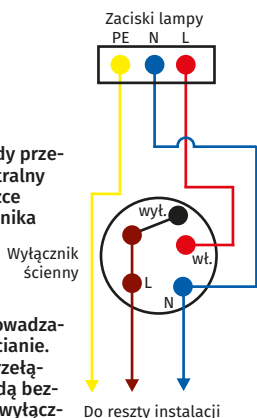
Rysunek 9a. Jest to typowe okablowanie wyłącznika światła dla pojedynczej lampy lub oprawy, tylko z parą przewodów (bez neutralnego), dochodzących z puszki instalacyjnej do wyłącznika ściennego. Należy zwrócić uwagę na to, że przewód ochronny (żółto-zielony) często nie jest używany w wielu starszych domach, ale w każdym przypadku, przewód ten nie odgrywa żadnej roli w konstrukcji ściemniacza.



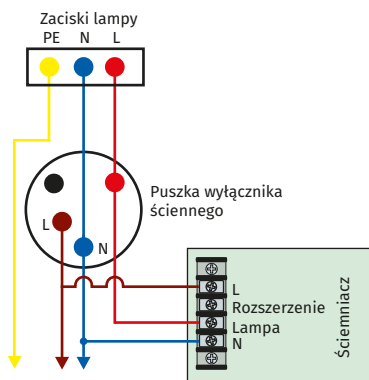
Rysunek 9b. Zastąpienie wyłącznika światła ściemniaczem w typowej instalacji z rysunku 9a jest tak proste jak tutaj pokazano: przewód fazowy (brązowy) instalacji elektrycznej idzie do zacisku „L” ściemniacza, podczas gdy zaciski „N” i „LAMP” ściemniacza są zwarte i połączone z przewodem prowadzącym, przez puszke instalacyjną, do zacisku fazowego w oprawie oświetleniowej. Przewód ten ma zwykle kolor ograniczony tylko fantazją instalatora i dostępnością kabli. Przed zainstalowaniem ściemniacza należy upewnić się, że wyłączono zasilanie na tablicy rozdzielczej!



Rysunek 9c. Niekiedy przewody: fazowy i neutralny są dostępne w puszcze montażowej wyłącznika ściennego, np. gdy puszka pełni również rolę puszki instalacyjnej rozprowadzającej przewody w ścianie. Przewód fazowy (przetaczany) i neutralny idą bezpośrednio z puszki wyłącznika ściennego do gniazda lampy lub oprawy oświetleniowej. Przewód ochronny (jeśli występuje) jest często podłączony bezpośrednio do gniazda lampy/oprawy oświetleniowej.



Rysunek 9d. Oto jak okablować ściemniacz w miejsce wymontowanego wyłącznika ściennego, gdy w puszcze wyłącznika dostępne są zarówno przewód fazowy, jak i neutralny. Takie okablowanie pozwoli na regulację jasności od zera do 100%.



się, że pozostanie ona na swoim miejscu. Aby to zrobić, nałóż kilka kropli uszczelnacza na spód płytki nośnej przed włożeniem jej do pustej obudowy Clipsal. Można też użyć (ostrożnie!) kleju momentalnego.

Włóż soczewkę Fresnela w otwór płytki nośnej, a następnie ustaw płytkę PCB ściemniacza nad płytką nośną i przełóż drut kontaktowy płytki dotykowej przez otwór w płycie nośnej i przez otwór w obudowie Clipsal. Pokazane jest to na trzech fotografiach na dole strony 25. Następnie przykręć płytkę PCB ściemniacza do płytki nośnej za pomocą śrubek M3×10.

Upewnij się, że przewód łączący płytkę dotykową wystaje teraz na zewnątrz pustej plastikowej obudowy Clipsal. Zegnij ten przewód o 90°, aby przylegał do zewnętrznej powierzchni płytki. Będzie się on stykał, po zamontowaniu, z ozdobną płytką aluminiową, zapewniając funkcję sterowania dotykowego. Wspomniane trzy fotografie dokładnie to ilustrują.

Rysunek 7 to etykieta ostrzegawcza, którą należy wydrukować i przykleić na zewnątrz plastikowej obudowy Clipsal. Dzięki temu w przypadku

usunięcia aluminiowej płytki dotykowej będzie widoczne ostrzeżenie o konieczności wyłączenia zasilania na tablicy bezpieczników.

Można również pobrać tę etykietę bezpłatnie ze strony internetowej SC jako plik PDF, wymieniony w sklepie internetowym SC w sekcji „Panele i elementy obudowy”.

Budowa płytki rozszerzeń

Potrzebujesz tej płytki tylko wtedy, gdy chcesz mieć więcej niż jedną płytkę dotykową do sterowania tym samym zestawem światła.

Układ płytki dodatkowej jest zbudowany na PCB o kodzie 10111193 i wymiarach 58,5×104 mm. Potrzebna jest również płytki nośna (o kodzie 10111192), aby przymocować płytkę dodatkową wewnątrz pustej obudowy Clipsal, również stosowanej z ozdobną aluminiową płytką czołową.

Podobnie jak w przypadku samej głównej płytki ściemniacza, w przypadku montażu płytki dodatkowej do metalowej puszki ściennego (np. w ścianie z cegły), niezbędna będzie izolacyjna wkładka dystansowa, utrzymująca odstęp co najmniej 30 mm pomiędzy elementami płytki, a metalowymi ściankami puszki.

Dodatkową płytkę sterującą można też zamontować bezpośrednio do kołków rozporowych w ścianie z płyt gipsowo-kartonowych przy użyciu standardowego sprzętu montażowego. Alternatywnie może być umieszczona w cienkiej lub standardowej głębokości puszce z tworzywa sztucznego do montażu powierzchniowego.

Rysunek 6 to schemat montażowy PCB płytki dodatkowej. Rezystory, diody Zenera, diodę i tranzystor należy zamontować wg schematu, w podanej kolejności. Zwróć uwagę, że niektóre rezystory nie będą normalnie dostępne z tolerancją 1%.

Dobłą praktyką jest również użycie multimetru cyfrowego do pomiaru wartości każdego rezystora. Zwróć uwagę, że dwa rezystory 4,7 MΩ na spodzie płytki są montowane później.

Dwie diody Zenera, ZD3 i ZD4, mają taką samą wartość, więc ich nie pomylisz i trzeba tylko uważać na ich polaryzację, jak i diody D3. Orientacja tranzystora Q3 również ma znaczenie, ale będzie ona prawidłowa, jeśli zamontujesz go płaską stroną w kierunku pokazanym

REKLAMA

Certyfikat Underwriters Laboratories

94V-0 E480148 TYPE 1

Zakład produkcyjny:

05-660 Warka
ul. M. Ropielewskiej 17
tel. 22 781 63 95
22 761 95 80
fax. 22 781 63 95 w 23
www.elmax.waw.pl
elmax@elmax.waw.pl

OBWODY DRUKOWANE

Produkcja, Projektowanie, Montaż

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

na rysunku. Prawdopodobnie będziesz musiał lekko wygiąć wyprowadzenia (np. używając małych szczypiec), aby dopasować je do pól lutowniczych na płytce PCB.

Listwę zacisków śrubowych CON2, dwa rezystory bezpieczeństwa VR37 firmy Vishay o wartości 4,7 MΩ oraz drut kontaktowy lutuj identycznie jak w przypadku głównej płytki ściemniacza.

Po raz kolejny przypominamy, że rezystorów bezpieczeństwa nie wolno zastępować innymi elementami. Są to rezystory wysokonapięciowe, których napięcie znamionowe wynosi 2,5 kV RMS. Są one koloru jasnoniebieskiego.

Procedura mocowania płytki dodatkowej do obudowy Clipsal za pośrednictwem płytki nośnej jest taka sama jak opisana dla głównej płytki ściemniacza. Wyjątkiem jest brak wiercenia otworu na soczewkę.

Instalacja

Do tej pory przetestowałeś ściemniacz zgodnie z procedurą opisaną w panelu bocznym i na schematach wg **rysunku 8**.

Wg tych procedur, w połączeniu ze schematami na **rysunkach 9a–d**, przeprowadź instalację ściemniacza w dwóch typowych wariantach – brak dostępnego przewodu neutralnego (bardziej typowy – rysunek 9a i 9b) i druga możliwość, przewód neutralny dostępny (rysunek 9c i 9d).

Na rysunkach 9 nie pokazano dodatkowych płytek sterujących ani przycisków włączających; sposób ich instalacji i okablowania pokazuje rysunek 8.

Ściemniacz i płytki rozszerzające muszą być pewnie zamocowane do ściany przed włączeniem zasilania. Oczywiście podczas instalacji urządzenia zasilanie w skrzynce bezpieczników lub na tablicy wyłączników musi być wyłączone.

Przed zainstalowaniem tych urządzeń przeprowadź konieczną kontrolę bezpieczeństwa. Przelicz multimetr na najwyższy zakres pomiaru rezystancji i sprawdź oporność pomiędzy zaciskiem fazowym złącza CON1 (CON2) i stykiem płytki dotykowej.

Wykonaj to zarówno dla głównej płytki ściemniacza, jak i płytki dodatkowej, jeśli jej używasz. Rezystancja powinna być zbliżona do 9,4 MΩ. W ten sposób sprawdzimy czy płytka dotykowa będzie bezpieczna.

Jeśli masz już starszy ściemniacz, który chcesz wymienić (może chcesz zamienić tradycyjne żarówki na LED-y?), nowy ściemniacz jest łatwy do zainstalowania, ponieważ okablowanie jest takie same. Po prostu łączysz odpowiednie zaciski z przychodzącymi przewodami: fazowym L (brązowy) i przewodem idącym do lampy przez puszkę instalacyjną.

Zwróć uwagę, że to również jest przewód fazowy; po włączeniu lampy jest on pod napięciem sieci! Stosowanie w tym miejscu przewodu oznaczonego kolorem niebieskim jest KARYGODNYM błędem! Niestety, ze względów praktycznych – trudności z zakupem odpowiednio oznaczonych kabli – jest często spotykany.

Okablowanie jest pokazane na dole **rysunku 8**. Ten przykład zawiera jedną płytkę dodatkową oraz oddzielny włącznik/wyłącznik chwilowy (podłączony do zasilania z sieci), ale te dodatkowe urządzenia są opcjonalne i mogą być pominięte, jeśli nie są potrzebne.

Jeśli instalujesz nowy ściemniacz i możesz poprowadzić przewód neutralny sieci zasilającej do miejsca montażu ściemniacza, to jesteś w idealnej sytuacji, ponieważ zapewni to pełny zakres ściemniania od wyłączenia aż do 100% (pełna jasność).

Jak pokazano na rysunku 8, dodatkowy moduł wymaga połączenia z przewodem fazowym L, i przewodu przedłużającego, który łączy się z wejściem EXTN (Rozszerz.) ściemniacza. Gorąco radzimy zlecić elektrykowi wykonanie nowego okablowania, jeśli go jeszcze nie ma.

Wolne zaciski pętli na płytce dodatkowej mogą być użyte do podłączenia końców dodatkowych przewodów, które trzeba doprowadzić.

Opcja przełącznika chwilowego, jak pokazano na rysunku 8, może być użyta w zestawie przełączników ściennych, co ułatwia instalację w miejscach o ograniczonej przestrzeni, jak np. w obramowaniu drzwi.

Kalibracja

Jeśli udało Ci się doprowadzić przewód neutralny „N” do modułu ściemniacza, to nie ma potrzeby przeprowadzania jakiegokolwiek kalibracji. Moduł jest wstępnie ustawiony tak, aby po pełnym włączeniu dostarczał pełne napięcie sieciowe do lampy.

Jeśli nie masz dostępnego przewodu neutralnego, ściemniacz będzie zasilany przez lampę. Ściemniacz będzie musiał być wyregulowany tak, aby zapewnić maksymalną jasność lampy bez migotania.

Regulację należy rozpocząć w ciągu 9 sekund od podłączenia zasilania do ściemniacza. W przeciwnym razie ściemniacz przejdzie w swój normalny tryb pracy.

Zasilanie ściemniacza polega na włączeniu obwodu oświetleniowego w rozdzielni elektrycznej. Tak szybko jak tylko możesz i przed upływem dziesięciu sekund, dotknij i przytrzymaj bez przerwy płytkę dotykową. Poczekaj, aż światło zacznie zwiększać swoją jasność. Zabierz rękę, gdy tylko światła zaczną migotać, co powinno nastąpić w pobliżu pełnej jasności.

Następnie należy nacisnąć i przytrzymać płytkę dotykową, aż światło przyciemni się do jasności nieco poniżej tego ustawienia, a więc do momentu, kiedy nie występuje migotanie. Ponownie zabierz rękę, a następnie szybko naciśnij płytkę dotykową, aby wyłączyć światło (światła). Ta czynność spowoduje ustawienie maksymalnej jasności lampy na ostatnio używanym poziomie jasności. Ściemniacz będzie od tej pory stosował ten poziom jako maksymalne ustawienie jasności, nawet w przypadku utraty zasilania sieciowego.

Ponowną kalibrację maksymalnej jasności można wykonać powtarzając powyższą procedurę, zaczynając od wyłączenia zasilania obwodu światła. Następnie można ustawić maksymalną jasność na wyższym lub niższym poziomie niż poprzednie ustawienie.

Należy pamiętać, że tempo wzrostu jasności lampy podczas tej procedury celowo jest wolne, aby można było ustawić jasność z rozsądną precyzją. Należy również pamiętać, że po rozpoczęciu procedury kalibracji poprzez dotknięcie płytki, mamy do pięciu sekund po odjęciu ręki na ponowne przyłożenie jej do płytki, aby rozpocząć zmniejszanie jasności.

Po odjęciu ręki przy zmniejszeniu jasności, mamy kolejne pięć sekund czasu, aby ponownie dotknąć płytki i wyłączyć lampę.

Jeśli nie dotkniesz płyty przed upływem tych pięciu sekund, kalibracja zostanie przerwana i będzie stosowana poprzednia nastawa maksymalnej jasności. Kalibrację trzeba będzie zacząć od nowa.

Należy pamiętać, że kalibracja powinna być przeprowadzona z lampami, które będą używane ze ściemniaczem. Jeśli używasz różnych lamp LED lub tradycyjnych żarówek, maksymalne ustawienie jasności bez migotania może być inne.

Użytkowanie

Należy pamiętać, że płytka ściemniacza zazwyczaj jest nieco ciepła w dotyku, ze względu na rozpraszanie w MOSFET-ach Q1 i Q2 łącznej mocy około 1 W.

Aby uzyskać niezawodne działanie, pilot musi być skierowany w stronę odbiornika na głównej płytce ściemniacza. Stwierdziliśmy, że nasz prototyp działa dobrze w odległości do 7 m od płytki ściennej, o ile pilot jest prawidłowo skierowany. ■

John Clarke

Adaptacja do wydania polskiego – Andrzej Nowicki



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://bit.ly/3HKJAN9>
Materiały dodatkowe są również dostępne na stronie edw.elportal.pl: <https://bit.ly/3BLTRAa>

Zdalna stacja nadzoru (monitoringu)

Jeśli masz drogi samochód, łódź, przyczepę kempingową, domek letniskowy, gospodarstwo ... w ogóle praktycznie wszystko... musisz wiedzieć, co się dzieje, gdy jesteś daleko. Czy akumulator się rozładowuje? Czy Twoja łódź nabiera wody? Czy Twoja pompa wodna pracuje bez przerwy? Musisz dowiedzieć się o tym ASAP (As Soon As Possible czyli tak szybko jak tylko można). Wszystko czego potrzebujesz do tego to kilka nakładek podłączonych do modułu Arduino i trochę oprogramowania. Możesz nawet zdalnie wyzwać działania, takie jak wyłączenie źle pracującej pompy, zanim spuści całą wodę!

Musimy przyznać: zastosowanie tego projektu nie miało pierwotnie nic wspólnego z monitorowaniem dróg samochodów lub łodzi, odległych domów letniskowych, zbiorników wody w gospodarstwie rolnym lub czymś innym tak przyziernym.

Wszystko to miało związek z wombatami. Dla naszych zagranicznych czytelników: wombaty, gatunek sporych torbaczy nieco zagrożony w naturze, to urocze, (zazwyczaj) wolno poruszające się zwierzęta futerkowe, które zamieszkują australijski busz (i, nawiasem

mówiąc, są wyjątkowe w tym, że ich odchody mają kształt sześcianu!).

Ale nawet to nie jest cała historia.

Brendan Akhurst, rysownik SC, mieszka w buszu i jest członkiem lokalnego stowarzyszenia ochrony wombatów.

Częścią ich działań jest ponowna introdukcja wombatów na obszarach, gdzie są narażone na mniejsze niebezpieczeństwo zaatakowania przez inne zwierzęta (np. psy). Robią to poprzez ich odłowienie i przeniesienie.

Problemem jest to, że wombaty bardzo łatwo się stresują i zginą, jeśli zostaną uwięzione zbyt długo.

Brendan chciał mieć możliwość poinformowania członków stowarzyszenia, że jedna z ich pułapek zadziałała, tak szybko jak to możliwe.

„Aha!” powiedzieliśmy. „Jest projekt, nad którym pracujemy od kilku miesięcy, który powiadomi Cię, za pośrednictwem telefonu komórkowego, o praktycznie każdym zdarzeniu”. „Obejmuje to sprężynową pułapkę na wombaty?” zapytał.

Powiedzieliśmy: „Praktycznie każde zdarzenie!”

Więc Brendan's Wonderful Wombat Warning Whatchamacallit (możesz to tłumaczyć jako: Magiczny Wihajster Brendana Alarmujący o Wombacie w Puddle) jest wynikiem...

Oczywiście to, do czego go użyjesz, zależy wyłącznie od Ciebie!

2G, 3G i 4G

Sieć komórkowa drugiej generacji 2G (GSM) została już w zasadzie wyłączona, a niektóre telkomy zaczynają planować wyłączenie sieci 3G.

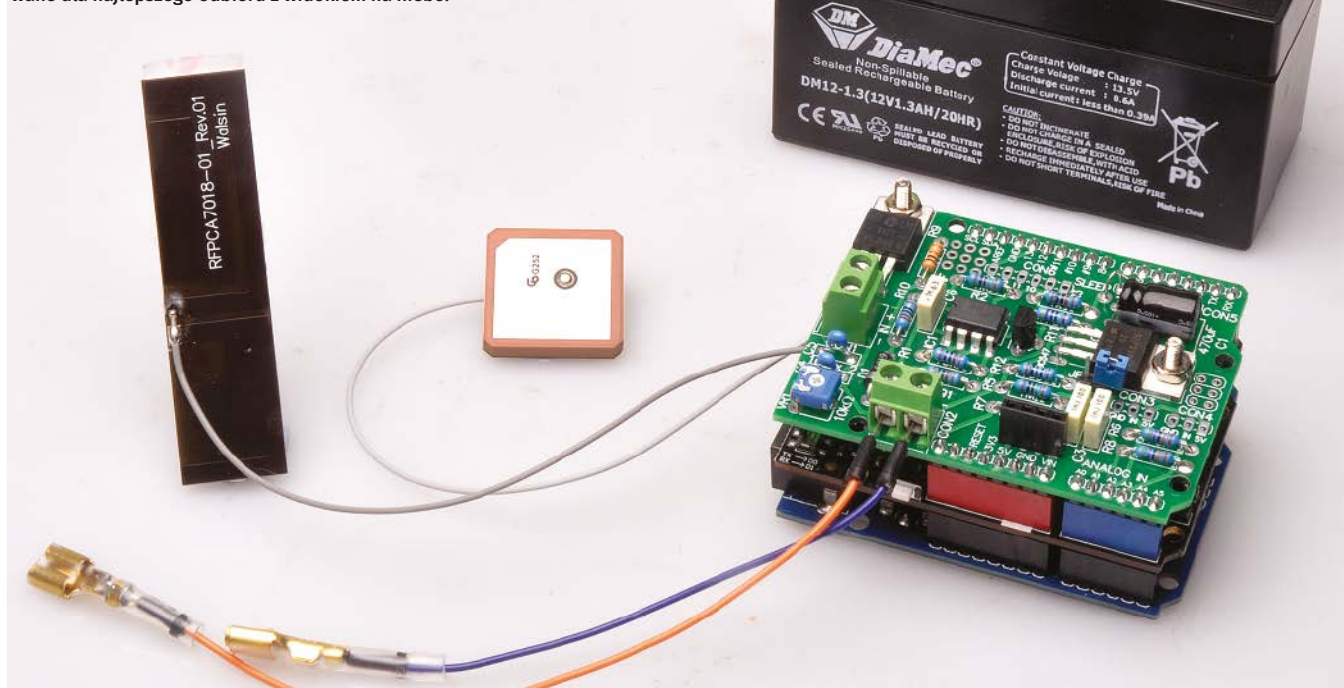
Tak więc, aby wykonywać tego rodzaju połączenia niezawodnie co najmniej przez kilka następnych lat, potrzebne jest urządzenie 4G.

Kiedy więc znaleźliśmy nakładkę Arduino 4G w rozsądnej cenie ok. 300 zł, skorzystaliśmy z okazji, aby zaprojektować wokół niej Zdalną Stację Monitoringu.

Ponieważ sercem tej stacji jest moduł Arduino, można ją łatwo zaprogramować, aby dostosować ją do konkretnych wymagań. Może monitorować stan przełączników, napięcia, czujników – po prostu wszystko, pod warunkiem, że istnieje biblioteka Arduino,



Gotowy zespół trzech modułów ma zwartą budowę, potrzebuje tylko kilku luźnych przewodów. Nawet najmniejszy akumulator 12 V SLA wystarczy do zasilania. Moduł Arduino (w tym przypadku Duinotech Leonardo) jest na dole, nakładka SIM7000E w środku, a nasza nakładka sterująca zasilaniem na górze. Anteny, zarówno dla modułu GPS jak i główna antena GSM (po lewej stronie zdjęcia) powinny być zamontowane dla najlepszego odbioru z widokiem na niebo.



obsługująca ten czujnik, przełącznik, pomiar napięcia etc. (i zwykle jest)!

Podobnie, możesz wysyłać polecenia do Arduino z telefonu komórkowego lub komputera, aby robić takie rzeczy jak zdalne włączanie lub wyłączanie zasilania sieciowego, używając prostego urządzenia dodatkowego, takiego jak nasz Opto-Isolated Mains Relay (przełącznik zasilania sieciowego z optoizolacją) (SC październik 2018 r., siliconchip.com.au/Article/11267).

Być może pamiętasz projekt Zdalnej Stacji Monitorowania GSM z numeru SC z marca 2014 roku, który również wykorzystywał moduł Arduino (siliconchip.com.au/Article/6743). To rozwiązanie jest już przestarzałe.

Jeśli zrealizowałeś kiedyś ten projekt, możesz być w stanie zaktualizować go, po prostu wymieniając nakładkę i wprowadzając kilka małych zmian w kodzie oprogramowania, ponieważ zestaw poleceń nowej nakładki 4G jest bardzo podobny.

Jedną drobną różnicą pomiędzy nakładką 2G a 4G, której tutaj używamy, jest to, że sygnał sterujący zasilaniem korzysta z innego portu Arduino. Nie testowaliśmy naszej nowszej nakładki 4G ze starszym projektem stacji monitoringu, ale prawdopodobnie będzie działał po dokonaniu pewnych zmian.

Nowa nakładka ma wystarczająco dużo dodatkowych możliwości, aby jej budowa była zasadna; i ta nowsza zdalna stacja monitoringu 4G dobrze wykorzystuje wiele przydatnych funkcji.

Ponieważ jest to projekt oparty na Arduino, musisz znać oprogramowanie Arduino IDE (zintegrowane środowisko programistyczne).

Można je pobrać bezpłatnie via link siliconchip.com.au/link/aatq lub którejs z lokalnych wersji stron opisujących środowisko Arduino, wraz z dodatkowymi bibliotekami. Najnowsza stabilna wersja to 2.0.

Nakładka 4G

Projekt ten wykorzystuje nakładkę 4G zaprojektowaną przez DFRobot. Sercem nakładki jest układ scalony SIM7000E firmy SIMCom, który zapewnia kompatybilność z europejskimi sieciami 4G.

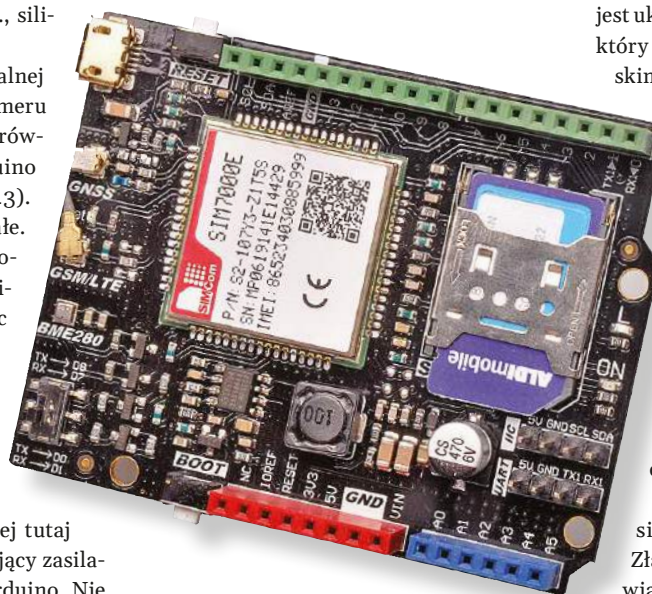
Schemat ideowy został przedstawiony na **rysunku 1**.

Moduł SIM7000E jest zasilany ze złącza VIN Arduino poprzez stabilizator MP2307. Wytwarza on stałe napięcie 3,3 V z wejściowego co najmniej 4,75 V.

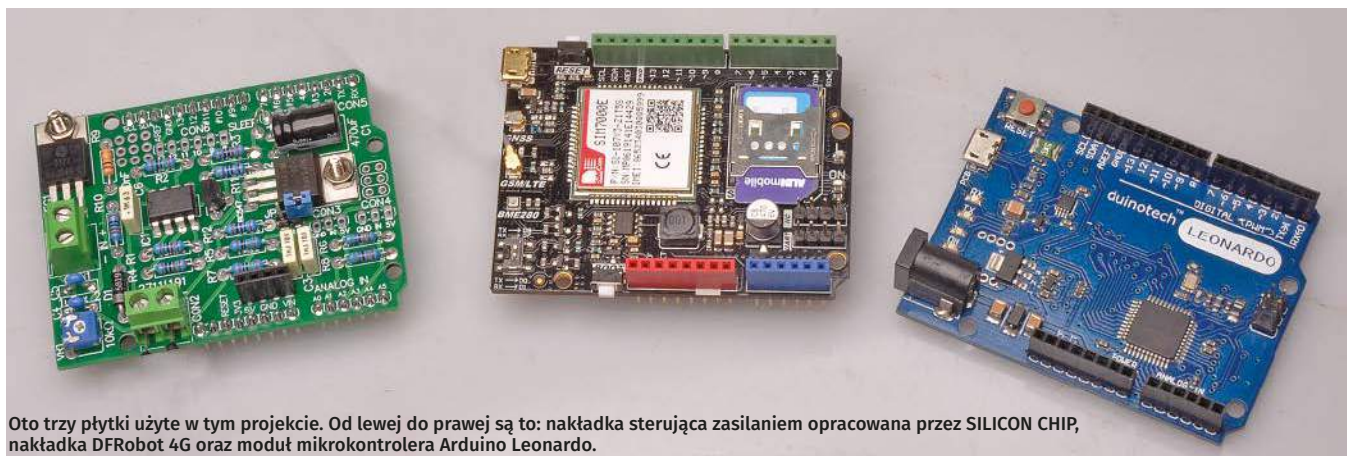
Chociaż eliminuje to możliwość zasilania nakładki z ogniwa litowo-jonowego lub LiPo o napięciu 3,7 V, to jednak będzie działał z większością aplikacji Arduino zasilanych ze złącza VIN lub gniazda DC.

Poniżej stabilizatora znajduje się sekcja sterowania zasilaniem. Złącze PWRKEY na SIM7000E jest ustawiane w stan niski, aby zasygnalizować, że ma ono włączyć lub wyłączyć zasilanie.

Przycisk chwilowy S1 łączy ten styk z masą, natomiast tranzystor NPN Q1 umożliwia osiągnięcie tego samego efektu wskutek



Użyliśmy karty SIM Aldi w sieci Telstra, aby przetestować naszą tarczę SIM7000E. Karta SIM kosztowała 5 dolarów i nie wymagała żadnych dopłat, nawet po dwóch miesiącach testów. Nakładka posiada również gniazda dla zewnętrznych anten sieci komórkowej i GNSS.



Oto trzy płytki użyte w tym projekcie. Od lewej do prawej są to: nakładka sterująca zasilaniem opracowana przez SILICON CHIP, nakładka DFRobot 4G oraz moduł mikrokontrolera Arduino Leonardo.

przejścia styku D12 modułu Arduino w stan wysoki.

Komunikacja pomiędzy modułem Arduino a nakładką odbywa się za pomocą pary łączących TX/RX, poprzez MOSFET-y z kanałami P: Q8 i Q9. Przełącznik S2 kieruje sygnały do D0/D1 (który jest zwykle sprzętowym portem szeregowym na płytce Arduino; sygnały dostępne są wtedy na złączu UART nakładki SIM7000E) lub D8/D7 na płycie Arduino.

Port USB SIM7000E jest wyprowadzony na złącze micro-USB.

Nie zasila ono nakładki, ale może być wykorzystane przez komputer do komunikacji z modułem SIM7000E.

Nie badaliśmy tego szczegółowo, ale wydaje się, że wiele funkcji modułu można wykorzystać poprzez wymienione złącze USB. Być może jest ono w stanie uruchomić nakładkę jako modem USB 4G.

Są też gniazda dla standardowej karty SIM, anteny 4G i anteny GNSS (Global Navigation Satellite System), która jest używana do odbioru GPS (USA) i GLONASS (Rosja).

Więcej informacji na temat tych i wielu innych istniejących systemów nawigacji satelitarnej można znaleźć w artykule „A look at SatNav systems” w wydaniu SC z listopada 2019 roku (siliconchip.com.au/Article/12075).

Jednym z częstych zastosowań zdalnej stacji monitorowania jest śledzenie pojazdów, a w tym przypadku odbiornik GNSS jest praktycznie obowiązkowy.

Nie musimy dodawać żadnego dodatkowego sprzętu, aby zaimplementować śledzenie pozycji w naszej Zdalnej Stacji Monitoringu 4G (a więc natychmiast dowiemy się, że ktoś, poza wombatami, zainteresował się naszą Stacją).

Przy zakupie nakładki do zestawu dołączone są dwie anteny. Antena sieci komórkowej to prosta, samoprzylepna antena typu PCB.

Niektóre zdjęcia tej nakładki pokazują małą antenę typu bicz, ale wygląda na to, że została ona zastąpiona przez antenę typu PCB. W zestawie znajduje się płaska ceramiczna antena mikropaskowa do zastosowań GNSS.

Na nakładce znajduje się również czujnik temperatury, ciśnienia i wilgotności BME280.

Red. EdW: co najmniej dziwna polityka firmy DFRobot spowodowała, że od wersji 2.0 nakładka SIM7000E nie posiada groźowego czujnika BME280. Oczywiście możesz zdobyć z zapasów magazynowych pełnowartościową nakładkę z czujnikiem, ale musisz się upewnić, że dystrybutor Ci to zagwarantuje. Zwróć uwagę, że do Twoich celów nadaje się TYLKO SIM7000E (europejski zakres częstotliwości), nakładki SIM7000A i SIM7000C przeznaczone są na inne

rynki i w Europie są nieprzydatne. Zachowaj ostrożność przy zamawianiu!

Wykorzystaliśmy moduły zbudowane z urządzeniem podobnych czujników jeszcze w 2017 roku (siliconchip.com.au/Article/10909). Jest to świetny bonus, ponieważ dodaje jeszcze więcej danych z czujników do naszej zdalnej stacji monitorującej 4G bez konieczności stosowania dodatkowego sprzętu.

Red. EdW: sugerujemy zapoznanie się z artykułem „Czujnik smogu IoT” w nr 5/6 z 2019 r. EdW autorstwa Pana Pawła Hoffmana.

Stwierdziliśmy, że temperatura odczytywana przez czujnik była wyższa od temperatury otoczenia, prawdopodobnie z powodu ciepła generowanego przez otaczające obwody; pomyśl, jak gorące są niektóre telefony komórkowe!

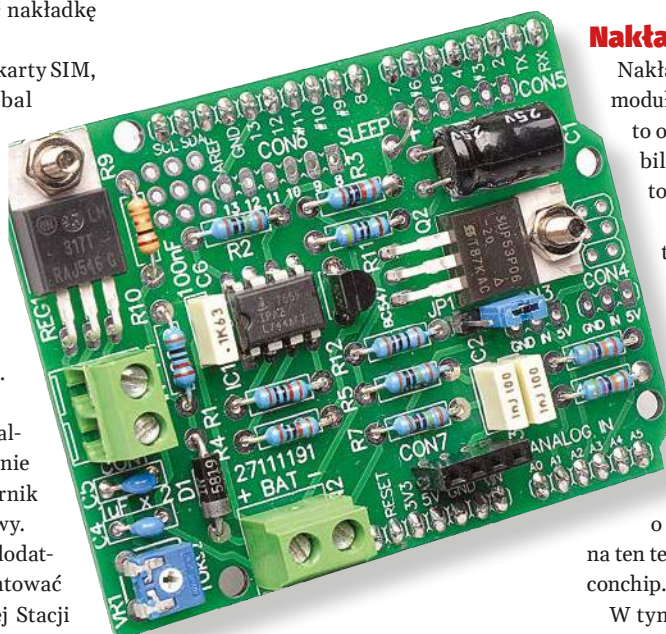
Nakładka SIM7000E

Nakładka SIM7000E jest określana jako moduł NB-IoT/LTE/GPRS/GPS. LTE i GPRS to od dawna stosowane technologie mobilnej transmisji danych, ale NB-IoT to nowszy standard.

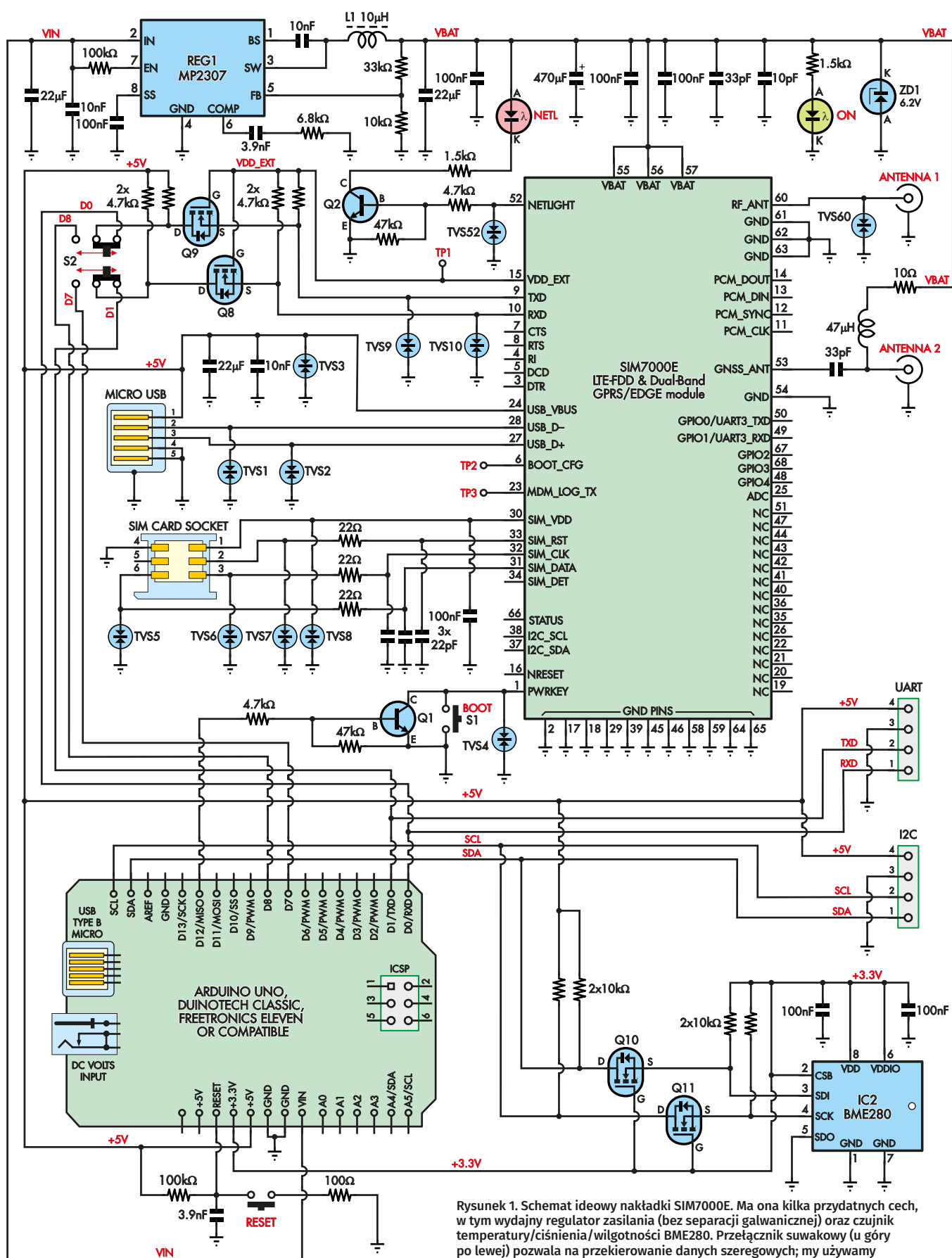
NB-IoT to wąskopasmowy wariant technologii telefonii komórkowej o niskim poborze mocy, zaprojektowany z myślą o wykorzystaniu przez urzędników IoT (Internet of Things, tłumaczone jako „Internet Rzeczy”, ale lepiej brzmi dwuznaczne „Internet do Rzeczy” albo „Internet do Wszystkiego” – pralki, lodówki czy szczoteczki do zębów). Więcej o IoT można dowiedzieć się z artykułu na ten temat w SC z listopada 2016 roku (siliconchip.com.au/Article/10425).

W tym projekcie nie korzystamy z funkcji NB-IoT; na tym etapie wydaje się, że technologia ta jest wciąż wdrażana w Australii (i w Polsce); wymagana jest karta SIM specyficzna dla NB-IoT.

Moduł SIM7000 występuje w kilku wariantach, które obsługują różne pasma



Kondensator elektrolityczny i timer 555 zastosowane na nakładce sterującej zasilaniem zostały starannie dobrane pod kątem niskiej upływności i niskiego prądu spoczynkowego, aby wydłużyć żywotność akumulatora. Należy zwrócić uwagę na zworkę łączącą port D7 Arduino z zaciskiem SLEEP.



Rysunek 1. Schemat ideowy nakładki SIM7000E. Ma ona kilka przydatnych cech, w tym wydajny regulator zasilania (bez separacji galwanicznej) oraz czujnik temperatury/ciśnienia/wilgotności BME280. Przełącznik suwakowy (u góry po lewej) pozwala na przekierowanie danych szeregowych; my używamy portów D0 i D1.

częstotliwości komórkowych. My korzystamy z SIM7000E (model sprzedawany przez Core Electronics lub Twojego lokalnego dystrybutora), który przeznaczony jest na rynek europejski i obsługuje pasma 3, 8, 20 i 28. Jest też SIM7000C, który jest przeznaczony dla częstotliwości używanych w Chinach.

W naszych testach w podmiejskim rejonie Sydney nie mogliśmy uzyskać połączenia stosując kartę SIM sieci Optus, chociaż sieć Optus używa niektórych z powyższych pasm. Natomiast sukcesem zakończyły się próby z kartą SIM sieci Telstra.

Ponieważ nie wszystkie częstotliwości są oferowane we wszystkich obszarach, Twoje doświadczenie może być inne. W Polsce najprawdopodobniej będziesz musiał odwiedzić salony kilku sieci, aby wybrać kartę SIM w najlepszej ofercie transmisji danych.

Sugerujemy, aby dokładnie sprawdzić, jakie częstotliwości są używane tam, gdzie planujesz wdrożyć Zdalną Stację Monitoringu 4G, aby upewnić się, że ta nakładka je obsługuje.

Ten moduł nie obsługuje połączeń głosowych. Większość stacji monitorujących zazwyczaj wykorzystuje do komunikacji SMS-y (wiadomości tekstowe) lub pakiety danych. Moduł SIM7000E obsługuje dane mobilne (dane przekazywane za pośrednictwem Internetu) i jest to świetny sposób na transmisję wielu małych pakietów informacji z monitoringu.

Nasz projekt wykorzystuje przy pomocy nakładki zarówno funkcje SMS-ów, jak i dane mobilne.

Rejestrowanie danych ThingSpeak

Opisany w SC miernik poziomu zbiornika wodnego z lutego 2018 roku (siliconchip.com.au/Article/10963) to zdalne urządzenie, które okresowo przysyłało dane na stronę ThingSpeak. Korzystało ono z połączenia WiFi z istniejącym punktem dostępu do Internetu, co ograniczało lokalizację, w której mogło być używane.

Takie ograniczenia można usunąć stosując nakładkę 4G SIM7000E, taką jak ta.

W tym projekcie również używamy strony ThingSpeak. Ma ona proste API (interfejs programowania aplikacji) do przysyłania danych, co jest idealne dla urządzeń o ograniczonych zasobach, takich jak mikrokontrolery Arduino. Zapewnia również prostą graficzną wizualizację zarejestrowanych danych. Dane można również pobrać jako plik CSV (comma-separated value – liczby oddzielone przecinkiem).

Pliki te można otworzyć w programie arkusza kalkulacyjnego, aby umożliwić bardziej zaawansowaną analizę. W wielu programach

arkuszy kalkulacyjnych opcją jest również tworzenie wykresów. Masz również do wyboru profesjonalne programy graficznej prezentacji danych.

Przesyłanie danych na stronę ThingSpeak wymaga transmisji danych mobilnych, więc używana karta SIM musi obsługiwać taką możliwość.

W przypadku tanich, przedpłaconych („pre-paid”) planów telefonicznych o dłuższym okresie ważności, które wypróbowaliśmy, wysyłanie przez wiele tygodni danych mobilnych na stronę ThingSpeak było zazwyczaj tańsze niż wysyłanie pojedynczych wiadomości tekstowych.

Nasza Zdalna Stacja Monitoringu 4G idealnie nadaje się do zapewnienia ciągłej rejestracji danych za pośrednictwem sieci 4G, jak również wysyłania wiadomości tekstowych w celu zaalarmowania o nietypowych sytuacjach, które wymagają natychmiastowego działania.

Nakładka sterująca zasilaniem

Oprócz firmowej nakładki 4G, nasza Zdalna Stacja Monitoringu wykorzystuje również specjalnie zaprojektowaną drugą nakładkę, która zapewnia zasilanie z akumulatora, słoneczne ładowanie akumulatora oraz niektóre procedury oszczędzania energii, aby zapewnić długi czas pracy przy użyciu akumulatora o niewielkiej pojemności.

Niestety, większość modułów Arduino ma niską sprawność energetyczną; po prostu nie projektowano ich z myślą o tym. Nawet z mikroprocesorem ustawionym w stan uśpienia, inne komponenty, takie jak liniowe regulatory napięcia i diody LED mają prądy spoczynkowe rzędu dziesiątek miliamperów.

Nasza nakładka redukuje pobór prądu z akumulatora w trybie czuwania do mikroamperów, co jest możliwe dzięki całkowitemu odłączeniu płytki Arduino (i nakładki SIM7000E) od zasilania za pomocą MOSFET-a, i tylko okresowemu zasilaniu tych komponentów.

Nakładka zapewnia dość wydajny sposób ładowania akumulatora, a także sprawdza jego napięcie i napięcie zasilania poprzez wejścia analogowe modułu Arduino.

Większość niewykorzystanych portów jest zdublowana na dodatkowe złącza, co pozwala na dołączenie innych czujników lub urządzeń peryferyjnych. Jest na niej nawet mały obszar do testowania dodatkowych komponentów.

Schemat nakładki

Schemat ideowy nakładki pokazany jest na **rysunku 2**. Układ jest prościutki i zawiera niewiele elementów. Niestabilizowane napięcie zasilania DC jest doprowadzane do złącza CON1 (z panelu słonecznego,

zasilacza wtyczkowego itp.), podczas gdy akumulator jest podłączony przez CON2.

Akumulator musi dostarczać napięcia w zakresie 7-15 V, więc odpowiedni jest 12 V akumulator kwasowo-ołowiowy, najlepiej bezobsługowy SLA (Sealed Lead-Acid). W naszym prototypie zastosowaliśmy mały akumulator SLA 12 V o pojemności 1,3 Ah.

Ze złącza CON1 zasilany jest REG1, regulowany stabilizator LM317. Rezystor 220 Ω i potencjometr nastawny 10 k Ω VR1 pozwalają na ustawienie jego napięcia wyjściowego.

Ponieważ REG1 utrzymuje różnicę napięć około 1,25 V pomiędzy końcówkami OUT i ADJ, przez rezystor 220 Ω przepływa prąd około 5 mA.

Ten prąd w większości przepływa również przez VR1, więc zmieniając jego rezystancję zmieniamy napięcie pomiędzy ADJ a GND.

Stąd można ustawić napięcie na VOUT, gdyż będzie to napięcie na VR1 plus 1,25 V.

Wyjście z regulatora jest filtrowane przez kondensator 1 μ F i podawane do akumulatora przez diodę Schottky'ego 1A, D1. Zapobiega to rozładowaniu akumulatora przez źródło zasilania, np. gdy jest to nieoświetlony panel słoneczny.

Rezystor 1 Ω pomiędzy wyjściem REG1 a anodą D1 zmniejsza napięcie wyjściowe wraz ze wzrostem prądu pobieranego z REG1.

Hipotetycznie, gdyby prąd płynący przez ten rezystor osiągnął wartość 1,25 A (co w praktyce nie byłoby możliwe), to napięcie na tym rezystorze wzrosłoby do 1,25 V, kompensując napięcie odniesienia REG1, a więc napięcie wyjściowe spadłoby do 0 V.

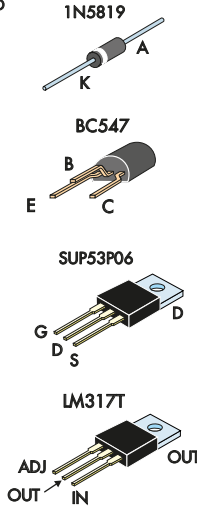
Tak więc napięcie wyjściowe stabilizatora REG1 spada o około 1 V na każde 100 mA prądu obciążenia.

Jeśli więc akumulator jest mocno rozładowany i napięcie na jego zaciskach jest niskie, prąd wyjściowy stabilizatora jest ograniczany do momentu, gdy jego napięcie wzrośnie do normalnego poziomu, w którym to momencie do akumulatora nie popłynie praktycznie żaden prąd.

W rzeczywistości prąd ładowania jest ograniczony przez rozpraszanie do około 160 mA dla ogniwa słonecznego 12 V (o nominalnym napięciu 18 V w obwodzie otwartym) zasilającego rozładowany akumulator 12 V.

Aczkolwiek potencjometr VR1 pozwala na ustawienie napięcia na zaciskach od 1,25 V do 56 V, jednak nie powinno się go ustawiać wyżej niż około 15 V, gdyż może to spowodować uszkodzenie stabilizatora niektórych modułów Arduino, a także układu IC1.

Jeśli nie musisz używać akumulatora, zasilanie może być doprowadzone bezpośrednio



lutowicznego SLEEP i nie jest zwykle używany do żadnych konkretnych celów.

Następnie zamontuj trzy prostokątne kondensatory MKT (nie są spolaryzowane). Kondensator 100 nF może być oznaczony kodem 104 lub ewentualnie 0,1 μ F, natomiast kondensatory 1 nF – kodem 102. Dalej lutujemy dwa kondensatory ceramiczne 1 μ F, które również nie są spolaryzowane.

Ostatnim jest kondensator elektrolityczny o małej upływności. Jest on spolaryzowany, a na płycie drukowanej zostawiliśmy miejsce na zamontowanie go poziomo na boku, aby można było ewentualnie zamontować nad nim kolejną nakładkę. Ujemne wyprowadzenie (zwykle krótsze i oznaczone paskiem na obudowie) trafia do pola bliżej Q2. Jeśli chcesz użyć większego kondensatora dla dłuższego opóźnienia (pomiędzy kolejnymi włączeniami urządzenia), zostawiliśmy trochę dodatkowego miejsca.

Możesz użyć niewielkiej ilości kleju termoplastycznego lub neutralnego samoutwardzalnego kitu silikonowego, aby unieruchomić kondensator w przypadku, gdy urządzenie będzie narażone na wibracje.

Następnie zamontuj półprzewodniki. D1 jest jedyną diodą i umieszczona jest w pobliżu CON2, z paskiem katody najbliżej CON2. Zamontuj tranzystor Q1 w pobliżu środka płytki, w kierunku pokazanym na rysunku. Ostrożnie wygnij jego wyprowadzenia, aby dopasować je do otworów w płycie, wcisnij mocno i przylutuj.

Q2 i REG1 mają obudowy TO-220, które są zamontowane płasko na płycie drukowanej, aby nie zwiększać wysokości modułu. Nie należy ich pomylić między sobą. Zegnij do tyłu wyprowadzenia każdego elementu, około 8 mm od miejsca, skąd końcówki wychodzą z obudowy.

Włóż wyprowadzenia w otwory PCB i upewnij się, że otwór radiatora pokrywa się z otworem w PCB, a następnie użyj śrubki M3 (od spodu) i nakrętki (od góry), aby unieruchomić stabilizator i MOSFET przed przylutowaniem ich wyprowadzeń.

Następnie należy zamontować układ scalony timera IC1. Może być konieczne delikatne wygięcie wyprowadzeń do środka, aby pasowały do płytki. Sprawdź orientację układu scalonego, aby upewnić się, że jest zgodna ze schematem montażowym, następnie przylutuj dwa przeciwległe wyprowadzenia po przekątnej i sprawdź, czy obudowa przylega płasko do płytki. Jeśli nie, przetop lutowanie, dociskając obudowę, a następnie przylutuj pozostałe wyprowadzenia.

Aby przylutować listwy kołkowe do połączenia z gniazdami na nakładce SIM7000E, wcisnij ich odcinki o odpowiedniej długości

do jakiejś płytki Arduino (np. Leonardo), aby były proste. Umieść nakładkę zasilania nad końcami szpilek i nasuń na nie jej otwory. Po upewnieniu się, że kołki ustawione są prosto, a korpus listwy przylega od dołu do nakładki, przylutuj każdą końcówkę szpilki do PCB. Zwykle listwy kołkowe są przylutowane na górze, natomiast listwy kołkowe z gniazdami do kołków, z możliwością układania w stos, są konieczne przylutowane pod spodem. Na fotografii na dole strony, po prawej, widać oba typy złączy. Zwykła listwa kołkowa z korpusem w kolorze czarnym jest przylutowana od góry nakładki zasilania, natomiast listwa kołkowa z gniazdami do szpilek (korpus w kolorze jasnozielonym) jest fabrycznie przylutowana od dołu nakładki SIM7000E. Gniazda łączą się z kołkami nakładki zasilania, natomiast kołki z gniazdami modułu Arduino. Następnie zdejmij nakładkę z modułu Arduino.

Teraz zamontuj CON1 i CON2, czyli złącza z zaciskami śrubowymi. Są one identyczne, przy czym należy zwrócić uwagę, aby otwory do wprowadzania przewodów były skierowane na zewnątrz płytki.

Możesz teraz zamontować dwuszpilekowy odcinek listwy kołkowej jako JP1 oraz złącza CON3-CON7, na które wyprowadzone są różne porty modułu Arduino. Nie są one potrzebne do najbardziej podstawowego zastosowania Zdalnej Stacji Monitoringu 4G, ale są przydatne do dodawania dodatkowych czujników, kiedy to konieczne. CON3-CON6 nie są lutowane w naszym prototypie.

Na płycie prototypowej, wg fotografii, wlutowane jest tylko złącze CON7. Jest to poczwórne żeńskie gniazdo do wtyków kołkowych, zamontowane, aby umożliwić łatwe monitorowanie napięcia VIN.

Aby sprawdzić, czy nakładka zasilania podaje napięcie na wejście VIN, zastosowaliśmy testową diodę LED poprzez przylutowanie rezystora 5 k Ω do jednego z wyprowadzeń diody. Następnie podłączyliśmy zestaw dioda/rezystor do pary gniazd VIN/GND na CON7, z anodą diody do VIN.

Testowanie

Możemy wykonać kilka podstawowych testów, aby sprawdzić czy wszystko działa jak należy. Nakładka podczas tych testów nie może być podłączony do żadnego innego modułu. Przed włączeniem zasilania należy skrócić potencjometr nastawny VR1 całkowicie w lewo.

Pierwszym krokiem jest ustawienie VR1 na właściwe napięcie ładowania akumulatora. Robi się to bez podłączonego akumulatora.

W tym celu podłącz do CON1 zasilacz, który dostarcza napięcie o co najmniej 3 V wyższe niż napięcie w pełni naładowanego akumulatora. Ustaw VR1 aż do osiągnięcia prawidłowego maksymalnego napięcia ładowania na CON2.

W przypadku akumulatora SLA lub podobnego typu o nominalnym napięciu 12 V należy ustawić je na około 14,4 V. W praktyce napięcie takie jest osiągnięte tylko przy zerowym

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

- 1 szt. moduł Arduino Leonardo lub płytka kompatybilna
- 1 szt. nakładka DFRobot SIM7000E [Digi-Key/Mouser (Cat DFR0505) lub bezpośrednio z www.dfrobot.com albo lokalnego dystrybutora] w wykonaniu wcześniejszym niż v.2.0.
- 1 szt. karta SIM 4G obsługująca SMS-y i dane internetowe
- 1 szt. nakładka sterująca zasilaniem (patrz poniżej)
- 1 szt. akumulator 12 V (bezoługowy) i odpowiednie źródło ładowania (np. mały panel słoneczny 12 V)

Części do nakładki sterującej zasilaniem

- 1 szt. dwustronna płytka drukowana o kodzie 27111191 i wymiarach 53,5×68,5 mm
- 2 szt. 2-drożne zaciski śrubowe do montażu na PCB, raster 5,08 mm [Jaycar HM3172, Altronics P2032B] (CON1, CON2).
- 1 szt. zestaw listew kołkowych do podłączenia do Arduino (1× 6-szpilek, 2× 8-szpilek, 1× 10-szpilek – patrz tekst) lub analogicznych listew męsko-żeńskich
- 1 szt. 2-szpilekowy odcinek listwy kołkowej ze zworką (JP1)
- 2 szt. 3-szpilekowe odcinki listwy kołkowej (CON3, CON4) – nie występuje na fotografii, opcja
- 2 szt. 6-szpilekowe odcinki listwy kołkowej (CON5, CON6) – opcja
- 1 szt. 4-drożne jednorzędowe gniazdo żeńskie (CON7)
- 2 kompl. śruby M3×6 i nakrętki (do montażu REG1 i Q2)

Półprzewodniki:

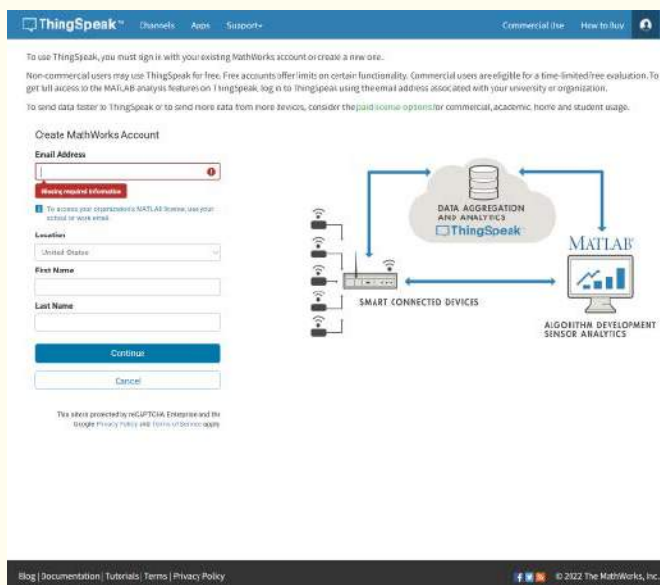
- 1 szt. 7555 timer IC CMOS, DIP-8 (IC1)
- 1 szt. LM317 regulowany stabilizator napięcia, TO-220 (REG1)
- 1 szt. BC547 tranzystor NPN, TO-92 (Q1)
- 1 szt. SUP53P06 MOSFET P-kanalowy, TO-220 (Q2)
- 1 szt. 1N5819 dioda Schottky'ego (D1)

Kondensatory:

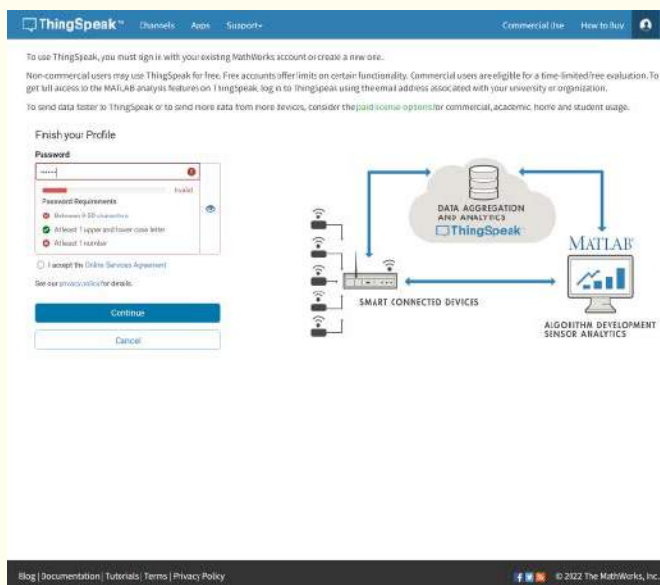
- 1 szt. 470 μ F/25 V kondensator elektrolityczny (low ESR)
- 2 szt. 1 μ F kondensatory ceramiczne MLC [Jaycar RC5499]
- 1 szt. 100 nF kondensator foliowy MKT
- 2 szt. 1 nF kondensatory foliowe MKT

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

- | | | |
|---|-----------------------|----------------------|
| 3 szt. 1 M Ω | 2 szt. 470 k Ω | 1 szt. 10 k Ω |
| 2 szt. 1 k Ω | 1 szt. 470 Ω | 1 szt. 220 Ω |
| 1 szt. 100 Ω | 1 szt. 1 Ω | |
| 1 szt. potencjometr nastawny 10 k Ω mini poziomy (VR1) | | |



Rysunek 4. Zarejestruj się i załóż konto ThingSpeak poprzez pokazaną tu stronę internetową. Jest to niezbędne do korzystania z napisanego przez nas oprogramowania, gdyż ThingSpeak pozwala na przesyłanie danych do „chmury”. Użytkownicy Matlab-a mogą wykorzystać swoje istniejące konto w ThingSpeak.



Rysunek 5. Jak w przypadku wielu usług online, musisz utworzyć na stronie ThingSpeak nazwę użytkownika i hasło. Strona ta informuje, czy wybrana przez Ciebie nazwa użytkownika jest wolna, oraz jak silne według niej jest Twoje hasło.

prądzie, więc rzeczywiste napięcie ładowania jest nieco niższe od tego.

Jeśli nie jesteś w stanie ustawić poprawnie napięcia, sprawdź elementy powiązane z REG1. Jeśli układ ładowania jest w porządku, podłącz akumulator i sprawdź, czy jest ładowany. Powinieneś zobaczyć, że napięcie akumulatora powoli rośnie.

Następnie sprawdź napięcie pomiędzy VIN i GND, używając diody LED, o której wspominaliśmy wcześniej, lub woltomierza. Te wyjścia są łatwo dostępne na CON7. Prawdopodobnie uzyskasz zerowy odczyt, co oznacza, że MOSFET Q2 jest wyłączony.

W takim przypadku, jeśli sprawdzisz napięcie na wyprowadzeniach kondensatora elektrolitycznego, powinieneś zauważyć, że powoli rośnie. Można je zmierzyć również na styku 6 IC1 w odniesieniu do GND. Zauważ, że obciążenie kondensatora multimetrem może wpłynąć na ten odczyt.

Teraz zewrzyj JP1, aby wymusić włączenie Q2 i ponownie sprawdź napięcie VIN. Powinno być ono zbliżone do napięcia akumulatora, a napięcie na styku 6 układu IC1 powinno być niskie.

Aby zasymulować żądanie wyłączenia zasilania przez moduł Arduino, podłącz na chwilę pole stykowe SLEEP za pomocą rezystora 1 kΩ do zacisku VBAT złącza CON2. VIN powinno spaść do zera, a kondensator elektrolityczny powinien zacząć się ładować. Jeśli tak nie jest, może to oznaczać, że Twój kondensator jest kiepskiej jakości.

Jeśli zależy Ci na dokładnym ustawieniu trwania okresu uśpienia, możesz zmierzyć czas i zmienić wartości elementów determinujących

stałą czasową, aby ten czas ustawić. Pamiętaj, że Arduino musi działać przez co najmniej 30 sekund, aby zaktualizować swój stan, więc okresy uśpienia krótsze niż dwie minuty nie są zbyt przydatne, ponieważ Arduino poświęci o wiele za dużo czasu na uruchamianie.

Po zakończeniu testów odłącz zasilacz i akumulator.

Budowa Zdalnej Stacji Monitoringu

Po zbudowaniu i sprawdzeniu nakładki zasilania, możemy teraz złożyć wszystko razem. Do naszego prototypu wybraliśmy moduł Arduino Leonardo. Używa on ATmega32U4 zamiast ATmega328 w wersji Uno. Litera „U” wskazuje, że ten układ scalony obsługuje USB.

Ich specyfikacja jest poza tym dość zbliżona, ale Leonardo ma tę przewagę, że sprzętowy port szeregowy na D0/D1 nie jest współdzielony z interfejsem szeregowym hosta (węzła) USB używanym do programowania. Możemy więc użyć go do komunikacji z SIM7000E. Leonardo ma również dodatkowe 512 bajtów pamięci RAM; może to być przydatne do zdalnego nadzoru, ponieważ musimy przechowywać i przetwarzać dane przed ich wysłaniem.

Gdybyśmy użyli Arduino Uno, byłibyśmy zmuszeni do wyboru między użyciem sprzętowego portu szeregowego (D0/D1) do komunikacji z SIM7000, co przeszkadzałoby w programowaniu i debugowaniu, a użyciem programowego portu szeregowego, który jest powolny i ma dużo ograniczeń.

Ustawiliśmy więc przełącznik linii komunikacyjnych na nakładce SIM7000 w pozycji

D0/D1 i jak wspomnieliśmy wyżej, wykorzystaliśmy D7 jako styk sterujący uśpieniem.

Aby skonfigurować nakładkę 4G, zamontuj dwie anteny i włóż aktywną kartę SIM. Jak w przypadku wielu tego typu aplikacji, preferowana jest karta SIM typu „prepaid”, na wypadek gdyby mikrokontroler „zwarował”. Z kartą „prepaid” jej limit uniemożliwia przypadkowe kosmiczne opłaty za przesłane dane lub wykonane połączenia.

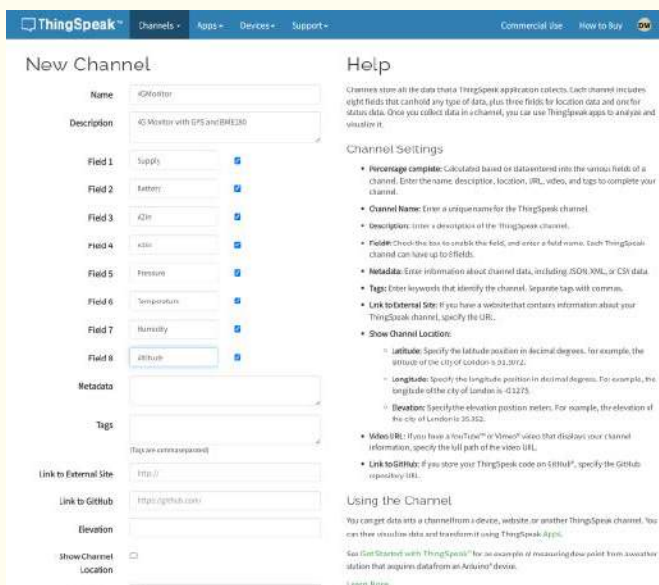
Teraz podłącz moduł 4G do Leonardo, a następnie podłącz nakładkę sterowania zasilaniem na górze. Sprawdź, czy nie ma żadnych zanieczyszczeń pomiędzy płytkami; jeśli nie przycięłeś dokładnie wszystkich przewodów, mogą się one zewrzeć ze sobą.

Nasze przykładowe oprogramowanie po prostu rejestruje dane z czujników nakładki 4G. Zaznaczyliśmy również kilka miejsc w kodzie, aby dodać własne testy lub działania. Na przykład, możesz monitorować napięcie i wysłać SMS, jeśli będzie zbyt niskie lub za wysokie. Lub podobnie, możesz wysłać SMS, jeśli przełącznik zostanie otwarty lub zamknięty. Aby w pełni wykorzystać nasz przykładowy kod, musisz założyć konto ThingSpeak.

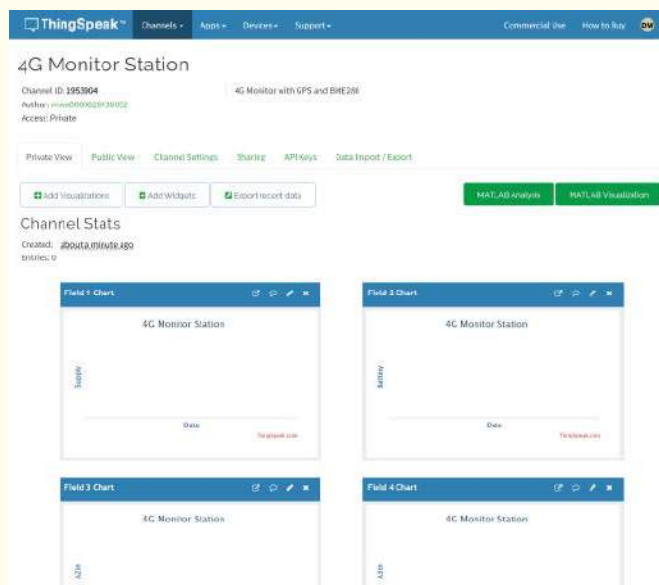
Zakładanie konta w ThingSpeak

ThingSpeak może być dostępny do użytku komercyjnego z ograniczonym czasowo okresem próbnym, ale darmowa licencja dostępna do użytku osobistego oferuje cztery „kanały” oraz do trzech milionów rejestracji danych rocznie.

Jeśli mielibyśmy wysłać aktualizacje stanu Stacji Monitoringu co dziesięć minut, wtedy



Rysunek 6. Zalecamy (przynajmniej na początku) utworzenie kanału ThingSpeak i skonfigurowanie jego pól w sposób pokazany tutaj. Pola te odpowiadają danym wytwarzanym przez oprogramowanie Zdalnej Stacji Monitoringu 4G. W razie potrzeby można je później zmienić.



Rysunek 7. Po utworzeniu kanału można przejść do jego przeglądania. Kanał domyślnie zawiera serię wykresów. Kolejne można dodać za pomocą przycisków „Dodaj”. Domyślnie dane kanału są prywatne, ale można je ustawić tak, aby były widoczne dla innych. Dla porównania można zapoznać się z kilkoma kanałami publicznymi.

potrzebowalibyśmy tylko około 50 000 pakietów danych rocznie.

Przejdź do strony https://thingspeak.com/users/sign_up [Sorry, that user cannot be found] czyli przejdź do strony <https://thingspeak.com/login?skipSSOCheck=true> i wprowadź informacje pokazane na rysunku 4. Możesz zostać poproszony o potwierdzenie, że chcesz użyć osobistego adresu e-mail, a także o kliknięcie na link wysłany w wiadomości e-mail w celu weryfikacji Twojego adresu e-mail.

Gdy to zrobisz, utwórz identyfikator użytkownika i hasło, zaakceptuj „Umowę o świadczenie usług online” i kliknij „Dalej”, jak na rysunku 5. Zostaniesz poproszony o wybranie sposobu korzystania z ThingSpeak. Aby móc korzystać z darmowej licencji, powinieneś wybrać opcję „Personal, non-commercial projects”.

Następnym krokiem jest utworzenie kanału. Każdy kanał składa się z maksymalnie ośmiu pól, więc teoretycznie mógłbyś mieć do czterech Zdalnych Stacji Monitoringu 4G, z których każda korzysta ze swojego niezależnego kanału.

Kliknij na „New Channel” i wypełnij informacje jak pokazano na rysunku 6. Nie musisz używać wszystkich ośmiu pól, ale ustawiliśmy oprogramowanie Arduino, aby używało wszystkich ośmiu, jak pokazano. Powinieneś użyć tych samych pól, chyba że planujesz zmodyfikować oprogramowanie.

Kliknij przycisk „Save”, a na stronie „Channel management” pojawią się dane kanału, jak widać na rysunku 7.

Zauważ, że nie utworzyliśmy pól dla szerokości i długości geograficznej. ThingSpeak

ma ukryte pola dla tych informacji. Nie widać ich na wykresach, ale są pobierane w danych CSV. Nasz kod Arduino wysyła szerokość i długość geograficzną do tych ukrytych pól.

Klucze API

Aby nasze urządzenie (i tylko nasze urządzenie) mogło przysyłać dane do naszych kanałów, potrzebujemy klucza API. Musi on zostać dołączony do kodu oprogramowania Arduino, aby Twoja Stacja Zdalnego Monitoringu 4G mogła współpracować z Twoim kanałem ThingSpeak.

Po zalogowaniu na stronie ThingSpeak przejdź do zakładki „API keys”. Skopiuuj 16-znakowy alfanumeryczny kod pod nagłówkiem „Write API Key” w bezpieczne miejsce; dodamy go wkrótce do kodu Arduino.

Możesz sprawdzić, czy Twój kanał działa, kopiując tekst po słowie „GET” w polu „Write a Channel Feed”. Wklej go do przeglądarki internetowej i naciśnij Enter; powinieneś zobaczyć pustą stronę z małą liczbą „1” w lewym górnym rogu.

Oznacza to, że jest to pierwsza aktualizacja i pokazuje jak Zdalna Stacja Monitoringu 4G przysyła dane do ThingSpeak. W ten sposób aktualizowane jest tylko jedno pole; jeśli jesteś zaznajomiony z protokołem HTTP, możesz z tym poeksperymentować.

Wróć do „Prywatnego widoku” („Private View”) utworzonego kanału i powinieneś zobaczyć jakąś aktywność w pierwszym polu; to są dane, które wysłałeś z przeglądarki internetowej. Możesz zostawić to okno otwarte podczas testowania, ponieważ będzie

ono aktualizowane w czasie prawie rzeczywistym i będziesz mógł zobaczyć wyniki. Możesz zapoznać się też z przykładowymi danymi wysyłanymi przez innych użytkowników do kanałów publicznych, przechodząc do zakładki „Dane publiczne” (Public View”)

Biblioteki Arduino

Potrzebne są cztery biblioteki niezbędne do działania napisanego przez nas oprogramowania; dwie są dołączone do większości dystrybucji Arduino IDE. My korzystaliśmy z wersji 1.8.5 Arduino IDE. Najnowsza wersja 2.0. oferuje przydatne narzędzia podczas konfigurowania i uruchamiania oprogramowania.

Biblioteki „avr/sleep” i „Wire” to dwie zwykle dołączane. Pierwsza biblioteka dostarcza funkcji dla trybów niskiego poboru energii, natomiast druga zapewnia interfejs I²C do komunikacji z czujnikiem BME280.

Trzecia biblioteka, którą stworzyliśmy, nosi nazwę „cwrite”. Pozwala ona na odczyt i zapis z tablicy znaków, tak jakby była ona obiektem strumienia, dzięki czemu możemy wykorzystać zdolność funkcji „print” do formatowania liczb zmiennoprzecinkowych w celu utworzenia adresu URL.

Wynikowy zestaw danych może być następnie wysłany za jednym zamachem do nakładki 4G.

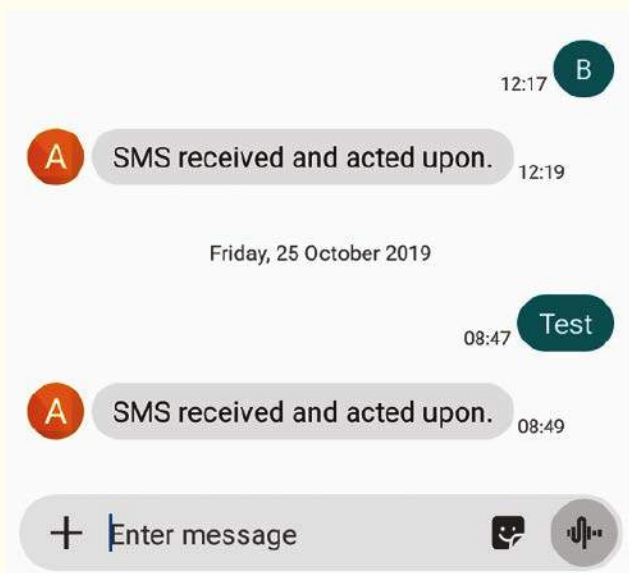
Biblioteka ta może być postrzegana jako dwie dodatkowe zakładki w projekcie Arduino. Jeśli zrobisz kopię projektu (używając File -> Save As...), to ta biblioteka również zostanie skopiowana.

```
COM34 - Tera Term VT
File Edit Setup Control Window Help

+CFUN: 1
+CPIN: READY
SMS Ready
DST: 1
*PSUTTZ: 19/10/24,22:53:33", "+44", 1
Format set
GNSS on
Success:
505 01 ALDImobile
Wait for GNSS
Wait for GNSS
Wait for GNSS
Wait for GNSS
GNSS fix fail
T=31.84C
H=34.66%
P=101088.72Pa

http://api.thingSpeak.com/update?api_key=
35&field3=758.00&field4=651.00&field5=101
200
ThingSpeak Success
No messages received.
Staying Powered On
```

Rysunek 8. Oto kilka przykładowych danych debugowania z portu szeregowego Zdalnej Stacji Monitoringu podczas normalnej pracy. Twoje dane mogą się różnić, jeśli na przykład korzystasz z usług innego operatora GSM.



Rysunek 9. Ponieważ Zdalna Stacja Monitoringu 4G jest zasilana okresowo, może nie odpowiadać na polecenia SMS natychmiast. Dwuminutowe opóźnienie pokazane tutaj wystąpiło podczas testów, kiedy mieliśmy nakładkę ustawioną na wyłączenie zasilania na około jedną minutę. Reakcja może być zmieniona poprzez edycję kodu Arduino.

Ostatnia biblioteka służy do wykonywania rzeczywistych odczytów temperatury, wilgotności i ciśnienia z czujnika BME280.

Jest ona napisana przez firmę SparkFun i może być zainstalowana poprzez narzędzie Library Manager w Arduino IDE. Wyszukaj „sparkfun BME280” w Library Manager i kliknij „Install”.

Dołączyliśmy tę bibliotekę do naszego pakietu oprogramowania dla tego projektu, na wypadek gdybyś nie mógł jej znaleźć.

Oprogramowanie Arduino

Skonfigurowaliśmy oprogramowanie Arduino do pracy z ośmioma polami, które właśnie utworzyliśmy, plus trzy ukryte pola szerokości i długości geograficznej oraz wysokości nad poziomem morza.

Te trzy pola pochodzą z odbiornika GNSS modułu SIM7000E oraz z czujnika ciśnienia atmosferycznego BME280 użytego w celu określenia wysokości n.p.m.

Oprogramowanie jest w postaci modułów, tak że biegli użytkownicy Arduino mogą je zmodyfikować, aby w razie potrzeby wykonywać inne operacje.

Na początek będziesz musiał edytować oprogramowanie, aby wkleić w nie swój klucz API. W okolicach linii 28 zamień tekst API_KEY_HERE na klucz API, który skopiowałeś wcześniej. Powinieneś otrzymać 16-znakowy ciąg alfanumeryczny otoczony cudzysłowami.

Poniżej, w liniach 29 i 30, znajdują się wpisy dotyczące numerów telefonów. Każda przychodząca wiadomość tekstowa ma porównywany

numer z ciągiem AUTH_NUMBER. Ciąg AUTH_NUMBER_HERE należy zastąpić końcowymi cyframi numeru telefonu.

Zrobiliśmy to w ten sposób, aby umożliwić dopasowanie zarówno numerów krajowych, jak i międzynarodowych, niezależnie od sformatowania. Tak więc dla australijskiego numeru telefonu komórkowego pierwszą cyfrą powinno być „4”, co oznacza, że początkowe „0” zostaje usunięte.

Szkie po prostu dopasowuje wszystkie cyfry, które są obecne. Jeśli więc zmieniono by to na „693”, to każdy numer kończący się na „693” zostałby zaakceptowany. Jeśli nie chcesz korzystać z tej funkcji, pozostaw ją jako domyślny ciąg znaków, ponieważ jest bardzo mało prawdopodobne, aby odpowiadał on istniejącemu numerowi telefonu.

Numer wychodzący powinien być w pełni kwalifikowanym międzynarodowym numerem telefonu komórkowego; np. polski numer telefonu komórkowego wraz z numerem kierunkowym kraju zaczynałby się od „+48”, po którym następuje dziewięć cyfr. Jest on używany do wysyłania alertów w postaci wiadomości tekstowych.

Wiele z pozostałych „definicji” jest ustawianych w celu określenia takich parametrów jak współczynniki pomiaru analogowego napięcia wejściowego oraz ilość dostępnej pamięci. Nie ma powodu, aby je zmieniać, chyba że modyfikujesz sprzęt.

Oprogramowanie wykonuje podstawową inicjalizację sprzętu w procedurze setup(). Więcej inicjalizacji ma miejsce w funkcji loop(), szczególnie w przypadku nakładki 4G.

Program wysyła do nakładki pewne dane, aby sprawdzić, czy jest ona zasilana, a jeśli nie, przełącza linię zasilania.

Wysyłany jest zestaw stałych sekwencji, aby wprowadzić nakładkę 4G w znany stan. Nakładka ma dziesięć sekund na ustalenie pozycji GNSS. Jeśli to się uda, sprawdzana jest prędkość urządzenia i wysyłany jest komunikat, jeśli jest ona większa niż 100 km/h. Jest to podstawowa demonstracja tego, jak łatwo można wykonać akcję na podstawie stanu czujnika.

Oprogramowanie wysyła następnie aktualizację do ThingSpeak, w postaci funkcji warunkowej, która sprawdza poprawność danych GNSS i wysyła te dane tylko wtedy, gdy są one prawidłowe.

Następnie Arduino sprawdza, czy nie ma wiadomości tekstowych z autoryzowanego numeru. Jeśli zostanie znaleziona, wysyłana jest standardowa odpowiedź.

Możesz zmodyfikować kod tak, aby sprawdzał zawartość wiadomości i wykonywał różne akcje (i dostarczał różne odpowiedzi) w zależności od otrzymanej treści.

Jeśli dane GNSS są niepoprawne, to zamiast się wyłączyć, Arduino przechodzi w stan uśpienia i pozostawia działającą nakładkę 4G, aby umożliwić jej uzyskanie poprawki. To nie zmniejsza zużycia energii tak bardzo jak wyłączenie Arduino, ale daje modułowi szansę na ustalenie pozycji.

Jeśli dane GNSS są prawidłowe, wtedy styk zasilania modemu jest przełączany (aby wykonać kontrolowane wyłączenie), a styk D7 jest ustawiany w stan wysoki, aby wyłączyć wszystko inne.

Gdy Arduino ponownie się włączy, sekwencja powtarza się, w wyniku czego ThingSpeak jest aktualizowany mniej więcej co 10 minut.

Każda aktualizacja zużywa około 2 kB danych, co według naszego planu komórkowego kosztuje około 0,0001\$. Podczas naszych testów wysłaliśmy około 6000 aktualizacji (ponad miesiąc aktualizacji) za łączną kwotę 1,08\$. Oczywiście Twój plan może się różnić.

Zakończenie

Podłącz płytke Arduino do komputera i wybierz płytkę Leonardo oraz odpowiadający jej port szeregowy z menu Arduino IDE.

Skompiluj i załaduj szkic „4G_Monitoring_Station.ino”. Odłącz Leonardo i podłącz obie nakładki.

Podłącz akumulator do CON2, a źródło zasilania do CON1. Zewrzyj na krótko JP1 i sprawdź czy cały zespół jest zasilany.

Przesyłanie do ThingSpeak powinno trwać mniej niż minutę. Jeśli tak się nie stanie, to być może trzeba będzie wykonać kilka razy debugowanie, aby sprawdzić, co szwankuje.

Podłącz ponownie płytkę Leonardo do komputera i otwórz program terminala szeregowego, aby monitorować wyjście. Powinno ono wyglądać jak jest to pokazane na Rysunku 8. Szukaj kodu „200” HTTP i komunikatu „ThingSpeak Success”.

Jeśli to uzyskasz, to znaczy, że przesyłanie danych do ThingSpeak przebiega prawidłowo.

Możesz zauważyć, że Arduino Serial Monitor nie zachowuje się dobrze, gdy Leonardo włącza się. My odnieśliśmy sukces używając programu o nazwie TeraTerm, ponieważ automatycznie łączy się on ponownie z portem szeregowym, jeśli ten port się rozłączy i ponownie połączy.

Niestety, kabel USB będzie również zasilał Leonardo, więc funkcje wyłączania mogą nie działać zgodnie z oczekiwaniami, gdy moduł Arduino jest podłączony do komputera.

Sztuczka do testowania

Aby przetestować nasz prototyp, potrzebowaliśmy sposobu, aby umożliwić komunikację USB, ale bez zasilania modułu Leonardo z komputera, ponieważ zakładka to pracę nakładki sterowania zasilaniem.

Aby to osiągnąć, użyliśmy jednej z naszych płytek PCB USB Port Protector opisanych w maju 2018 roku (siliconchip.com.au/Article/11065).

Jeśli płytka Port Protector jest okablowana bez żadnych elementów poza wtyczką i gniazdem USB (CON1 i CON2), to łączy GND, D+ i D–, ale nie 5 V.

Tak więc ta przejściówka może być użyta do podłączenia urządzenia USB, aby umożliwić transfer danych, ale nie zasilania. Leonardo jest zasilany napięciem 5 V poprzez wbudowany stabilizator, z napięcia podawanego na końcówkę VIN.

Należy zwrócić uwagę, aby masa komputera nie miała innego potencjału niż masa stacji zdalnego monitoringu 4G; na przykład, jeśli zasilasz ją z zasilacza sieciowego lub podobnego, upewnij się, że wyjścia zasilacza są odizolowane od jakiegokolwiek innego potencjału.

Jeśli nie jesteś tego pewien, możesz użyć do testów komputera zasilanego z własnego akumulatora.

Debugowanie

Chociaż kod jest dość skomplikowany, nie napotkaliśmy z nim wielu problemów.

Ale na wypadek gdyby wystąpiły, prześleliśmy niektóre komunikaty o błędach, które może wyświetlić Arduino.

Jeśli nie widzisz komunikatów „GNSS on” lub „Format set”, Twoje Arduino prawdopodobnie nie komunikuje się z nakładką 4G.

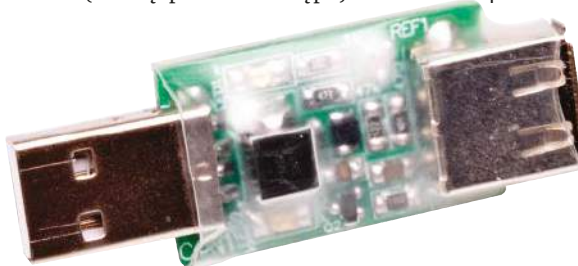
Zgodnie z arkuszem danych nakładki, komunikuje się ona z prędkością 115,200 bodów, ale nasze urządzenie było ustawione na 19,200 bodów. Możesz zmienić to ustawienie w linii 5 kodu Arduino.

Po komunikacie „GNSS on” powinieneś zobaczyć „Success” i nazwę sieci komórkowej.

Jeśli widzisz tutaj „Fail”, to znaczy, że nakładka 4G nie rejestruje się w sieci. Dzieje się tak zazwyczaj, gdy nakładka 4G nie odbiera stacji GSM. Może to być spowodowane tym, że nie obsługuje ona pasma częstotliwości Twojego operatora, lub możesz być poza zasięgiem nadajnika. Sprawdź, czy antena sieci komórkowej jest prawidłowo podłączona.

Możesz czasami zobaczyć „GNSS fix fail”, kiedy nakładce 4G nie udaje się pobrać danych niezbędnych do ustalenia położenia ze względu na ustawienia planu oszczędzania energii.

Oprogramowanie próbuje pobrać APN (nazwę punktu dostępu) z nakładki 4G



Oto USB Port Protector z numeru SC z maja 2018 – podczas testów używaliśmy jednej z przejściówek bez wlutowanych komponentów, poza wtyczką i gniazdem USB. Pozwala to na przesyłanie danych, ale nie zasilania.

i użyć go do powiązania z danymi mobilnymi. Jeśli widzisz komunikat odnoszący się do APN, CSST lub „bearer failing”, to znaczy, że nie jest to skonfigurowane prawidłowo. Sprawdź nazwę APN, której używa Twój dostawca sieci komórkowej.

Powinien zostać wyświetlony adres URL, którego używa Zdalna Stacja Monitoringu 4G, a następnie kod wynikowy http, jak na Rysunku 8.

Jeśli nie jest to „200” (sukces HTTP), sprawdź <https://www.mathworks.com/help/thingspeak/error-codes.html>, aby zobaczyć, co oznaczają inne kody błędów.

Jeśli nadal są problemy, liczne dodatkowe linie debugowania zostały w kodzie skomentowane (poprzez dodanie „//” na początku). Możesz włączyć dodatkowe komunikaty usuwając je, kompilując i przysyłając kod ponownie.

Możesz również spróbować naszego szkicu „Leo_FTDI_with_passthrough.ino”. Konfiguruje on Leonardo, aby umożliwić bezpośrednią komunikację pomiędzy portem szeregowym a nakładką 4G.

Możesz spróbować różnych szybkości transmisji, aby zobaczyć, która z nich działa i wysłać polecenia bezpośrednio do nakładki 4G.

Załaduj skompilowany kod do modułu Leonardo i zewrzyj JP1 na nakładce sterowania zasilaniem.

Być może będziesz musiał nacisnąć przycisk BOOT nakładki 4G, aby włączyć ją ręcznie.

Po potwierdzeniu prawidłowej prędkości transmisji, prześlij „4G_Monitoring_Station.ino” do Leonardo ponownie.

Podsumowanie

Celowo pozostawiliśmy ten projekt jako otwarty; oczekujemy, że czytelnicy dostosują sprzęt i kod do różnych zastosowań.

Do użytku na zewnątrz, zalecamy umieszczenie wszystkiego w plastikowej obudowie klasy IP-65, z obiema antenami zamontowanymi na spodzie pokrywy.

Niektóre otwory wentylacyjne, skierowane w dół, mogą pomóc odprowadzić wszelkie mogące się tworzyć kondensaty pary wodnej, i pozwalają na zewnętrzny dostęp powietrza do próbkowania przez czujnik BME280. ■

Tim Blythman

Adaptacja do wydania polskiego – Andrzej Nowicki

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

Szkoła Konstruktorów



Rozwiązanie zadania głównego 318

Temat wrześnieowego zadania 318 brzmiał: **Zaproponuj nietypowe wykorzystanie multimetru.**

Piotr Grzegorzczak z Siedlec zaproponował: (...)

1. Z multimetru można zrobić... stoper. Wystarczy ustawić pomiar rezystancji na odpowiednio duży zakres i do zacisków miernika wpiąć kondensator. W moim przypadku ustawienie zakresu 2000 k Ω i zastosowanie kondensatora o pojemności 1 mF pozwoliło na zrobienie stopera – co ciekawe z możliwością zatrzymywania i wznawiania. Wadą jest słaba precyzja stopera i w praktyce pozwala ona na mierzenie czasu rzędu 2 minut.

2. Natomiast wpięcie czujnika np. PT100 umożliwia wyznaczenie temperatury w tańszych miernikach. Należy odczytać opór i obliczyć temperaturę ze wzoru albo skorzystać z przygotowanej wcześniej tabelki lub wykresu.

3. Multimetrem można testować również MOSFET-y (zwłaszcza te poziomu-logicznego). Wystarczy ustawić pomiar przewodności. Najpierw należy naładować bramkę testując przewodność między bramką a źródłem. Złącze dren-źródło powinno przewodzić. Krótkie zwarcie bramki do źródła (najlepiej zewrzeć wszystkie nóżki) powinno spowodować zmianę stanu przewodzenia. Ponadto należy sprawdzić czy bramka jest odizolowana od pozostałych elektrod i wewnętrzną diodę zabezpieczającą wpięty zaporowo na wyjściu tranzystora.

Piotr przypomniał trzy bardzo interesujące możliwości. Multimetr istotnie może być stoperem, ale być może warto w tym samym układzie pracy wykorzystać go w roli miernika pojemności dużych kondensatorów. Czas ładowania jest przecież wprost proporcjonalny do pojemności kondensatora.

Jeżeli chodzi o pomiar temperatury, to czujniki rezystancyjne PT100 (a w tym przypadku coraz częściej spotykane Pt1000 i Pt500) mogą zapewnić zdecydowanie większą dokładność, niż dołączane do niektórych multimetrów termopary. Problemem może być, a praktycznie zawsze jest, dokładność pomiaru rezystancji, która na zakresach omomierza często jest gorsza niż $\pm 1\%$. Warto pamiętać, że dokładniejsze termometry o dokładność rzędu 0,1 stopnia a nawet dużo lepszej, można w warunkach amatorskich zrealizować właśnie z użyciem platynowych czujników RTD. Układ elektroniczny nie jest największym kłopotem. Największym jest precyzyjna kalibracja, ale to odrębny i szeroki temat.

Testowanie MOSFET-ów za pomocą multimetru to kolejne bardzo interesujące zastosowanie. Obecnie, w związku z brakami półprzewodników na rynku, elektronika

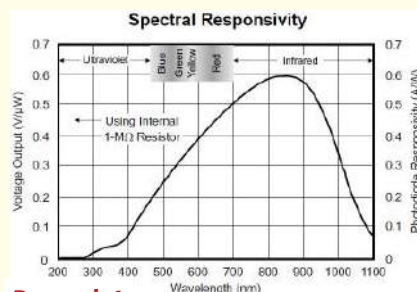
światowa cierpi z powodu fali podróbek. Już napisałem kilka artykułów na temat podróbek, między innymi właśnie MOSFET-ów. Na rynku oferowane są chińskie tanie podróbki. Aby je zidentyfikować, na pewno trzeba zmierzyć rezystancję $R_{DS(on)}$, co można zrealizować w opisany właśnie sposób. Ale to nie wszystko. Oprócz rezystancji, w większości zastosowań ważne są też parametry dynamiczne, w tym kluczowy parametr: ładunek bramki Q_g . To też można zmierzyć, a przynajmniej z dobrym przybliżeniem oszacować za pomocą większości multimetrów, nawet tych tanich (ale nie najtańszych).

Interesująca propozycję przedstawił **Michał Stach** z Kamionki Małej: *Celem instalacji fotowoltaicznej jest pozyskiwanie jak największej ilości energii. Niestety z tym bywa różnie. Posiadacze czują się oszukani kiedy instalacja nie pracuje jak należy. Jednak nie zawsze winna leżeć po stronie instalacji, bowiem w grę wchodzi pogoda, a dokładnie jej parametry określone jako*

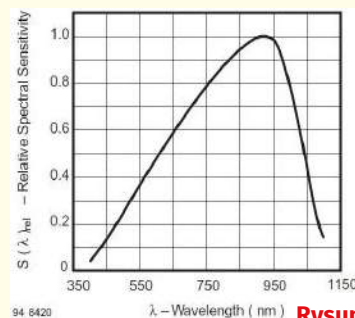
Fotografia 1



Fotografia 2



Rysunek 1



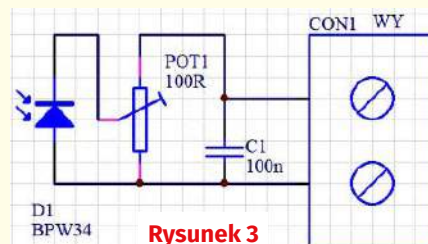
Rysunek 2



Fotografia 3



Fotografia 4



Rysunek 3

uśłonecznienie i nasłonecznienie. To drugie (nasłonecznienie) mierzy się w W/m^2 przy pomocy pyranometru. Przyrządu o dość abstrakcyjnej nazwie, budowie i wysokiej cenie. Na szczęście zwykły multimetr może stać się w pełni użytecznym pyranometrem po zastosowaniu prostego adaptera.

Przykładowy i profesjonalny pyranometr pokazany jest na fotografii 1. Tak naprawdę mierzy on ciepło dostarczane/pozyskiwane przez ciało idealnie czarne. Nie tylko bezpośrednio ze Słońca ale również rozproszone i odbite. (...) widmo słoneczne wykracza poza zakres widzialny zarówno w paśmie UV jak i podczerwieni. Zatem czy pomiar luksomierzem ma jakąś wartość? Fotografia 2 pokazuje że tak, ale tylko orientacyjny i przy dobrej pogodzie. Przybliżone wartości zawiera znaleziona w Internecie tabelka:

Intensity in different conditions	
Illuminance	Example
120,000 lux	Brightest sunlight
111,000 lux	Bright sunlight
109,870 lux	AM 1.5 global solar spectrum sunlight
20,000 lux	Shade illuminated by entire clear blue sky, midday

Dlatego do przystawki pyranometrycznej musimy poszukać elementu fotoczulęgo z szerszym pasmem, najlepiej obejmującym zakres promieniowania emitowanego przez Słońce. Pierwszym elementem który przychodzi na myśl jest OPT101. Zakres przetwarzania OPT101 pokazany jest na rysunku 1. W zadowalającym zakresie pokrywa się on z emisją słoneczną. Ten optymistyczny fakt zachęcił do rozpoczęcia prób właśnie od OPT101. Niestety wzmacniacz transimpedancyjny OPT101 zdecydowanie lepiej radzi sobie z minimalnym oświetleniem niż z ekspozycją wprost na Słońce. Testy pokazały

że rezystor Rext ustalający wzmacniacza OPT101 musiałby mieć wartość poniżej 10 k Ω , a już przy 13 k Ω OPT101 wpada w oscylacje. Zakładanie (optycznych) filtrów szarych w celu stłumienia światła nie wchodzi w grę, gdyż ich pasmo z reguły zawężone jest do zakresu widzialnego. W trakcie poszukiwania komponentu przewijały się: PT202C, RPM075, LTR323, BPW21R aż w końcu rozwiązaniem okazała się fotodioda BPW34. Jej zakres czułości przedstawia rysunek 2.

Budowę przystawki pyranometrycznej przedstawia fotografia 3, a jej schemat pokazany jest na rysunku 3. Montaż rozpoczynamy od przyklejenia diody BPW34 do ścianki obudowy. Następnie przyklejamy osłonę od warunków atmosferycznych. Wystarczająco dobra okazuje się bańka pozyskana z żarówki/żarówki LED. Współczynnik konwersji BPW34 to 47 μA przy 1 mW/cm^2 (co odpowiada 10 W/m^2). Zatem aby odczytów nasłonecznienia dokonywać multimetrem na zakresie 199,9 mV (wtedy 100,0 mV odpowiada 1000 W/m^2) potencjometr P1 należy ustawić około 25 Ω .

Jeżeli chodzi o jeszcze inne przykłady nietypowego wykorzystania multimetru, ja zainteresowałem się niedawno możliwościami pomiaru z ich pomocą skrajnie małych i skrajnie dużych rezystancji. Zwykle zakres pomiaru rezystancji w popularnych

multimetrach rozciąga się od 0,1 Ω do 20 lub 200 megaomów, a w bardzo niewielu zakres ten jest szerszy. Tymczasem każdy multimetr pozwala w bardzo prosty sposób mierzyć małe oporności, co najmniej od 0,1 milioma (0,0001 Ω), ale nie na zakresie omomierza. Wykorzystałem to w mikroomomierzu, który może mierzyć rezystancje już od 1 mikrooma! Fotografia 4 pokazuje pomiar centymetrowego kawałka drutu instalacyjnego o przekroju 2,5 mm^2 , którego rezystancja między punktami pomiaru wynosi 92 mikroomy, czyli 0,000092 Ω . Dolna granica i rozdzielczość takich pomiarów to 1 mikroom!

Multimetrami można też stosunkowo łatwo mierzyć dużo większe oporności, na pewno powyżej 1 gigaoma. Już przeprowadziłem wstępne próby z pomiarami rezystancji rzędu gigaomów, co prawdopodobnie umożliwi też pomiar rezystancji o wartości nawet teraomów. Przy bardzo dużych rezystancjach pomiary są dużo trudniejsze i w grę wchodzi dodatkowe czynniki i poważne problemy. Także kwestie właściwości używanego do takich pomiarów multimetru. W związku z tym kupiłem dwa dodatkowe multimetry i przetestowałem wszystkie, które mam w swojej pracowni, których większość pokazana jest na fotografii 5. Zachęcam do zainteresowania się kwestią pomiarów takich skrajnych rezystancji! ■

Fotografia 5





Rozwiązanie zadania głównego 319

Temat październikowego zadania 319 brzmiał: **Zaproponuj układ elektroniczny związany z recyklingiem, ponownym wykorzystaniem odpadów lub przedmiotów zużytych i zbędnych.**

Do czasu oddania materiałów do druku dotarło do mnie tylko jedno rozwiązanie. **Michał Stach** z Kamionki Małej napisał:

Oświetlenie LED-owe staje się standardem. Niestety oszczędności wprowadzane przez producentów skracają żywotność tych źródeł światła. Czasem niedługo po zakupie stajemy się posiadaczami złomu. Nie ma

przesłanek do zmiany trendu w najbliższych latach więc naprawa i usunięcie kluczowych wad konstrukcyjnych ma uzasadnienie ekonomiczne i jest działaniem proekologicznym. Podstawowym problemem jest ciepło. Nie tyle jego obecność (utrata sprawności LED) co kłopoty z odprowadzaniem. Struktura diody LED jest delikatna; jak pokazują karty katalogowe szanujących się producentów, zalecana temperatura pracy leży poniżej +100°C. W praktyce oznacza to stosowanie dużych radiatorów, a tymczasem w strukturach COB diody ściśnięte są na małej powierzchni. Kończy się to tak, jak pokazuje **fotografia 1**. Uszkodzenie jednej z diod wyłącza cały sznurek diod. Z reguły jest ich kilkanaście, więc na chwilę pozostałe przejmują prąd przewidziany dla całości. Jednak po wypadnięciu jednego wiersza usterka kolejnych postępuje lawinowo. Naprawa polega na wymianie diody, a czasem również sterowania. **Fotografia 2** przedstawia klasyczną „żarówkę LED” z gwintem E27 po wklejeniu nowej diody. Operacja nie wymagała wymiany zasilacza. Zastosowana dioda – Citizen charakteryzuje się bardzo dobrym współczynnikiem oddawania barw: CRI=97. Za jednym posunięciem udało się przywrócić

sprawność „żarówce” i zamienić trupi odcień na barwę Słońca w czerwcowe popołudnie. **Fotografie 3 i 4** przedstawiają kolejne etapy regeneracji halogena LED. W tym przypadku odzyskana została obudowa; nie było jaka, bo ciśnieniowy odlew aluminium z dodatkową komorą na zasilacz. Ten zamontowany przez producenta uległ zwarciu. Zalany w całości żywicą nie przedstawiał żadnej wartości. Sama dioda COB przetrwała, jednak jej barwa: 11000 K była nie do zaakceptowania. Z pomocą przyszła dioda CLU036-120840H7 generująca 2500 lm przyjemnego światła o barwie 4000 K wystęrowana z regulowanego źródła prądowego, którego schemat pokazany jest na **rysunku 1**. Regulację jasności zapewnia potencjometr cyfrowy XR9511 włączony w obwód sprzężenia zwrotnego TL431. Całość zasilają zasilacz MeanWell do zastosowań PoE (48 VDC). Nowy – zregenerowany halogen



Fotografia 1



Fotografia 2

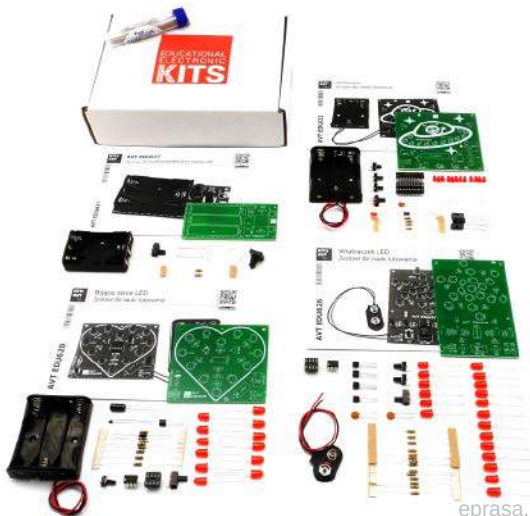


Fotografia 3



Fotografia 4

REKLAMA



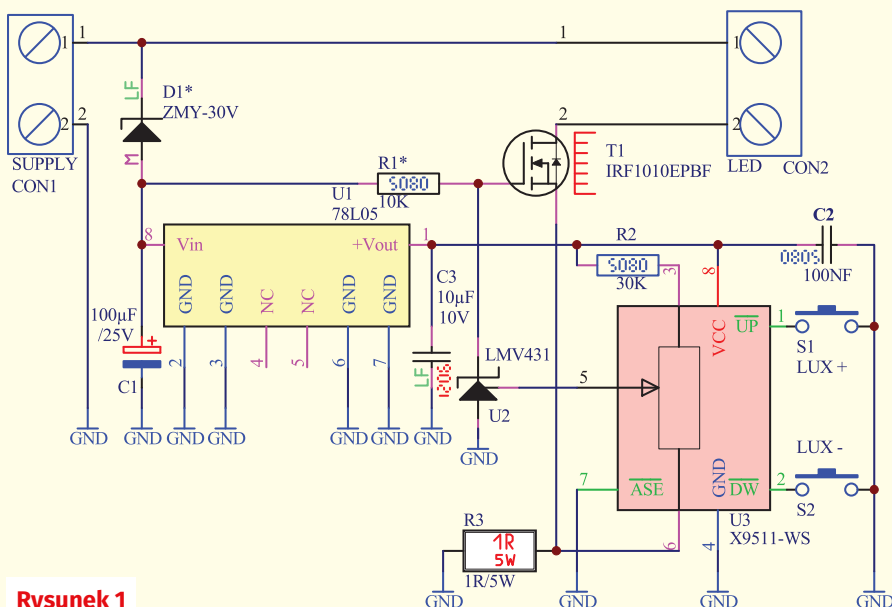
AVTEDU4PAKIET

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



Rysunek 1

Imię i nazwisko	Miejscowość	Punkty	Talon AVT [PLN]
Jarosław Węgliński	Warszawa	6	200
Rafał Orodziński	Białystok	8	300
Jacek Konieczny	Poznań	3	100
Andrzej Nowicki	Warszawa	6	100
Zygmunt Flisak	Opole	6	200
Piotr Grzegorzcyk	Siedlce	4	100
Michał Stach	Kamionka Mała	6	350

tak spodobał się właścicielowi, że dostarczył drugi na modernizację.

Opisany przykład pokazuje, że w wielu przypadkach zakup tanich „żarówek LED” wcale nie jest wyrazem oszczędności, a wprost przeciwnie – naraża na dodatkowe koszty. Zagadnienie jest skomplikowane, a w wielu sytuacjach, zamiast samodzielnej modernizacji, lepszym rozwiązaniem jest zakup droższych i znacząco lepszych wersji. Problem w tym, jak rozpoznać te lepsze wersje? Czy tylko po wyższej cenie? To odrębny, szeroki i niełatwy temat

Zygmunt Flisak z Opola napisał: *Dzień dobry. Wdrażany obecnie model gospodarki obiegu zamkniętego polega na maksymalnym wydłużeniu czasu życia produktów, zachowaniu wartości surowców i materiałów oraz ograniczeniu ilości wytwarzanych odpadów. W elektronice jednym ze sposobów realizacji tego modelu mogłoby być ponowne*

wykorzystanie elementów pochodzących z demontażu starych urządzeń.

Nadawanie drugiego życia łatwo dostępnym w handlu i tanim elementom biernym, takim jak rezystory małej mocy czy kondensatory elektrolityczne, nie jest optymalne. Jednak warto czasem zastanowić się nad możliwością pozyskania niektórych elementów czynnych, np. lamp elektronowych, nieprodukowanych już układów scalonych i tranzystorów germanowych. Te ostatnie (**fotografia 5**) mogą bowiem posłużyć do budowy wiernych replik klasycznych efektów gitarowych lub rozmaitych wariacji układu typu Joule thief.

Zastąpienie tranzystora krzemowego jego germanowym odpowiednikiem prowadzi do obniżenia minimalnego napięcia wymaganego do startu takiej przetwornicy. W moim przypadku zamiana tranzystora BC547C na ASY29 (też NPN!) doprowadziła



Fotografia 5

do tego, że dioda LED zaczynała świecić już przy napięciu wejściowym wynoszącym 200 mV. Oczywiście nic nie stoi na przeszkodzie, aby po zamianie biegunów zasilania i polaryzacji diody zastosować bardziej rozpowszechnione tranzystory germanowe PNP. Należy jednak mieć świadomość, że współczesne tranzystory polowe i specjalizowane układy scalone (np. LTC3108) oferują jeszcze lepsze możliwości w tym zakresie.

Aktualne informacje nagród za zadania o numerach 316...319 podane są w tabelce. Osoby nagrodzone kuponami otrzymują z AVT stosowny e-mail z informacją i wskazówkami, a dopiero potem zamawiają w sklepie AVT (wrzucając do koszyka pod adresem www.sklep.avt.pl) towary za przydzieloną sumę, a w uwagach pisząc, że jest to talon ze Szkoły Konstruktorów. Talonów za zadania **nie można** sumować. Istnieje możliwość dopłaty w przypadku zamówienia o większej wartości niż przydzielony talon. Ale uwaga: **talon ważny jest tylko 30 dni**. Po tym terminie traci ważność i przepada.

Zadanie 319 było ostatnim zadaniem Szkoły. Chciałbym serdecznie podziękować wszystkim liczным uczestnikom Szkoły Konstruktorów, którzy brali w niej udział od 1996 roku oraz sympatykom, którzy śledzili zadania i ich rozwiązania. A tych, którzy nadal chcą obserwować moje działania, zapraszam na nową stronę: <https://piotr-gorecki.pl>. ■

Piotr Górecki

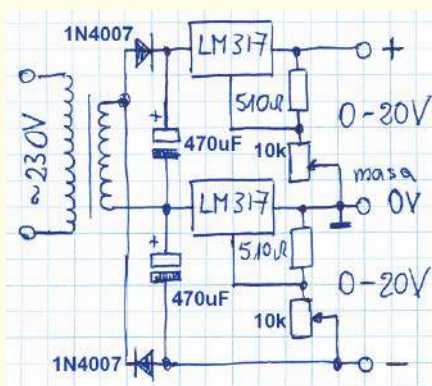
Co tu nie gra? Rozwiązanie zadania 318

Na **rysunku A** pokazany jest zamieszczony w EdW 9/2022 schemat symetrycznego zasilacza uniwersalnego $\pm(0...20\text{ V})$, który chcemy wykorzystywać między innymi

do testów modułów audio wymagających symetrycznego zasilania.

Zadanie było z jednej strony dość łatwe, ponieważ bez problemu można było wskazać wymagany

jeden błąd lub usterkę. Z drugiej strony było trudne, bowiem pełna analiza tego, tylko na pozór prostego układu, wcale nie jest taka łatwa, jak mogłoby się w pierwszej chwili wydawać.



Rysunek A

Schemat z rysunku A to nie jest wcale teoretyczny wymysł. Hobbyści nadal chętnie realizują różnego rodzaju zasilacze i stabilizatory. I bardzo słusznie! Tylko często brak wiedzy i doświadczenia prowadzi do „przedobrzania” i „przekombinowania”. I właśnie na pozór sprytny i oszczędny schemat z rysunku A jest przykładem „przekombinowania” i braku wiedzy. Tak tak to zwykle jest z zadaniami NieGra, to ja narysowałem schemat i celowo podałem takie a nie inne wartości elementów, jednak inspiracją tego zadania były kłopoty pewnego młodego elektronika, który podobny zasilacz zrealizował i dziwił się, dlaczego nie działa on według oczekiwań.

Szczegółowa analiza wszystkich aspektów byłaby zbyt obszerna i wymagałaby kilku artykułów. Spróbujmy omówić, a przynajmniej zasygnalizować najważniejsze.

Może na początek przytoczę zwięzły i zgrabny opis problemów, nadesłany przez jednego ze stałych uczestników: (...) Pierwszym mankamentem tego zasilacza jest próba uzyskania napięć symetrycznych z transformatora sieciowego bez odczepu za pomocą podwajacza napięcia z prostownikiem jednopółkowym. Owszem, takie rozwiązanie jest poprawne, jednak cechuje się istotnymi wadami. W porównaniu z prostownikiem dwupółkowym należy się liczyć z większym napięciem tętnień i mniejszą ich częstotliwością, a także mniejszą wartością średnią wyprostowanego napięcia. Wymienione niedoskonałości stawiają pod znakiem zapytania przydatność takiego zasilacza do aplikacji audio. Można próbować ograniczyć ich wpływ odpowiednio dużymi kondensatorami filtrującymi. Nie określono wprowadzie oczekiwanego poboru prądu z zasilacza, jednak 470 μF wydaje się tutaj za mało. Niezależnie od tego, transformator z odczepem (lub dwoma identycznymi uzwojeniami wtórnymi połączonymi szeregowo) i mostek Graetza to zdecydowanie lepszy wybór.

Drugi problem polega na wykorzystaniu układu LM317 do stabilizacji napięcia

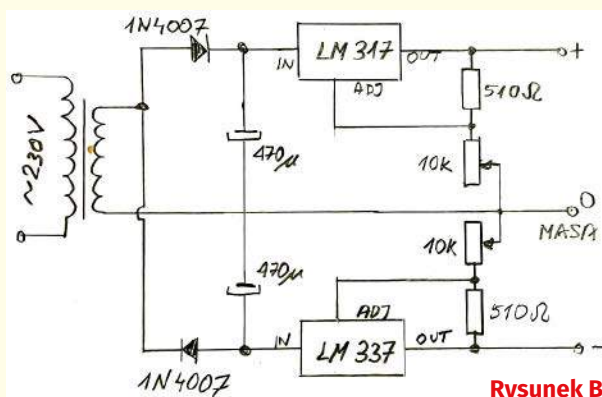
ujemnego. Zasilacz przedstawiony na schemacie faktycznie generuje dwa regulowane napięcia dodatnie. W tej konfiguracji szyna oznaczona minusem stanowi napięcie 0 V, względem którego pierwszy regulator wytwarza stabilizowane napięcie dodatnie dostępne na zacisku opisanym jako „masa 0 V”. Z kolei na wyprowadzeniu oznaczonym plusem uzyskamy również stabilizowane napięcie dodatnie będące sumą napięć wyjściowych obydwóch układów LM317 względem 0 V. Aby konstrukcja pełniła rolę zasilacza symetrycznego, pierwszy z układów scalonych [dolny na schemacie] powinien być regulatorem napięcia ujemnego (np. LM337) w aplikacji zgodnej z kartą katalogową, tj. z wejściem podłączonym do diody prostowniczej, rezystorami dzielnika zamienionymi kolejnością i wyjściem podłączonym do zacisku oznaczonego minusem.

Odnosnie wspomnianego dzielnika napięcia: wartości przedstawione na schemacie umożliwiają regulację napięcia w zakresie od 1,25 do 25,76 V, a nie od 0 do 20 V. Karta katalogowa LM317 przedstawia sposób wytwarzania napięć niższych od 1,25 V.

Ponadto charakterystykę impulsową oraz impedancję wyjściową zasilacza może poprawić kondensator wyjściowy o pojemności rzędu 1 μF. Warto też dodać kondensator wejściowy o niewielkiej pojemności, np. 100 nF oraz kolejny kondensator w obwodzie wyprowadzenia ADJ regulatora wraz z diodą zabezpieczającą celem ograniczenia tętnień.

Należy również mieć świadomość znacznych strat mocy na układzie LM317 dla niskich napięć wyjściowych. Napięcie wejściowe układu LM317 powinno być wyższe o ok. 3 V od maksymalnego napięcia wyjściowego, zatem w rozpatrywanym układzie powinno wynosić co najmniej 23 V. Jeżeli na wyjściu ustawimy 1,25 V, to przy założonym prądzie wyjściowym 1 A straty mocy przekroczą 20 W. Pozdrawiam

Tak wstępnie poprawiona, ale nadal niedoskonała wersja,



Rysunek B

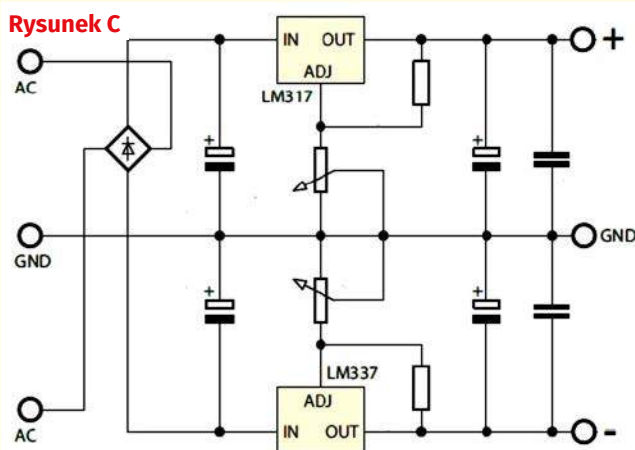
nadesłana przez kolejnego ze stałych uczestników, pokazana jest na **rysunku B**. Usunięty jest tylko podstawowy problem, ale pozostało wiele wad wersji pierwotnej. Pisali o nich inni uczestnicy konkursu. Jeden z nich przysłał schemat pokazany na **rysunku C**. Oto uwagi dotyczące prostownika i głównego filtra tętnień sieci:

(...) Trzeba zwrócić uwagę na to, że każdy z LM317 zasilany jest „jednopółkowo”. Górna część przewodzi „dodatnią” półokwę sinusoidy a dolna „ujemną” półokwę. To wymusza zastosowanie większej pojemności niż 470 μF celem ograniczenia wartości tętnień napięcia. To determinuje zastosowanie układu dla małych obciążeń.

(...) Nie jest to najlepsze rozwiązanie, ale możliwe. Lepiej zastosować inne rozwiązania z prostownikiem dwupółkowym i transformatorem z uzwojeniem dzielonym.

(...) W układzie źle wykonana jest masa i zasilanie układu. Należałoby dać transformator z napięciami symetrycznymi i mostek prostowniczy.

(...) Nie jest podany prąd wyjściowy układu, ale teoretycznie będzie on ograniczony przez wbudowane w LM317 układy zabezpieczeń do 1,5 A. To z kolei oznacza, że źle dobrano diody prostownicze (1 A). Jeszcze bardziej wątpliwym wydaje się także użycie prostowników 1-półkowych i relatywnie małych (jak na zastosowania audio) pojemności



Rysunek C



wejściowych kondensatorów filtrujących. Spowoduje to, przy większych prądach wyjściowych, powstawanie dużych tętnień napięcia wejściowego.

(...) można wspomnieć że taki typ prostownika sprawdza się tylko przy niewielkich mocach i symetrycznym obciążeniu. Przy większych prądach zaczynamy mieć problem z filtracją, a przy niesymetrycznym obciążeniu – z podmagnesowywaniem rdzenia transformatora składową stałą.

A oto nadesłane przez uczestników uwagi dotyczące innych kondensatorów, których na schemacie brak:

(...) Brakuje kondensatorów filtrujących (typowo 100 nF) umieszczonych blisko wejść i wyjść układów LM317 (szczególnie ważny jest ten wyjściowy, który istotnie poprawia stabilność ich pracy, choć niektóre noty podają, że nie jest wymagany). Typowo, noty katalogowe zalecają też użycie kondensatora 10 μ F równolegle podłączonego do potencjometru. Brakuje też kondensatorów elektrolitycznych na wyjściach układu, co jest szczególnie ważne w, proponowanej w tekście, aplikacji audio.

(...) Potencjometri z reguły bocznikujemy kondensatorem C_{Adj} 10 μ F celem stłumienia tętnień napięcia jak podają noty aplikacyjne. To narzuca stosowanie diod zabezpieczających układy LM. Jednej diody od rozładowania C_{Adj} i drugiej od C_{wy} . Elementów tych brak na rysunku.

(...) Brak na wejściu i wyjściu kondensatorów o wartości 100 nF, a na wyjściu każdej „połówki” brak kondensatorów elektrolitycznych np. kilka μ F, brak diod zabezpieczających przed odwrotną polaryzacją.

A teraz uwagi dotyczące zakresu regulacji napięcia:

(...) Nie jest możliwe uzyskanie regulacji napięcia od 0 V w tej konfiguracji jak podano na rysunku.

(...) Dobór rezystora i potencjometru również nie najlepszy. Wystarczy rezystor 121 Ω i potencjometr 2 k Ω dla uzyskania górnej wartości napięcia 20 V. Ze względu na prawidłową pracę układów LM rezystor 510 Ω jest za duży.

(...) w takim układzie stabilizator może stabilizować napięcie począwszy od 1,25 V. Stabilizacja od „zera” jest możliwa,

ale przy odpowiednim spolaryzowaniu wejścia ADJUST (patrz karta katalogowa)

(...) Rezystor 510 Ω [trzeba] zmienić na 680 Ω , to na wyjściu uzyskamy max. 20,1 V, ale napięcie minimalne będzie 1,25 V dla LM3x7. Można zastosować jeszcze diody przy LM.

Te uwagi nie wyczerpują problemu regulacji napięcia. Pozostaje kwestia minimalnego prądu obciążenia LM317/337, o czym jak widać powyżej, wspomniał tylko jeden z uczestników. Przegapienie tego szczegółu spowoduje, że przy braku zewnętrznego obciążenia napięcie wyjściowe stabilizatora będzie wyższe, niż wynikające z wartości współpracujących rezystancji.

No i wreszcie nadesłane uwagi dotyczące podstawowego problemu:

(...) Największą wadą układu jest to, że „dolny” LM317 miałby przewodzić prąd w obu kierunkach. Dla dodatniej połowy przebiegu obwód zamyka się przez dolny LM317. LM317 nie toleruje napięć wstecznych.

(...) w układzie prąd „dolnego” zasilacza musiałby być zawsze większy od „górnego” w przeciwnym przypadku „górny” w ogóle nie mógłby pracować, a na dolnym próbowałibyśmy wymuszać przepływ prądu z wyjścia na wejście stabilizatora, co jest niedopuszczalne i może skutkować jego uszkodzeniem. By się przed tym zabezpieczyć dodajmy diodę „antyrównoległą” do dolnego stabilizatora.

(...) Użyty układ LM317 jest zaprojektowany do stabilizacji napięcia dodatniego, a więc jest w stanie wydawać prąd, a nie pobierać (wewnętrzny tranzystor NPN w układzie Darlingtona). Dlatego prawdopodobnie autor schematu umieścił drugi układ w gałęzi „0 V”, ale w związku z tym przepływ prądu przez gałąź „0 V” będzie możliwy tylko „od układu do wyjścia”. Układ ma szansę pracować tylko w przypadku gdy prądy gałęzi + i – będą równe albo gdy prąd gałęzi – będzie większy niż +. Przy próbie pobierania prądu tylko z gałęzi + układ nie zadziała zgodnie z założeniami (brak drogi powrotu prądu). Przeznaczenie układu do testów układów, a także niezależna regulacja obu torów, w zasadzie dyskwalifikują tę konfigurację. Wystarczyło użyć układu LM337 w gałęzi ujemnej.

Podstawowa kwestia jest dość prosta: w zasilaczu symetrycznym, zwłaszcza z niezależnie

regulowanymi napięciami, prądy „dodatni” i „ujemny” nie muszą być równe. W omawianym układzie problem powstaje, gdy prąd „dodatni” miałby być większy od „ujemnego”. Wtedy środkowy zacisk, oznaczony MASA 0V musiałby przyjąć, pochłoniąc (ang. sink) różnicę prądu „dodatniego i „ujemnego”. A „dolny” stabilizator LM317 nie daje takiej możliwości. Właśnie dlatego produkowane są też stabilizatory „ujemne”, jak choćby właśnie LM337.

Ten podstawowy wniosek jest dość prosty, ale w grę wchodzi też inne szczegóły, które zaciemniają sytuację. Nie będziemy tego szczegółowo analizować, ale najbardziej dociekliwi mogą samodzielnie zastanowić się, czy prawdziwe są poniższe stwierdzenie, pochodzące z dwóch nadesłanych rozwiązań:

(...) Aby taki układ działał, trzeba go zasilić dwoma oddzielnymi napięciami z dwóch oddzielnych uzwojeń i prostowników.

(...) W układzie jak na rysunku A stabilizatory „przeszkadzają” sobie wzajemnie – napięcie i prąd jednego jest głównym zakłóceniem dla drugiego („interakcja” lub „zakłócenia skrośne”). Wykorzystanie obu, niezależnie od siebie w pełnym zakresie regulacji (tzn. 1,25...20 V) jest w tym układzie oczywiście niemożliwe, ale można znaleźć takie wartości napięć i prądów dla których stabilizatory – o dziwo! – jednak będą pracować.

Na koniec jeszcze jeden błąd, który celowo umieściłem na rysunku A, nieprzypadkowo nie rysując i nie podłączając trzeciej nóżki potencjometrów. To na pozór nieznaczący drobiazg, ale ważny w praktyce. I tylko jeden z uczestników napisał o tym bezpośrednio:

(...) Teoretycznie brak połączenia trzeciej nóżki potencjometru z suwakiem nie jest błędem, ale w praktyce zawsze zaleca się dokonanie tego połączenia. Jest to szczególnie ważne w momencie utraty styku. W tym przypadku rezystancja potencjometru w układzie wzrośnie tylko do rezystancji nominalnej potencjometru i nie powstaje przerwa w obwodzie (jak w przypadku braku wspomnianego połączenia). Przykładowo, w układach regulacyjnych audio powoduje to zdecydowanie mniej słyszalne trzaski i zmniejsza prawdopodobieństwo destabilizacji układu.

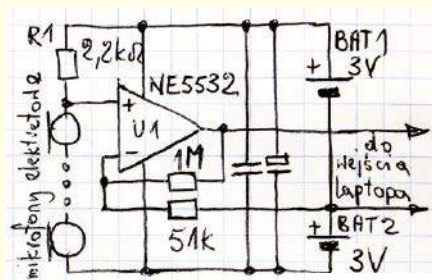
Tyle o błędach z zadania NieGra318. ■

Co tu nie gra? Rozwiązanie zadania 319

W EdW 10/2022 zamieszczony był schemat do zadania NieGra319, które brzmiało: W zapasach nzebierało się nam kilkanaście

mikrofonów elektretowych o różnego typu urządzeń. Teraz zachęteni informacjami z rozwiązania zadania NieGra314 na temat

zestawów mikrofonowych i możliwości kształtowania charakterystyki kierunkowej (beamforming) chcemy przeprowadzić



Rysunek A

eksperymenty z zestawem mikrofonów. Chcemy zastosować oszczędne rozwiązanie przedwzmacniacza z bezpośrednim sumowaniem sygnałów z mikrofonów według rysunku A.

Zadanie nie było trudne, ponieważ bez problemu można było wskazać jeden z kilku błędów zawartych na schemacie. Schemacie, który jak zwykle sam przygotowałem na potrzeby tego zadania.

Uczestnicy słusznie zwrócili uwagę głównie na mikrofony. Oto fragmenty dwóch rozwiązań:

(...) W układzie nieprawidłowo podłączone są mikrofony. Są to mikrofony elektretowe posiadające w środku tranzystor JFET, który musi być zasilany, więcłączenie szeregowo nie wchodzi w grę. W związku z tym, że tranzystory stosowane w mikrofonach mają różne parametry, nie można zastosować wspólnego rezystora podłączonego do zasilania, ale każdy mikrofon należy podłączyć oddzielnym rezystorem do plusa, a do wejścia wzmacniacza podłączyć poprzez kondensator.

(...) Jeżeli mikrofon elektretowy to kapsuła ze zintegrowanym przedwzmacniaczem na tranzystorze JFET, na co wskazuje rezystor polaryzujący R1, to należałoby połączyć mikrofony równolegle, a nie szeregowo. Taki tranzystor działa w przybliżeniu jak źródło prądowe i w przypadku połączenia równoległego prądy będą się sumować na R1. Wartość R1 przy równoległym łączeniu wielu mikrofonów trzeba by zmniejszyć.

Ponieważ według zadania chodzi o mikrofony elektretowe „od różnego typu urządzeń”, należałoby się zastanowić, czy kilka, a tym bardziej kilkanaście mikrofonów o znacząco różnych parametrach, może być połączonych równolegle i zasilanych przez jeden wspólny rezystor. Może się bowiem okazać, że któreś tranzystory FET wymuszają na mikrofonach niskie napięcie stałe, które będzie niekorzystne dla innych mikrofonów. To szeroka kwestia.

Należałoby też sprawdzić dokładniej, jakie właściwości statyczne i dynamiczne mają mikrofony elektretowe. Czy także dla

sygnałów akustycznych są one dobrymi źródłami prądowymi o dużej rezystancji dynamicznej? A może w zakresie częstotliwości akustycznych impedancja dynamiczna jest mała? Jak skuteczność mikrofonu zależy od panującego na nim napięcia stałego? Czy nie należałoby zasilać ich za pomocą źródła prądowego albo zestawu źródeł prądowych? To są trudne kwestie, bardzo słabo znane i słabo rozumiane. Sygnalizuję je, ale szczególnie zdecydowanie wykraczają poza ramy zadania NieGra319.

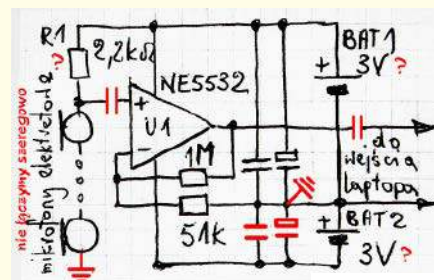
W każdym razie **szeregowe połączenie mikrofonów elektretowych jest ewidentnym błędem.**

W układzie jest też kolejny dyskwalifikujący błąd. Otóż pracujący w konfiguracji nieodwracającej wzmacniacz operacyjny ma wzmocnienie około 20x, ale co ważne, wzmacnia nie tylko napięcia zmienne, ale też i stałe. Potencjałem odniesienia jest masa – punkt połączenia dwóch baterii zasilających. A jakie będzie napięcie na wejściu nieodwracającym („plusowym”) wzmacniacza operacyjnego?

Tego nie wiadomo. Zależać to będzie od parametrów mikrofonów. Na pewno nie będzie równe potencjałowi masy, więc wzmacniacz operacyjny zapewne zostanie nasycony. Jeden z uczestników ujął to tak: (...) Brak odcięcia składowej stałej przez kondensator na wejściu nieodwracającym wzmacniacza. W rezultacie wzmacniacz usiłując około 20-krotnie wzmocnić napięcie stałe, najprawdopodobniej nasyci się w pobliżu jednego albo drugiego napięcia zasilania. Którego – nie jestem w stanie powiedzieć, bo mikrofony nie są dołączone do masy. (...) Być może nie jest to błąd, ale bez kondensatora odcinającego składową stałą na wyjściu, nawet przy zasilaniu symetrycznym, pojawi się wzmocnione około 20 razy napięcie niezrównoważenia według karty katalogowej max. 5 mV*20=0,1 V.

Zwróciliście też uwagę na kolejny problem. Jeden z uczestników słusznie napisał: (...) Karta katalogowa układu NE5532 i strona internetowa firmy TI podają minimalne całkowite napięcie zasilania 10 V, podczas gdy na schemacie jest 6 V. Podejrzewam, że wzmacniacze NE5532 mogą działać przy napięciu mniejszym od 10 V, ale 6 V to może już być za mało.

Co prawda w karcie katalogowej Philipsa (Signetics) można znaleźć informację, że najniższe napięcie zasilania to ±3 V, co jeden z uczestników skomentował tak: (...) Układ NE5532 jak wynika z danych katalogowych może być zasilany od ±3 do ±20 V, więc zasilanie go minimalnym napięciem z baterii 3 V może być za słabe, dlatego proponuję



Rysunek B

źródło napięcia stabilizowanego 3 V lub wstawić stabilizatory na 3 V, a całość zasilić np. baterią 4,5 V.

Nie ulega wątpliwości, że jest problem z niskim napięciem zasilania. A jeżeli chodzi o zasilanie, to jest też problem z kondensatorami. Oto fragmenty rozwiązań:

(...) Użycie dwóch kondensatorów odprężających to dobry pomysł, należałoby jedną użyć dwóch par: jednej pomiędzy + zasilania i masą, a drugiej pomiędzy minusem a masą.

(...) W zasilaniu symetrycznym kondensatory filtrujące na wejściach zasilania należy podłączyć jeden zestaw pomiędzy +BAT1 i masą (-BAT1 i +BAT2), a drugi pomiędzy -BAT2 i masą

(...) O ile wartość wzmocnienia, określona ilorazem rezystancji, równa 20 może wystarczyć do wzmocnienia sygnału do wymaganego poziomu, to już w tym przypadku zastosowanie tak dużej rezystancji 1 MΩ nie jest dobrym wyborem. Należy dobrać rezystory poniżej 1 MΩ.

Wspomniana w ostatnim fragmencie kwestia wzmocnienia równego 20x przy niskim napięciu zasilania i jakimś poziomie szumów, wymagałaby szerszej analizy i dyskusji. Intuicja podpowiada, że także i z innych względów dobrze byłoby zastosować wyższe napięcie zasilania, żeby silniejsze sygnały nie przesterowały wzmacniacza operacyjnego.

Ogólnie biorąc, zadanie dotyczyło eksperymentów z mikrofonami elektretowymi. Można stanowczo stwierdzić, że schemat z rysunku A do tego się nie nadaje. Na pewno napięcie zasilania powinno być większe, trzeba zmodyfikować obwody mikrofonów, a także układ pracy wzmacniacza. Pozostaje też bardzo ważna w takich zastosowaniach kwestia szumów. Należałoby się zastanowić, czy poziom szumów całkowitych określają szумы własne wzmacniacza operacyjnego, czy szумы własne tranzystorów JFET w mikrofonach, czy może same kapsułki elektretowe. To kolejny bardzo obszerny i trudny temat. Kwestia poprawy schematu zdecydowanie wykracza jednak poza ramy zadania NieGra319. Nadesłany **rysunek B** zawiera



proponując poprawy tylko niektórych błędów, natomiast wprowadza problem z szeregowym kondensatorem wejściowym. Oto lista uczestników nagrodzonych upominkami:

Marcin Kartowicz – Bolechowo,
Marek Smolarczyk – Częstochowa,

Paweł Pawłowski – Poznań,
Andrzej Kubiak – Rumia,
Zdzisław Landowski – Gdańsk,
Marcin Witkowski – Bełchatów,
Jerzy Kornaszewski – Radom,
Wojtek Kwedło – Białystok.

Ponieważ są to ostatnie zadania konkursu Co tu nie gra?, jest to zbiorcza lista osób nagrodzonych upominkami z omówionych czterech zadań 316...319. Pozdrawiam. ■

Piotr Górecki

Policz – rozwiązanie zadania 318

W EdW 9/2022 przedstawione było zadanie Policz318, które brzmiało: (...) chcemy zrobić układ do pomiaru bardzo małych rezystancji według rysunku A. Pomysł jest taki, że każdy multimetr na zakresie 200 mV (199,9 mV) ma rozdzielczość 0,1 miliwolta. Będziemy mierzyć spadek napięcia wprost na badanej rezystancji R_x . Jeżeli przez tę mierzoną rezystancję R_x przepuścimy prąd o wartości 1 A, to za pomocą dowolnego multimetru cyfrowego będziemy mogli mierzyć rezystancje już od 0,1 milioma. Do takich pomiarów potrzebne będzie możliwie dobre źródło prądowe o niezmiennym prądzie wyjściowym 1 A. W ramach zadania Policz318 należy: **zapropionować schemat i wartości elementów źródła prądowego 1 A.**

Jeden ze stałych uczestników napisał krótko: (...) Moja propozycja źródła prądowego 1 A jest na rysunku B. Zestaw 6 rezystorów, 7,5 Ω każdy z serii E24 (5%), połączonych równolegle w rezultacie daje wymaganą rezystancję 1,25 Ohma, konieczną do wytworzenia prądu 1 A. W każdym z rezystorów wydzielą się moc 1,25 W/6=0,2 W, zatem wystarczą rezystory o obciążalności 0,5 W. Pozdrawiam serdecznie

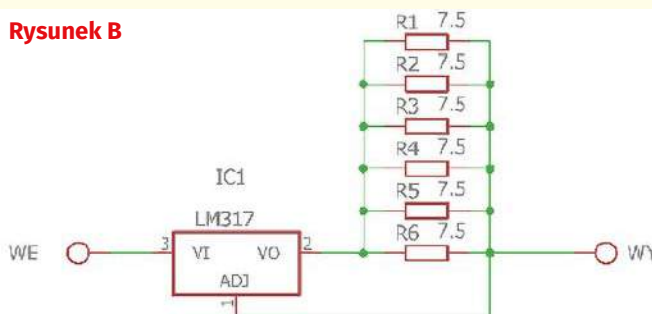
Inny stały uczestnik zaproponował inne rozwiązanie: (...) Na rysunku C pokazany jest układ z zastosowaniem rzeczywistych elementów zgodnie z założoną w zadaniu konfiguracją. (...) Ze względu na wydzielaną w rezystorze R_x moc, układ sprawdzi się do pomiaru małych rezystancji. Wydaje się, że w zakresie rezystancji 0,1 Ω ...0,0001 Ω układ spełni swoje zadanie. Na rysunku D

pokazany jest ten sam układ, ale z zastosowaniem tranzystora Darlingtona. Wartość V_2 w tym przypadku należy zwiększyć do 4 V.

Układ doświadczalny potwierdził poprawność przyjętego rozwiązania (...)

W układzie doświadczalnym użyłem rezystorów R_3 o wartości 1 Ω /1% – 10 W oraz 0,5 Ω /5% – 10 W. Na dokładność pomiaru ma wpływ wartość prądu, która powinna wynosić dokładnie 1 A. W układzie między tranzystor T_1 a rezystor R_x został włączony amperomierz celem kontroli wartości prądu. (...) konieczność wyregulowania układu za pomocą potencjometru P_1 (...) wieloobrotowego. (...) W tym wypadku pin 7 jest masą. Do masy układu dołączony jest „+” zasilania V_1 . Również dla źródła V_2 „+” jest na masie układu. Wykonałem szereg pomiarów z trzema rodzajami tranzystorów (...) miały przykręcony mały radiator. Radiator był niewielki, co powodowało mocne grzanie się tranzystorów. Na pewno w takich warunkach nie powinno się wykonywać pomiarów. Zwłaszcza kiedy tranzystor po prostu zmienia swoją temperaturę. Pozwoliło mi to stwierdzić, które tranzystory w tych warunkach najlepiej się sprawdzają. Najlepsze okazały się tranzystory typu MOSFET P (...) [Później] zmniejszyłem wartość napięcia V_2 na 3 V (dla tranzystorów typu MOSFET), a rezystor R_3 zmniejszyłem do wartości 0,5 Ω . W tych warunkach moc na tranzystorze wyniesie ok. 2,5 W, a na rezystorze 0,5 W. Rezystor R_x miał wartość podaną 0,1 Ω /5% a więc moc wydzielona na nim 0,1 W. Tranzystory

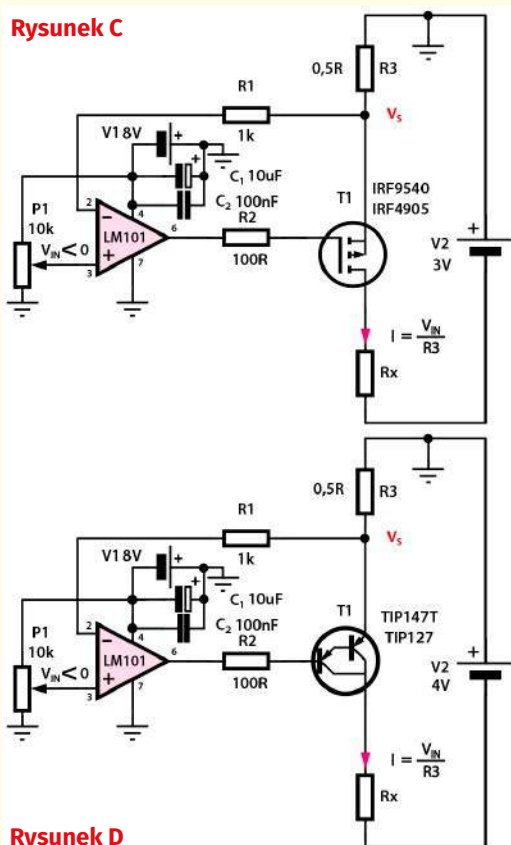
Rysunek B



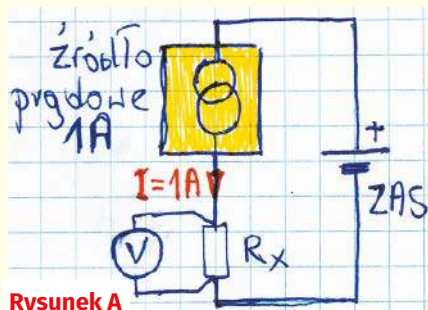
były wyposażone w ten sam radiator co we wcześniejszych pomiarach. Stwierdziłem, że tranzystory się nie grzeją i warunki pomiarów są stabilniejsze.

To samo przeprowadziłem dla tranzystora TIP147T (Darlington) przy czym stwierdziłem, że należy zwiększyć wartość napięcia V_2 do 4 V aby uzyskać prąd w obwodzie

Rysunek C



Rysunek D



Rysunek A

pomiarowym 1 A. W tym przypadku stwierdziłem grzanie się tranzystora, co utrudnia pomiar ze względu na niestabilność cieplną. Na podstawie pomiarów stwierdzam, że najlepsze w zastosowaniu będą tranzystory typu MOSFET P. (...) Układ doświadczalny zmontowany był na płytce stykowej. Na dokładność pomiaru ma wpływ

również użyty przyrząd pomiarowy i stąd może wynikać różnica między uzyskanym wynikiem a wartością podaną na rezystorze. Ale nie badałem tego problemu aby dokładnie sprawdzić, który z posiadanych mierników jest najdokładniejszy i najlepiej się sprawdzi w pomiarze tak małych napięć (rezystancji). To oddzielny problem.

Rzeczywiście, nawet w na pozór prostych zadaniach ujawniają się dodatkowe okoliczności, które finalnie trzeba uwzględnić podczas praktycznej realizacji. Jeżeli jednak takie problemy zostaną skutecznie rozwiązane, analizowany tu układ źródła prądowego może służyć do precyzyjnych pomiarów bardzo małych rezystancji, nawet rzędu mikroomów. ■

Policz – rozwiązanie zadania 319

W EdW 10/2022 przedstawione było zadanie, które brzmiało: **Zadanie Policz319 jest kontynuacją zadania 314, które pokazało, że nie tak łatwo jest zrobić sensowny przedwzmacniacz na jednym tranzystorze. Dlatego w następnym kroku nadal chcemy zaprojektować przedwzmacniacz mikrofonowy, tym razem zawierający dwa dowolne tranzystory (niekoniecznie bipolarne) według ogólnej idei z rysunku A. W ramach zadania Policz319 należy: zaproponować schemat przedwzmacniacza mikrofonowego na dwóch dowolnych tranzystorach.**

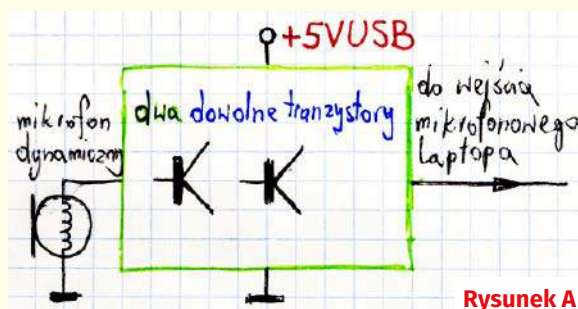
Do zadania można było podejść w różny sposób. Na przykład stosując dwa oddzielne, niezależne stopnie. Przykład takiego rozwiązania pokazany jest na **rysunku B**. W opisie można przeczytać: (...) Wykorzystam (...) schemat jednostopniowego przedwzmacniacza wykorzystując go w układzie dwustopniowym. (...) napięcie zasilania wynosi tylko 5 V, co determinuje maksymalna wartość napięcia sygnału na wyjściu, który chcemy wzmacnić (a tym samym wzmacnienie całego układu). Idealem byłoby ustawić napięcie DC kolektora tranzystorów na poziomie 2,5 V. Głównie w stopniu II. Rozwiązanie zadania polega na ustaleniu odpowiedniego wypadkowego wzmacnienia układu oraz dobraniu

punktów pracy tranzystorów tak, aby suma napięć na kolektorach tranzystorów DC + AC nie przekraczała 5 V i 0 V. W przeciwnym przypadku sygnał zmienny na wyjściu będzie zniekształcony. Będzie obcinany. Układ z zadania POLICZ314 miał wzmocnienie $\times 10$ (20 dB). Teoretycznie połączenie dwóch stopni powinno skutkować wzmocnieniem wypadkowym $\times 100$ (iloczyn wzmocnień obu stopni). (...) Policzymy jakie faktycznie jest wzmocnienie układu. Obliczymy napięcie na rezystorze R_{B4}

$$U_{RB4} = \frac{V_{CC}}{R_{B3} + R_{B4}} \cdot R_{B4} = \frac{5}{11,3 + 2,43} \cdot 2,43 = 0,885 \text{ V}$$

Napięcie na rezystorach emiterowych $R_{E3}, R_{E4}: U_{RE3,RE4} = 0,885 - 0,7 = 0,185 \text{ V}$. Prąd emitera tranzystora $T: I_{ET3} = 0,185/12 = 15,4 \text{ mA}$

Rezystancja emitera r_{eT2} II stopnia:
 $r_{eT2} = 26 \text{ mV} / I_{ET2} = 26 / 15,4 = 1,69 \Omega$. Prąd
 emitera tranzystora T_1 : $I_{ET1} = 0,185 / 30 =$
 $6,17 \text{ mA}$. **Rezystancja emitera r_{eT1}**
I stopnia: $r_{eT1} = 26 \text{ mV} / I_{ET1} = 26 / 6,17 =$
 $4,21 \Omega$. Rezystancja dla II stopnia wi-
 dziana od strony bazy tranzystora



Rysunek A

$$R_{\text{WEBASE}}:R_{\text{WEBASE}}=B \times r_{\epsilon T2}=200 \times (1,69+10)=2338 \, \Omega.$$

Rezystancja wejściowa dla II stopnia: $R_{\text{WE}}=R_{B3} \parallel R_{B4} \parallel R_{\text{WEBASE}}=11,3 \, \text{k}\Omega \parallel 2,43 \, \text{k}\Omega \parallel 2,338 \, \text{k}\Omega=1077,9 \, \Omega.$

Efektywne obciążenie kolektora dla I stopnia: $R_{\text{AC}}=R_{\text{C1}} \parallel R_{\text{WE}}=200 \parallel 1077,9=168,7 \, \Omega.$

Wzmocnienie napięciowe I stopnia

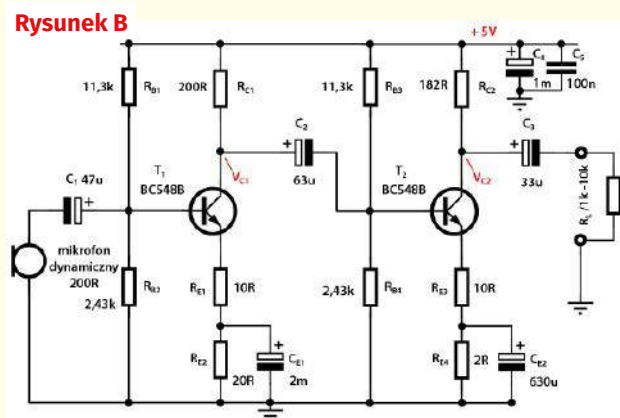
$$A_{V1} = \frac{R_{AC}}{r_{eT1} + R_{E1}} = \frac{168,7}{4,21 + 10} = 11,87$$

Efektywne wzmocnienie napięciowe II stopnia dla $R_L=1\text{ k}$. Ponieważ:
 $R_{A_2}=R_{C_2}\parallel R_L=182\parallel 1000=154\ \Omega$, stąd:

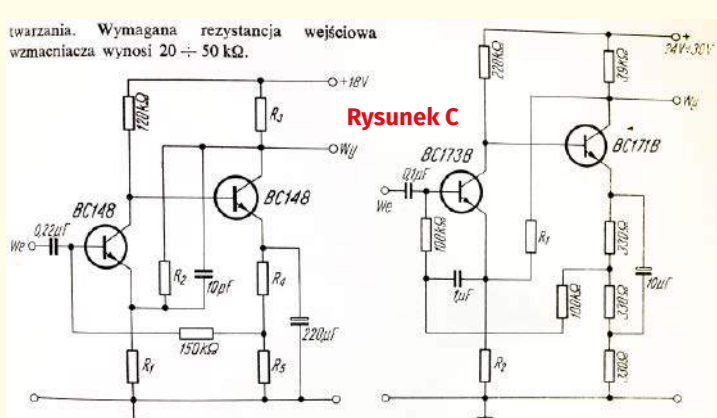
$$A_{V2} = \frac{R_{AC}}{r_{eT2} + R_{E2}} = \frac{154}{1,69 + 10} = 13,17$$

Całkowite wzm. napięciowe jest iloczynem wzm. napięciowego poszczególnych stopni

$$A_{V1,2} = A_{V1} \times A_{V2} = 11,87 \times 13,17 = 156,3 \text{ (43,9 dB)}$$



Rysunek B



Rysunek



(...) W rozważaniach przyjęto wsp. $\beta = 200$. Przy doborze tranzystorów należy zadbać aby ten współczynnik nie był mniejszy niż 200, a lepiej więcej. Poprawi to rezystancję wejściową każdego ze stopni.

Potencjały na kolektorach tranzystorów:
 $V_{C1} = 5 - (I_{C1} \times R_{C1}) = 5 - (200 \times 0,00617) = 3,766 \text{ V}$,
 $V_{C2} = 5 - (I_{C2} \times R_{C2}) = 5 - (182 \times 0,0154) = 2,197 \text{ V}$

Przy tak ustawionych punktach pracy amplituda sygnału AC na wyjściu układu może osiągnąć 1 V ($V_{pp} = 2 \text{ V}$) bez obawy o przekroczenie dopuszczalnych granic. Opis doboru pojemności pominąłem. Dla tych przyjętych wartości pojemności dolny zakres pasma przenoszenia jest poniżej 20 Hz, a górny zdecydowanie powyżej 100 kHz (pasma 3 dB). Przynajmniej w teorii. Rezystancje wejściowe I i II stopnia: $R_{WE1} = 1173,9 \Omega$, $R_{WE2} = 1077,9 \Omega$. Uzyskane wyniki są punktem wyjścia do budowy układu praktycznego. W tych rozważaniach teoretycznych (uproszczonych) nie został uwzględniony prąd bazy tranzystora T_2 . Prąd ten spowoduje większy spadek napięcia na R_{B3} , a tym samym zmniejszy się napięcie na R_{B4} . To przełoży się na mniejszy prąd I_{ET2} , a tym samym na mniejszy prąd kolektora T_2 . Mniejszy prąd emitera to większa wartość r_{eT2} . To przełoży się na mniejsze wzmocnienie II stopnia. Również potencjał V_{C2} tranzystora przesunie się w „górze” ze względu na mniejszy prąd kolektora I_{C2} . Ten sam problem dotyczy I stopnia, ale w mniejszym stopniu ze względu na mniejsze prądy. Ogólnie wzmocnienie się zmniejszy dążąc w kierunku wartości 40 dB.

Jeden z uczestników napisał: (...) nie potrafię policzyć [wartości] elementów (...) przesyłam schematy trzech wzmacniaczy z 2 tranzystorami ze starej książki (...), [gdzie] podane są gotowe schematy z elementami (...) Schematy te pokazane są na rysunkach C, D.

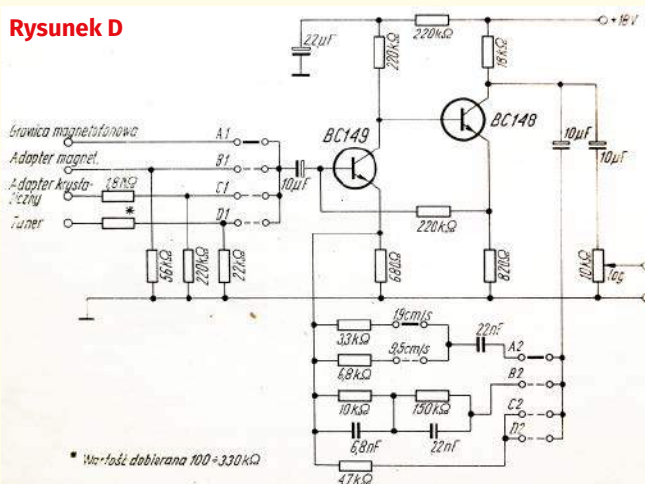
Natomiast pewien młody uczestników wyraził taką opinię: (...) zaprojektowanie

[przedwzmacniacza] na 2-óch tranzystorach jest jeszcze trudniejsze. Szkoda czasu na grzebanie się w tranzystorach, ponieważ teraz wszystko robi się na układach scalonych (...) taki przedwzmacniacz jak w zadaniu można łatwo zrobić na wzmacniaczu operacyjnym i nie trzeba długich obliczeń, a tylko dobrać dwa oporniki ustawiające wzmocnienie (...).

Tu sprawa jest mocno skomplikowana. Nieprzypadkowo postawiłem takie właśnie zadanie *Policz319*, ponieważ od dawna planuję, a ostatnio podjąłem konkretne kroki, żeby w sposób praktyczny sprawdzić jakość dźwięku (przed)wzmacniaczy tranzystorowych i porównać z dźwiękiem (przed)wzmacniaczy scalonych. Otóż to prawda, że dziś dominują układy scalone, ale w technice audio nadal wykorzystywane są układy budowane z pojedynczych tranzystorów. Mało tego! Wzmacniacze budowane z elementów dyskretnych uważane są za lepsze od tych, które zawierają wzmacniacze scalone. Po części jest to kolejny audiofiliński mit, ale tylko po części.

Otóż nie ma powodów do generalizowania, że pojedyncze tranzystory z zasady dają lepszy dźwięk niż wzmacniacze scalone, w sumie też przecież tranzystorowe. Można przedstawić wiele argumentów, że nie zawsze dźwięk „tranzystorowy” jest lepszy od „scalonego”, a już na pewno, że nie zawsze tak musi być. Wzmacniacze scalone mają poważne zalety, a jedyną kwestią dyskusyjną może być (mikro) opóźnienie i duże wzmocnienie z otwartą pętlą

Rysunek D



w rozbudowanym torze wzmacniającym. To są słabo rozumiane tematy, które ogólnie wrzuca się do kosza nazywanego „głębokością sprzężenia zwrotnego”. I właśnie z uwagi na problem „głębokości sprzężenia zwrotnego” można z powodzeniem bronić koncepcji dobrze zrealizowanego przedwzmacniacza na dwóch tranzystorach, który ma słabsze ujemne sprzężenie zwrotne i mniejsze opóźnienia, niż układ na wzmacniaczu operacyjnym. To są jednak mocno skomplikowane i dyskusyjne kwestie, które należy sprawdzić praktycznie i które na pewno wykraczają poza ramy zadania *Policz319*.

Oto lista uczestników nagrodzonych upominkami:

Maciek Blim – Łódź,
Tadeusz Suszał – Warszawa,
Andrzej Nowicki – Warszawa,
Jakub Jakubczyk – Kluczbork,
Ryszard Magdycz – Wrocław.

Ponieważ są to ostatnie zadania konkursu *Policz*, jest to zbiorcza lista osób nagrodzonych upominkami z omówionych czterech zadań 316...319. Pozdrawiam. ■

Piotr Górecki

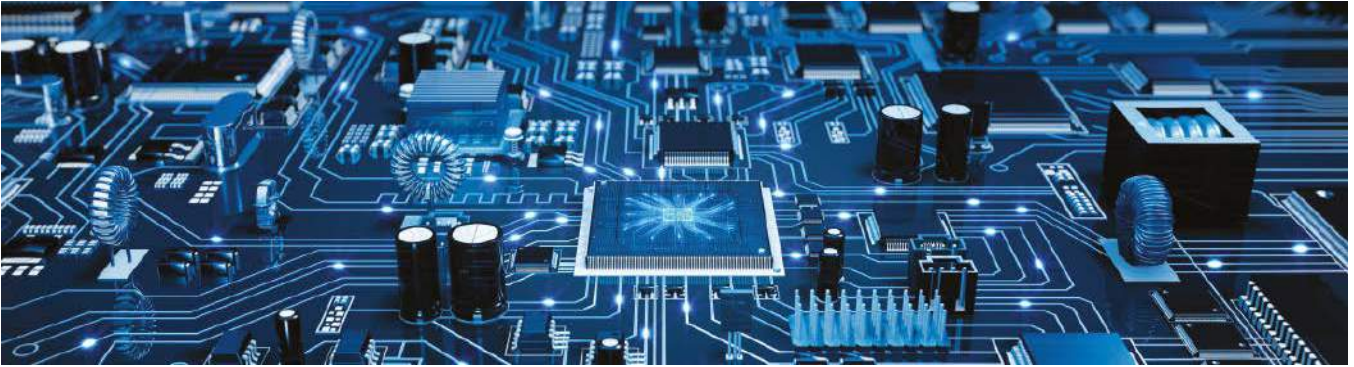
REKLAMA

Kursy w Ulubionym Kiosku

IT i Hi-tech • Muzyka i Dźwięk

Pełna oferta na stronie

www.ulubionykiosk.pl



Poziomy logiczne, część 1

Tematyka tego odcinka została zasugerowana przez redaktora PE, Matta Pulzera – a są nią poziomy logiczne i związane z nimi problemy we współpracujących układach.

Termin „poziom logiczny” odnosi się do napięć używanych do reprezentowania logicznych jedynek i zer w obwodach cyfrowych. Zazwyczaj projektujemy układy logiczne, rozważając tylko jedynki i zera danych oraz sygnałów sterujących w kombinacjach lub sekwencjach. Jednakże bramki logiczne i przerzutniki, z których zbudowane są takie obwody, są z kolei realizowane za pomocą tranzystorów – i obwody te mają właściwości analogowe, jeśli przyjrzymy się relacjom między ich wejściami, wyjściem i napięciami zasilającymi na wystarczającym poziomie szczegółowości. W większości przypadków nie trzeba się tym przejmować – możliwe jest projektowanie procesorów z dziesięcioma miliardami tranzystorów, ponieważ możemy pominąć analogowe aspekty obwodów i zajmować się jedynie cyfrową reprezentacją sygnałów. Czasami jednak musimy brać pod uwagę rzeczywiste napięcia i inne „analogowe” właściwości obwodów logicznych, na przykład gdy łączymy urządzenia z różnych rodzin technologicznych, pracujące przy różnych napięciach zasilania lub gdy rozważamy wpływ zakłóceń elektrycznych na działanie obwodu. Aby podkreślić potencjalne różnice pomiędzy obwodami, przyjrzymy się pokrótce dwóm kluczowym typom układów logicznych – CMOS i TTL – a następnie przyjrzymy

się szczegółowo napięciom logicznym i praktycznemu przykładowi ich łączenia. W tym artykule używamy terminów „technologie logiczne” i „rodziny układów logicznych”. „Technologie” odnoszą się do procesów wytwarzania stosowanych do produkcji wszystkich typów układów scalonych, w tym na przykład mikroprocesorów wysokiej klasy. Technologia określa rodzaj i wielkość stosowanych tranzystorów oraz wiele charakterystyk, takich jak szybkość i zużycie energii. „Rodzina układów logicznych” odnosi się bardziej szczegółowo do różnych zakresów podstawowych elementów logicznych (takich jak bramki, przerzutniki, rejestry i liczniki) produkowanych przez producentów takich jak Texas Instruments. Każda rodzina układów logicznych jest realizowana w określonej technologii i wykorzystuje określone podejście do projektowania bramek logicznych na poziomie pojedynczego tranzystora.

Definiowanie poziomów

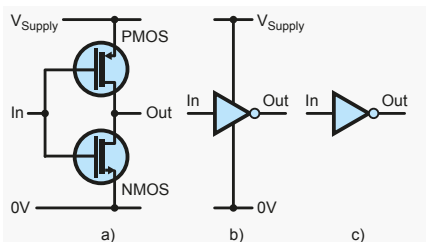
Systemy logiczne i obliczeniowe, od najprostszych funkcji po zaawansowane procesy arytmetyczne i algorytmiczne, mogą być projektowane i analizowane w kategoriach matematycznych i abstrakcyjnych reprezentacji opartych na algebrze Boole’a, liczbach binarnych i różnych systemach kodowania. Aby faktycznie zbudować obwody do realizacji tych koncepcji, musimy zdecydować, w jaki sposób reprezentować logiczne jedynki i zera w rzeczywistym obwodzie elektronicznym. Moglibyśmy użyć dwóch napięć, powiedzmy +5 V dla logicznej 1 i 0 V dla logicznego 0, lecz jest to kwestia umowna. Równie dobrze mogłoby to być –2 V dla 0 i +2 V dla 1, lub 0 V dla 1 i +5 V dla 0. Ogólnie rzecz biorąc, jeśli wyższe napięcie jest używane dla logicznej 1, mamy do czynienia z logiką pozytywną, a jeśli niższe napięcie jest używane dla 1, mamy logikę ujemną. Sprawy mogą się

skomplikować – przykładowo ten sam fizyczny obwód może działać jako bramka NAND dla logiki dodatniej, ale jako bramka NOR dla logiki ujemnej. Na szczęście zdecydowana większość obwodów, z którymi prawdopodobnie się zetkniesz, będzie używać logiki dodatniej. Jeśli wybierzemy 5 V dla logicznej 1 i 0 V dla logicznego 0, to co reprezentuje 4,9 V? W prawdziwych obwodach musimy zdefiniować zakres napięć, które reprezentują prawidłowy poziom logiczny, powiedzmy 0 V do 2 V dla 0 i 3 V do 5 V dla 1. Musimy to robić, ponieważ nie możemy budować obwodów, które działają dla dokładnie ustalonych napięć w zmiennych warunkach obciążenia, zasilania i temperatury, a także z różnicami w poszczególnych produkowanych elementach. Napięcia używane do reprezentowania jedynek i zer będą zależały od napięcia zasilania używanego do zasilania obwodu.

Zazwyczaj jedna z szyn zasilania jest masą (0 V), a druga ma napięcie dodatnie względem masy (napięcie zasilania). Dla logiki dodatniej idealny poziom jedynki będzie równy napięciu zasilania, a idealny poziom zera będzie równy 0 V. Jak przed chwilą wspomniano, musimy ustalić zakres napięć, więc poziom zera będzie mieścił się w zakresie od 0 V do stosunkowo małego napięcia dodatniego, a poziom jedynki będzie od większego napięcia dodatniego aż do napięcia zasilania. Przyjrzymy się temu nieco bardziej szczegółowo w dalszej części.

Zasilanie logiki

Wspomniałszy o zasilaniu, warto zwrócić uwagę, że schematy układów logicznych są zazwyczaj rysowane bez pokazywania połączeń zasilających. Jest to spowodowane tym, że gdybyśmy narysowali wszystkie przewody zasilające, to zapchalibyśmy schemat dodatkowymi

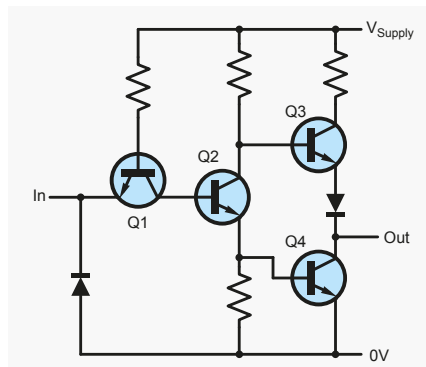


Rysunek 1. Bramka CMOS NOT: a) obwód tranzystorowy bramki CMOS NOT, b) bramka NOT z symbolem bramki i zasilaniem, c) bramka NOT z symbolem bramki bez zasilania.

liniami, co utrudniłoby zrozumienie struktury układu logicznego. **Rysunek 1** pokazuje schemat bramki NOT zaimplementowanej w jej najbardziej podstawowej formie w technologii CMOS (więcej na ten temat wkrótce) – zwykle pokazujemy zasilanie, gdy patrzymy na obwody logiczne na odpowiednim poziomie szczegółowości. Rysunek pokazuje również ten sam obwód narysowany przy użyciu znanego symbolu bramki NOT, zarówno z zasilaniem, jak i bez niego – ten drugi wyraźnie skupia się tylko na funkcji logicznej. Rzeczywiste napięcia używane do zasilania układów cyfrowych stały się bardziej zróżnicowane w miarę postępu technologicznego. Był czas, kiedy wiele cyfrowych układów scalonych używało wyłącznie zasilania 5 V lub preferowało pracę przy 5 V (w większym możliwym zakresie). Od tego czasu nastąpił trend spadkowy napięć zasilających, a typowe wartości to 3,3 V, 2,5 V, 1,8 V, 1,5 V, 1,2 V i 0,8 V. Dlatego też nierzadko zdarza się, że dwa kluczowe układy lub moduły, które realizują istotne funkcje dla danego projektu, wymagają różnych napięć zasilania, a co za tym idzie, mają potencjalnie niekompatybilne poziomy logiczne. Tendencja spadkowa i związana z nią dywersyfikacja napięć zasilających dostępnych układów jest związana z fizyką skalowania tranzystorów do małych rozmiarów. Jest to siła napędowa naszej coraz późniejszej technologii cyfrowej, którą opisuje prawo Moore’a (obserwacja, że liczba tranzystorów w układzie scalonym podwaja się mniej więcej co dwa lata). Wpływ skalowania rozmiarów tranzystorów na właściwości technologii układów scalonych jest złożony. Obniżenie napięć zasilania jest związane z rozpraszaniem mocy, prądami upływu (które wpływają na moc w stanie czuwania) oraz natężeniem pola elektrycznego wewnątrz układów. Ta tendencja spadkowa napięć zasilających mniej więcej podążała za trendem zmniejszania rozmiarów tranzystorów od lat 90. do połowy lat 2000, gdy ustabilizowała się, podczas gdy tranzystory nadal stawały się mniejsze, mniej więcej zgodnie z prawem Moore’a.

Logika oparta na MOSFET-ach

Bramka NOT (inwerter) z rysunku 1 wykorzystuje technologię CMOS. CMOS to skrót od „complementary metal-oxide semiconductor”, gdzie słowo „complementary” odnosi się do faktu, że układ implementuje oba typy tranzystorów MOSFET (PMOS i NMOS). CMOS jest obecnie najpowszechniej stosowaną technologią w układach cyfrowych. Na najbardziej podstawowym poziomie, tranzystory MOS w logice CMOS działają jak elektrycznie sterowane przełączniki, z napięciem bramki kontrolującym przewodzenie między źródłem a drenem. Tranzystory PMOS i NMOS wymagają przeciwnych polaryzacji bramka-źródło, aby się włączyć. W bramce NOT



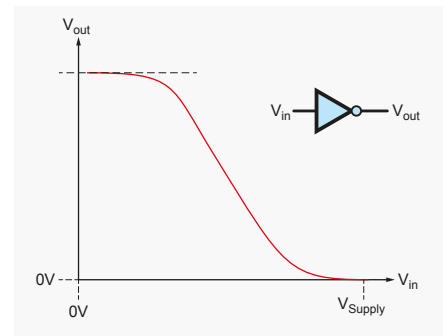
Rysunek 2. Schemat bramki (inwertera) TTL NOT.

pokazanej na rysunku 1, wejście logiczne 1 (bliskie V_{Supply}) powoduje włączenie tranzystora NMOS i wyłączenie PMOS, co powoduje połączenie wyjścia z masą (0 V) dając logiczne 0 na wyjściu – zgodnie z wymaganiami logiki NOT. Dla wejścia logicznego 0 (0 V) występują przeciwne warunki (PMOS włączony, NMOS wyłączony), więc wyjście jest podłączone do zasilania dając logiczną 1 na wyjściu – zgodnie z wymaganiami. CMOS nie był pierwszą technologią logiczną wykorzystującą tranzystory MOS. Rozwój rozpoczął się we wczesnych latach 60., kiedy to pierwsze układy scalone były wykonane wyłącznie w technologii PMOS, ponieważ tranzystory NMOS były wówczas zbyt trudne do wyprodukowania. W połowie i pod koniec lat 60. dostępne były również układy wyłącznie w technologii NMOS i CMOS, a firma RCA wprowadziła dobrze znaną serię układów 4000. Tranzystory NMOS były szybsze od PMOS, a wraz z poprawą technik produkcji, pod koniec lat 70. zdominowały logikę opartą na MOS. Wraz z dalszym postępowaniem, CMOS, który ma podobną szybkość, ale znacznie mniejszą moc niż NMOS, stał się najbardziej rozpowszechnioną technologią logiczną (np. w procesorach) w latach 80. Na przestrzeni lat pojawiło się wiele rodzin układów logicznych CMOS, o nazwach takich jak High-Speed CMOS (HC/HCT), Advanced CMOS (AC/ACT) i Advanced Ultra-Low-Voltage CMOS (AUC). Poszczególne rodziny CMOS pracują z różnymi napięciami zasilania i mogą pracować z różnymi maksymalnymi szybkościami.

Logika oparta na tranzystorach bipolarnych

CMOS nie jest jedyną technologią logiczną, zwłaszcza gdy spojrzeć na ogólną historię rozwoju zintegrowanej elektroniki cyfrowej. Podobnie jak w przypadku logiki opartej na MOS, we wczesnych latach 60. rozpoczyna się rozwój różnych układów scalonych opartych na tranzystorach bipolarnych:

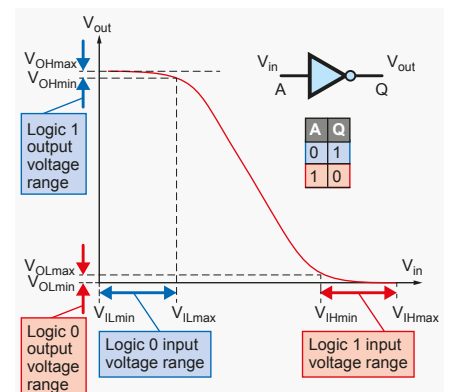
- Resistor Transistor Logic (RTL) – logika rezystorowo-tranzystorowa,
- Diode Transistor Logic (DTL) – logika diodowo-tranzystorowa,



Rysunek 3. Szkic charakterystyki przenoszenia inwertera.

- Transistor Transistor Logic (TTL) – logika tranzystorowo-tranzystorowa.

TTL stał się dominującą wersją spośród tych technologii, szczególnie wraz z wypuszczeniem serii układów 7400 przez Texas Instruments w połowie lat 60. TTL nadal się rozwijał, z wykorzystaniem tranzystorów Schottky’ego i innych osiągnięć prowadzących do szybszych urządzeń o mniejszej mocy. Przez lata wypuszczano różne rodziny urządzeń, w szczególności: Low-power Schottky (LS), Advanced Low-power Schottky (ALS) oraz FAST (F) – Fairchild Advanced Schottky TTL. TTL jest obecnie nieco przestarzałą technologią, której użycie spada od końca lat 90, pomimo tego, że jest dobrze znana. Układy w tej technologii są jednak nadal dostępne, a poziomy logiczne „zgodne z TTL” są powszechnie stosowane. Niektóre rodziny logiczne mają dwie wersje (np. AC i ACT) z wersją „T” wskazującą poziomą kompatybilność z TTL. CMOS i TTL to nie jedyne główne typy układów logicznych. Istnieją również układy scalone, które łączą oba typy tranzystorów w tak zwanej „technologii BiCMOS”. **Rysunek 2** pokazuje schemat podstawowego inwertera TTL. Podobnie jak obwód CMOS z rysunku 1, może on być również przedstawiony jako symbol bramki NOT podczas rysowania schematu na poziomie bramek. Układ posiada trzy stopnie: sterowanie prądem (Q1), podział fazy (Q2) oraz sterownik wyjścia z szeregowo połączonymi tranzystorami



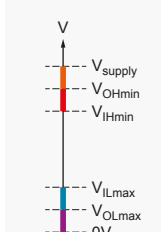
Rysunek 4. Charakterystyka przenoszenia z rysunku 3 pokazująca zakresy napięć dla operacji logicznej NOT.

bipolarnymi o jednakowej polaryzacji („totem-pole”) (Q3 i Q4). Wejście steruje prądem bazy Q2 poprzez Q1. Q2 z kolei steruje bazami Q3 i Q4 w przeciwnych kierunkach. Jest to wymagane, ponieważ w przeciwieństwie do układu CMOS, oba tranzystory wyjściowe są tego samego typu (NPN), więc ich bazy muszą byćysterowane w różny sposób, aby wyłączyć jeden z nich, podczas gdy drugi pozostaje włączony. Dla obwodu z rysunku 2, dla 0 V na wejściu, Q1 powoduje przepływ prądu od bazy Q2, wyłączając Q2, co powoduje powstanie napięcia bliskiego zasilaniu na bazie Q3 i 0 V na bazie Q4. To powoduje włączenie Q3 i wyłączenie Q4, dając na wyjściu logiczną 1. Jeśli wejście nie jest podciągnięte do niskiego napięcia, Q1 dostarcza prąd do bazy Q2, włączając go. Powoduje to obniżenie napięcia na bazie Q4 na tyle, aby go wyłączyć, ale spadek napięcia na rezystorze na emiterze Q2 jest wystarczający do włączenia Q4, co daje logiczne 0 na wyjściu.

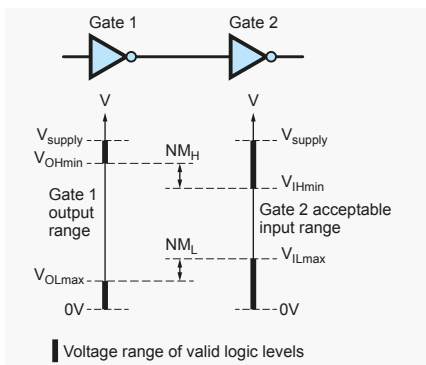
Napięcia logiczne

Różne typy tranzystorów, różne schematy bramek NOT (rysunki 1 i 2) i odmienne zasady ich działania sprawiają, iż prawdopodobnym jest, że gdybyśmy szczegółowo zbadali analogowe zachowanie bramek NOT CMOS i TTL, to zobaczmy różnice, na przykład w zakresie napięcia wejściowego, przy którym można z całą pewnością powiedzieć, że układ reaguje jako wejście logicznej 1. Jak już wskazano, dostępne są rodziny układów CMOS, które pracują przy różnych napięciach zasilania – więc nawet układy wykorzystujące tę samą podstawową technologię mogą mieć różne napięcia logiczne – w szczególności jedne z podanej wcześniej listy napięć. Z drugiej strony, TTL jest dobrze znany z działania tylko przy zasilaniu 5 V, ale oparte na BiCMOS układy „niskiego napięcia TTL” (LVTTTL), takie jak rodzina LVT, działają przy napięciu zasilania 3,3 V. Możemy potraktować bramkę logiczną jak obwód analogowy i przyjrzyć się zależnościom napięcia wyjściowego od wejściowego (charakterystyka przenoszenia). Najprościej będzie nam to zrobić dla inwertera, ponieważ ma on tylko jedno wejście i da ogólny obraz zachowania innych bramek zbudowanych w tej samej technologii. Wykres charakterystyki przenoszenia inwertera pokazano na rysunku 3 – stanowi on tylko szkic ilustracyjny – nie przedstawia on rzeczywistego układu ani danych symulacyjnych. Krzywa ma dwa poziome obszary, gdzie napięcie wyjściowe nie zmienia się zbyt i jest albo blisko napięcia zasilania, albo masy.

Do pracy logicznej chcemy, aby wyjście znajdowało się



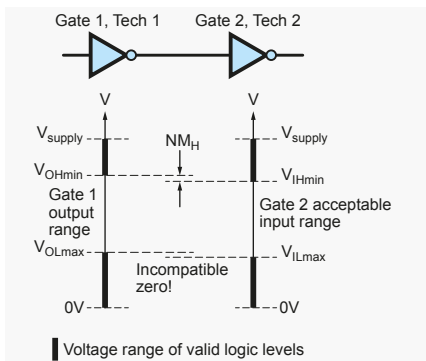
Rysunek 5 Graficzne przedstawienie napięć bramek logicznej.



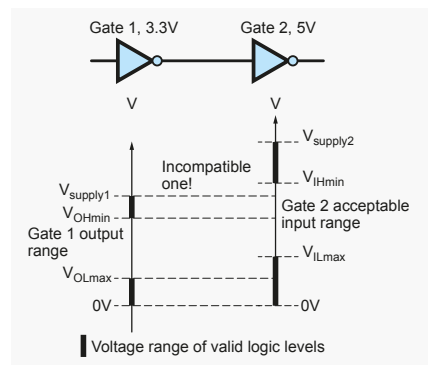
Rysunek 6. Poziomy logiczne i marginesy szumów. Zakres wyjściowy bramki dla logicznych zera i jedynki jest wyższy od dopuszczalnego zakresu wejściowego, co powoduje zawężenie poziomu napięcia w kierunku idealnego.

w jednej z tych dwóch części krzywej. Pomiędzy nimi znajduje się obszar nachylony (w prawdziwych bramkach może być on znacznie bardziej stromy niż pokazany). Dla operacji logicznych chcemy uniknąć pracy w tym obszarze, a także przejść z obszarów „poziomych” w część nachyloną. Nachylenie, przynajmniej w części prostoliniowej, jest jak charakterystyka wzmacniacza liniowego i rzeczywiście można stosować niektóre bramki NOT jako „analogowe” wzmacniacze. Rysunek 4 pokazuje sytuację z rysunku 3 z zaznaczonymi zakresami napięć, dla których moglibyśmy użyć tego inwertera jako bramki NOT. Oznaczono tutaj osiem różnych napięć w celu określenia granic czterech pokazanych zakresów. Połowa z tych granic jest równa zasilaniu lub masie, więc istnieją cztery kluczowe napięcia, które (wraz z napięciem zasilania) określają działanie przełączające bramki:

- Maksymalne napięcie wejściowe logicznego 0 (V_{ILmax})
- Maksymalne napięcie wyjściowe logicznego 0 (V_{OLmax})
- Minimalne napięcie wejściowe logicznej 1 (V_{IHmin})



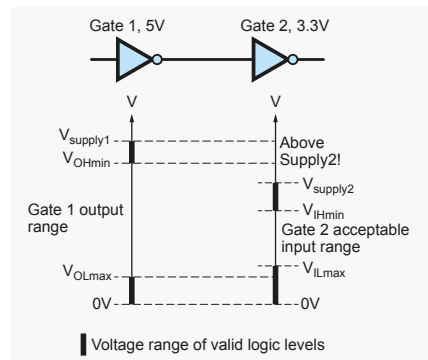
Rysunek 7. Przykład niekompatybilnych układów: różne technologie o tym samym zasilaniu. Zakres wyjściowy logicznego zera dla technologii 1 jest zbyt szeroki dla technologii 2 – nie wszystkie ważne sygnały 0 z bramki 1 zostaną rozpoznane przez bramkę 2. Logika dla 1 jest poprawna, ale margines logiczny jest bardzo mały, poziom logiki 1 będzie wrażliwy na zakłócenia.



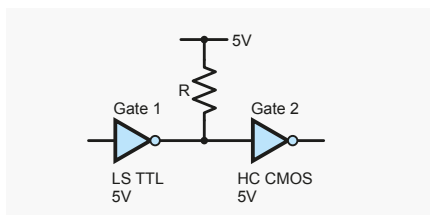
Rysunek 8. Przykład niekompatybilnych układów: bramka niskiego napięcia steruje bramką wyższego napięcia. Logika 0 jest poprawna, ale napięcie wyjściowe logiki 1 jest niewystarczające do rozpoznania przez drugą bramkę.

- Minimalne napięcie wyjściowe logicznej 1 (V_{OHmin})

Ogólnie rzecz biorąc, niezależnie od technologii, bramki logiczne powinny akceptować dany zakres napięć wejściowych jako 1 lub 0 i zapewniać tworzenie węższego zakresu możliwych napięć wyjściowych, bliższych odpowiedniemu idealnemu napięciu dla 1 lub 0 (zasilanie lub masa). Widać to na przykładzie charakterystyki przenoszenia z rysunku 4. Oznacza to, że każda bramka ma tendencję do formowania napięcia w kierunku idealnego dla danego poziomu logicznego. Chociaż krzywa przenoszenia jest użyteczna jako szczegółowa reprezentacja charakterystyki bramki, pełna krzywa nie jest nam potrzebna, jeśli chcemy uzyskać jedynie graficzną reprezentację napięć logicznych. Do tego celu tworzy się powszechnie schematy za pomocą pionowego paska lub linii, na której zaznaczone są pozycje odpowiednich napięć – patrz rysunek 5. Reprezentacja na rysunku 5 daje również natychmiastowe wskazanie różnicy pomiędzy najgorszym przypadkiem poziomu wyjściowego a najgorszym dopuszczalnym poziomem wejściowym zarówno dla logicznej 1 jak i logicznego 0. Wartość liczbową



Rysunek 9. Przykład niekompatybilnych układów: bramka wysokiego napięcia steruje bramką niższego napięcia. Logika 0 jest poprawna. Napięcie wyjściowe logiki 1 jest powyżej minimalnego wymogu wejścia logiki 1, ale napięcie to może być zbyt wysokie dla drugiej bramki.



Rysunek 10. Połączenie LSTTL z HC-CMOS.

tej różnicy nazywana jest marginesem szumów. Dokładniej, możemy zdefiniować marginesy szumów dla logicznego 0 (NM_L) i logicznej 1 (NM_H) w następujący sposób:

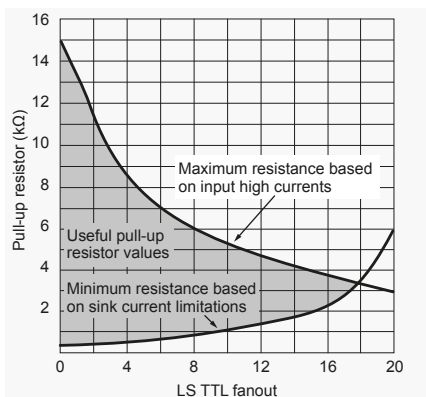
$$NM_L = V_{ILmax} - V_{OLmax}$$

$$NM_H = V_{OHmin} - V_{IHmin}$$

Zobrazowano to na **rysunku 6**. Marginesy szumów wskazują, jak dobrze połączenie w układzie cyfrowym radzi sobie z zakłóceniami elektrycznymi (przesunięciami napięcia) bez utraty prawidłowej wartości danych wejściowych. Jeśli napięcie znajduje się pomiędzy zdefiniowanymi poziomami logicznymi z jakiegokolwiek powodu (z wyjątkiem krótkiego czasu podczas przełączania pomiędzy poziomami) wtedy mamy niezdefiniowaną wartość logiczną i obwód może zachowywać się nieprzewidywalnie lub nawet ulec uszkodzeniu.

Niekompatybilne układy logiczne

Poza kwestiami związanymi z szumem, o czym wspomnieliśmy, istnieją dwie kluczowe cechy, które mogą się różnić, a co za tym idzie powodować potencjalne problemy z łączeniem ze sobą układów cyfrowych. Są nimi technologia logiczna (np. CMOS i TTL) oraz napięcie zasilania. Różne technologie mogą stwarzać problemy z niekompatybilnością nawet przy tym samym napięciu zasilania. Ilustruje to **rysunek 7**, na którym widać dwa możliwe w tej sytuacji problemy – wąski margines szumów i niedopasowane zakresy logiczne. **Rysunki 8 i 9** ilustrują możliwe problemy w przypadku połączenia ze sobą bramek logicznych pracujących na różnych napięciach. Na **rysunku 8** bramka o niższym napięciu steruje bramką o wyższym napięciu. Logiczne zera są zgodne, ale jedynki już nie. W tym przypadku wyjście logicznej 1 z pierwszej bramki może utrzymywać drugą bramkę w pośrednim zakresie napięcia, co oprócz tego, że nie jest rozpoznawane jako prawidłowa 1, może powodować inne niepożądane efekty, takie jak powodowanie większego niż normalnie rozpraszania mocy w drugiej bramce. Na **rysunku 9** bramka o wyższym napięciu steruje bramką o niższym napięciu. Logiczne 0 są ze sobą zgodne. Dla poziomu logicznego 1 wyjście pierwszej bramki jest zdecydowanie powyżej minimum 1 dla drugiej bramki, więc w tym elemencie są ze sobą kompatybilne; jednak wysokie napięcie wyjściowe

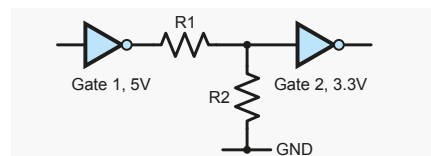


Rysunek 11. Wybór rezystora podciągającego dla połączenia LSTTL/HC-CMOS (źródło: Fairchild Semiconductor application note AN-314 via ON Semiconductor).

z bramki 1 może uszkodzić bramkę 2. Taka sytuacja nie jest wcale wyjątkowa – niektóre wejścia cyfrowe są specjalnie zaprojektowane do radzenia sobie z napięciami wejściowymi powyżej napięcia zasilania. Jest to tzw. tolerancja napięciowa. Na przykład sytuacja z **rysunku 9** byłaby poprawna, gdyby druga bramka miała tolerancję 5 V. Informacje o tolerancji napięcia wejściowego można znaleźć w kartach katalogowych urządzeń. Łącząc ze sobą bramki wykonane w tej samej technologii i o tym samym napięciu zasilania w zasadzie musimy się martwić wyłącznie o obciążenie. Zakresy napięć powinny być automatycznie kompatybilne. Jednak dla każdej technologii istnieje ograniczenie liczby wejść, które wyjście bramki może poprawnie wysterować. Efekty obciążenia będą stosunkowo łatwe do oceny dla pojedynczej technologii (szczegóły są łatwo dostępne w kartach katalogowych), ale przyłączeniu dwóch różnych technologii sytuacja może być trudniejsza do oceny.

Łączenie rodzin – przykład

Istnieje bardzo wiele rodzin układów logicznych, które potencjalnie mogą pracować na tym samym napięciu, więc nie możemy omówić każdej możliwej kombinacji związanej z ich łączeniem. Przyjrzyjmy się tylko jednemu przykładowi – wysterowaniu HC-CMOS z LSTTL. Są to nieco przestarzałe technologie, jakkolwiek wciąż dostępne. Ponadto dyskusja posłuży jako ilustracja niektórych problemów występujących w przypadku interfejsów logicznych. Specyfikacja LSTTL gwarantuje poziom wyjściowy 2,7 V dla logicznej 1, lecz HC-CMOS wymaga wejścia 3,5 V dla logicznej 1 przy zasilaniu 5 V. W praktyce wyjście LSTTL będzie prawdopodobnie wystarczające, ale kompatybilność nie jest gwarantowana przez charakterystykę najgorszego przypadku. Aby zapewnić kompatybilność, wyjście logicznej 1 z układu LSTTL może być podniesione poprzez umieszczenie rezystora podciągającego z wyjścia do napięcia zasilania (patrz **rysunek 10**). Gdy wyjście przechodzi



Rysunek 12. Interfejs logiczny dzielnika potencjału.

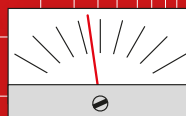
w stan wysoki, rezystor podciąga napięcie bardzo blisko napięcia zasilania. Wartość rezystora należy dobrać w zależności od liczby innych wejść LSTTL, które bramka obsługuje, oprócz wejścia (wejść) CMOS (określany jako „fan-out” LSTTL), korzystając z wykresu na **rysunku 11**. Na przykład, jeśli urządzenie LSTTL napędza tylko obwody CMOS, to fan-out LSTTL wynosi 0, więc wartość rezystora jest wybierana z lewej osi wykresu. Dane te pochodzą ze szczegółowej noty aplikacyjnej dotyczącej interfejsów logicznych firmy Fairchild Semiconductor (obecnie część ON Semiconductor) (<http://bit.ly/pe-jan20-ttl>), która analizuje problemy związane z interfejsami dla szeregu rodzin układów logicznych dostępnych w późnych latach 90. i jest użytecznym źródłem informacji na ten temat dla każdego, kto jest zainteresowany łączeniem tych starszych technologii. Kiedy obwody logiczne pracują na różnych napięciach, bezpośrednie połączenie nie zawsze jest możliwe i potrzebujemy specjalnego układu interfejsu. Najprostszym rozwiązaniem jest użycie dzielnika potencjału, aby zredukować poziom wyjściowy obwodu o wyższym napięciu i uczynić go kompatybilnym z obwodem o niższym (patrz **rysunek 12**). Typowymi wartościami mogą być $R1=18\text{ k}\Omega$ i $R2=33\text{ k}\Omega$ dla 5 V do 3,3 V. Zastosowanie dzielnika potencjału zwiększy pobór mocy i może spowodować wolniejsze przełączenie.

Układy scalone do translacji poziomów

W tym artykule omówiliśmy podstawy poziomów napięć logicznych oraz niektóre problemy, które mogą wystąpić podczas interakcji układów logicznych wykorzystujących różne technologie i źródła zasilania. Jest to powszechny problem w projektowaniu układów cyfrowych i dlatego dostępne są różne układy scalone, które są specjalnie przeznaczone do translacji napięć logicznych – obejmują one translację sygnałów dwukierunkowych, jak również jednokierunkowych, które omówiliśmy. Użycie takich układów scalonych jest często najlepszym rozwiązaniem. ■

Ian Bell

AUDIO OUT



Zwrotnica do mini monitora PE, część 1

Witam w kolejnym etapie naszej podróży po LS3/5a. W tym odcinku mieliśmy zamiar zająć się obudową, a następnie wyborem zwrotnic, ale w rzeczywistości poczyniliśmy większe postępy w tej drugiej dziedzinie, więc najpierw zajmiemy się tym tematem, a w następnej kolejności obudowami. To zaskakujące, jak skomplikowana może być prosta drewniana skrzynka, ale chcemy, aby była ona wystarczająco elastyczna dla całej gamy przyszłych konstrukcji, również dla LS3/5a, więc warto poświęcić trochę więcej czasu, aby zrobić to dobrze.

Wprowadzenie

Ten projekt nie polega tylko na „zbudowaniu LS3/5a”, ale na dostarczeniu Ci bazy do wykorzystania różnych głośników i związanych z nimi zwrotnic w oparciu o jedną obudowę i (w większości) jedną płytkę drukowaną zwrotnicy. **Rysunek 1** pokazuje możliwe opcje zwrotnic. Wersje 1 i 2 stanowią wyjątek od powyższej zasady uniwersalności. Wersja 1 to po prostu oficjalna konstrukcja zwrotnicy LS3/5a firmy Falcon Acoustics. Jest ona oczywiście doskonała, ale droga i przeznaczona tylko dla głośników LS3/5a. Wersja 2 to budżetowa, chińska wersja z eBay’a. Wykorzystuje ona płytkę drukowaną – z eBay’a i cewki – także z eBay’a. Rezultaty są dobre, ale niestety niezgodność komponentów skutkuje wersją „budżetową”. Następne są wersje 3 do 7 (plus przyszłe potencjalne odmiany). Wykorzystują one naszą własną uniwersalną płytkę zwrotnicy, bardzo elastyczną konstrukcję z mnożeniem możliwości dostrojenia i modyfikacji. Pozwala ona na uniknięcie drogich cewek (‘oficjalnych BBC’) poprzez zastosowanie wysokiej jakości elementów z rdzeniem ferrytowym. W tym miesiącu zbudujemy wersje 2 i 3 (wersja 1 pochodzi bezpośrednio

od Falcon Acoustics). W następnym miesiącu zajmiemy się konstrukcją Wavectora. Istnieje więc kilka podstawowych konstrukcji zwrotnic dla projektów mini-monitorów PE. Pierwszą, którą omówimy, jest klasyczny układ 15 Ω LS3/5a w standardzie BBC, pokazany na **rysunku 2** (to zredagowana wersja z rysunku 33, Audio Out, PE z października 2019).

Wersja 2: klasyczna zwrotnica BBC

Wersja ta wykorzystuje płytkę LS3/5a 15 Ω dostępną na eBayu, pokazaną na **rysunku 3**. Wykorzystuje ona głośniki B110A i T27 z Falcon Acoustics lub z dużego rynku komponentów używanych (jak zawsze, patrz eBay). Na eBayu dostępny jest także pasujący zestaw cewek. Te chińskie cewki nie do końca spełniają specyfikację tolerancji

komponentów BBC, wynoszącą $\pm 5\%$ (raczej -7%), ale na pewno spełniają swoje zadanie. Są one pokazane na **rysunku 4**.

Komponenty – zwrotnica, wersja 2

Wszystkie tolerancje komponentów to $\pm 5\%$.

PCB. ZeBay – np. eBay part 191278013041, lub szukaj „LS3/5a crossover 15ohm”.

Rezystory. (WW = wire-wound, rezystor drutowy)

R1, R3 82 Ω 2 W, WW lub tlenek metalu

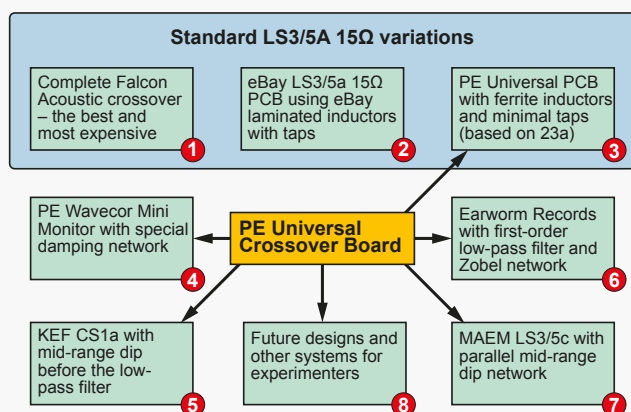
R2 22 Ω 6 W, WW (patrz kolejny odcinek)

R4 8,2 Ω 6 W, WW

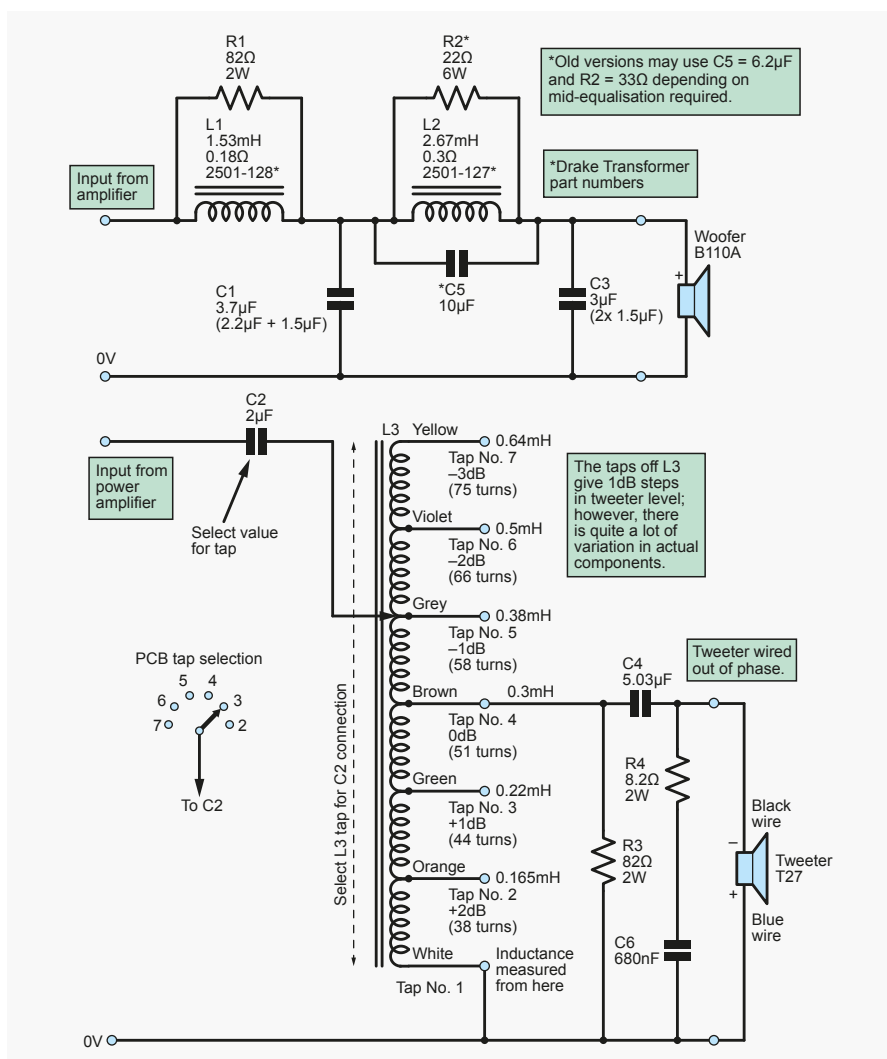
Kondensatory. Oryginalny projekt BBC używał kondensatorów typu MKC, takich jak STC/ITT PMC, Wima MKB3, i Advance Filmcap, z których wszystkie nie są już produkowane (choć Falcon miał kilka kondensatorów poliwęglanowych specjalnie wyprodukowanych). Ja używam poliestrowych lub MKT, które są im najbliższe w jakości dźwięku, ale nie są tak stabilne. (Rogers zawsze używał Rifa PHE w przezroczystych obudowach, które były typu MKT.) Wszystkie kondensatory mają obudowy cylindryczne, minimalne napięcie 63 V, chociaż użycie kondensatorów na wyższe napięcia 100 V i 160 V daje nieco niższe zniekształcenia.

C1 3,7 μF – użyj 2,2 μF (C1a) + 1,5 μF (C1b)

C2 Wartość zależy od odczepu cewki dobrego dla poziomu głośnika wysokotonowego, aby dopasować się do głośnika niskotonowego. Zwykle 1,72 μF z odczepem nr 6 (fioletowym) dla głośników wysokotonowych T27 produkowanych po 1984 roku. Składają się na nią C2a 1,5 μF i C2b 220 nF. (Dla innych



Rysunek 1. Schemat przedstawiający możliwe konstrukcje zwrotnic do zbudowania w naszej serii artykułów. Zwróć uwagę na mnogość zwrotnic dwudrożnych, które można zbudować na płytce uniwersalnej.



Rysunek 2. Oryginalny układ BBC 15 Ω LS3/5a. Cewka z odczepami ma za zadanie zmieniać poziom głośnika wysokotonowego z krokiem 1 dB. (Zauważ, że na rysunku 33 artykułu z października 2019 roku ustawiliśmy krok odczepu na 0,5 dB; jednak Malcolm Jones z Falcon Acoustics doradził, że odczepy powinny być w rzeczywistości krokami 1 dB. Rzeczywiście, w moich testach licznych cewek w różnych odmianach LS3/5a jest bardzo dużo odchytek – cewki istotnie mają ogromne wartości tolerancji).

odczepów wartość C2 różni się – patrz tabela w dalszej części.)

C3 3 μF (2 × 1,5 μF)

C4 5,03 μF (4,7 μF + 330 nF)

C5 10 μF (może wymagać zmniejszenia dla niektórych B110A – patrz następny odcinek)

C6 680 nF

Cewki. Zestaw sześciu (dla pary zwrotnic) przez eBay z Hifikits – eBay part 323926563397, lub szukaj 'Inductor for LS3/5a 15ohm Crossover/Frequency divider One Pair(6 Pieces)' – tak, 'ivider' jest poprawnym słowem! Cena to około 65 funtów z dostawą.

L1 1,53 mH, rezystancja przy prądzie stałym (DCR) 0,18 Ω lub mniej. Laminowany Radiometal, lub podobny 45 do 49% stop niklu i stali. Rozmiar rdzenia 40 mm długości × 35 mm × 10 mm szerokości. Środki mocowań na 46 mm, otwory 3,8 mm.

L2 2,67 mH, DCR 0,3 Ω lub mniej, ten sam rdzeń co L1.

L3 0,64 mH z odczepami na 0,5, 0,38, 0,3, 0,22 i 0,165 mH, DCR 0,13 Ω lub mniej, ten sam rdzeń co L1.

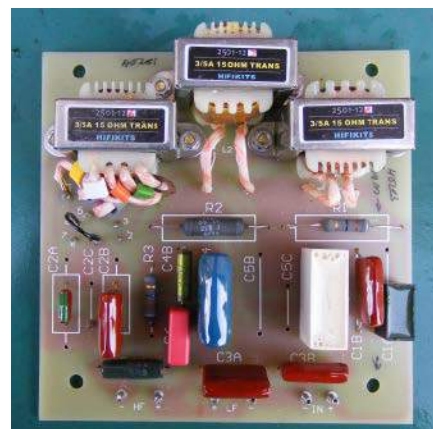
Akcesoria. Do zamocowania laminowanych cewek indukcyjnych użyj sześciu śrub z łbem M3,5 typu Pozi z nakrętkami i podkładkami zabezpieczającymi.

Siedem 0,15-calowych i siedem 0,1-calowych pinów jest potrzebnych do końcówek, plus sześć 0,2-calowych pinów

Do montażu płytki drukowanej użyj śrub z łbem stożkowym M5 × 40 mm, aby przymocować płytkę drukowaną do przedniej przegrody głośnika; potrzebne jest także 12 nakrętek i podkładek zabezpieczających.

Budowa – wersja 2

Montaż jest bardzo prosty, w zasadzie dużo łatwiejszy niż w przypadku większości płytek drukowanych, ponieważ elementy są duże – bardziej przypominają radio



Rysunek 3. Płytkę zwrotnicy LS3/5a 15 Ω z eBaya, ale na jak długo? Opisy elementów odpowiadają układowi BBC.

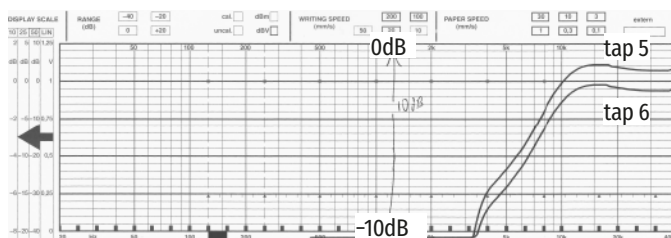


Rysunek 4. Laminowane cewki LS3/5a są dostępne na eBayu. Producent nie podaje, czy zastosowano w nich Radiometal czy stop żelaza z krzemem, ale brzmią dobrze.

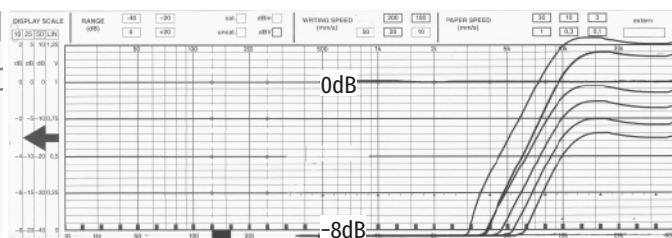
z lat 60. Najpierw wkładamy piny, potem montujemy i lutujemy cewki, a na końcu rezystory i kondensatory. Grube przewody odczepowe na cewce L3 składają się z dwóch kawałków drutu i nie przejdą one przez otwory w płytce, więc są przylutowane do pinów na górze płytki. Nie dostarczyliśmy schematu warstwy opisowej, ponieważ płytka jest bardzo wyraźnie oznaczona i można zobaczyć jej duże zdjęcia na stronie 51 Audio Out, PE z października 2019.

Tłumienie głośnika wysokotonowego

W większości systemów głośnikowych, głośnik niskotonowy jest mniej czuły niż wysokotonowy, przy czym dla danej elektrycznej mocy wejściowej emitowana jest mniejsza moc akustyczna. Czułość jest zwykle określana w dB na wat w odległości 1 m. Głośniki wysokotonowe charakteryzują się szeroką tolerancją tego parametru, ponieważ w większym stopniu niż w przypadku głośników niskotonowych wpływają na nie zmiany siły magnesu i masy membrany. Musi to być uwzględnione w projekcie zwrotnicy, aby zapewnić spójność pasma przenoszenia

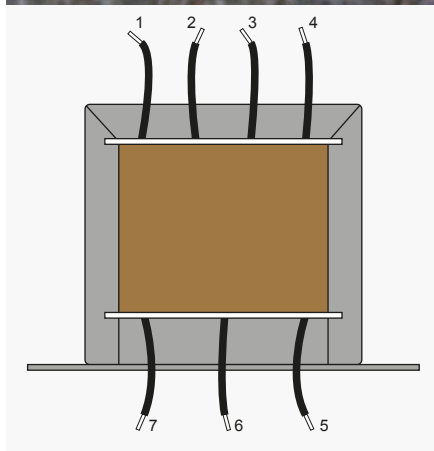


Rysunek 5. Krzywa zwrotnicy wysokotonowej z wykorzystaniem cewki Spendor. Zwróć uwagę, że obie krzywe są tłumione poniżej osi 0 dB w górnej części wykresu.



Rysunek 6. Poziome przełączanie głośnika wysokotonowego przy użyciu cewki eBay. Od góry do dołu, krzywe dla odczepów 2 do 7.

gotowych głośników. Zwykle robi się to poprzez zmianę rezystora zasilającego filtr górnoprzepustowy. W BBC LS3/5a i innych głośnikach zaprojektowanych przez BBC zastosowano bardziej złożone rozwiązanie, mianowicie autotransformator z odczepami. Transformator ten pełnił również rolę cewki indukcyjnej filtra górnoprzepustowego. W fabryce poziom głośnika wysokotonowego jest ustawiany poprzez pomiar pasma przenoszenia. W przypadku konstruktorów domowych, zwykle trzeba go ustawić ze słuchu. Dla większości ludzi o w miarę muzycznym słuchu jest zwykle dość oczywiste, czy głośnik wysokotonowy jest zbyt głośny czy zbyt cichy (w stosunku do głośnika niskotonowego). Oczywiście trzeba użyć dobrze zbalansowanego źródła muzyki.



Rysunek 7. Połączenia do cewki z odczepami z eBay.

Tłumienie głośnika wysokotonowego w zwrotnicy BBC

W oryginalnej zwrotnicy BBC LS3/5a tłumienie głośnika wysokotonowego jest ustawiane za pomocą odczepów na L3; widać to na rysunku 2. Na oryginalnych płytkach drukowanych BBC i eBay znajduje się małe półkole z sześcioma pinami o numerach od 2 do 7, z których jeden musi być podłączony za pomocą łącznika do pinu w środku półkole, aby ustawić poziom tłumienia. Na oryginalnej cewce BBC i Falcona te odczepy są również oznaczone kolorami – niestety nie są zgodne z kodami kolorów rezystorów. Kolory te są opisane poniżej. W miarę zwiększania tłumienia od odczepu 2 do 7, impedancja wejściowa rośnie, ponieważ L3 jest transformatorem, więc aby zachować tę samą krzywą filtru, kondensator wejściowy C2 musi być odpowiednio zmniejszony. Ważne jest, aby zrozumieć, że wartość C2 zależy od rzeczywistego odczepu zastosowanego na L3. Dlatego najlepiej jest mieć pod ręką wszystkie możliwe kondensatory

wymagane dla odczepów tłumiących, aby dopasować głośnik wysokotonowy do zwrotnicy. Wartości kondensatorów wymagane dla każdego odczepu to:

Odczep 2 – pomarańczowy (maksymalne podbicie) 5,5 μF (3,3 μF i 2,2 μF)

Odczep 3 – zielony 3,98 μF (3,3 μF i 680 nF)

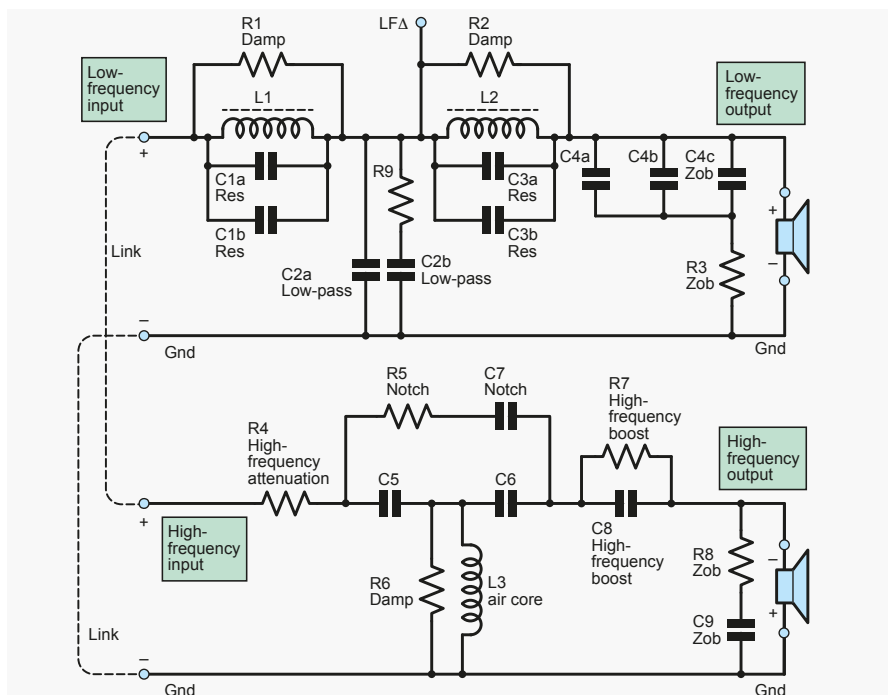
Odczep 4 – brązowy (teoretyczne wzmocnienie równe jedności) 3 μF (1,5 μF $\times$ 2)

Odczep 5 – szary 2,3 μF (2,2 μF i 100 nF)

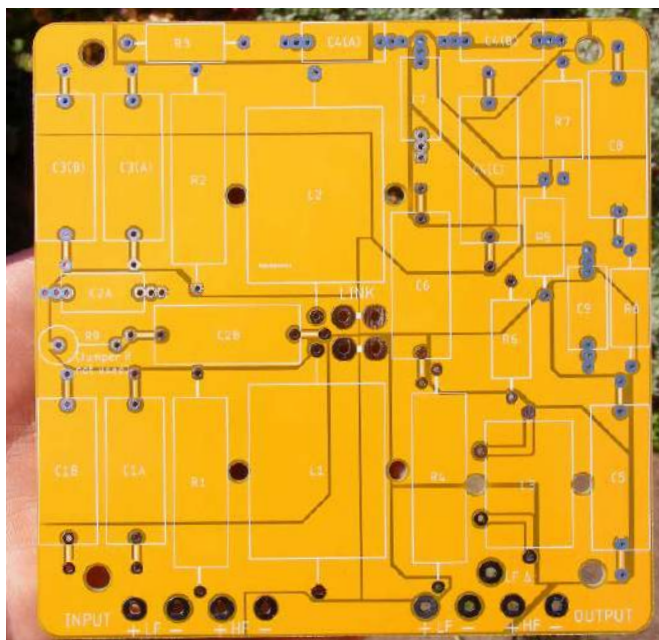
Odczep 6 – fioletowy 1,72 μF (1,5 μF i 220 nF)

Odczep 7 – żółty (minimalna głośność) 1,36 μF (680 nF $\times$ 2)

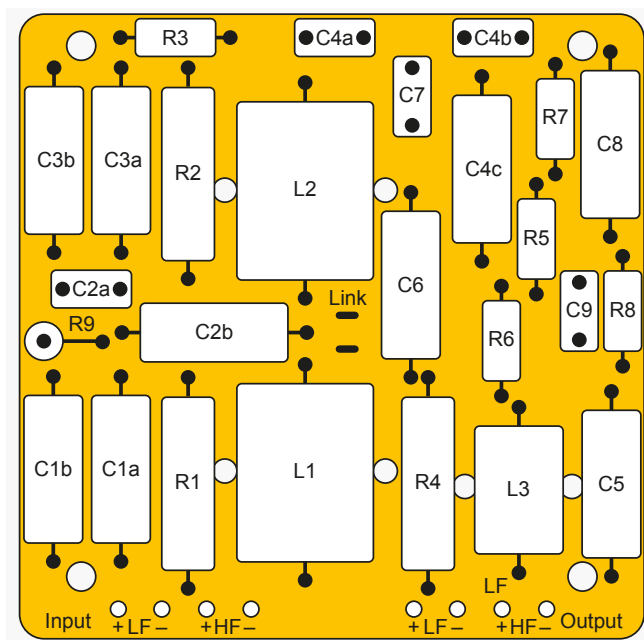
W rzeczywistości jest bardzo mało prawdopodobne, że głośnik wysokotonowy będzie miał większą moc wyjściową w punkcie podziału zwrotnicy niż głośnik niskotonowy, tak więc odczepy numer 2 i 3 prawie nigdy nie są potrzebne. Pozostaje więc tylko opcja tłumienia, a zwykle potrzeba tylko kilku dB. Mamy już prawie 10 dB tłumienia średnich częstotliwości/podbicia niskich częstotliwości z sekcji dolnoprzepustowej zwrotnicy



Rysunek 9. Pełne możliwości płytki uniwersalnej. Nie wszystkie elementy są wykorzystywane w każdym projekcie.



Rysunek 8. Płytki uniwersalnej zwrotnicy PE do mini-monitora.



Rysunek 10. Ułożenie elementów na płytce zwrotnicy.

napędzającej głośnik niskotonowy, co jest również brane pod uwagę przy dopasowywaniu poziomów względnych. W praktyce w większości zwrotnic LS3/5a najczęściej stosowane są odczepy 5 lub 6. Najlepiej zacząć od odczepu 6, co oznacza, że C2 to 1,5 μF + 220 nF połączone równolegle. Zobaczymy jak możemy uprościć sprawę na płytce uniwersalnej zwrotnicy, gdzie zastosowano znacznie prostszą (i tańszą) cewkę.

Pomiar cewek

Cewki są najmniej doskonałym elementem elektrycznym, a ich indukcyjność zmienia się odczuwalnie wraz z częstotliwością. W projektowaniu i budowie zwrotnic najlepiej traktować podane wartości z pewną dozą ostrożności i faktycznie mierzyć to co się ma. Specyfikacja BBC podaje, że indukcyjność powinna być mierzona przy 3 kHz, czyli mniej więcej przy częstotliwości podziału zwrotnicy. Jeśli masz stary mostek indukcyjny, częstotliwość może być ustawiona na 3 kHz lub też na cokolwiek innego. Próbowałem użyć cyfrowego miernika LCR Tema 72-6634, który był ustawiony na 200 Hz, co było wartością zbyt niską. Analizator impedancji Peak LCR45, który ma trzy przełączane częstotliwości, dał dobre wyniki przy 1 kHz, ale zaniżone przy 15 kHz i 200 kHz z powodu pasożytniczej pojemności uzwojenia cewki. Zmierzyłem starą dwuodczepową cewkę Spondora (mającą tylko odczepy 5 i 6) oraz cewkę z eBay'a i wygenerowałem krzywe pokazane na **rysunku 5**. Produkowane przez Volta cewki Falcona są najlepiej trafione, ale czekam na próbkę, żeby to potwierdzić. Na stronie Volta jest ciekawy artykuł o tej cewce: <https://voltloudspeakers.co.uk/not-your-average-inductor>. Zmierzyłem też cewki z eBay'a i ich wartości indukcyjności prezentuję poniżej, mierzone przy 1 kHz. Miałem tylko po dwie sztuki z każdej, co stanowi mało reprezentatywną próbkę:

L1 1,62 i 1,58 mH, 0,15 Ω , a nie 1,53 mH
L2 2,53 i 2,68 mH, 0,28 Ω , a nie 2,67 mH
L3 Cewka z odczepami 0,586 i 0,597 mH, 0,04 Ω , a nie 0,64 mH.

Otrzymane w ten sposób krzywe zwrotnicy głośnika wysokotonowego przedstawiono na **rysunku 6**. Zauważmy, że pomiary dla dwóch odczepów wychodzą poza skalę, ale to nie ma znaczenia, i tak nikt ich nie używa. Połączenia z cewką pokazano na **rysunku 7**. Na moich wyprowadzeniach umieściłem kolorowe tulejki z numerami przewodów, aby uniknąć pomyłki przed przylutowaniem do płytki.

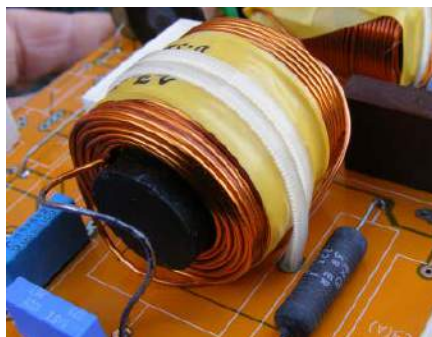
Teraz przechodzimy do naszej super uniwersalnej płytki zwrotnicy. Płytkę eBay LS3/5a została zaprojektowana specjalnie dla zwrotnicy BBC LS3/5a, ale aby wdrożyć różne wersje, które chcemy zaoferować, zaprojektowaliśmy nową płytkę PCB, która jest znacznie bardziej elastyczna w konfiguracji i wyborze komponentów – szczególnie w odniesieniu do cewek. Ponadto wszyscy wiemy, jak niepewne mogą być dostawy z eBay – dziś są, jutro ich nie ma – więc ze wszystkich powyższych powodów zleciłem Mike'owi Grindle'owi stworzenie uniwersalnej płytki PCB do zwrotnicy według mojego projektu, pokazanego na **rysunku 8**.

Wersja 3: 15 Ω LS3/5a z filtrem górnoprzepustowym 23a przy użyciu uniwersalnej płytki zwrotnicy

Będzie ona dostępna w serwisie PE PCB. Płytkę ta wykorzystuje cewki ferrytowe, które są dostępne u wielu dostawców, takich jak Volt, Falcon, Wilmslow Audio i Cyber Market. Jest to uniwersalna płytkę, która może obsłużyć wiele standardowych układów zwrotnic dwudrożnych, podsumowanych na **rysunku 1**. (Zastępuje ona również płytkę Falcon PCB7 używaną w wersji 23a, która nie jest już wymieniana). Uniwersalna płytkę zwrotnicy może być ustawiona dla wszystkich standardowych konfiguracji dwudrożnych. Dostępne są opcje dla filtrów do trzeciego rzędu, równoległych i szeregowo dostrojonych sieci „mid-dip”, rezystorów tłumiących i sieci Zolba. Wykorzystuje ona cewki ferrytowe dużej mocy w sekcji niskoczęstotliwościowej oraz małe ferrytowe (lub opcjonalnie powietrzne) w sekcji wysokoczęstotliwościowej. Jest to również idealne rozwiązanie do opracowywania własnych projektów. Podam szczegóły budowy kilku różnych wersji, wszystkie z wykorzystaniem tej samej, uniwersalnej płytki. Wymiary i otwory mocujące są określone według specyfikacji BBC, gdyż muszą być odpowiednie w celu zamontowania z tyłu zdejmowanej przedniej przegrody lub tylnego panelu. W przypadku uniwersalnej płytki zwrotnicy, jeśli stosowane są duże cewki z rdzeniem powietrznym lub laminowane, będą one musiały być zamontowane z wyprowadzeniami poza płytkę. Ogólny układ pokazany jest na **rysunku 9**, a ogólne ułożenie na **rysunku 10**. Przewidziano możliwość bi-wiringu, tak aby oddzielne przewody mogły być użyte do zasilania sekcji basów i wysokich tonów, oznaczonych na płytce odpowiednio jako „LF” i „HF”. Jeśli bi-wiring nie jest pożądanym, obie sekcje można połączyć za pomocą

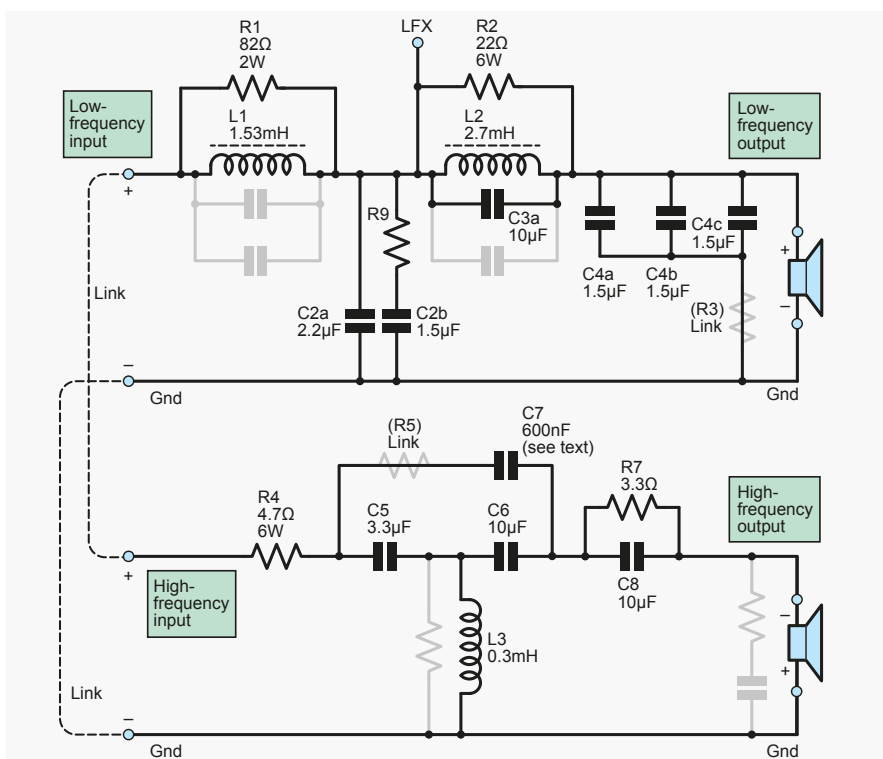


Rysunek 11. Płytkę posiada opcję biwiringu. Jeśli opcja ta nie jest używana to sekcje niskotonowa i wysokotonowa są połączone jak na rysunku powyżej.



Rysunek 12. Ciężkie cewki należy złączać opaskami zaciskowymi, w przeciwnym razie przewody mogą się złamać.

zworek na środku płytki (rysunek 11). Na płytce drukowanej, wiele konturów pozwala na zastosowanie różnych rozmiarów kondensatorów, w tym typów osiowych. Przewidziano również możliwość zrównoleglenia kondensatorów w celu uzyskania dużych/niestandardowych wartości. Na przykład, pojemność C4 może tworzyć do trzech kondensatorów: C4a, C4b i C4c.

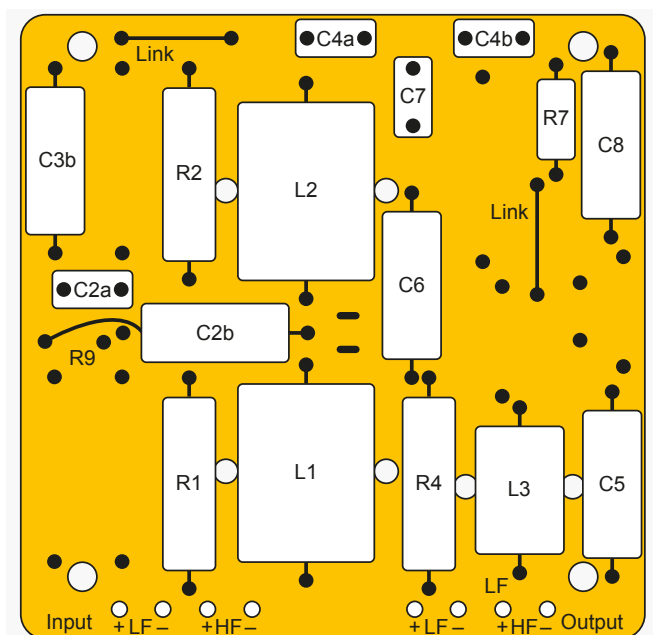


Rysunek 13. Elementy (wyróżnione) używane w zwrotnicy LS3/5a 23a zbudowanej na uniwersalnej płytce zwrotnicy. Uwaga: numery części odpowiadają układowi na z rysunku 9.

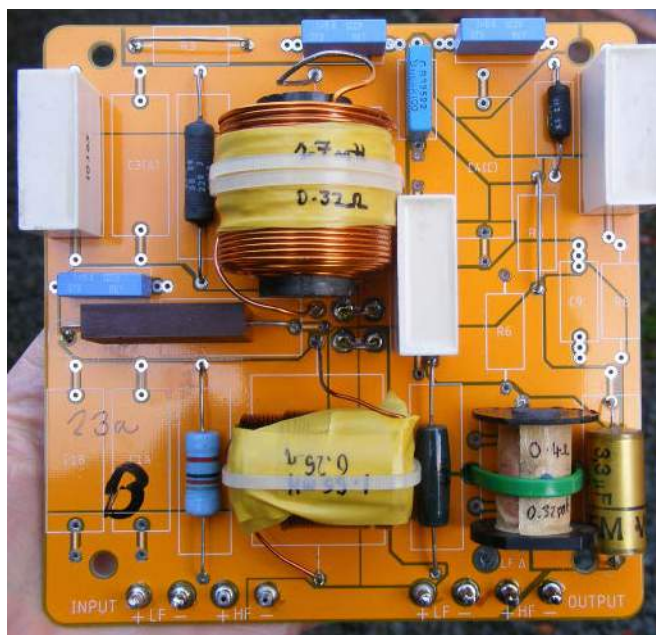
Tanie cewki

Niskoczęstotliwościowe cewki z rdzeniem ferrytowym (L1 i L2) powinny być dostosowane do dużej mocy, 50 mm × 12,5 mm i nawinięte drutem 1 mm, aby uzyskać rezystancję DC tak małą jak w cewkach z rdzeniem laminowanym. Mniejszy 25 mm rdzeń może być użyty dla cewki głośnika wysokotonowego (L3). Ze względu na swoją wagę, cewki te muszą

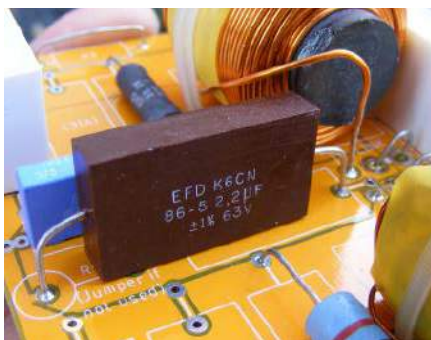
być przymocowane parami opasek kablowych przewleczonych przez 3 mm otwory w płytce (rysunek 12). Cewki o dokładnej wartości dla LS3/5a można kupić w Falconie, ale biorąc pod uwagę tolerancje, standardowe wartości cewek z serii E12, 1,5 mH i 2,7 mH są akceptowalne. Aby wyeliminować wymóg stosowania drogiej cewki (L3), sekcja górnoprzepustowa wykorzystuje układ Falcona 23a, który ma



Rysunek 14. Warstwa opisowa dla wersji 3: 15 Ω LS3/5a z filtrem górnoprzepustowym 23a z wykorzystaniem płytki uniwersalnej zwrotnicy.



Rysunek 15. Ukończona wersja 3: 15 Ω LS3/5a z filtrem górnoprzepustowym 23a z wykorzystaniem płytki uniwersalnej zwrotnicy.



Rysunek 16. Uniknięcie stosowania zworki R9 poprzez przedłużenie przewodu C2 do otworu R9.

tylko zwykłą cewkę 0,3 mH i rezystor (R4) do tłumienia głośnika wysokotonowego.

Komponenty – zwrotnica w wersji 3

Do budowy układu LS3/5a 15 Ω 23a na płytce uniwersalnej zwrotnicy nie wszystkie pozycje komponentów są wykorzystywane (pamiętaj, że jest to płytka oferująca wiele opcji), dlatego numeracja komponentów jest zmieniona w stosunku do standardu BBC w wersji 2 opisanej powyżej. Wersja 3 odnosi się do schematu obwodu uniwersalnego na rysunku 9, a użyte komponenty są zaznaczone na **rysunku 13**. Zwróć uwagę, że niektóre pozycje komponentów (R3, R5 i R9) muszą być połączone zwarte miedzianym ocynowanym 20swg. Elementy wersji 3 (ale NIE ich numery katalogowe) są takie same jak w filtrze dolno-przepustowym dla wersji 2, poza zastosowaniem cewek ferrytowych. Warstwę opisową płytki pokazano na **rysunku 14**, a gotową płytkę 15 Ω LS3/5a 23a na **rysunku 15**. Dla

C2b można rozciągnąć kondensator osiowy 2,2 μ F na tyle daleko, aby połączyć R9, unikając połączenia pokazanego na **rysunku 16**. Poziom głośnika wysokotonowego jest teraz ustawiany przez R4 i może mieć wartość od 2,7 Ω (najgłośniejszy) do 6,8 Ω (najcichszy).

Tolerancje wszystkich elementów wynoszą $\pm 5\%$. Uwaga: numery elementów odnoszą się do rysunku 9, a nie rysunku 2.

PCB. Dostępne w PE PCB Service. Zestawy komponentów będą również dostępne po ukazaniu się zestawów obudów.

Rezystory.

R1 82 Ω 2 W WW lub tlenek metalu

R2 22 Ω 6W WW

R3 zwarte

R4 4,7 Ω 3 W WW (dobrać do poziomu głośnika wysokotonowego)

R5 zwarte

R6 nieużywane

R7 3,3 Ω 3 W WW

R8 nieużywane

R9 zwarte

Kondensatory. Wszystkie 63 V minimum, 5% poliestrowe

C1a, C1b nieużywane

C2a, C4a, C4b 1,5 μ F

C2b 2,2 μ F

C3a, C6, C8 10 μ F

C3b nieużywane

C5 3,3 μ F

C7 600 nF (użyj 680 nF)

Cewki.

L1 1,5 mH 50 mm $\times$ 12,5 mm rdzeń ferrytowy, DCR 0,3 Ω

L2 2,7 mH 50 mm $\times$ 12,5 mm rdzeń ferrytowy, DCR 0,43 Ω

L3 0,3 mH rdzeń powietrzny, jeśli to możliwe dla najniższych zniekształceń, ale można użyć ferrytowego, DCR 0,4 Ω – 0,5 Ω www.voltlo-udspeakers.co.uk

Opstręt.

100 mm $\times$ 3 mm – opaski kablowe, 6 sztuk
12 $\times$ 2 mm – kołki

100 mm 20swg – cynowany drut miedziany do połączeń

Budowa – Wersja 3

Montaż jest prosty; komponenty są duże jak na współczesne standardy. Upewnij się, że są one sztywno zamocowane, aby zapobiec ich przemieszczaniu się. Włóż najpierw wyjścia kołki i zworki. Stopiona obudowa kondensatora wygląda nieprofesjonalnie a łatwo może się to zdarzyć, gdy próbujesz dostać się do niedostępnych pinów. Potrzebna będzie wyższa moc lutownicy (co najmniej 45 W) ze względu na dużą ilość miedzi na płytce drukowanej w porównaniu ze standardowymi płytkami drukowanymi.

Kolejny odcinek

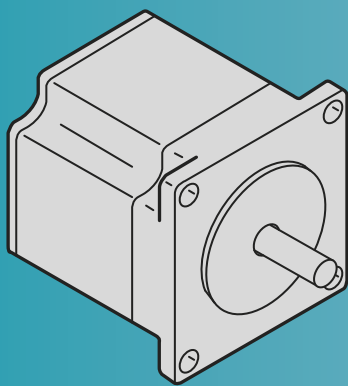
W następnym odcinku opiszemy użycie uniwersalnej płytki zwrotnicy z głośnikami Wavector, kilka dalszych usprawnień dla LS3/5a oraz kilka wskazówek pomiarowych. ■

Jake Rothman

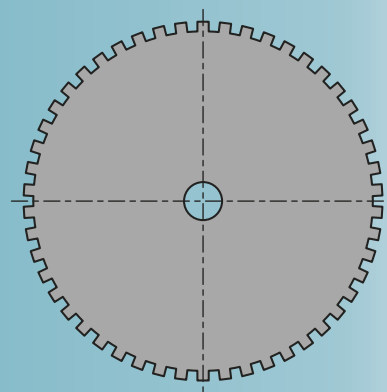
REKLAMA

m.technik
Ciekawi świata są zawsze młodzi

przejrzyj i kupisz na
www.ulubionykiosk.pl



Silniki krokowe w praktyce



Część 2: Wybór i identyfikacja silników krokowych

W tym miesiącu przyjrzymy się kwestii wyboru silnika krokowego i przedstawimy kompletny przewodnik, w jaki sposób zidentyfikować przewody w silnikach hybrydowych o 4, 5, 6 i 8 wyprowadzeniach.

Wybór silnika krokowego

Istnieje szereg parametrów, których wartości należy założyć przed wybraniem odpowiedniego silnika krokowego dla swojego projektu. Niestety, niektóre z tych parametrów mogą być trudne do dokładnego oszacowania lub obliczenia bez zagłębiania się w matematykę czy symulacje komputerowe. Wyspecyfikowanie konwencjonalnego silnika szczotkowego DC lub silnika z przekładnią zazwyczaj sprowadza się do określenia napięcia, prędkości obrotowej i momentu obrotowego. Moment obrotowy dla silnika DC będzie określony przy ustalonej liczbie obrotów na minutę i może nawet zawierać wartość momentu silnika w stanie utknięcia. Zablockowanie silnika szczotkowego DC może szybko doprowadzić do przegrzania uzwojeń, dymu i awarii. Silniki krokowe są w tej kwestii inne, gdyż są przeznaczone do pracy w zakresie od zera do maksymalnych obrotów. Maksymalny moment przy zerowych obrotach jest określany jako moment trzymania. Wraz ze wzrostem obrotów, dostępny moment obrotowy silnika spada. Spróbuj

przekręcić silnik krokowy zbyt szybko do punktu, do którego fizycznie nie może nadążyć, a może zdarzyć się wypadnięcie kroku lub zablockowanie silnika. Może to również zdarzyć się przy próbie zbyt szybkiego przyspieszenia silnika. Wracając zaś do parametrów, które jako hobbyści powinniśmy brać pod uwagę – są one następujące:

- moment trzymania,
- wymagana maksymalna ilość obrotów na minutę,
- moment krytyczny – obciążenie, jakie silnik może poruszać przy danej maksymalnej ilości obrotów na minutę,
- wielkość kroku w stopniach lub w krokach na obrót.

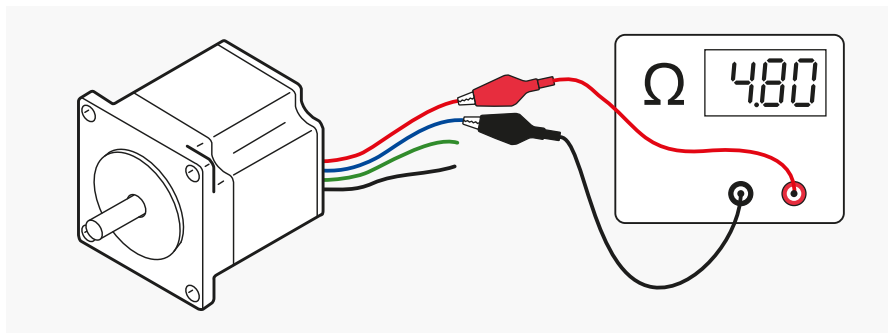
Oprócz powyższych, może być ważne także przyspieszenie silnika, w przypadku którego indukcyjność silnika jest również do rozważenia – więcej na ten temat później. Po znalezieniu silnika spełniającego potrzeby, otrzymamy wartość prądu fazowego przy określonym napięciu, która jest wykorzystywana do doboru parametrów sterownika silnika.

Kompletny przewodnik po identyfikacji przewodów w silniku krokowym

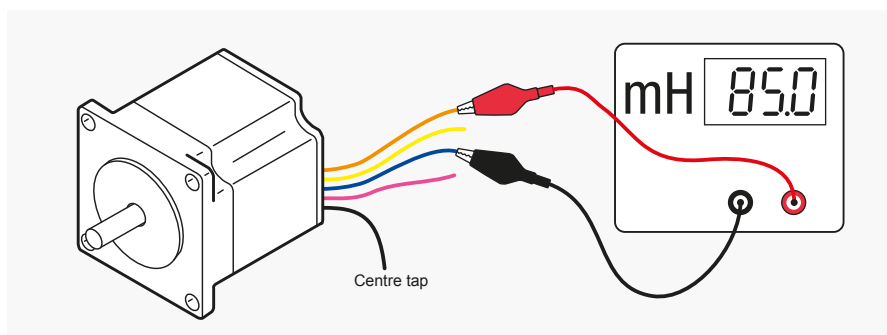
Nie masz karty katalogowej pod ręką lub masz uszkodzony silnik bez numeru modelu producenta? Nie ma problemu – przy pomocy multimetru, możesz zidentyfikować przewody w silniku z 4 i 6 wyprowadzeniami przed podłączeniem do sterownika. Z silnikiem z 5 wyprowadzeniami nie jest już tak łatwo, a dla 8 wyprowadzeń będziesz potrzebował dwukanałowego oscyloskopu!

Silniki krokowe z 4 wyprowadzeniami

Taka sytuacja zwykle oznacza, że mamy do czynienia z 2-fazowym silnikiem bipolarnym, jak na rysunku 7a. Jest to mało prawdopodobne, ale może to być również silnik o zmiennej reluktancji. Sprawdźmy to podczas testów. Używając multimetru ustawionego na zakres pomiaru rezystancji 200 Ω , podepnij czerwoną sondę testową do dowolnego z czterech przewodów. Następnie używając drugiej sondy, sprawdź rezystancję na każdym z trzech pozostałych przewodów (rysunek 11). Jeden przewód powinien wskazywać niską rezystancję, pozostałe zaś dwa przewody brak odczytu/otwarty obwód. Przewody dające niską rezystancję to jedna faza, pozostałe dwa przewody to druga faza, ale należy ten fakt potwierdzić i sprawdzić ich rezystancję, gdyż powinny mieć identyczną rezystancję jak pierwsza para. Sterownik może mieć oznaczenia 1A 1B (dla fazy 1) i 2A 2B (dla fazy 2) lub analogicznie A+ A-, B+ B-. Tak długo jak podłączasz fazy we właściwych parach, biegunowość fazy lub czy jest to faza 1



Rysunek 11. W przypadku silnika z 4 wyprowadzeniami należy najpierw znaleźć parę, która daje małą rezystancję.



Rysunek 12. Pomiar indukcyjności 4 przewodów w celu identyfikacji tych, które znajdują się na tej samej ścieżce strumienia magnetycznego.

czy 2 nie ma znaczenia. Jedyne co się stanie, to możesz uzyskać niewłaściwy kierunek obrotów, co można łatwo skorygować zamieniając miejscami dwa przewody na jednej z faz. Jeśli uzyskasz niski odczyt oporu na wszystkich czterech przewodach, wtedy masz do czynienia z silnikiem o zmiennej reluktancji.

Silniki krokowe z 5 wyprowadzeniami

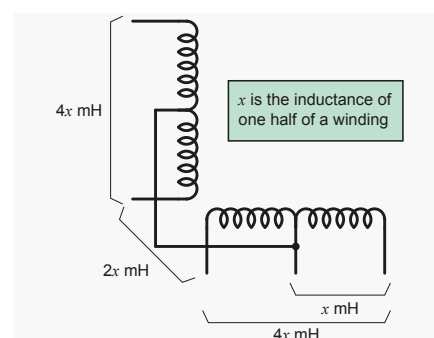
Zwykle taka sytuacja oznacza, że jest to 4-fazowy silnik unipolarny, ale może być

to też 2-fazowy silnik bipolarny z dodatkowym przewodem do obudowy. Wyeliminujmy najpierw tę drugą możliwość, podłączając czerwoną sondę testową do obudowy silnika i sprawdzając rezystancję w stosunku do wszystkich 5 przewodów. Jeśli jeden przewód daje prawie zero omów a reszta wskazuje na obwód otwarty, to masz do czynienia z silnikiem 2-fazowym, dwubiegunowym z przewodem ekranującym. Zignoruj więc przewód ekranowy i sprawdź pozostałe przewody tak, jak opisano w części dla „4 wyprowadzeń”.



Rysunek 13. Zastosowanie miernika RLC do pomiaru indukcyjności uzwojeń silnika krokowego.

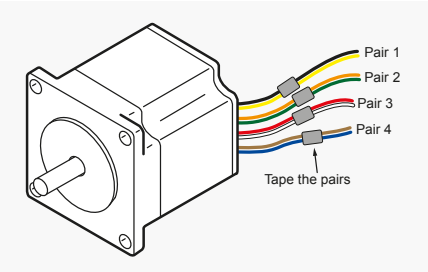
Bardziej prawdopodobne jest jednak, że jest to silnik unipolarny, jak na rysunku 7b). Ponieważ obie pary uzwojeń są połączone ze sobą za pomocą środkowego odczepu, jest to nieco trudniejsze do zidentyfikowania i będzie wymagało jak najdokładniejszego pomiaru i zanotowania wartości rezystancji. Istnieje 10 kombinacji par przewodów, które musisz zmierzyć; cztery z nich będą miały wartość o połowę niższą niż pozostałe sześć. Cztery wyniki o wartości połowy pozostałych będą miały jeden wspólny kolor przewodu – jest to odczep środkowy. Jako przykład, tabela 2 przedstawia reprezentatywny zestaw wyników. Identyfikacja środkowego odczepu jest najważniejsza dla podłączenia do sterownika, reszta może zostać znaleziona metodą prób i błędów, ponieważ tylko jedna kombinacja spowoduje, że silnik będzie prawidłowo pracował. Podczas wypróbowywania każdej z możliwości nie dojdzie do uszkodzenia sterownika lub silnika. Czytałem o technice wykorzystującej baterię lub zasilacz, aby zobaczyć, w którą stronę silnik „drży”, gdy zasilanie DC jest przyłożone między środkowym odczepem i każdym z pozostałych przewodów końcowych cewki po kolei, ale nie znalazłem wiarygodnych wyników. Uzwojenia mają charakter nie tylko rezystancyjny, ale również indukcyjny, więc zastanawiałem się, czy można zastosować także pomiar indukcyjności. Jeśli masz miernik RLC (patrz rysunek 13) to nie ma z tym problemu. Jeśli nie, to można użyć oscyloskopu z generatorem sygnału. Możesz kupić multimetr cyfrowy z zakresami indukcyjności za mniej niż 20 funtów. Warto zainwestować, jeśli masz kilka silników krokowych do zidentyfikowania. Zmierz indukcyjność sześciu kombinacji pozostałych czterech przewodów (pomiń środkowy odczep) (rysunek 12). Patrz tabela 3 dla zmierzonych indukcyjności z przykładowego silnika. Zauważ, że dwa odczyty były dwukrotnie wyższe od pozostałych. Wiemy, że każdy pomiar dotyczy dwóch połówek uzwojenia w szeregu, a cewki w szeregu



Rysunek 14. Podczas gdy rezystancje uzwojeń będą takie same, pomiar indukcyjności pozwoli nam określić wyprowadzenia.

Tabela 2. W tym przykładzie, cztery odczyty rezystancji o wartości połowy reszty mają wspólny czerwony przewód, więc czerwony przewód jest środkowym odczepem

Kolor przewodu	Kolor przewodu	Rezystancja (Ω)
Różowy	Pomarańczowy	60
Różowy	Żółty	60
Różowy	Czerwony	30
Różowy	Niebieski	60
Pomarańczowy	Żółty	60
Pomarańczowy	Czerwony	30
Pomarańczowy	Niebieski	60
Żółty	Czerwony	30
Żółty	Niebieski	60
Czerwony	Niebieski	30



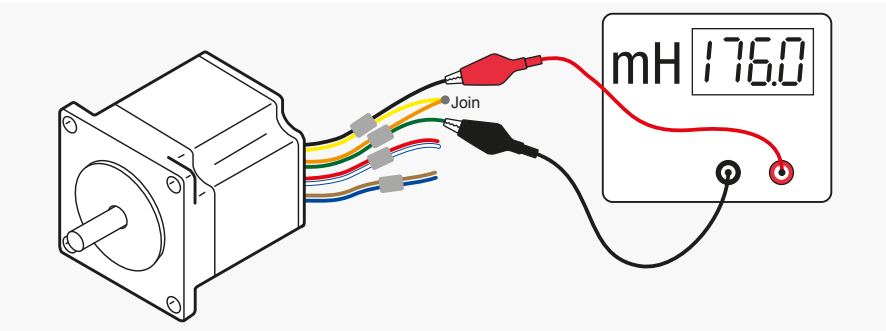
Rysunek 15. Znajdź i pogrupuj 4 pary wyprowadzeń, które dają niski odczyt rezystancji.

sumują swoje wartości. Jak więc dwa odczyty mogą być 4 razy większe od połowy uzwojenia? Odpowiedź wydaje się być właściwie prosta. W silniku krokowym, dwa poszczególne uzwojenia lub fazy nie mają wspólnej ścieżki strumienia magnetycznego, ale każde pojedyncze uzwojenie ze wspólnym odczepem

środkowym ma. Każda połowa uzwojenia testowanego silnika miała indukcyjność 42 mH, więc dwie połowy uzwojenia połączone szeregowo dają 85 mH, a jeśli są częściami tego samego uzwojenia, wartość ta jest ponownie podwojona ze względu na wspólną ścieżkę strumienia magnetycznego. To podwojenie z kolei nie występuje przy osobnych uzwojeniach (rysunek 14). Wiemy teraz, który przewód podłączony jest do środkowego uzwojenia. W naszym testowanym silniku, czerwony to środkowy odczep, różowy/pomarańczowy to końcówki jednego uzwojenia, a żółty/niebieski to końcówki drugiego. Teraz, gdy podłączysz silnik do sterownika unipolarnego, będzie on albo działał, albo nie. Jeśli nie działa, po prostu zamień jedną z par końców uzwojenia, na przykład przewód różowy i pomarańczowy.

Tabela 3. W tym przykładzie dwie wartości indukcyjności są dwukrotnie większe od pozostałych

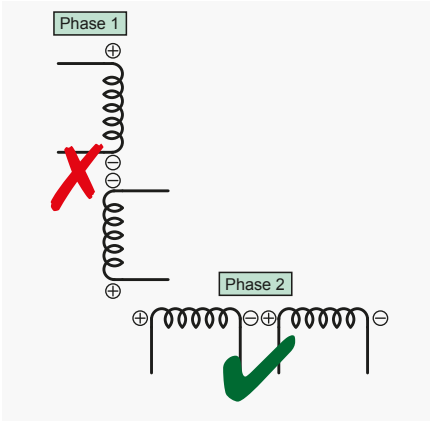
Kolor przewodu	Kolor przewodu	Induktancja (mH)
Różowy	Pomarańczowy	167
Różowy	Żółty	85
Różowy	Niebieski	85
Pomarańczowy	Żółty	85
Pomarańczowy	Żółty	85
Żółty	Niebieski	167



Rysunek 16. Zmierz indukcyjność, aby zobaczyć, które dwie pary są na tej samej ścieżce strumienia magnetycznego.

Silniki krokowe z 6 wyprowadzeniami

W przeciwieństwie do silników z 5 wyprowadzeniami, silnik z 6 nie stanowi problemu, ponieważ środkowe odczepty są wyprowadzone oddzielnie, jak to widać na rysunku 7c. Ponieważ oba uzwojenia nie są ze sobą połączone elektrycznie, poszukujemy dwóch identycznych grup odczytów rezystancji. Przypnij czerwoną sondę multimetru do któregoś z przewodów i sprawdź rezystancję pozostałych 5 przewodów – trzy z nich pokażą rozwarcie, więc odłóż je na chwilę. Dwa, które dały odczyty, będą miały albo tę samą wartość, albo jeden z nich będzie dwa razy większy od drugiego. Jeśli oba odczyty były takie same, to czerwona sonda jest podłączona do środkowego odczepu, a pozostałe dwa przewody są końcami uzwojenia. Jeśli wartość jest dwa razy większa, należy przełączyć czerwoną sondę do jednego z dwóch badanych przewodów i powtórzyć pomiary. Jeśli odczyty są teraz takie same, to czerwona sonda jest podłączona do środkowego odczepu. Jeśli nadal wartość jest dwukrotnie wyższa, wówczas pozostałe wyprowadzenie musi być środkowym odczepem – zmierz je i upewnij się. Powtórz te czynności dla pozostałych trzech przewodów odłożonych wcześniej, aż zidentyfikujesz ich środkowy odczep. Teraz znasz już wyprowadzenia dla każdego uzwojenia i wiesz, które z nich są środkowe. Gdy podłączysz silnik do sterownika unipolarnego i nie będzie on prawidłowo pracował, jak w przypadku opisanego powyżej silnika 5-przewodowego, zamień jedną parę końcówek uzwojenia. Ponieważ jest to silnik 6-cio przewodowy, możesz go sterować sterownikiem silnika bipolarnego. W przypadku bipolarnym, nie podłączaj środkowych odczepów (odetnij je lub zaizoluj indywidualnie), tylko końcówki uzwojeń. Jeśli silnik będzie chodził w złym kierunku, jak poprzednio wystarczy zamienić jedną z par.

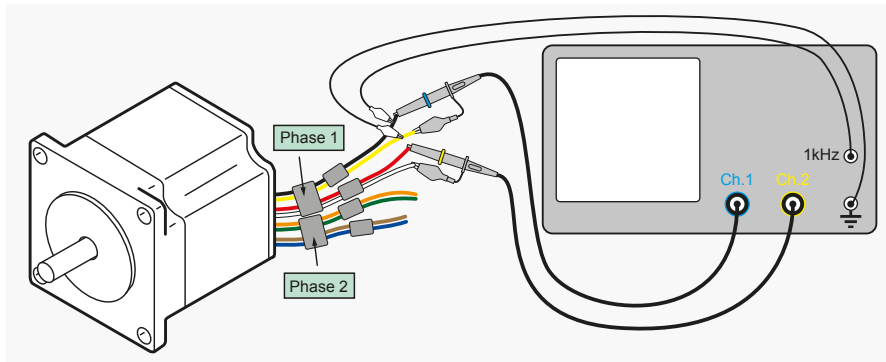


Rysunek 17. Biegunowość uzwojeń musi być prawidłowa, tutaj faza 1 jest pokazana nieprawidłowo, musi być taka jak faza 2.

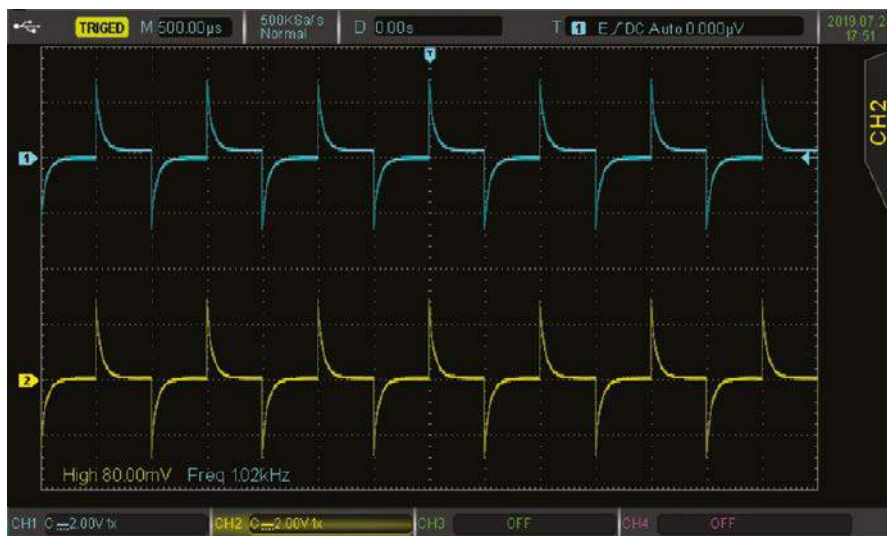
Silniki krokowe z 8 wyprowadzeniami

Teraz naprawdę mamy problem, gdyż mamy 4 uzwojenia i 28 możliwych pomiarów.

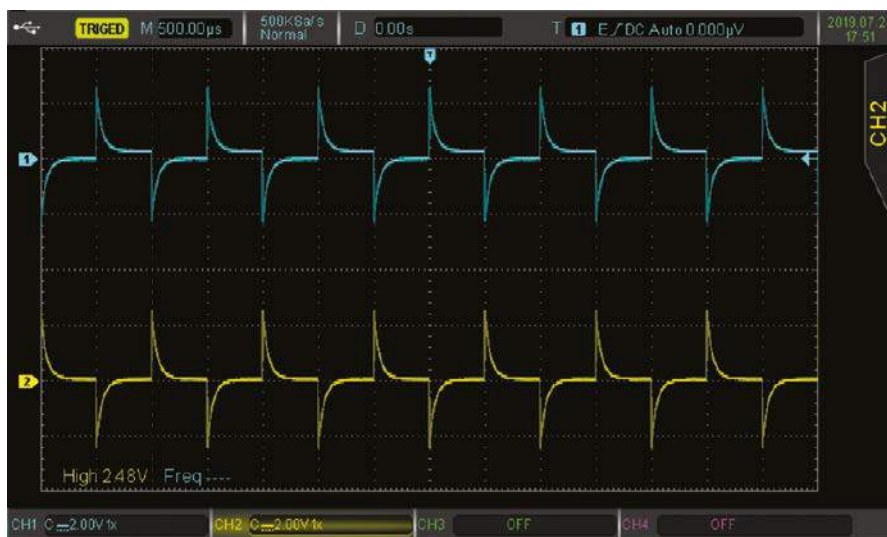
Naprawdę przyda się karta katalogowa producenta, ale zobaczmy jak wiele możemy się dowiedzieć testując! Wykonaj testy rezystancji, aby zidentyfikować i oznaczyć 4 pary,



Rysunek 18. Podaj sygnał o częstotliwości 1 kHz do pierwszej pary i zmierz na oscyloskopie napięcie indukowane w drugiej parze o tej samej fazie.



Rysunek 19. Sygnał oscyloskopowy fali prostokątnej (kanał 1) podany jest na jedną parę w fazie 1, kanał 2 pokazuje sygnał indukowany w drugiej parze fazy 1. Obie pary uzwojeń są zsynchronizowane co jest zjawiskiem prawidłowym.



Rysunek 20. Jeżeli obie pary uzwojeń są poza fazą, to mierzony sygnał na kanale 2 jest przesunięty w fazie o 180°.

z których każda daje niski odczyt rezystancji, a które możemy nazwać parami od 1 do 4 (rysunek 15). Zmierz indukcyjność 1 pary, którą możemy wykorzystać do sprawdzenia, które pary znajdują się na tej samej wspólnej ścieżce strumienia magnetycznego, stosując pomiary indukcyjności, jak opisano w testach silnika z 5 wyprowadzeniami. Połącz pary 1 i 2 szeregowo i zmierz łączną indukcyjność par 1 i 2 (rysunek 16). Jeśli indukcyjność jest 4 razy większa od indukcyjności pojedynczej pary, którą mierzyliśmy, to wiemy, że oba te uzwojenia są na tej samej ścieżce strumienia magnetycznego, którą możemy nazwać fazą 1, pary 3 i 4 muszą być wtedy fazą 2. Jeśli odczytana łączna indukcyjność jest tylko dwa razy większa od pojedynczej pary to połącz szeregowo pary 1 i 3 i ponownie zmierz łączną indukcyjność. Jeśli ponownie otrzymasz dwukrotność odczytu, spróbuj par 1 i 4. Jedna z tych kombinacji powinna dać czterokrotnie większy odczyt, którego szukamy. Wiemy teraz, które pary są w której fazie, więc najlepiej pogrupować pary na ich fazy. To czego nie wiemy, to jak zorientowane są dwa uzwojenia w każdej fazie (rysunek 17). Jest to ważne, ponieważ gdy podłączamy silnik z 8 wyprowadzeniami do sterownika to uzwojenia fazowe są łączone szeregowo lub równolegle. Pomiary rezystancyjne czy indukcyjne nie pomogą w określeniu biegunowości uzwojeń, do tego potrzebny jest oscyloskop. Oscyloskopy mają zwykle źródło sygnału 1 kHz używane do kalibracji, ale możemy je wykorzystać jako źródło sygnału w naszych próbach. Podłącz źródło sygnału 1 kHz do jednej z par uzwojeń w fazie 1 i do kanału 1 na oscyloskopie. Teraz podłącz drugą parę uzwojeń fazy 1 do kanału 2 na oscyloskopie. Jeśli oba przebiegi są zsynchronizowane, jak na rysunku 19, to można wywnioskować, że połączenia masy źródła sygnału i masy kanału 2 idą do tych samych końcówek uzwojeń, co lepiej widać na rysunku 18. Jeśli oba przebiegi są przesunięte o 180° (rysunek 20), to należy zamienić miejscami przewody kanału 2 a wtedy przebiegi będą zsynchronizowane. Powtórz ten sam test dla uzwojeń fazy 2 i zidentyfikuj uzwojenia w stopniu wystarczającym do podłączenia do sterownika.

W kolejnym odcinku

W części 3 omówimy temat sterowania silnikami krokowymi. ■

Paul Cooper

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

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.

Wykład 2

Układy scalone

A. Historia, czyli lektura lekka na rozgrzewkę

Kto wynalazł układ scalony? W podręcznikach z prawniczą poprawnością wymienia się dwa nazwiska: Jack Kilby (1923–2005) i Robert Noyce (1927–1990), gdyż po wieloletnim sporze prawnym o pierwszeństwo wynalazku układu scalonego, toczonym przez firmy Texas Instruments (Jack Kilby) i Fairchild (Robert Noyce), obie firmy zgodziły się podzielić tym wynalazkiem. Gdybym miał wskazać jedną konkretną datę narodzin mikroelektroniki, tj. elektroniki opartej na układach scalonych, to bez wahania podałbym rok 1959.

Annus Mirabilis 1959

Rok 1905, w którym Albert Einstein opublikował cztery epokowe prace (ramka obok), jest uznawany przez fizyków jako rok narodzin współczesnej fizyki i nazywany rokiem cudownym (Annus Mirabilis). Dla elektroników rok 1959 zasługuje na miano Annus Mirabilis, choć nie zdarzył się wtedy cudowny akt twórczy jednego genialnego mózgu, ale pojawiły się aż trzy wynalazki fundamentalne dla narodzin i rozwoju mikroelektroniki. Są to:

- technologia planarna, w szczególności odkrycie właściwości pasywacyjnych i maskujących oraz izolacyjnych warstwy SiO_2 na powierzchni płytki krzemowej (J. Hoerni),
- pierwszy rzeczywiście monolityczny układ scalony wykonany w płytce krzemowej technologią planarną (R. Noyce),
- pierwszy tranzystor MOS (Mohamed M. Atalla, Dawon Kahng), choć najpierw układy scalone były wytwarzane z tranzystorów bipolarnych, a tranzystory MOS dopiero dekadę później zaczęły swój triumfalny marsz jako elementy składowe układów scalonych coraz większej i jeszcze większej, wreszcie gigantycznej skali integracji.

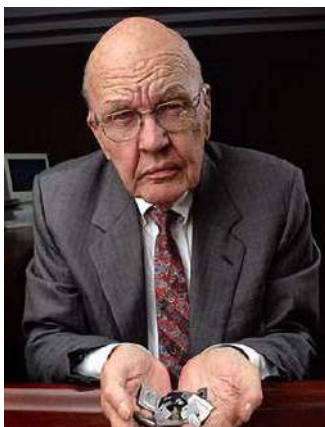
Zaraz, zaraz, a gdzie się podział wynalazek Jacka Kilby'ego? Świadomie go pominąłem, choć zgłoszenie patentowe swojego wynalazku J. Kilby zrobił też w 1959 roku, a działający model układu scalonego zademonstrował latem 1958 roku. Jednak jego układ scalony (rysunek 1) nic nie wniósł do technologii produkcji



Albert Einstein

Annus Mirabilis 1905 zaowocował czterema epokowymi publikacjami Alberta Einsteina:

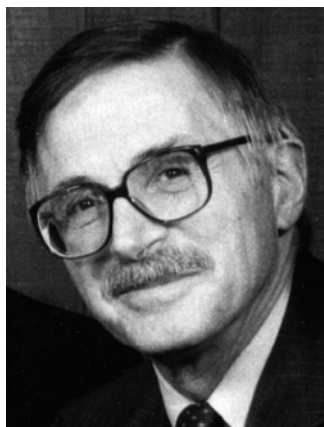
- efekt fotoelektryczny (nagroda Nobla w 1921 r.),
- model ruchów Browna,
- szczególna teoria względności,
- $E = mc^2$.



Jack St. Clair Kilby



Robert Noyce

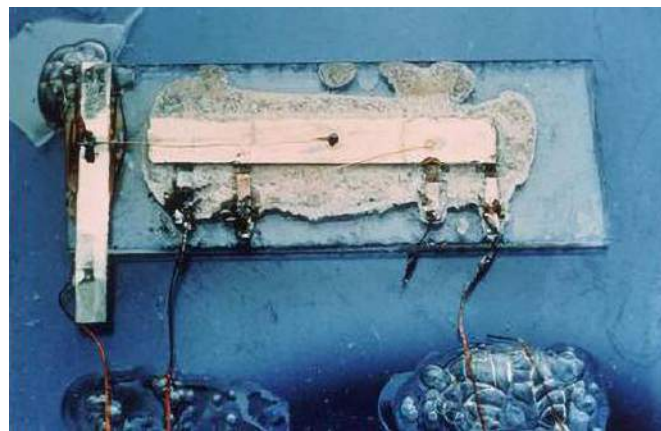


Jean Amédée Hoerni



Mohamed Martin Atalla

układów scalonych. Po pierwsze, był to model wykonany na płytce z germanu. Wprawdzie w tamtych latach była masowa produkcja tranzystorów i diod germanowych, ale produkcja układów scalonych od samego początku (od roku 1960) była oparta na płytkach krzemowych. Nigdy nie próbowano wytwarzać układów scalonych na płytkach germanowych. Nie dlatego, że Si mając znacznie większą szerokość pasma zabronionego (1,12 eV wobec 0,7 eV dla Ge) zapewnia możliwość pracy przy wyższych temperaturach (do ok. 150°C dla Si, a dla Ge do 90°C), co pozwala na rozpraszanie większych mocy, ale przede wszystkim ze względu na „nietechnologiczność” germanu. Zastosowanie technologii planarnej dla Ge byłoby bardzo trudne, bo brakuje łatwo wytwarzanej warstwy o takich właściwościach jak SiO_2 na powierzchni Si. Poza tym, układ scalony Jacka Kilby’ego nie był w istocie monolitycznym, gdyż połączenia poszczególnych elementów wytworzonych na wyizolowanych wyspach Ge były wykonane drutem złotym. Zatem wynalazek Jacka Kilby’ego nic nie wniósł do rozwoju technologii układów scalonych, ale mógł odegrać pewną rolę psychologiczną jako dowód, że możliwe jest wykonanie w jednym krysztale półprzewodnika gotowego układu (generatora przebiegu sinusoidalnego, multiwibratora lub przesuwника fazowego). Jack Kilby udowodnił eksperymentalnie słuszność idei ogłoszonej w roku 1952 na konferencji, w której Jack Kilby też uczestniczył.



Rysunek 1. Układ scalony Jacka Kilby

Na ogół uważa się, że G.W.A. Dummer na tej konferencji sformułował jedną z pierwszych idei układu scalonego. Prognozując rozwój technologii półprzewodnikowej stwierdził on: „Wraz z odkryciem tranzystora i ogólnym rozwojem technologii półprzewodnikowej zaistniały przesłanki do rozważenia możliwości realizacji urządzeń elektronicznych w postaci bloków (brył) ciała stałego, pozbawionych połączeń drutowych. Bryła ciała stałego może się składać z warstw spełniających funkcje izolujące, przewodzące, prostownicze i wzmacniające, przy czym występują bezpośrednie połączenia funkcjonalne wydzielonych obszarów ciała stałego, zawierających różne warstwy”.

Po wielu latach, w roku 2000, wynalazek Jacka Kilby’ego został doceniony przyznaniem połowy udziału w Nagrodzie Nobla w dziedzinie fizyki. Odbierając nagrodę J. Kilby zwrócił uwagę, że Robert Noyce z pewnością miałby udział w tej nagrodzie, gdyby nie zmarł na atak serca dekadę wcześniej.

Oddaliśmy część wynalazcom układu scalonego i powinniśmy wrócić do Annus Mirabilis 1959, tj. do trzech wynalazków, które zapoczątkowały erę mikroelektroniki, rozwijanej na fundamencie technologii planarnej. Zajmiemy się tym już w części merytorycznej wykładu (B. Meritum), a teraz dokończymy krótką opowieść o historii układów scalonych.

Trudno oddzielić historię rozwoju układów scalonych od historii komputerów. Nie byłoby układów scalonych, gdyby nie było komputerów. I vice versa. To trochę przesadzone twierdzenia. Może bez komputerów byłyby układy scalone, ale na pewno nie osiągnęłyby tak gigantycznej skali integracji, tj. miliardów tranzystorów w jednym chipie. Potrzeby rozwojowe komputerów stymulowały rozwój technologii półprzewodnikowej. Z kolei osiągnięcia tej



Dawon Kahng



Geoffrey William Arnold Dummer

technologii otwierały nowe, fantastyczne możliwości dla rozwoju komputerów. Spróbujmy więc obie historie, układów scalonych i komputerów, sprzęgnąć w jedną opowieść. By nie rozwlekać tematu skupimy uwagę na pięciu momentach przełomowych.

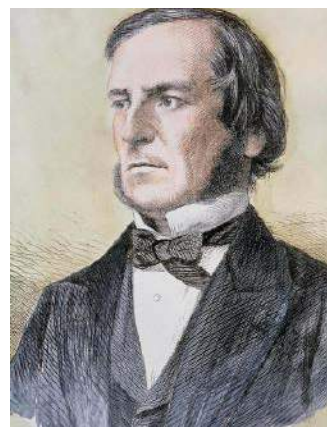
Przełom pierwszy

Historię komputerów można wyprowadzić od liczydła koralikowego, znanego na Bliskim Wschodzie już 500 lat przed naszą erą. Przez ponad 2000 lat był to najlepszy kalkulator, w Polsce dość powszechnie stosowany w biurach i sklepach jeszcze długo po drugiej wojnie światowej. Od XVII wieku triumfowały trybiki, tj. kalkulatory mechaniczne, jeszcze 50 lat temu uważane za szczyt techniki biurowo-skłepowej, szczególnie, gdy zamiast korbki miały wspomaganie silnikiem elektrycznym. To skrótowne kalendarium – ponad 2000 lat panowania liczydła, 300 lat rozwoju kalkulatorów mechanicznych i to, co wydarzyło się w ostatnich 50 latach – obrazuje niezwykle przyspieszenie tempa zmian cywilizacji technicznej. Zadajmy sobie jednak pytanie, co wspólnego mają kalkulatory z komputerem? Niewiele mają wspólnego. Nawet współczesne kalkulatory to tylko silniki przyspieszające liczenie. W komputerze chodzi o coś więcej niż tylko szybkie rachowanie, chodzi o uniwersalność w rozwiązywaniu dowolnych problemów wyrażonych przez algorytm i program, operujący poprzez instrukcje na zbiorach wprowadzonych danych. Dlatego rozwój coraz doskonalszych kalkulatorów trybikowych w znikomym stopniu przyczynił się do powstania komputerów. O wiele ważniejszym wydarzeniem historycznym na ścieżce do ery komputerów była idea **kołu binarnego**, czyli sposobu reprezentowania dowolnej liczby dziesiętnej przy użyciu tylko dwóch cyfr zero i jeden, zaproponowana przez Wilhelma Leibniza (1646–1719).

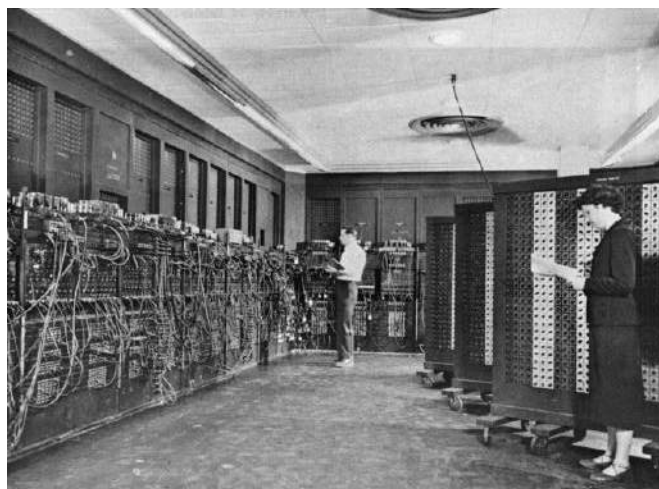
Sto pięćdziesiąt lat później George Boole (1815–1864) rozwinął ideę Leibniza w nową gałąź matematyki, zwaną algebrą Boole’a, która jest teoretycznym fundamentem całej współczesnej elektroniki cyfrowej. Sto lat później, w pierwszej połowie XX wieku sięgnięto do algebry Boole’a i zaczęły powstawać kalkulatory działające na elementach dwustanowych – najpierw na przekaźnikach, a potem na lampach elektronowych. Wreszcie w roku 1945 w USA powstał pierwszy komputer, znany pod akronimem ENIAC (Electronic Numerical Integrator And Computer). Twórcy tego komputera – naukowcy z Uniwersytetu Pensylwanii, John Mauchly i J. Presper Eckert współpracowali z genialnym matematykiem Johnem von Neumannem (Uniwersytet Princeton), który zdefiniował zasady funkcjonowania komputera, tj. sposób przechowywania programów i danych oraz ich przesyłania i przetwarzania. Tak powstała słynna **architektura von Neumanna**, stanowiąca podwaliny wszystkich nowoczesnych komputerów (oczywiście z późniejszymi modyfikacjami). **Sformułowanie architektury von Neumanna należy więc uznać za pierwszy przełom.** Natomiast pierwszy komputer ENIAC był dużym osiągnięciem konstrukcyjnym, ale nie przełomowym. Przeciwnie, demonstrował kres możliwości istniejącej wówczas technologii. ENIAC zawierał prawie 18000 lamp próżniowych, zajmował pomieszczenie o wymiarach 12 m na 6 m i ważył 27 ton. Gabaryty i potężny pobór mocy można by jakoś wytrzymać, ale nie do zaakceptowania była zawodność takiego kolosa. Jeśli czas życia lampy wynosi 10.000 godzin, to średnio co pół godziny trzeba w którymś miejscu wymienić lampę. Wtedy powstało w literaturze amerykańskiej określenie „tyranny of numbers”. Ta tyrania liczb, coraz większych liczb, wołała o nową technologię. Niebawem, bo już w roku 1947 wynaleziono tranzystor i pojawiła się nadzieja, że zamiana lamp na tranzystory pozwoli pokonać barierę rozwojową pod nazwą tyrania liczb.

Drugi przełom

Piłka znalazła się po stronie technologii. Młodym Czytelnikom wyjaśniam, że słowo technologia kiedyś odnosiło się tylko do wytwarzania produktów materialnych, w tym wypadku – podzespołów elektronicznych. Lata pięćdziesiąte to okres zamiany lamp elektronowych na tranzystory – elementy znacznie mniejsze, bardziej niezawodne i pobierające mniej mocy niż lampy. Konstruowano komputery o coraz większych zdolnościach obliczeniowych, ale rosnąca liczba połączeń między elementami nieuchronnie pogarszała wskaźniki niezawodności i ograniczała szybkość przetwarzania informacji ze względu na transmisję sygnału przewodami lub ścieżkami, czyli liniami o rozproszonych pasożytniczych parametrach L, C. W latach sześćdziesiątych, po uruchomieniu masowej produkcji układów scalonych serii TTL (Transistor-Transistor Logic), konstruktorzy komputerów dostali do dyspozycji nowe cegielki – zamiast z tranzystorów, czyli kluczy dwustanowych, mogli teraz projektować komputery z bramek logicznych. Najpierw



George Boole (1815–1864) stworzył nową gałąź matematyki, zwaną algebrą Boole’a, która jest teoretycznym fundamentem całej współczesnej elektroniki cyfrowej



Rysunek 2. ENIAC – pierwszy komputer z roku 1945



John von Neumann, węgierski matematyk pracujący w USA – twórca architektury komputerów

były to układy zawierające kilka tranzystorów (mała skala integracji – SSI), z czasem pojawiły się układy (m.in. liczniki, rejestry przesuwne) zawierające setki tranzystorów (średnia skala integracji – MSI). Jednak ciągle były to układy działające na tranzystorach bipolarnych, które z istoty ich działania pobierają znaczne prądy. Dlatego w latach pięćdziesiątych i sześćdziesiątych komputery nadal miały gabaryty sporych szaf z kosztownymi systemami zasilania i chłodzenia. Pojawienie się układów scalonych po roku 1960 nie spowodowało od razu przełomu w rozwoju komputerów. Ciągłe komputer był drogim urządzeniem, przeznaczonym tylko do rozwiązywania problemów naukowych, a poza ośrodkami naukowymi stosowanym tylko

przez instytucje rządowe, militarne lub wielkie korporacje. Zadaniom projektowania takich urządzeń i wdrażania ich do produkcji mogły podołać tylko duże zespoły konstrukcyjne realizujące programy rządowe. Na przykład w RWPG (Rada Wzajemnej Pomocy Gospodarczej – zrzeszająca kraje obozu komunistycznego, na czele z ZSRR) istniał program pod nazwą RIAD, który obejmował sześć typów komputerów, różniących się mocą obliczeniową. Dodajmy, że w tamtym okresie opóźnienie technologiczne ZSRR w mikroelektronice nie było bardzo duże, można je szacować na ok. 5 lat. Dopiero rok 1970 – opanowanie technologii MOS (znanej teoretycznie od dziesięciu lat) i rozpoczęcie produkcji mikroprocesorów i mikrokomputerów (dziś nazywanych mikrokontrolerami) jednoukładowych otworzyło drogę do nowego etapu rozwoju komputerów. **To był drugi przełom.** W latach siedemdziesiątych jak grzyby po deszczu powstawały dziesiątki różnych konstrukcji komputerów, dostępnych dla użytkownika indywidualnego. Sukces technologów – masowa produkcja układów scalonych dużej skali integracji (LSI – duża skala integracji, tysiące tranzystorów w jednym chipie) – zrewolucjonizował informatykę. Komputery trafiły pod strzechy. Może nie dosłownie, bo posiadanie komputera w gospodarstwie domowym nie wydawało się wtedy tak niezbędne jak posiadanie telewizora, pralki czy lodówki. Był to raczej gadżet dla „wykształciuchów”. Ale rozpoczął się proces osławiania społeczeństwa z pojęciem „komputer osobisty”.

Trzeci przełom

W połowie lat siedemdziesiątych, gdy pojawił się niezwykle popularny procesor Intela 8080, zaczęły powstawać hobbistyczne konstrukcje komputerów, a niektóre z nich rozwinęły się w duże serie produkcyjne. Pierwszym sukcesem rynkowym był Altair 8080, zbudowany przez Eda Roberta. Największym beneficjentem tych nowych możliwości biznesowych, jakie otwarły się przed konstruktorami garażowymi, stał się duet Steve Wozniak i Steve Jobs, których komputer Apple I, skonstruowany przez Wozniaka na bazie mikroprocesora 6502 zrobił furorę. A kolejna wersja Apple II wydzignęła firmę garażową Apple do poziomu jednej z największych firm komputerowych na świecie. W tym czasie powstało kilkadziesiąt różnych konstrukcji komputerów osobistych (rysunek 3). Została zagrożona pozycja IBM, firmy dominującej dotychczas w przemyśle komputerowym. Reakcją IBM było wejście na rynek z IBM Personal Computer, opartym na mikroprocesorze Intel 8080. Był początek lat osiemdziesiątych. Rynek komputerów osobistych rozwijał się chaotycznie i rozbieżnie w sferze oprogramowania. Dziesiątki różnych modeli komputerów były niezgodne programistycznie, każdy producent stosował swoją wersję oprogramowania (na ogół różne warianty języka BASIC) ściśle związaną z konstrukcją komputera. Programy pisane przez użytkowników określonego typu komputera nie działały na komputerach innej konstrukcji bez uciążliwej konwersji.

Dojrzała potrzeba zmiany koncepcji software'u. Rozwiązanie wymyślił Gary Kindall, twórca programu operacyjnego o nazwie CP/M, który działał jako pośrednik między programami użytkownika a sprzętem maszyny. Gdyby wszystkie komputery wyposażać w identyczny program operacyjny, to mogłyby uruchamiać identyczne programy użytkownika bez żadnych modyfikacji. Ten rewolucyjny krok wykonał IBM, jednak nie skorzystał z programu Kindalla, tylko kupił licencję na program operacyjny DOS od maleńkiej firmy Microsoft, należącej do Billa Gatesa. Sprytny biznesmen Bill Gates zachował prawa do innej wersji programu operacyjnego (MS-DOS) i inni producenci komputerów, którzy uruchamiali produkcję „klonów” komputera IBM, kupowali licencję na program



Rysunek 3. Komputery osobiste drugiej połowy lat siedemdziesiątych – od lewej: Commodore PET, Apple II, TRS-80

Skala integracji

Ze względu na liczbę tranzystorów w pojedynczym układzie scalonym przyjęto następujące określenia skali integracji:

Akronim	Liczba tranzystorów
SSI – Small Scale Integration	1÷100
MSI – Medium Scale Integration	100÷1000
LSI – Large Scale Integration	1000÷10000
VLSI – Very Large Scale Integration	10000÷1 milion
ULSI – Ultra Large Scale Integration	1 milion÷10 milionów
GSI – Giant Scale Integration	ponad 10 milionów

operacyjny MS-DOS z firmy Microsoft. Na tym wyrosła potęga Microsoftu. Był to **trzeci przełom** w historii komputerów. Dotyczył software'u. Otworzył możliwości potężnego rozwoju rynku programów użytkownika, które mogły teraz działać na dowolnych komputerach wyposażonych w system operacyjny DOS.

Przełom czwarty i piąty

Musimy skrócić tę opowieść o historii komputerów, bo zdaje się, że oddala nas od tematu wykładu, czyli układów scalonych. **Przełom czwarty** to wprowadzenie ikon obrazkowych, które można było przesuwac za pomocą myszy. Najpierw dokonał tego Steve Jobs w komputerze Mac Intosh, w roku 1984 (właściwie rok wcześniej w komputerze Lisa, który jednak okazał się komercyjną kląpą). Zaraz potem Microsoft zrobił to samo, wprowadzając Windows jako ulepszoną wersję programu MS-DOS. Komunikacja człowieka z komputerem przy pomocy obrazków, a nie przez wpisywanie poleceń tekstowych, otworzyła drogę do totalnego upowszechnienia komputera jako urządzenia obsługiwanego niemal intuicyjnie, nawet przez dziecko nie umiejące pisać. Trzeba było tylko zamienić myszkę ekranem dotykowym. Trzeba było też sprzęgnąć komputery w jedną sieć. Trzeba było wynaleźć Internet i sieć WWW. I to był **piąty przełom**, który wraz z przełomem czwartym wprowadził nas w cywilizację smartfona – komputera udającego telefon, bez którego nie potrafią żyć miliardy ludzi na świecie. Pozornie przełomy trzeci (DOS), czwarty (ikony) i piąty (Internet) nie mają związku z układami scalonymi. Nastąpiła też zmiana pojęcia technologia. Teraz Google, Facebook, Twitter to firmy technologiczne. Pamiętajmy jednak, że zaplecze rozwoju tych firm są osiągnięcia mikroelektroniki, czyli technologii produkcji układów scalonych. Nie byłoby smartfona, gdyby nie działało prawo Moore'a, gdyby nie było układów scalonych, zawierających miliony tranzystorów. A czym jest „chmura”, w której z tak bezbrzeżnym zaufaniem przechowujemy wszystkie nasze zasoby ważnych informacji, zdjęcia rodzinne itp. Pięknie wymyślono z tą „chmurą” – wyobraźnia podsuwa nam bezpieczną opiekę potęgi Niebios, a to potęga serwerów Google'a, których nie byłoby bez układów scalonych zawierających miliardy tranzystorów.

Średnica płytek Si osiągnięta w miarę rozwoju technologii

50 mm – 1965 r.
100 mm – 1975 r.
125 mm – 1981 r.
150 mm – 1987 r.
200 mm – 1992 r.
300 mm – 2000 r.
450 mm – obecnie

B. Meritum, czyli sedno tematu

Zajmiemy się najpierw technologią wytwarzania układów scalonych, oczywiście tylko w zakresie elementarnym, by zaspokoić naturalną ciekawość Czytelników, którzy przecież w EdW oczekują wiedzy o aplikacjach układów scalonych, a nie o ich wytwarzaniu. Dlatego dalej zajmiemy się klasyfikacją układów scalonych pod względem ich funkcji aplikacyjnych.

Proces produkcji układów scalonych składa się z trzech etapów, realizowanych na różnych liniach technologicznych, a nawet w różnych firmach. Są to:

- wytwarzanie płytek krzemowych,
- wytwarzanie chipów w płytkach krzemowych,
- montaż chipów w obudowach.

Wytwarzanie płytek Si

Płytki krzemowe stosowane do produkcji układów scalonych muszą spełniać wiele parametrów, przede wszystkim muszą być wolne od zanieczyszczeń i muszą mieć idealną, ciągłą budowę krystaliczną, czyli muszą mieć strukturę monokrystaliczną. Poziom wymaganej czystości określa się jako „jedenaście dziewiątek”, tj. zawartość atomów krzemu ma wynosić 99,999999999%, czyli jeden atom pierwiastka obcego przypada na 100 miliardów atomów Si. Krzemu każdy kraj ma pod dostatkiem – występuje w ilości 28% w skorupie ziemskiej jako składnik pospolitego piasku i skał (krzemionka). Mimo powszechnej dostępności surowca wyjściowego krzem czysty jest materiałem drogim ze względu na koszt złożonego procesu oczyszczania. Czysty materiał polikrystaliczny jest przetapiany i przetwarzany w monokryształ. Najczęściej monokryształy krzemu są wytwarzane metodą wyciągania z fazy ciekłej, znaną na całym świecie jako **metoda Czochralskiego**. (Jan Czochralski, polski uczone, profesor Politechniki Warszawskiej, zastosował tę metodę w 1916 r. i opublikował ją w 1918 r.). Nazwisko Czochralski jest znane w każdym zakątku świata, gdzie są elektronicy. Ma najwięcej cytowań spośród wszystkich polskich uczonych. Jest w każdym podręczniku o technologii półprzewodników (widziałem również w podręczniku chińskim pisanym hieroglifami).

Proces wyciągania monokryształów metodą Czochralskiego przedstawia **rysunek 4**. Zarodek kryształu jest podnoszony powoli i jednocześnie obraca się kilkanaście razy na minutę. Na styku zarodka z roztopionym krzemem narastają stopniowo kolejne warstwy kryształu. W ramce obok podano jak z biegiem lat zwiększano średnicę monokryształów Si wytwarzanych metodą Czochralskiego. Po pocięciu walca na płytki o grubości 0,5÷2 mm następuje szlifowanie i polerowanie powierzchni płytek – mechaniczne i chemiczne, do osiągnięcia lustrzanej gładkości (**rysunek 5**). W produkcji płytek Si specjalizuje się na świecie kilkadziesiąt firm. Produkcją chipów zajmują się inne firmy, które kupują gotowe płytki Si o określonych parametrach.



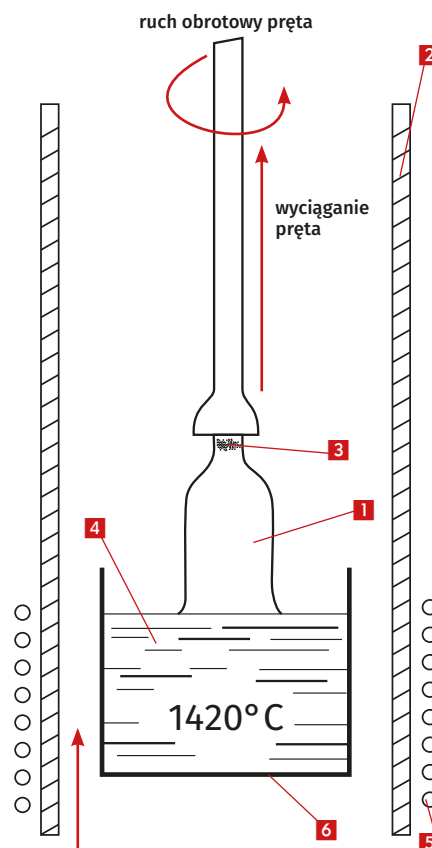
Jan Czochralski (1885–1953) – polski uczone, wynalazca powszechnie stosowanej na świecie metody otrzymywania monokryształów krzemu

Produkcja chipów

Swego rodzaju fenomenem w historii rozwoju technologii półprzewodnikowej jest długowieczność technologii planarnej, wynalezionej w 1959 roku. Oczywiście, linia technologiczna produkcji współczesnych układów scalonych o ścieżkach szerokości kilku nanometrów (**rysunek 6**) nie przypomina w niczym hali produkcyjnej sprzed 50 lat, w której powstawały układy scalone o ścieżkach szerokości kilku mikrometrów. Zmieniły się urządzenia technologiczne, wymogi czystości atmosfery, w której są prowadzone poszczególne procesy doprowadziły do całkowitego wyeliminowania możliwości kontaktu człowieka z płytkami podlegającymi obróbce. A dodajmy, że cały proces wytwarzania chipów w płytce krzemowej dla najbardziej złożonych układów, wymaga wykonania ok. 400 do 500 operacji technologicznych, co trwa do dwóch – trzech miesięcy. W tym czasie płytki automatycznie wędrują od urządzenia do urządzenia, cały czas w idealnie czystej atmosferze ochronnej. Jednak podstawowe zasady technologii planarnej zostały zachowane do dziś. Istotą tej technologii jest zastosowanie fotolitografii oraz warstwy SiO_2 na powierzchni Si do maskowania i lokalnego domieszkowania płytki podłożowej. W powtarzających się sekwencjach procesów fotolitografii, nanoszenia warstw, trawienia okien i domieszkowania lokalnego, w cienkiej warstwie przypowierzchniowej o głębokości kilku mikrometrów (ok. 1% grubości płytki) powstają struktury chipów, zawierające wiele warstw metalizacji, tj. ścieżek połączeń między właściwymi obszarami domieszkowanego Si. Wszystkie procesy są wykonywane tylko z jednej strony płytki Si, której 99% grubości pełni rolę nośnika mechanicznego, a tylko 1% przy jednej powierzchni zawiera niezwykle skomplikowany układ warstw tworzących strukturę chipa.

Skoro idea technologii planarnej ciągle pozostaje ta sama, to spróbujmy prześledzić pojedynczy proces lokalnego domieszkowania według rysunku pokazującego zasadnicze fazy klasycznego procesu planarnego (**rysunek 7**).

Tak to wyglądało przed laty, gdy ścieżki miały rozmiary mikrometrowe. W miarę zmniejszania wymiaru charakterystycznego do setek, a później dziesiątków i wreszcie kilku nanometrów niezbędne były przede wszystkim zmiany sposobów wykonywania fotolitografii. Zaczniemy od naświetlania. W latach 60. stosowano światło widzialne o długości fali $\lambda = 435$ nm (fiolet). Próba wykonywania ścieżek o wymiarze np. 130 nm przy stosowaniu światła o $\lambda = 435$ nm byłaby podobna do próby namalowania linii o szerokości 1 mm przy pomocy pędzla o średnicy włosa 5 mm. W kolejnych latach wprowadzano standardy coraz krótszych fal zakresu ultrafioletowego (UV): 365 nm, 248 nm, 193 nm, aż wreszcie osiągnięto $\lambda = 13,5$ nm (EUV – extreme ultraviolet). Stosowanie coraz krótszych fal UV wiąże się z szeregiem problemów optycznych związanych z soczewkami i szkłem fotomasek. Piętrzą się też problemy z centrowaniem fotomasek przy kolejnych procesach fotolitografii. Wreszcie zniknęły fotomaski i pojawiło się urządzenie o nazwie stepper, które bezpośrednio naświetla chip za chipem, skanując krokowo całą powierzchnię płytki Si. Pojawił się problem wydajności steppera, który na jednej płytce musi wykonać setki lub nawet tysiące kroków i naświetleń. Najlepiej radzi sobie holenderska firma ASML, produkująca steppery na EUV o wydajności 200 płytek/godzinę. Kilka lat temu uważano, że brak możliwości dalszego postępu z rozdzielczością litografii zatrzyma ciągły marsz ku coraz większej skali integracji według prawa Moore’a. Na razie to się nie stało, choć chyba osiągnięto kres możliwości litografii optycznej. Poza tym, na horyzoncie jest następna generacja litografii z zastosowaniem promieni X o długości fali poniżej 1 nm. Prowadzone są też prace nad litografią promieniami elektronowymi i jonowymi. Jednak od demonstracji laboratoryjnych do opłacalnej ekonomicznie, czyli wydajnej litografii prowadzi trudna droga. Wyraźnie nad dalszym postępowaniem i nad prawem Moore’a wiszą ciemne chmury. Może jednak uda się jeszcze trochę poprawić rozdzielczość litografii i osiągniemy wymiary tranzystora tak małe, że będzie składał się z 3 atomów, po jednym na dren, bramkę i źródło (myślę,



gaz obojętny

Rysunek 4. Urządzenie do wyciągania krystalizatorów z fazy ciekłej metodą Czochralskiego: 1 – wałek monokrystaliczny; 2 – rura kwarcowa; 3 – zarodek krystalu; 4 – roztopiony krzem; 5 – zwojnica indukcyjna; 6 – tygiel grafitowy lub kwarcowy



Rysunek 5. Walce i otrzymane z nich płytki monokrystalicznego Si (źródło: <https://anyisilicon.com/silicon-wafer/>)



Rysunek 6. Clean room – obecność człowieka sprowadzona do minimum (źródło: wikipedia.org)

że zażartowałem, ale pewny tego nie jestem). Wspomnijmy jeszcze, że postęp w rozwoju technologii planarnej wiąże się nie tylko z litografią, ale też z wprowadzeniem takich procesów jak suche trawienie anizotropowe, czy implantacja jonów. Wróćmy do opisu procesów technologicznych od momentu, gdy płytki Si przeszły wszystkie operacje wsadowe (wykonywane dla setek lub tysięcy chipów na całej płytce) i mamy płytkę z gotowymi chipami (rysunek 8).

Testowanie, cięcie i montaż chipów w obudowie

Pozostaje pociąć płytkę na oddzielne chipy, umieścić każdy chip w obudowie i wykonać testowanie. Gdybyśmy tak zrobili, to kilkadziesiąt procent gotowych, zamkniętych w obudowach układów trafiłoby do odpadów. Żeby uniknąć strat związanych z montażem w obudowie nie działających układów, jeszcze na płytce wykonuje się tzw. testowanie ostrzowe każdego chipa i układy wadliwe są zaznaczane, by po pocięciu płytki zostały odrzucone (rysunek 9). Procent działających poprawnie chipów, odniesiony do całkowitej liczby chipów na płytce, jest nazywany **uzyskiem**. Uzysk zwykle zawiera się w przedziale od 30% do 80%.

Po testowaniu ostrzowym następuje dzielenie płytki na chipy. Stosuje się następujące metody:

- zarysowanie linii diamentem i przełamanie (jak przy cięciu szyby),
- cięcie piłą mechaniczną,
- cięcie laserem,
- cięcie plazmą.

Ostatnie operacje technologiczne polegają na umieszczeniu struktury układu scalonego w obudowie. Jest to tzw. mikromontaż, przebiegający na różne sposoby, w zależności od rodzaju obudowy, a jest ich dużo, jak widać na rysunku 10. W każdym przypadku wykonywane są połączenia pól kontaktowych struktury układu scalonego z wyprowadzeniami obudowy i zamknięcie hermetyczne obudowy. (W języku potocznym to zdanie mogłoby brzmieć: „... wykonywane są połączenia padów chipa z pinami obudowy...” Muszę przyznać, że obrona przed slangiem angielskim w języku polskim chyba nie ma szans. Przez długie lata nie używaliśmy słowa chip, tylko ten mały kryształek układu scalonego na płytce był nazywany strukturą lub mikroplątką. Podałem się i używam słowa chip [lub czip], które stało się zresztą wieloznaczne i w języku potocznym odnosi się na ogół do gotowych układów scalonych w obudowach, a struktura, czy mikroplątką ma też angielską nazwę die).

Najogólniej obudowy układów scalonych można podzielić na przeznaczone do montażu przewlekane (Through Hole) i przeznaczone do montażu powierzchniowego (Surface Mount). W praktyce hobbista najczęściej spotyka się z układami w obudowach DIP (Dual Inline Packages).

Klasyfikacja układów scalonych

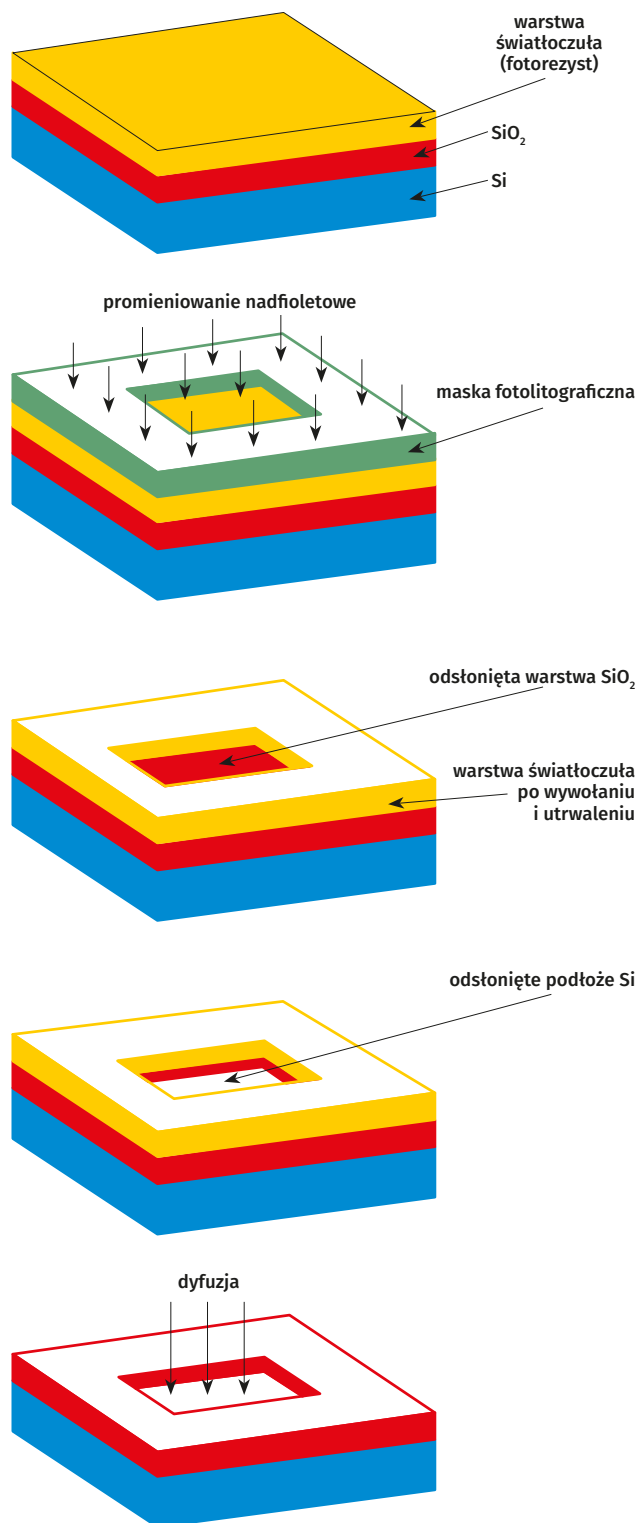
Ograniczmy się tylko do dwóch kryteriów klasyfikacji:

- rodzaj technologii,
- funkcje aplikacyjne.

Rodzaj technologii

Dotychczas zajmowaliśmy się wyłącznie technologią układów scalonych monolitycznych. Istnieją jeszcze trzy rodzaje technologii:

- cienkowarstwowa,
- grubowarstwowa,
- hybrydowa.



Rysunek 7. Zasadnicze fazy procesu planarnego

Może się wydawać dziwne, dlaczego klasyfikację rodzajów technologii układów scalonych nie podałem na początku wykładu. Powód jest bardzo prosty, otóż ponad 99% układów scalonych na rynku to układy monolityczne. Zatem informacje o pozostałych technologiach mają charakter suplementu.

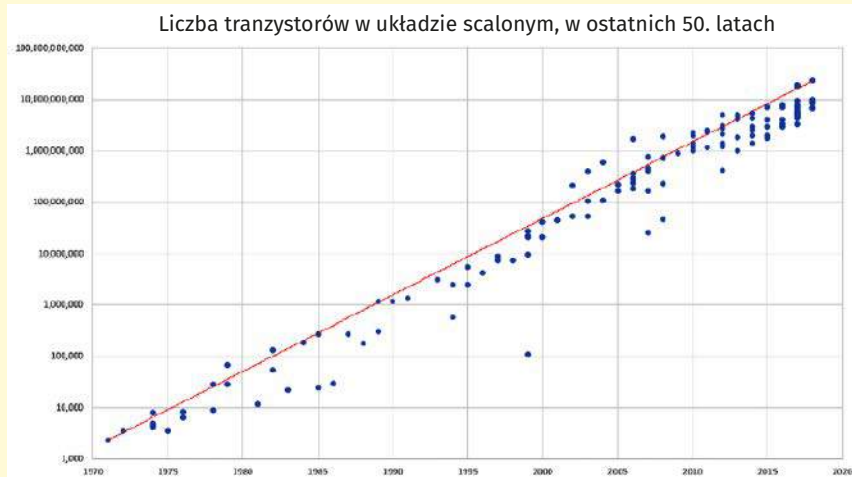
Pół wieku temu nie było to oczywiście i technologie cienkowarstwowa oraz grubowarstwowa były brane pod uwagę jako alternatywne kierunki rozwoju mikroelektroniki w stosunku do układów monolitycznych, których niektóre ograniczenia wydawały się wtedy istotne. Jakże to ograniczenia? Otóż w kryształach Si łatwo można wytwarzać tranzystory i diody, ale trudno sobie poradzić z elementami biernymi – R, L, C. O indukcyjności można w ogóle zapomnieć, nie ma sposobu na wykonanie cewek, poza niewielkimi spiralkami o znikomym małej indukcyjności. Z kondensatorami można sobie radzić w zakresie do kilkudziesięciu pF, można też w niektórych przypadkach wykorzystać pojemność złącza p-n. Jako rezystory w bardzo ograniczonym zakresie można wykorzystywać warstwy półprzewodnika, co jednak zajmuje dużo miejsca w strukturze i rezystancja nie dość, że ma słabą tolerancję, to jeszcze silnie zależy od temperatury. Jednak kreatywne rozwiązania schematowe, szczególnie w układach cyfrowych, pozwoliły niemal całkowicie wyeliminować potrzebę stosowania elementów R, L, C w układach monolitycznych.

W technologiach warstwowych, zarówno cienkowarstwowej jak i grubowarstwowej sytuacja jest odwrotna. Dość łatwo wykonuje się elementy bierne (poza indukcyjnością), ale tranzystory i diody wytwarza się w oddzielnym procesie i montuje się indywidualnie w strukturze układu warstwowego. Zatem w istocie są to układy hybrydowe.

W technologii cienkowarstwowej podstawowym procesem jest nanoszenie warstw w próżni na podłożu szklanym lub ceramicznym. Metodą naparowywania próżniowego otrzymuje się warstwy przewodzące, oporowe i dielektryczne takich materiałów jak aluminium, nichrom, złoto, tlenek krzemu.

W technologii grubowarstwowej metodą sitodruku nanosi się na podłoże izolacyjne warstwy przewodzące, oporowe i dielektryczne. Podobnie jak w technologii cienkowarstwowej otrzymuje się w ten sposób dobrej jakości rezystory i kondensatory, tranzystory i diody są natomiast wytwarzane w oddzielnym procesie i montowane indywidualnie w strukturze układu grubowarstwowego.

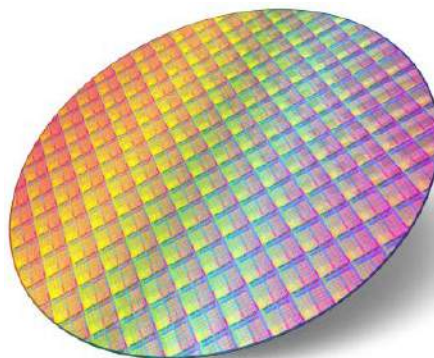
Obie wyżej wymienione technologie znajdują obecnie zastosowanie w wybranych, szczególnych aplikacjach, jako technologie hybrydowe, nazywane „multi chip”. Jedną z takich aplikacji są wzmacniacze mocy od 5 W nawet do 50 W. Innym zastosowaniem są układy SoC (System on Chip), w których do podłoża izolacyjnego z warstwami połączeń montuje się chipy monolitycznych układów scalonych.



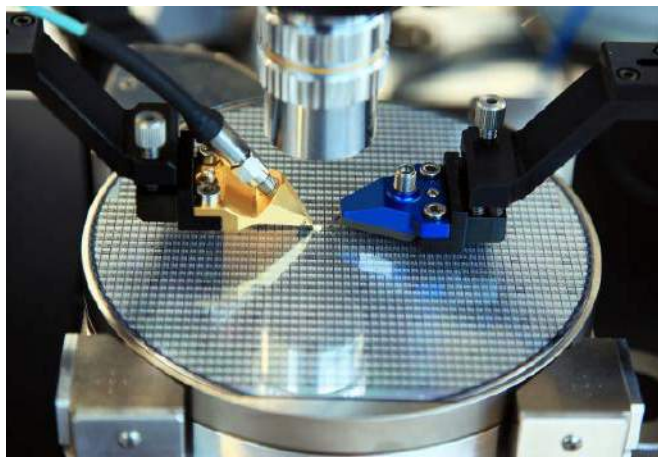
Prawo Moore'a

Wykres w skali logarytmiczno (y) – liniowej (x), pokazujący 2-krotny wzrost skali integracji „sztandarowych” układów scalonych co półtora – dwa lata.

Nie jest to żadne prawo przyrody, tylko spostrzeżenie prognostyczne Gordona Moore'a, jednego z założycieli firmy Fairchild. W połowie lat sześćdziesiątych zauważył on, że co rok pojawiają się układy scalone o coraz większej skali integracji, przy czym liczba tranzystorów w nich zawartych jest dwukrotnie większa niż przed rokiem. Potem osłabiono to tempo wzrostu do dwóch razy co półtora-dwa lata. W wywiadzie z roku 2015 Gordon Moore powiedział, że prognozował takie tempo rozwoju skali integracji w następnych 10 latach i jest zdumiony, że ta prognoza sprawdza się już pół wieku.



Rysunek 8. Płytki Si ze strukturami wytworzonymi w niej chipów (źródło: waferworld.com/post/silicon-wafer-manufacturers-materials)



Rysunek 9. Testowanie chipów na płytce Si (źródło: waferpro.com/silicon-wafer-manufacturing)



Rysunek 10. Różne rodzaje obudów układów scalonych (źródło: components101.com/articles/different-ic-package-types-and-which-one-should-you-select)

Funkcje aplikacyjne

Ze względu na funkcje aplikacyjne najogólniej dzielimy układy scalone na **analogowe** i **cyfrowe**. W **układach analogowych** obróbce podlegają sygnały ciągłe, w **układach cyfrowych** – sygnały dyskretne, charakteryzujące się zwykle dwoma poziomami napięcia utożsamianymi z „zerem” i „jedyneką” logiczną. Układy analogowe są tradycyjnie nazywane **liniowymi**, chociaż często służą do nieliniowej obróbki sygnałów ciągłych, jak na przykład w przypadku kształtowania przebiegów logarytmicznych, stabilizacji napięcia lub prądu, itp. Pierwsze układy analogowe scalone pojawiły się już na początku lat sześćdziesiątych. Niezwykłym faktem jest ciągła, ponad 50-letnia popularność dwóch układów:

- wzmacniacza operacyjnego $\mu A741$, opracowanego przez Davida Fullagara (Fairchild) w roku 1968. Dla sprawiedliwości dodajmy, że jest to nieco ulepszona wersja układu $\mu A709$, opracowanego kilka lat wcześniej przez Boba Widlara
- układu czasowego 555, opracowanego przez Hansa Camenzinda (Signetics) w roku 1971. Układ 555 to prawdziwy koń roboczy elektroniki, szczególnie amatorskiej, któremu poświęcono wiele wydań książkowych.

Najbardziej popularne rodziny układów analogowych przedstawia schematycznie **rysunek 11**.

W rozwoju układów analogowych nie chodziło o upakowanie coraz większej liczby tranzystorów w chipie. Na ogół zawierają one kilkadziesiąt tranzystorów. Ulepszenia dokonywane z biegiem lat dotyczyły poprawy parametrów funkcjonalnych, co łączyło się ze zmianą

Quiz: Układy scalone – historia

Kto jest uznawany za wynalazcę układu scalonego?

- ☐ W. Shockley
- ☐ J. Kilby
- ☐ wspólnie J. Kilby i R. Noyce

W której firmie powstał pierwszy, w pełni monolityczny układ scalony?

- ☐ Texas Instruments
- ☐ Fairchild
- ☐ Intel

Dla jakiego półprzewodnika opracowano technologię planarną?

- ☐ German (Ge)
- ☐ Arsenek galu (Ga As)
- ☐ Krzem (Si)

Pierwsze układy scalone produkowano z Ge, czy z Si?

- ☐ zarówno z Ge jak i Si
- ☐ Si
- ☐ Ge

Kto jest autorem idei kodu binarnego?

- ☐ W. Leibniz
- ☐ G. Boole
- ☐ I. Newton

W którym roku powstał pierwszy komputer ENIAC?

- ☐ 1940
- ☐ 1945
- ☐ 1952

John von Neumann (a właściwie János Lajos Neumann) był uczonym węgierskim pracującym w USA, który zasłynął jako twórca...

- ☐ teorii złącza p-n
- ☐ architektury komputera
- ☐ algebry działań na liczbach binarnych

W którym roku wynaleziono układ scalony?

- ☐ 1952
- ☐ 1957
- ☐ 1959

W którym roku rozpoczęto produkcję układów MOS LSI na dużej skali?

- ☐ 1960
- ☐ 1970
- ☐ 1980

Na bazie którego mikroprocesora Steve Wozniak skonstruował pierwszy komputer Apple?

- ☐ 8080
- ☐ Z80
- ☐ 6520

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 30.12.2022.

Skalowanie

Od początku lat siedemdziesiątych, gdy pojawiły się pierwsze układy scalone MOS LSI (zawierające ponad 1000 tranzystorów), powstało pojęcie skalowania jako ciągłego dążenia do zmniejszania wymiarów charakterystycznych ścieżek w układzie scalonym, odnoszone głównie do długości kanału tranzystora MOS. W istocie, skalowaniu zawdzięczamy ciągły wzrost skali integracji, tak zgrabnie sformułowany w postaci prawa Moore'a.

Oto jak z biegiem lat wprowadzano kolejne generacje technologiczne, wyrażone wymiarem charakterystycznym:

10 μm	6 μm	3 μm	1,5 μm	1 μm
1971	1974	1977	1981	1984

800 nm	600 nm	350 nm	250 nm	180 nm	130 nm
1987	1990	1993	1996	1999	2001

90 nm	65 nm	45 nm	32 nm	22 nm	14 nm
2003	2005	2007	2009	2012	2014

10 nm	7 nm	5 nm	3 nm	2 nm
2016	2018	2020	2022	2024 (?)

Jak widać, co 2-3 lata pojawia się nowy standard technologiczny. Dla 10-krotnego zmniejszenia wymiaru charakterystycznego potrzeba trochę ponad dekadę:

od 10 μm do 1 μm – (1971–1984)
 od 800 nm do 130 nm – (1987–2001)
 od 90 nm do 14 nm – (2003–2014)
 od 10 nm do 1 nm – (2016–?)

Czy rzeczywiście za kilka lat będzie technologia o wymiarze charakterystycznym 1 nm?

Wydaje się, że zbliżyliśmy się do kresu możliwości skalowania. Wystarczy sobie uświadomić, że w sieci krystalicznej Si odstęp między atomami wynosi 0,235 nm, czyli wymiar charakterystyczny 1 nm oznacza ścieżki o szerokości 4 atomów Si. Często utożsamia się wymiar charakterystyczny z długością kanału tranzystora. Trudno wyobrazić sobie działanie tranzystora, w którego kanale jest parę atomów. Jednak w najnowszych technologiach wymiar charakterystyczny nie jest tożsamy z długością kanału, bo wykorzystuje się trójwymiarową konfigurację struktury tranzystora. Dla przykładu technologią, w której sięga się po trzeci wymiar jest FinFET (Fin od kształtu płetwy) – cienka prostopadłościenna fosa, na której zboczach znajduje się tranzystor.

technologii bipolarnej na technologię MOS. Przykładem może być wzmacniacz operacyjny CA3130, konkurent $\mu\text{A}741$ wykonany w technologii BiMOS.

Klasyfikacja **układów scalonych cyfrowych** zależy od przyjętego kryterium. Według kryterium akademickiego, biorącego pod uwagę sposób przetwarzania informacji, rozróżnia się

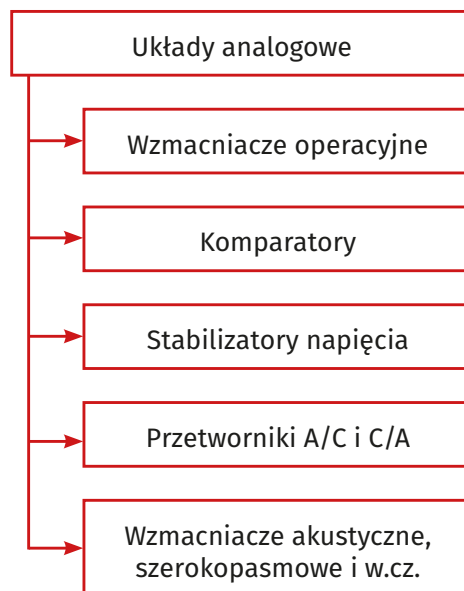
- układy kombinacyjne (bez „pamięci”),
- układy sekwencyjne (z „pamięcią”).

Ze względu na technologię można przyjąć klasyfikację jak na **rysunku 12**.

Technologia bipolarna schodzi do przeszłości. Jeszcze w konstrukcjach amatorskich pojawiają się układy TTL. W technologiach unipolarnych PMOS ma znaczenie już tylko historyczne, a współcześnie najważniejsza jest technologia CMOS.

Najważniejsza w praktyce wydaje się być klasyfikacja układów scalonych odpowiadająca kategoryzacji produktów w ofertach dystrybutorów. Przy takim podejściu wyróżnia się:

- układy logiczne (takie jak mikroprocesory i mikrokontrolery),
- pamięci (takie jak RAM, ROM, Flash, itp.),
- układy interfejsów,
- układy zarządzania mocą,
- układy programowalne,
- układy analogowe, na ogół dzielone na liniowe i układy RF,
- układy scalone z sygnałem mieszanym (takie jak układy akwizycji danych, przetworniki C/A i A/C, itp.),
- układy 3D, tj. pakiety płytek lub chipów Si łączonych metodą przelotek przez otwory w płytkach krzemowych (TSV – Through-Silicon Vias) lub metodą połączeń Cu-Cu,
- inne układy według specyficznych kategoryzacji handlowych.



Rysunek 11. Klasyfikacja układów analogowych

Zakończenie

Temat „układy scalone” jest tak rozległy, że wiele zagadnień kwalifikuje się na odrębny wykład. Liczę też na to, że na różne pytania szczegółowe będę mógł odpowiedzieć w ramach „Konsultacji”. Na pewno po 60 latach rozwoju krzemowej technologii planarnej widać kres jej

Quiz: Układy scalone – technologia

Polski uczyń Jan Czołchalski zasłynął w świecie opracowaniem...

- ☐ metody wytwarzania monokryształów
- ☐ metody wydobycia krzemu ze skał w Sudetach
- ☐ metody utleniania krzemu

Jakie są obecnie największe średnice wytwarzanych płytek krzemowych?

- ☐ 5 cali
- ☐ 30 cm
- ☐ 45 cm

W którym roku wynaleziono technologię planarną?

- ☐ 1957
- ☐ 1959
- ☐ 1962

Ile operacji technologicznych składa się na cały proces wytwarzania chipów w płytce krzemowej dla najbardziej złożonych układów?

- ☐ ok. 90 do 100
- ☐ ok. 150 do 200
- ☐ ok. 400 do 500

Jaka jest rola warstwy SiO₂ wytwarzanej w technologii planarnej na powierzchni płytki Si?

- ☐ Służy jako warstwa przewodząca
- ☐ Służy do polerowania chemicznego płytki Si
- ☐ Służy jako warstwa pasywnująca, maskująca i dielektryczna

Do czego służy stepper?

- ☐ Do naświetlania krok po kroku płytki Si
- ☐ Do trawienia lokalnego warstwy SiO₂
- ☐ Do testowania chipów krok po kroku na płytce Si

Jaki jest wymiar charakterystyczny ścieżek w najbardziej zaawansowanej obecnie technologii?

- ☐ 12 nm
- ☐ 7 nm
- ☐ 2 nm

Ile wynosi na ogół uzysk w produkcji chipów?

- ☐ 1÷2%
- ☐ 30÷80%
- ☐ 95÷99%

Czy testowanie ostrzowe chipów wykonuje się przed pocięciem płytki krzemowej na chipy?

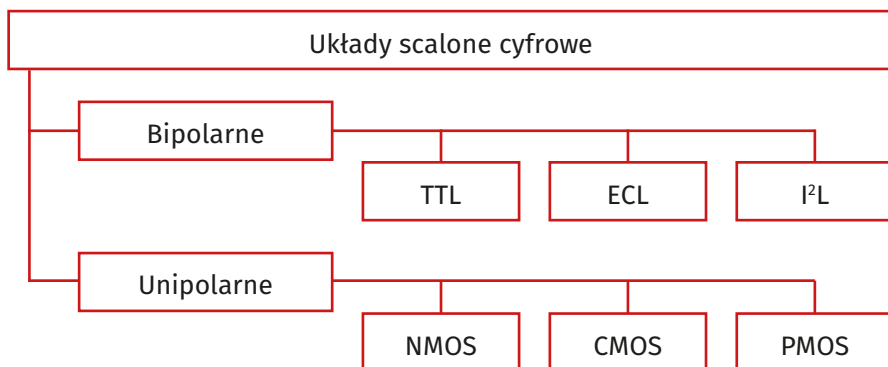
- ☐ TAK, zawsze
- ☐ TAK, niekiedy
- ☐ NIE

Która technologia wytwarzania układów scalonych dominuje?

- ☐ Cienkowarstwowa
- ☐ Grubowarstwowa
- ☐ Monolityczna

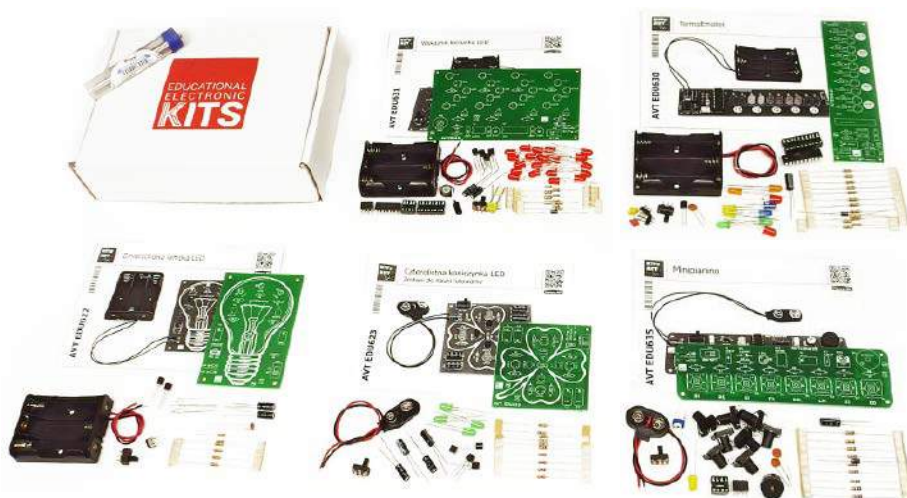
Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 6.01.2023.

możliwości. Na pytanie co dalej nikt nie ma pewnej odpowiedzi, bo chodzi o zmianę jakościową, a nie kolejne lata doskonalenia tej samej koncepcji technologicznej. Najczęściej tę szansę zmiany jakościowej upatruje się w komputerach kwantowych, które obiecują 1000-krotny wzrost mocy obliczeniowej. Do czego może to się nam przydać? Futuryści na czele z Raymondem Kurzweilem z utęsknieniem wypatrują osiągnięcia możliwości przeniesienia zawartości mózgu człowieka do komputera, twierdząc, że będzie to równoznaczne z osiągnięciem nieśmiertelności, jeśli uznać za mało istotną cielesną powłokę mózgu człowieka. Łącząc też to zdarzenie z punktem „singularity”, czyli osobliwości, polegającej na tym, że sztuczna inteligencja przewyższy intelektualnie ludzi i będzie gwałtownie zwiększać swoje zasoby wiedzy aż do nieskończoności. Ma to nastąpić już niedługo – Kurzweil i wyznawcy jego przewidywań czekają niecierpliwie. Karpie też nie mogły się doczekać Świąt Bożego Narodzenia, o czym napisałem we wstępie. ■



Rysunek 12. Klasyfikacja układów scalonych cyfrowych ze względu na technologię

REKLAMA



AVTEDU5PAKIET – 119,00 zł

AVTEDU5PAKIET – to zestaw 5 kitów DIY do nauki lutowania:

- AVTEDU622 – Zmierzwowa lampka LED
- AVTEDU623 – Czterolistna koniczyzna LED
- AVTEDU630 – TermoEmotek
- AVTEDU631 – Wskaźnik kierunku LED
- AVTEDU635 – Minipianino

sklep.avt.pl / Allegro Sklep-AVT
lub 03-197 Warszawa, ul. Leszczynowa 11

eprasa.pl 8b6fc5b24e

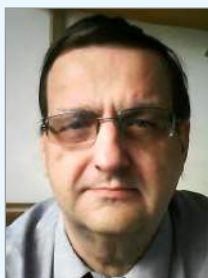
Uczmy się na cudzych błędach

Celem tej rubryki jest kształtowanie u Czytelników EdW umiejętności krytycznego czytania schematów i opisów projektów autorskich. Wszyscy jesteśmy omylni. Konstruktorzy projektów elektronicznych też. W projektach publikowanych w Internecie, ale też w artykułach drukowanych zdarzają się błędy różnej wagi, w tym też takie, które sprawiają, że układ nie może działać prawidłowo. Uczmy się wykrywać te błędy na przykładach projektów sprawdzonych w naszym redakcyjnym Pokoju Nauczycielskim.

Pamiętajmy! Nie oceniamy Autorów, tylko uczymy się na cudzych błędach.

Zapraszamy Czytelników do współpracy z naszym Pokojem Nauczycielskim. Jeśli natrafiliście w Internecie lub źródłach drukowanych na opis projektu z poważnymi Waszym zdaniem błędami, to przysyłajcie takie opisy do naszej redakcji (redakcja@elportal.pl w tytule wiadomości: Pokój Nauczycielski) wraz z Waszymi uwagami.

Projekt sprawdza i poprawia Paweł Sujko



Mgr inż. elektronik po Politechnice Warszawskiej, specjalność aparatura elektroniczna. Od 1992 roku pracownik Polskiego Radia SA jako inżynier serwisowy. Największa forma w jakiej „maczałem” palce, to nadajnik w Solcu Kujawskim (1 MW), najmniejsza, to pendrive (<1 W).



Prosty generator potrójnej fali sinusoidalnej

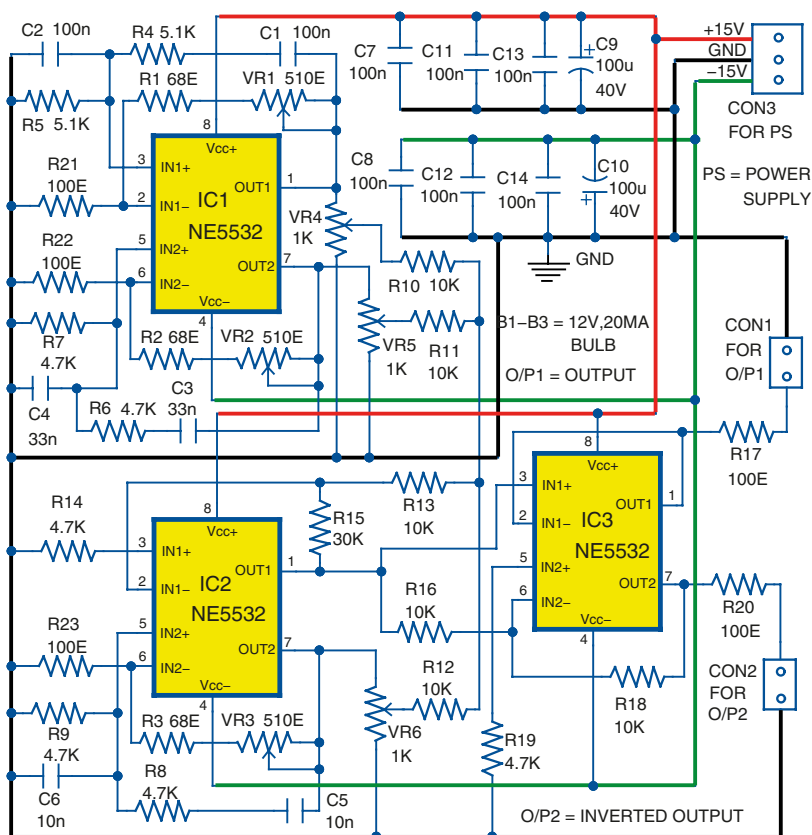
Opisany generator wytwarza sygnały dźwiękowe do testowania torów mikrofonowych i kodeków używanych w sprzęcie wzmacniającym audio. Może on być także używany do testowania innych systemów dźwiękowych, w tym sprzętu VoIP (Voice over Internet Protocol). Sprzęt do transmisji sygnałów mowy jest zwykle przeznaczony do przesyłania sygnałów audio w zakresie od około 300 Hz do 3400 Hz. Testowanie takiego sprzętu odbywa się zwykle na trzech częstotliwościach – 300 Hz, 1000 Hz i 3400 Hz. Podczas testów zwykle używamy oddzielnych sygnałów lub ich kombinacji.

Układ i jego działanie

Schemat obwodu prostego potrójnego generatora fal sinusoidalnych pokazano na rysunku 1. Jest on zbudowany z użyciem trzech podwójnych wzmacniaczy operacyjnych typu NE5532 (od IC1 do IC3), zasilacza symetrycznego i kilku innych komponentów. Możesz też użyć wzmacniacza operacyjnego, takiego jak NE5532A, RC4560 lub dowolnego podobnego ewentualnie lepszego, zdolnego do wysterowania obciążeń od 600 omów.

W układzie są trzy generatory przebiegów sinusoidalnych z mostkami Wiena oraz trzy dodatkowe wzmacniacze operacyjne pracujące jako: sumator, bufor wyjściowy nieodwracający i bufor wyjściowy odwracający i bufor wyjściowy odwracający fazę sygnału o 180 stopni. Wyliczone częstotliwości oscylatorów, dla podanych wartości elementów, są podane w tabeli 1. Lepiej jest zastosować rezystory od R4 do R9 z tolerancją 1% i kondensatory od C1 do C6 z tolerancją $\pm 2\%$ lub lepszą. Częstotliwości pracy możemy obliczyć, korzystając ze wzoru na standardowy oscylator z mostkiem Wiena.

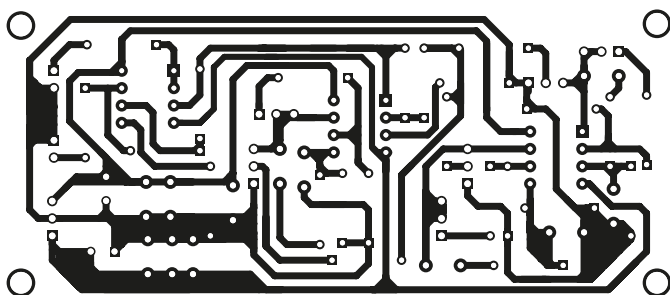
Końcówki 1 i 7 układu IC1 są wyjściami odpowiednio pierwszego i drugiego generatora, podczas gdy sygnał trzeciego oscylatora pochodzi z końcówki 7 układu IC2. Wszystkie trzy sygnały



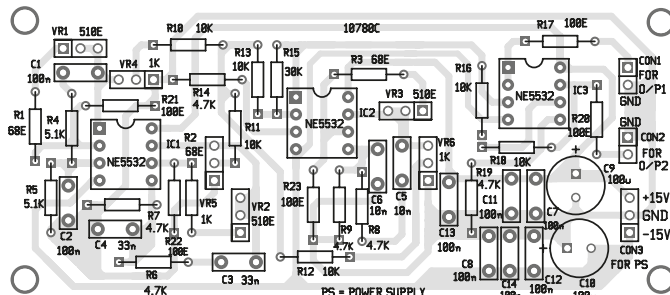
Rysunek 1.

Tabela 1. Częstotliwości generatorów sinusoidalnych w układzie

R4 = R5 = 5,1 k	C1 = C2 = 100 nF	312 Hz
R6 = R7 = 4,7 k	C3 = C4 = 33 nF	1026 Hz
R8 = R9 = 4,7 k	C5 = C6 = 10 nF	3386 Hz



Rysunek 2.



Rysunek 3.

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

Półprzewodniki:

IC1-IC3 – NE5532, podwójne wzmacniacze operacyjne

Rezystory: (wszystkie 1/4 W, $\pm 5\%$ węglowe)

R1-R3 – 68 Ω
R4, R5 – 5,1 k Ω
R6-R9, R14, R19 – 4,7 k Ω
R10-R13, R16, R18 – 10 k Ω
R15 – 30 k Ω
R17, R20-R23 – 100 Ω
VR1-VR3 – 510 Ω
VR4-VR6 – 1 k Ω

Kondensatory:

C1, C2, C7, C8, C11-C14 – ceramiczne 100 nF
C3, C4 – ceramiczne 33 nF
C5, C6 – ceramiczne 10 nF
C9-C10 – elektrolityczny 100 μ F/40 V

Inne:

CON1-CON2 – złącze 2-pinowe
CON3 – złącze 3-pinowe
- podstawa 8-pinowa (3 cyfry)
- przewody potężeniowe

z generatorów są sumowane w drugim wzmacniaczu układu IC2. Suma sygnałów sinusoidalnych jest dostępna na końcówce 1 układu IC2.

Amplitudy sygnałów składowych pojawiających się na końcówce 1 układu IC2 są regulowane za pomocą potencjometrów VR4, VR5 i VR6. Potrójny generator fali sinusoidalnych ma dwa wyjścia. Wyjście proste jest zrealizowane z użyciem pierwszego wzmacniacza operacyjnego układu IC3 (końcówki: 2, 3 i 1 wyjście), a wyjście odwrócone zrealizowano na drugim wzmacniaczu IC3 (końcówki: 6, 5 i 7 – wyjście). Wyjście sinusoidalne (O/P1) jest dostępne na CON1, a jego odwrócona wersja (O/P2) jest dostępne na CON2. Generator potrójnej sinusoidy pracuje w zakresie napięcia zasilania ± 5 V do ± 15 V, ale lepiej jest go używać z zasilaniem z zakresu ± 7 V do ± 15 V.

Budowa i testowanie

Układ ścieżek na płytce drukowanej generatora potrójnej sinusoidy pokazano na rysunku 2, a rozmieszczenie jego elementów na rysunku 3. Zmontuj obwód na PCB. Podłącz symetryczne zasilanie ± 15 V przez CON3 i obwód jest gotowy do użycia. Dany obwód może generować jednocześnie trzy sygnały sinusoidalne (312 Hz, 1026 Hz i 3386 Hz jeżeli wartości elementów są dokładnie takie jak podano), zgodnie z wykazem w tabeli. Sygnały te są dostępne na CON1, a w postaci odwróconej na złączu CON2. Obwód wymaga prostej regulacji amplitudy dla każdego oscylatora. Potencjometry VR1, VR2 i VR3 służą do regulacji amplitudy wyjściowej oscylatorów wokół IC1 i IC2. Możemy zastąpić te potencjometry odpowiednimi stałymi rezystorami. ■

Petre Tzv Petrov

Uwagi i poprawki

Generator z mostkiem Wiena wymaga obwodu stabilizacji amplitudy, w przeciwnym razie sygnał zanika lub, wprost przeciwnie, narasta aż do przesterowania wzmacniacza co daje silne zniekształcenia i raczej trapez niż sinusoidę. Najprostszą metodą jest użycie żarówki w obwodzie ujemnego sprzężenia zwrotnego wzmacniacza. Przy zachowaniu symetrii elementów mostka, wzmocnienie wzmacniacza musi wynosić 3 V/V ale aby generator wystartował musi być większe od tej wartości. Jeżeli jednak, to wzmocnienie pozostanie większe niż 3 V/V, to amplituda sygnału będzie narastać, aż do przesterowania wzmacniacza. Przy zastosowaniu żarówki, jej rezystancja zimna jest ok. 10x mniejsza niż przy zasilaniu napięciem nominalnym. Po włączeniu zasilania włókno żarówki jest zimne i ma małą rezystancję, co z pozostałymi rezystorami ujemnego sprzężenia zwrotnego (np. R1+VR1 dla pierwszego generatora) daje potrzebne do startu generatora wzmocnienie większe niż 3 V/V. W miarę narastania amplitudy sygnału, włókno żarówki rozgrzewa się, jego rezystancja rośnie aż do ustalenia wzmocnienia układu wynoszącego 3 V/V. Gdyby amplituda dalej rosła, to rezystancja włókna wzrośnie, wzmocnienie wzmacniacza zmaleje poniżej 3 V/V co spowoduje malenie amplitudy i powrót układu do stanu stabilnego. Na schemacie oryginalnym zamiast żarówek zastosowano rezystory stałe (R21, R22 i R23) mimo, że w środku schematu napisano parametry żarówek B1 do B3, 12 V, 20 mA.

Drugi problem, to układ sumowania sygnałów wyjściowych. Normalnie, to sygnały powinny być sumowane prądowo w punkcie masy pozornej na wejściu odwracającym wzmacniacza w IC2, końcówka 2. W układzie jest jednak rezystor R13, którego istnienie spowoduje, że regulacje amplitud

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

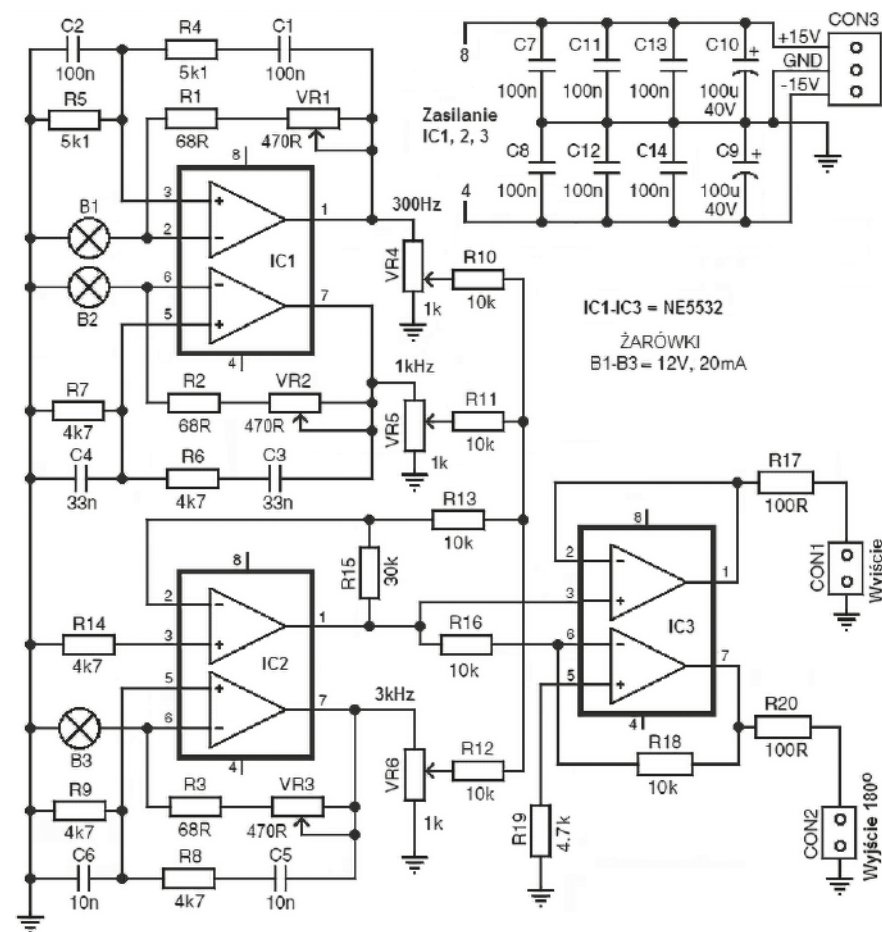
generatorów będą częściowo współzależne i ustawianie całości będzie kłopotliwe. R13 powinien być zastąpiony zworą.

Przyjmując pracę układu ze stabilizacją żarówkową np. z napięciem na żarówce wynoszącym 1 Vsk i ~5 mAsk (dla żarówki 12 V, 20 mA), to na wyjściu pojedynczego generatora amplituda będzie wynosiła ok. 3 Vsk*1,41 ~4,23 V, to być może konieczna była by zmiana rezystora R15 na mniejszy by wzmacniacz sumujący nie był przesterowywany przy niektórych ustawieniach potencjometrów VR4 do VR6,

W założeniu układ ma służyć jako źródło sygnału pomiarowego dla torów mikrofonowych i liniowych ale z powodu opisanego w punkcie pierwszym (sygnały niesinusoidalne) i dużych amplitud wyjściowych układu, nie nadaje się on w tej wersji do wymienionych pomiarów (szczególnie w przypadku torów mikrofonowych gdzie sygnały wejściowe są na poziomie miliwoltów).

W układzie brakuje możliwości dostrajania częstotliwości generatorów składowych co może być przydatne tak ze względu na same pomiary, jak i łatwości uruchamiania układu przy nieuchronnym rozrzucie parametrów elementów określających częstotliwość oraz zmianach ich parametrów w czasie.

Podane wartości potencjometrów VR1 do VR3 wynoszące 510 Ω nie są produkowane. Potencjometry, zależnie od strony świata, są masowo produkowane w seriach



wartości: 100 Ω, 200 Ω, 500 Ω lub 100 Ω, 220 Ω, 470 Ω.

W układzie do celów pomiarowych należy raczej używać metalizowanych

rezystorów precyzyjnych zamiast standardowych węglowych o tolerancji 5%. ■

Quiz: Poziomy logiczne

Jakie rodziny układów logicznych są najczęściej stosowane współcześnie?

- ☐ DTL i RTL
- ☐ ECL
- ☐ CMOS i TTL

Co to jest technologia BiCMOS?

- ☐ Technologia łącząca tranzystory bipolarne i MOSFET
- ☐ Technologia CMOS z bramkami z bizmutu
- ☐ Technologia CMOS z zasilaniem dwupolarnym

Która z tych technologii nie jest już stosowana?

- ☐ PMOS
- ☐ NMOS
- ☐ CMOS

Który z trzech przypadków napięć logicznych odpowiada logice dodatniej?

- ☐ 0 V dla poziomu 1, +5 V dla poziomu 0
- ☐ 0 V dla poziomu 0, +5 V dla poziomu 1
- ☐ +2 V dla poziomu 0, -2 V dla poziomu 1

Który z trzech przypadków napięć logicznych odpowiada logice ujemnej?

- ☐ 0 V dla poziomu 1, +5 V dla poziomu 0
- ☐ 0 V dla poziomu 0, +5 V dla poziomu 1
- ☐ -2 V dla poziomu 0, +2 V dla poziomu 1

Który ciąg napięć zasilania układów logicznych jest stosowany w rzeczywistości?

- ☐ +12 V, +8 V, +5 V, +3,3 V, +1,5 V, +1,2 V
- ☐ +5 V, +3,3 V, +2,5 V, +1,8 V, +1,5 V, +1,2 V, +0,8 V
- ☐ +15 V, +5 V, +2,5 V, +1,25 V

Bramka ma maksymalny poziom zera napięcia wyjściowego równy 0,4 V i steruje kolejną bramką o maksymalnym poziomie zera napięcia wejściowego równym 0,8 V. Jaki jest margines zaktóceń?

- ☐ -0,4 V
- ☐ +0,4 V
- ☐ +1,2 V

Układy logiczne nazywamy kompatybilnymi, gdy spełniają warunki:

- ☐ mają jednakowe szybkości działania
- ☐ są wykonane w tej samej technologii
- ☐ mają zgodne poziomy logiczne 0 i 1

Po co stosuje się rezystor podciągający w układach logicznych?

- ☐ Dla zapewnienia kompatybilności poziomów logicznych 1 dwóch współpracujących układów
- ☐ Dla zmniejszenia napięcia zasilania
- ☐ Dla zmniejszenia zaktóceń

Układ translacji poziomów służy do:

- ☐ zamiany logiki dodatniej na ujemną
- ☐ inwersji poziomów logicznych
- ☐ współpracy niekompatybilnych układów logicznych

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 13.01.2023.

Sterowanie natężenia światła diod LED przy pomocy techniki PWM

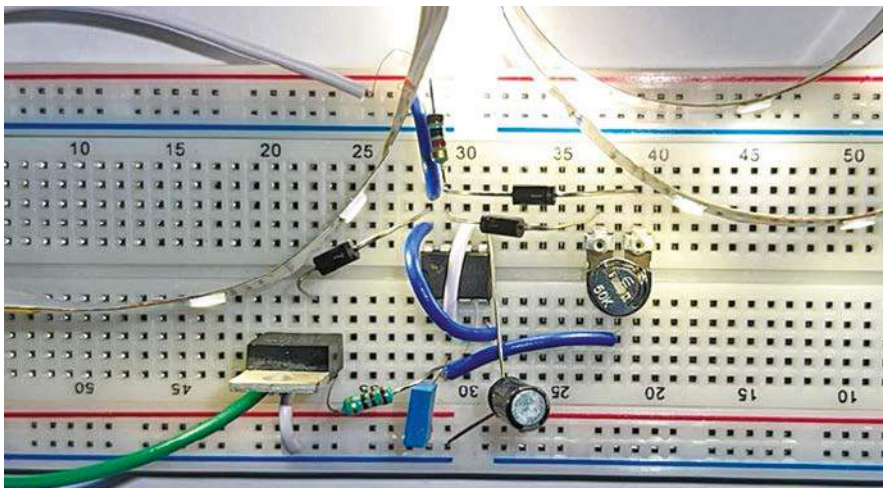


W artykule opisano, w jaki sposób można wykonać prosty obwód pozwalający na regulację intensywności świecenia taśmy z diodami LED. Obwód jest oparty na układzie scalonym timera 555, który realizuje technikę modulacji szerokości impulsu (PWM) do sterowania napięciem zasilającym LED-y. Układ 555 pracuje w trybie astabilnym. Ten sam obwód można również wykorzystać do sterowania prędkością silnika prądu stałego. Układ 555 steruje tranzystorem MOSFET pracującym jako klucz i zwiększającym obciążalność układu. Komponenty wymagane do projektu są wymienione w poniższej tabeli.

Układ sterownika jest bardzo prosty i wymaga tylko kilku elementów. Ze schematu ideowego pokazanego na **rysunku 1** widać, że układ wykorzystuje timer 555 pracujący jako generator o mniej więcej stałej częstotliwości (ok. 13 kHz) z możliwością regulacji

szerokości impulsu. W związku z tym możemy sterować średnią wartością prądu zasilającego diody LED, a tym samym ich jasnością, poprzez zmianę wartości elementów w obwodzie generatora. Timer 555 to powszechnie używany układ scalony (IC) w obudowie

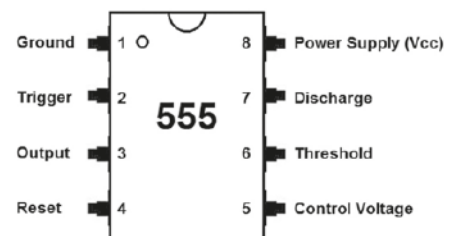
ośmionóżkowej, mający wielorakie zastosowanie w elektronice, w tym jako generator impulsów czy układ odmierzenia czasu (timer). Napięcie zasilania układu scalonego może wynosić od 4,5 V do 16 V. **Rysunek 2** pokazuje rozkład wyprowadzeń układu 555.



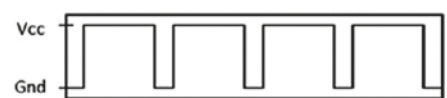
Fotografia 1. Obwód na płytce prototypowej

Sterowanie PWM za pomocą układu 555

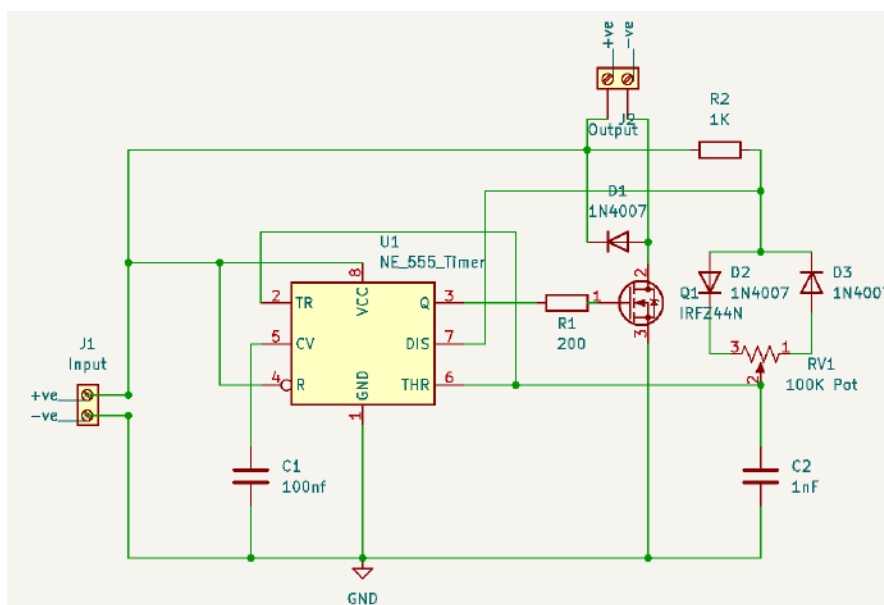
Zanim zagłębimy się w projektowany układ, poznamy technikę PWM, znaną również jako modulacja czasu trwania impulsu. Jest to metoda regulowania średniej mocy dostarczanej do obciążenia poprzez odpowiednie przerywanie przepływu prądu stałego. Taką



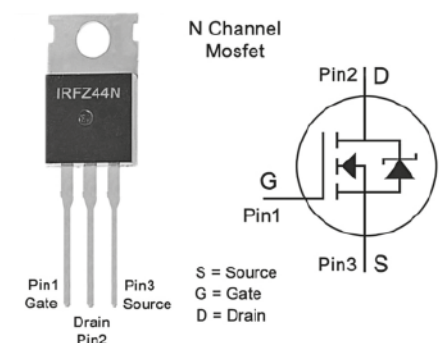
Rysunek 2. Opis wyprowadzeń timera 555



Rysunek 3. Przykład sygnału PWM



Rysunek 1. Schemat ideowy



Rysunek 4. Opis wyprowadzeń tranzystora IRFZ44N

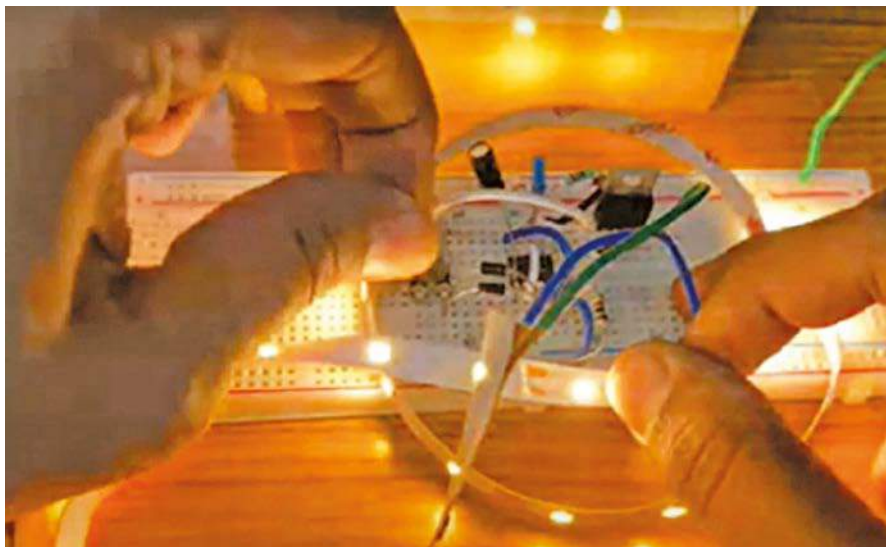
metodę zastosowano w opisywanym układzie. Aby uzyskać sygnał PWM, timer 555 pracuje w układzie multiwibratora astabilnego. Jak sama nazwa wskazuje, timer 555 nie ma stanu stabilnego w tym trybie, ale raczej dwa stany quasi-stabilne, których suma czasów trwania jest stała ale zmieniają się ich proporcje. Mówiąc prościej, w tym trybie timer 555 włącza się i wyłącza z bardzo dużą prędkością, którą można uznać za 0 i 1 w kontekście binarnym, wytwarzając w ten sposób sygnał prostokątny, jak pokazano na **rysunku 3**. Różne są tylko czasy trwania obu stanów.

Podłączenia timera 555 w układzie

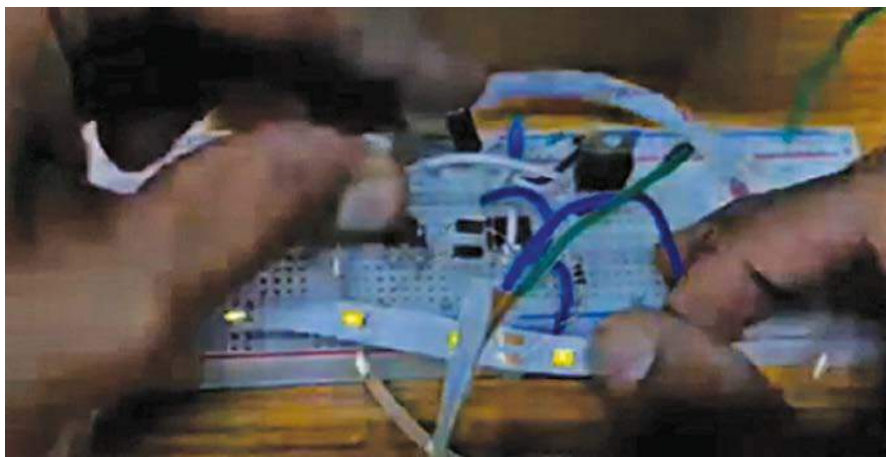
- Styk 1 jest stykiem masy i jest podłączony do ujemnego zacisku zasilania (–ve).
- Pin 8 jest stykiem zasilania VCC i jest podłączony do szyny (+ve).
- Styki 2 i 6, to odpowiednio wejścia wyzwalające i progowe, które są połączone z masą przez kondensator C2 (1 nF), wyznaczający stałą czasu multiwibratora.
- Pin 3 to wyjście multiwibratora sterujące bramką tranzystora MOSFET.
- Pin 4 to końcówka zerująca stan przelotnika w 555. W naszym układzie nie wykorzystujemy jej, więc podłączamy ją do szyny +ve.
- Pin 5 pozwala na sterowanie pracą multiwibratora przez zmianę napięć odniesień wewnętrznych komparatorów układu 555. Ponieważ nie wykorzystujemy tej funkcji, to podłączamy tę końcówkę przez kondensator 100 nF do masy (–Ve).
- Pin 7 końcówka kolektora wewnętrznego tranzystora układu 555 rozładowującego kondensator C2 przez diodę D3 i przez część potencjometru RV1.

Tranzystor IRFZ44N do szybkiego przełączania

Układ scalony timera 555 ma sporą ale ograniczoną wydajność prądową końcówki wyjściowej (3), którą można bezpośrednio wykorzystać doysterowania tylko dla kilku diod LED. Aby jednak sterować długim paskiem diod LED lub sterować prędkością silnika, to taki prąd nie wystarczy. Aby rozwiązać ten problem, w obwodzie zastosowano tranzystor MOSFET, którego bramka jest podłączona do wyjścia 3 układu 555 przez rezystor R1 (ogranicza on prędkość przełączania tranzystora co zmniejsza przepięcia na indukcyjności gdy układ wykorzystujemy do sterowania silnikiem). Dioda D2 zwiera przepięcia występujące na indukcyjności silnika. Zastosowany MOSFET typu n, zapewnia sterowanie obciążeniem od strony masy. W naszym układzie używamy tranzystora IRFZ44N, który posiada następujące parametry:



Rysunek 5. Autorski prototyp przedstawiający maksymalne oświetlenie paska LED



Rysunek 6. Zmniejszanie jasności paska LED za pomocą potencjometru

$V_{DSS} = 55 \text{ V}$ – maksymalne napięcie drenu-
źródło przy bramce zwartej ze źródłem
 $R_{DS(on)} = 17,5 \text{ m}\Omega$ – rezystancja przewodzą-
cego kanału, przy $V_{GS}=10 \text{ V}$ i $I_D=25 \text{ A}$
 $I_D = 49 \text{ A}$ – prąd maksymalny drenu w tem-
peraturze 25°C

Sterowanie intensywnością diody LED

Jak wspomniano wcześniej, elementy pasywne w układzie (R2, RV1 i C2), służą do ustalenia parametrów pracy multiwibratora 555, a tym samym jasności diod LED.

D – wypełnienie impulsu, Pw – szerokość impulsu (stanu wysokiego), T – całkowity okres przebiegu.

Wypełnienie impulsu reprezentuje stosunek czasu, w którym obciążenie jest podłączone do zasilania, do całego okresu przebiegu sterującego. Wyższe wypełnienie oznacza, że obciążenie jest podłączone do źródła zasilania przez dłuższy czas w całym okresie powtarzania.

Ta technika może być wykorzystana do sterowania jasnością intensywności diod LED

lub prędkości silnika. Gdy wypełnienie przebiegu jest duże, to dioda LED świeci jaśniej lub prędkość silnika będzie większa. Aby zmniejszyć intensywność świecenia diod LED lub zmniejszyć prędkość obrotową silnika, należy zmniejszyć wypełnienie przebiegu multiwibratora.

Intensywność diod = (czas włączenia/całkowity czas) $\times$ maksymalna intensywność

Prędkość obrotowa = (czas włączenia/całkowity czas) $\times$ maksymalna prędkość

Testowanie

Aby sprawdzić, czy układ działa prawidłowo, można użyć pakietu baterii o sumarycznym napięciu wynoszącym 12 V, tak jak to zrobił autor. Ponieważ używamy tranzystora MOSFET z kanałem typu N i sterujemy obciążeniem od strony masy, to dodatni zacisk obciążenia jest zawsze podłączony, podczas gdy ujemny jest podłączony do drenu tranzystora MOSFET, który jest sterowany z układu timera 555. Na **rysunku 5** pokazano pasek LED świecący z maksymalną jasnością.

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

Układ scalony timera 555
MOSFET – IRFZ44N
Rezystor 1 kΩ, 5%
Rezystor 200 Ω, 5%
Kondensator 100 nF
Kondensator 1 nF
Dioda 3×1N4001
Potencjometr 100 kΩ, lin.

Aby zmniejszyć intensywność świecenia diod, należy obrócić pokrętko potencjometru w kierunku przeciwnym do ruchu wskazówek zegara, uzyskując efekt pokazany na rysunku 6.

Tę samą konfigurację można przetestować też z silnikiem 12 VDC. Powinniśmy być w stanie kontrolować prędkość obrotową silnika.

Sharad Bhowmick

Dodatek. Film instruktażowy dotyczący tego projektu DIY można obejrzeć pod adresem <https://bit.ly/3VCKsaP>.

Ten artykuł był wcześniej opublikowany na łamach „EFY”, lipiec 2022 (efymag.com)

Od Red. EdW: uwagi dla dociekliwych

Jako D2 i D3 lepiej użyć diod impulsowych typu 1N4148 lub nawet Schottky’ego, np. BAT85.

Podane przez autora, w filmie z linku, wzory na czasy ładowania i rozładowania kondensatora C2 w tym układzie, nie są do końca prawdziwe, bo nie uwzględniają spadku napięcia na diodach D2 i D3. Pozornie ten spadek jest niewielki, ale... Autor podaje (film z linku, czas pomiędzy 6.40 a 8.00) – przyjmijmy środkowe położenie potencjometru RV1:

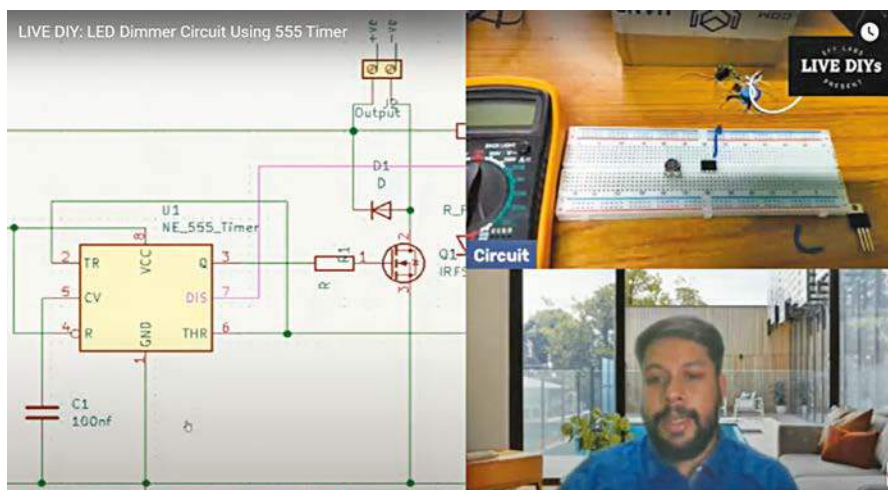
$$t_L = 0,693 \cdot (R_2 + 0,5 \cdot R_{V1}) \cdot C_2$$

oraz

$$t_R = 0,693 \cdot 0,5 \cdot R_{V1} \cdot C_2$$

gdzie $0,693 \approx \ln 2$

W rzeczywistości kondensator C2 jest ładowany, nie od napięcia U_{cc} , tylko od $V_{cc} - U_{D2}$ i rozładowywany od napięcia 0 V (pomijając napięcie nasycenia tranzystora rozładowującego w 555 końcówka 7) tylko napięcia U_{D3} , przy nie zmienionych progach przełączania komparatorów w układzie 555, wynoszących: $2/3 U_{cc}$ i $1/3 U_{cc}$. Zmienia to dwie sprawy w takiej strukturze multiwibratora: wzory na czasy t_L i t_R oraz wprowadza wrażliwość układu na zmiany



Rysunek 7. Zrzut ekranu z nagranego wideo demonstrującego projekt

napięcia zasilania, co może być istotne w innych zastosowaniach tej struktury multiwibratora.

$$t_A = (R_2 + 0,5 \cdot R_{V1}) \cdot C_2 \cdot \ln \frac{2 \cdot U_{cc} - 3 \cdot U_d}{U_{cc} - 3 \cdot U_d}$$

oraz

$$t_R = 0,5 \cdot R_{V1} \cdot C_2 \cdot \ln \frac{2 \cdot U_{cc} - 3 \cdot U_d}{U_{cc} - 3 \cdot U_d}$$

W przypadku gdybyśmy mieli idealną diodę i $U_d = 0$ V, to wzór sprowadza się do postaci podanej w filmie przez autora. Jeżeli postawimy wartości: $R_2 = 1$ kΩ, $0,5 \cdot R_{V1} = 50$ kΩ, $C_2 = 1$ nF, $U_{cc} = 12$ V i $U_d = 0,47$ V (wartość dla środkowego położenia RV1), to z wzorów uproszczonych otrzymamy okres pracy multiwibratora:

$$T_1 = t_{L1} + t_{R1} = (51 \text{ k}\Omega + 50 \text{ k}\Omega) \cdot 1 \text{ nF} \cdot \ln 2 \approx 70 \mu\text{s}$$

co daje częstotliwość pracy $f_1 = 1/T_1 \approx 14\,286 \text{ Hz}$

W przypadku uwzględnienia nieidealności diod:

$$T_2 = t_{A2} + t_{R2} =$$

$$(51 \text{ k}\Omega + 50 \text{ k}\Omega) \cdot 1 \text{ nF} \cdot \ln \frac{2 \cdot 12 \text{ V} - 3 \cdot 0,47 \text{ V}}{12 \text{ V} - 3 \cdot 0,47 \text{ V}} =$$

$$76,52 \mu\text{s}$$

$$f_2 = \frac{1}{T_2} \approx 13068 \text{ Hz}$$

Widać, że dla $U_{cc} = 12$ V współczynnik nie wynosi 0,693 tylko 0,758 i jest jeszcze zależny od wartości U_{cc} . Dodatkowo spadki napięcia na diodach D_2 i D_3 są zależne od prądów jakie przez nie płyną czyli będą się one zmieniać

zależnie od ustawienia potencjometru R_{V1} (mniej więcej od 0,45 V do 0,77 V)

Nawet w tym układzie ma to pewne praktyczne znaczenie, które ujawni się np. przy próbie filmowania kamerą przy oświetleniu LED-ami sterowanymi tym układem. Aby po obrazie z kamery nie pływały linie (mora), to częstotliwość kluczkowania diod oświetlających powinna być całkowitą wielokrotnością częstotliwości ramki (50 lub 60 Hz). Dla obu tych wartości najbliższą optymalną częstotliwością migotania diod LED jest 300 Hz lub wielokrotność tej wartości. Jeżeli policzymy układ z wzorów uproszczonych, to nie dość, że nie wyjdzie nam założona częstotliwość, to jeszcze będzie ona pływała jeżeli napięcie zasilania nie będzie stabilizowane... a jeszcze trzeba uwzględnić tolerancje wartości użytych elementów.

Dodatkową sprawą jest, to że przy zmianie wypełnienia impulsu zmienia się też częstotliwość pracy naszego multiwibratora, bo przy ładowaniu C2 w obwodzie jest rezystor R2, a przy rozładowywaniu go nie ma, więc regulacje czasów włączenia i wyłączenia nie są dokładnie współbieżne. Ten problem można ominąć dając szeregowo z D3 dodatkowy rezystor o wartości 1 kΩ.

Oczywiście w podstawowym zastosowaniu opisane problemy nie będą grały specjalnie roli ale w pewnych specyficznych zastosowaniach oświetleniowych lub innych zastosowaniach tej struktury multiwibratora warto je mieć na uwadze.

REKLAMA

Już ponad rok publikujemy dla projektantów
i programistów dla elektroniki. Odwiedź

ELPORTAL.pl

Lampa-sygnalizator przelotu Międzynarodowej Stacji Kosmicznej (ISS)

Jeśli jesteś entuzjastą kosmosu i lubisz patrzeć w gwiazdy, prawdopodobnie jesteś już zaznajomiony z ISS, czyli Międzynarodową Stacją Kosmiczną, która okrąża Ziemię sześć razy na dobę. Jeśli interesujesz się ISS, tak jak autor, który jest studentem inżynierii lotniczej, to spodoba ci się ten projekt.

To łatwe do wykonania urządzenie, to mała półkuliśta lampa, która daje ładną poświatę po podłączeniu do źródła zasilania. Kiedy jednak ISS przelatuje nad twoją lokalizacją, lampka zaczyna migać, od około 30 sekund do kilku minut. Prototyp autora pokazano na rysunku 1.

Wszystko, czego potrzebujesz do wykonania tego projektu, to mikrokontroler węzłowy (nodeMCU), dioda LED Ø 5 mm, kolor wedle uznania, kawałek izolowanego przewodu, rezystor 330 Ω, gruby karton do wykonania podstawki lampy i matowy klosz w kształcie półkuli.



Rysunek 1. Prototyp autora

Red. EdW: Jako klosz sygnalizatora można wykorzystać matową osłonę z uszkodzonej lampy LED, należy ją delikatnie oddzielić od podstawy z gwintem (jest

przyklejona). Zasilanie NodeMCU realizujemy przez USB. Przy wykonaniu podstawy należy uwzględnić wycięcie na złącze USB.

Do przetestowania tego projektu potrzebne są dwie platformy programowe: Adafruit IO i IFTTT. Adafruit IO, to usługa internetowa przeznaczona dla użytkowników, do wyświetlania, reagowania i interakcji z ich danymi w pewny i bezpieczny sposób w chmurze. IFTTT wywodzi swoją nazwę od wyrażenia warunkowego programowania „If This Then That” (jeśli stało się to, to zrób tamto) i jest platformą oprogramowania, która łączy aplikacje, urządzenia i usługi różnych programistów, aby uruchomić jedną lub więcej automatyzacji z ich udziałem.

Aby wykonać projekt, najpierw skonfiguruj Adafruit IO, a następnie zainstaluj

Rysunek 3. Drugi krok tworzenia kanału na Adafruit

Rysunek 4. Tworzenie pulpitu nawigacyjnego

Rysunek 2. Pierwszy krok tworzenia kanału na Adafruit

w smartfonie aplikację IFTTT, która jest łatwo dostępna, za darmo (w wersji ograniczonej) w Internecie. Aby skonfigurować Adafruit IO, otwórz stronę io.adafruit.com, utwórz tam konto i zaloguj się. Następnie wykonaj czynności wymienione poniżej:

1. Kliknij „Feeds” (kanały) na górnym pasku i wybierz „New Feed” (nowy kanał), jak to pokazano na rysunku 2. Nadaj nazwę swojemu kanałowi. Możesz użyć tej samej nazwy, która została użyta dla projektu (rysunek 3), czyli ISS_LED, ponieważ uprości to programowanie.

2. Kliknij „Dashboards” (panele) na górnym pasku i wybierz „New Dashboard” (nowy panel), jak to pokazano na rysunku 4. Nadaj mu taką samą nazwę jak kanałowi, czyli ISS_LED.
3. Kliknij utworzony panel nawigacyjny i wybierz ikonę koła zębatego po prawej stronie (rysunek 5). Następnie wybierz „Create a new block” (utwórz nowy blok) i kliknij przełącznik (patrz ON na rysunku 6).
4. Wybierz nazwę swojego kanału (w tym przypadku ISS_LED) w wyskakującym

okienku i kliknij „Next step” (następny krok). W następnym oknie nic nie zmieniaj, po prostu kliknij „Create block” (utwórz blok).

5. Po utworzeniu bloku kliknij „My key” (mój klucz) na górnym pasku (patrz rysunek 2) i zanotuj swoją nazwę użytkownika oraz aktywny klucz, który jest unikalny dla twojego projektu. Nie dziel się nimi z nikim, ponieważ może to narobić bałaganu w twoim projekcie (patrz rysunek 8).

Teraz, gdy konfiguracja Adafruit IO została zakończona, możesz skonfigurować IFTTT, po pobraniu aplikacji na swój smartfon (lub możesz po prostu skorzystać ze strony aplikacji na komputery). Aby skonfigurować usługę IFTTT, wykonaj poniższe czynności:

Kliknij „Create” (Utwórz) (na dole po lewej na rysunek 9) i wybierz pole „If This” (Jeśli To). Wyszukaj „Space” (przestrzeń) w pasku wyszukiwania i kliknij ją. Kliknij pole „ISS passes over a specific location” (ISS przelatuje nad określoną lokalizacją) i wybierz swoją lokalizację.

Teraz kliknij blok „Then That” (wtedy tamto) i wyszukaj Adafruit na pasku wyszukiwania. Kliknij „Send data to Adafruit IO” (patrz rysunek 10), wybierz nazwę kanału (ISS_LED) i wpisz ON w polu „Data to save” (dane do zapisania).

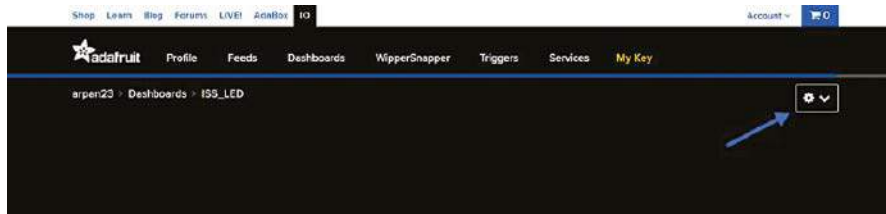
Po skonfigurowaniu Adafruit i IFTTT pozostaje tylko wgrać kod do NodeMCU. Podsumowując, stworzyliśmy przełącznik na Adafruit, który może sterować twoim NodeMCU przez Wi-Fi. Stworzyliśmy aplet na IFTTT, który przełącza przełącznik Adafruit za każdym razem, gdy ISS przelatuje nad wybraną lokalizacją.

Kod

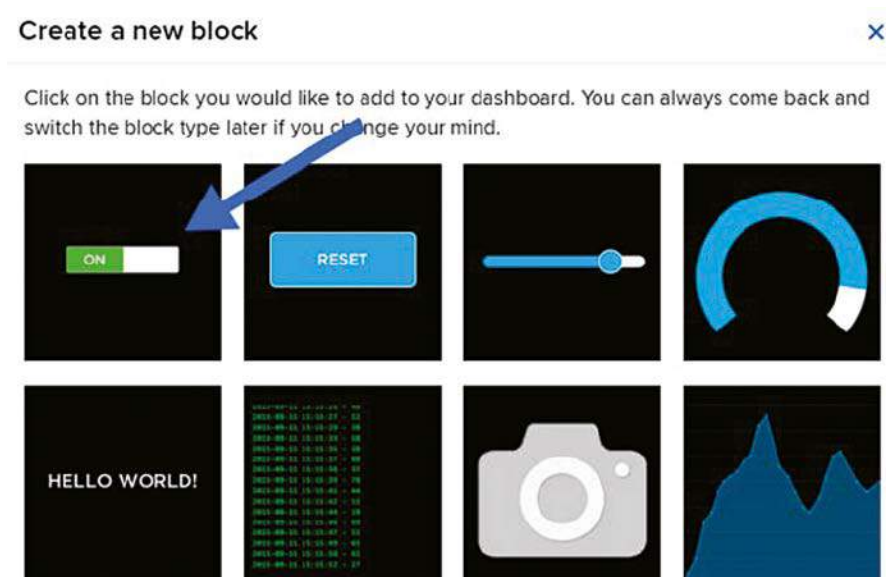
Sercem lampy sygnalizatora ISS jest NodeMCU, który musi zostać zaprogramowany, aby działał tak, jak chcemy. Upewnij się, że masz, na swoim Arduino IDE, zainstalowane biblioteki **NodeMCU** i **Adafruit mqtt**. Kod źródłowy w pliku **ISS_LED_SAMPLE.ino** pokazany na rysunku 11, należy nieco zmodyfikować, aby działał poprawnie.

Wprowadź następujące zmiany w przykładowym kodzie:

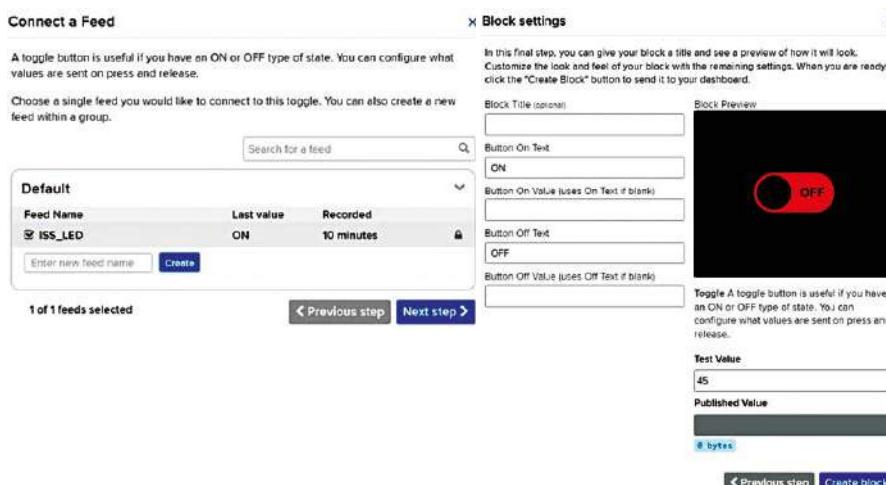
- Zastąp SSID nazwą twojego Wi-Fi (lub mobilnego punktu dostępu).
- Zastąp PASSWORD hasłem swojego Wi-Fi (lub mobilnego punktu dostępu).
- Zastąp identyfikator Adafruit swoją nazwą użytkownika Adafruit IO.



Rysunek 5. Tworzenie nowego bloku



Rysunek 6. Wybór przełącznika dwustabilnego jako bloku



Rysunek 7. Konfiguracja przełącznika dwustabilnego

Sięgnij po archiwalne wydania

- Zastąp AIO KEY aktywnym kluczem, który uzyskałeś po utworzeniu przełącznika Adafruit.

Jeśli nazwałeś swój kanał Adafruit jako ISS_LED, to skończyłeś z edycją kodu. Ale jeśli użyłeś innej nazwy, to zastąp ISS_LED wszędzie w kodzie tą inną nazwą, której użyłeś.

Podany kod programu sprawia, że dioda LED jest domyślnie wyłączona. Kiedy ISS nadleci, dioda zacznie migać. Jeśli chcesz, aby dioda LED pozostała domyślnie włączona, możesz zastąpić **digitalWrite(led,LOW);** linią **digitalWrite(led,HIGH);** w pierwszym wierszu funkcji **void loop()**.

Teraz możesz iść dalej i przesłać kod do swojego NodeMCU. Należy pamiętać, że podczas testów EFY Lab, nazwa użytkownika i hasło były ustawione w kodzie programu ISS_LAMP_SAMPLE.ino:
`#define MQTT_NAME "efytech2"`
`#define MQTT_PASS "aio_sXIY05Iri6o3ueuWYLLY5KnZlCaa"`
i należy je zastąpić własnymi danymi. Do lokalizacji ISS wykorzystano stronę internetową „<https://www.astroviewer.net/iss/en/>”.

Układ

W tym ostatnim kroku uruchamiania projektu podłącz diodę LED do pinu D7 swojego NodeMCU. Rezystor 330 omów połączony szeregowo (jak pokazano na rysunku 12) służy do ograniczenia prądu płynącego przez tę diodę.

Aby zrobić podstawę lampy, możesz wyciąć pasek czarnej tektury i zwinąć go w cylinder. Wysokość cylindra (szerokość paska) powinna być nieco większa niż wysokość NodeMCU, gdy spoczywa na sworzniach. Średnica cylindra powinna być taka sama jak średnica podstawy dyfuzora (klosza rozpraszającego).

Wytnij kolejny kawałek kartonu w kształcie koła o średnicy równej średnicy kartonowego cylindra i przyklej go klejem do cylindra. Umieść NodeMCU i diodę LED w tej tekturowej podstawie i umieść dyfuzor na górze (patrz rysunek 13 do rysunku 15).

Twój własny sygnalizator przelotu ISS jest teraz gotowy. Za każdym razem, gdy ISS będzie przelatywała nad twoją lokalizacją, lampa poinformuje cię o tym.

ISS przecina daną lokalizację raz lub dwa razy dziennie (czasami w środku nocy, kiedy może nie zostać zauważona). Są dni, kiedy w ogóle się nie pojawia, ale to rzadkość.

YOUR ADAFRUIT IO KEY

Your Adafruit IO Key should be kept in a safe place and treated with the same care as your Adafruit username and password. People who have access to your Adafruit IO Key can view all of your data, create new feeds for your account, and manipulate your active feeds.

If you need to regenerate a new Adafruit IO Key, all of your existing programs and scripts will need to be manually changed to the new key.

Username

Active Key

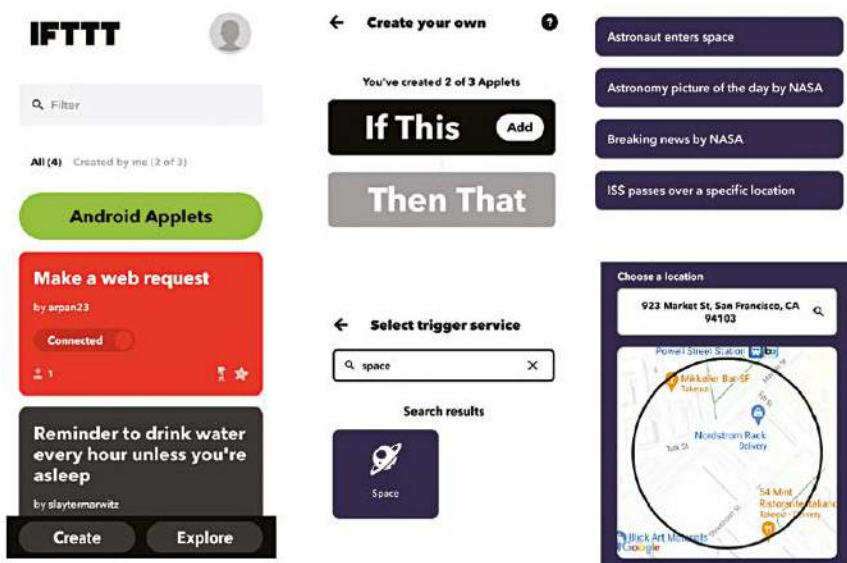
REGENERATE KEY

Hide Code Samples

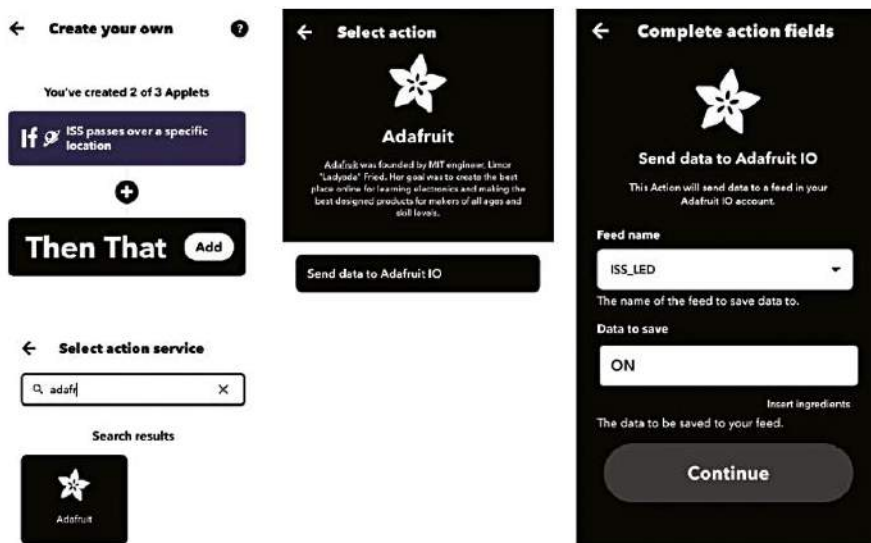
Arduino

```
#define IO_USERNAME "adafruit"
#define IO_KEY "adafruit"
```

Rysunek 8. Wprowadzanie aktywnego klucza Adafruit



Rysunek 9. Konfiguracja bloku „Jeżeli to” w IFTTT



Rysunek 10. Konfiguracja bloku „Wtedy Tamto” w IFTTT

ELEKTRONIKI dla WSZYSTKICH

ISS_LED_SAMPLE | Arduino 1.8.9

File Edit Sketch Tools Help

```
ISS_LED_SAMPLE

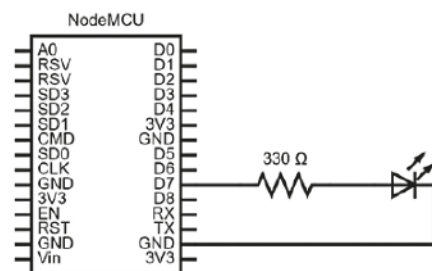
#include <ESP8266WiFi.h>
#include "Adafruit_MQTT.h"
#include "Adafruit_MQTT_Client.h"

#define WIFI_SSID ".....SSID....."//Your wifi name
#define WIFI_PASS ".....PASSWORD....."//your wifi password

#define MQTT_SERV "io.adafruit.com"
#define MQTT_PORT 1883
#define MQTT_NAME ".....Adafruit ID....." //Your adafruit ID
#define MQTT_PASS ".....AIO KEY....." //Your adafruit AIO key

int led = D7;
```

Rysunek 11. Przykładowy kod dla NodeMCU, który należy zmodyfikować

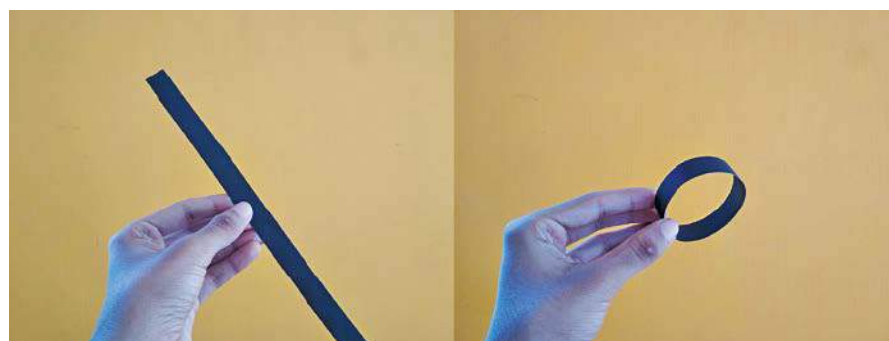


Rysunek 12. Połączenia NodeMCU i diody LED

Dobrze jest wiedzieć, że w ISS pracują ludzie i ekscytujące jest to, że są tuż nad głową. Kiedy ISS przybędzie po zachodzie słońca, możesz nawet wybiec i popatrzeć, jak porusza się po niebie i może pomachać jej na powitanie! ■

Arpan Mondal

Ten artykuł był wcześniej opublikowany na łamach „EFY”, lipiec 2022 (efymag.com)



Rysunek 13. Wykonanie tekturowej podstawy lampy



Rysunek 14. NodeMCU i dioda LED umieszczone na dole



Rysunek 15. Osłona rozpraszająca umieszczona na tekturowej podstawie



Przesyłka
GRATIS

Zamów wygodnie na
www.UlubionyKiosk.pl

Prosta ładowarka bezprzewodowa do smartfonów



Baterie smartfonów muszą być ładowane często, zwykle częściej niż raz dziennie. Niestety ich kable ładujące zwykle ulegają po pewnym czasie uszkodzeniu z powodu nadmiernego użytkowania, co może być denerwujące – zwłaszcza jeśli zawiodą, gdy są najbardziej potrzebne. Ładowarka bezprzewodowa może być wygodną alternatywą; po prostu ustawiasz ją i zapominasz o problemie. Ale jak to dokładnie działa?

Chociaż ładowanie bezprzewodowe może wydawać się niedawnym wynalazkiem, to istnieje ono od ponad stu lat. Jego powstanie to wynik doświadczeń słynnego serbsko-amerykańskiego wynalazcy, Nikoli Tesli.

Zasada działania

Pod koniec XIX wieku Nikola Tesla z powodzeniem przysyłał energię elektryczną drogą powietrzną. Użył do tego, procesu zwanego sprzężeniem rezonansowo-indukcyjnym,

które działało poprzez wytworzenie pola elektromagnetycznego przez nadajnik (do wysyłania energii elektrycznej) i odbieranie go przez odbiornik, wtedy aby zasilić żarówki w nowojorskim laboratorium Tesli. Kilka lat później opatentował cewkę transformator rezonansowy, nazywaną obecnie cewką Tesli – wieżę z kopułową elektrodą na szczycie, z której strzelały wyładowania elektryczne. Tesla miał znacznie szerszą wizję bezprzewodowej sieci energetycznej, ale jego marzenia nigdy

się spełniły. Teraz ta sama podstawowa zasada ładowania indukcyjnego może być wykorzystana do bezprzewodowego ładowania smartfona.

Budowa

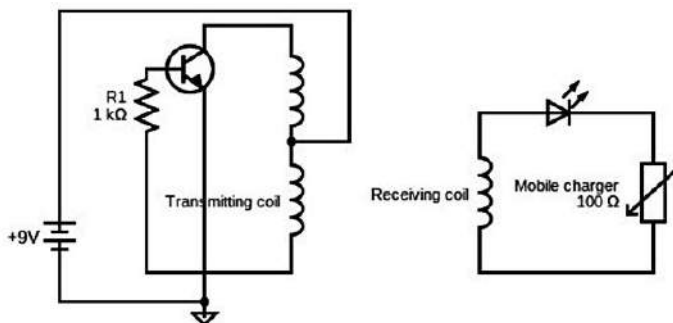
Do wykonania cewek możesz użyć, plastikowej butelki o średnicy około 8 cm, tak jak zrobił to autor. Możesz też do tego celu użyć rury PCV lub innego cylindrycznego przedmiotu. Nawin



Rysunek 1.



Rysunek 2.



Rysunek 3.



Rysunek 4.

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

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

Drut miedziany emaliowany o średnicy 0,3 – 0,4 mm – 5 metrów (do cewek nadajnika i odbiornika)

Rezystor 1 kΩ – 1 szt. do ograniczania prądu bazy tranzystora

Tranzystor BC547 – 1 szt.

LED – 1 szt. jako prostownik jednopółprzewodnikowy

Kabel USB – 1 szt. do ładowania telefonu komórkowego

Bateria 9 V ze złączem – 1 szt. do zasilania nadajnika

15 zwojów emaliowanego drutu miedzianego, wykonaj wyprowadzenie i nawiń kolejne 15 zwojów, coś jak transformator z centralnym odczepem. Podobnie nawiń 15 zwojów emaliowanego drutu miedzianego, aby utworzyć cewkę odbiorczą. Jako elementu czynnego nadajnika autor użył tranzystora npn, typu BC547, ale zamiast tego dla lepszych rezultatów można zastosować tranzystor mocy TTC5200. Podłącz rezystor 1 k Ω do bazy tranzystora, aby ograniczyć jej prąd. Cewkę z odczepem środkowym należy podłączyć do bieguna dodatniego akumulatora, którego biegun ujemny należy

połączyć z emiterem tranzystora. Pozostałe dwa końce cewki nadajnika należy podłączyć do zacisków kolektora i bazy tranzystora, jak pokazano na **rysunku 3**, poprzez rezystor ograniczający prąd 1 k Ω . Końcówki cewki odbiorczej połączone są z przewodem USB poprzez diodę LED, która służy tu do sygnalizacji oraz jako prostownik jednopółkowy.

Zasada działania

Cewka nadajnika połączona z baterią przez tranzystor tworzy obwód drgający. Prąd oscylacyjny wewnątrz cewki nadawczej powoduje, że emituje ona pole elektromagnetyczne

o częstotliwości ok. 1 MHz. To oscylujące pole elektromagnetyczne indukuje prąd elektryczny w cewce odbiornika. Indukowany prąd jest prądem przemiennym, dlatego należy go przekształcić na prąd stały w celu ładowania mobilnego. Dioda LED zastosowana w obwodzie prostuje prąd na prąd stały, poza tym, że służy jako wskaźnik. ■

Sakthivignesh r.

Ten artykuł był wcześniej opublikowany na łamach „EFY”, sierpień 2022 (efymag.com)

Quiz: Zrozumieć tranzystory bipolarne (cykl artykułów w EdW 7, 8, 9, 10, 11/2022)

Ile jest konfiguracji, czyli sposobów włączania tranzystorów bipolarnych?

- ☐ 2
- ☐ 3
- ☐ 4

Wzmocnienie prądowe w układzie ze wspólną bazą (WB) oznaczone literą α wynosi:

- ☐ ok. 0,99
- ☐ ok. 10
- ☐ ok. 100

Wzmocnienie prądowe w układzie ze wspólnym emiterem (WE) oznaczone literą β wynosi:

- ☐ ok. 1
- ☐ ok. 5
- ☐ ok. 100 i więcej

Największą rezystancję wejściową ma tranzystor włączony w układzie:

- ☐ WB
- ☐ WE
- ☐ WC

Największą rezystancję wyjściową ma tranzystor włączony w układzie:

- ☐ WB
- ☐ WE
- ☐ WC

Największą częstotliwość graniczną pracy tranzystora osiąga się w układzie:

- ☐ WB
- ☐ WE
- ☐ WC

Efekt Millera, tj. multiplikacja pojemności wejściowej występuje w układzie:

- ☐ WB
- ☐ WE
- ☐ WC

W układzie nazywanym wtórnikiem emiterowym tranzystor jest włączony w konfiguracji:

- ☐ WB
- ☐ WE
- ☐ WC

Wzmocnienie napięciowe wtórnika emiterowego wynosi:

- ☐ ok. 1
- ☐ ok. 10
- ☐ ok. 100

Bootstrapping jest rozwiązaniem minimalizującym wpływ obwodu polaryzacji na zmniejszenie impedancji wejściowej w układzie:

- ☐ WB
- ☐ WE
- ☐ WC

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 20.01.2023.

REKLAMA

świat radio

Magazyn wszystkich użytkowników eteru
KRÓTKOFALARSTWO CB RADIO TECHNIKA

przejrzyj i kup na
www.ulubionykiosk.pl



The advertisement shows a black ICOM radio with a digital display showing 7.070.50, resting on a wooden stump. In the background, there's a scenic view of a mountain range. To the right, a stack of 'Świat Radio' magazines is shown, with the top issue featuring a radio and the title 'Księga G106'.

Wykrywanie i klasyfikacja obiektów za pomocą Raspberry Pi i uczenia maszynowego wykorzystującego platformę Edge Impulse

Gdyby maszyny potrafiły rozpoznawać przedmioty tak, jak robią to ludzie, to byłoby to całkiem interesujące. Wykrywanie obiektów w obrazie jest obecnie bardzo popularnym tematem. Stwórzmy więc kamerę do wykrywania obiektów, która może na żywo klasyfikować przedmioty i ludzi, a każda aktywność jest wyświetlana na żywo w sieci przy użyciu adresu IP.

Wykorzystamy platformę programistyczną Edge Impulse do trenowania modelu ML (uczenia maszynowego) i integracji go z modułem Raspberry Pi oraz pozyskiwania na żywo, wideo i obrazów, za pośrednictwem interfejsu kamery. Projekt będzie mógł pracować z Internetem i bez niego. Wykorzystując Internet możemy użyć projektu do takich działań jak:

- Monitorowanie drzwi na żywo w celu ostrzegania, gdy wejdzie nieznana osoba.
- W przemyśle, do klasyfikowania i segregacji obiektów za pomocą ramion robotycznych.
- Liczenie owoców na drzewie lub w separatorze.

Istnieją różne wersje i odmiany systemu operacyjnego Raspberry Pi, ale my musimy przygotować kartę SD z najnowszym systemem operacyjnym. Aby załadować system operacyjny na kartę SD, wykonaj następujące czynności wymienione poniżej:

1. Pobierz Raspberry Pi Desktop Imager na komputer.
2. Uruchom Raspberry Pi Imager.
3. Wybierz system operacyjny jako Raspberry Pi OS (32-bitowy).
4. Wybierz kartę SD.
5. Wybierz Zapisz.
6. Włóż kartę SD do gniazda w Raspberry Pi.
7. Podłącz Raspberry Pi do zasilania oraz klawiaturę, mysz i monitor.
8. Jeśli system operacyjny zostanie poprawnie zainstalowany, zobaczysz komunikat „Welcome to Raspberry Pi Desktop”.

Teraz, gdy system operacyjny jest gotowy, podłącz kamerę USB lub kamerę Raspberry Pi. Jeśli używana jest kamera Raspberry Pi (podłączona taśmą przewodową), to najpierw włącz interfejs kamery w ustawieniach konfiguracji Raspberry Pi. Następnie zainstaluj na Raspberry Pi platformę Edge Impulse, uruchamiając następujące polecenia na terminalu:

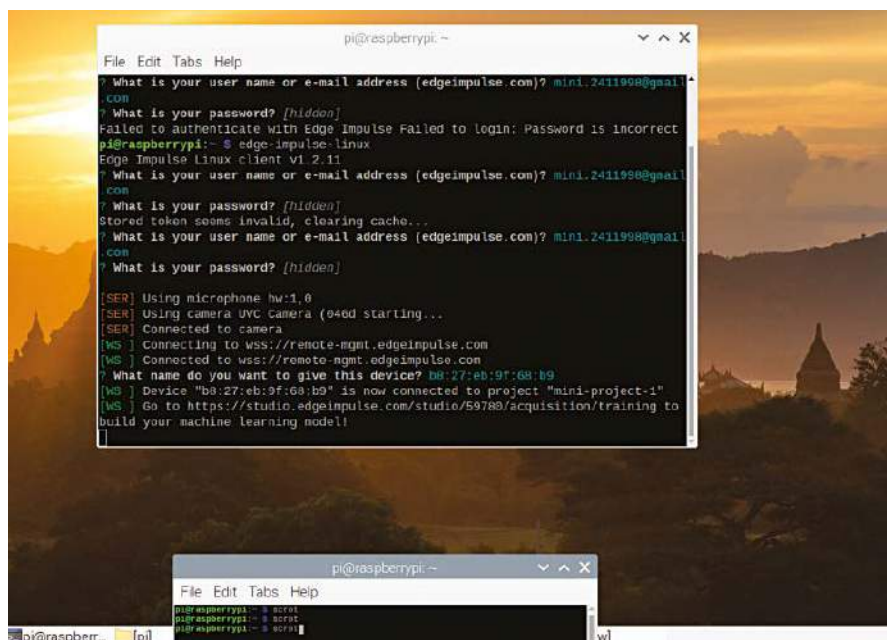
```
curl -sL https://deb.nodesource.com/
setup_12.x | sudo bash -
```

```
sudo apt install -y gcc g++
make build-essential nodejs sox
gststreamer1.0-tools gststreamer1.0-
plugins-good gststreamer1.0-plugins-base
gststreamer1.0-plugins-base-apps
sudo npm install edge-impulse-linux -g
--unsafe-perm
```

Po skonfigurowaniu Edge Impulse utwórz i wytrenuj model ML. W tym celu otwórz stronę Edge Impulse, a następnie wprowadź swoje imię, nazwisko oraz identyfikator e-mail i zarejestruj się. Utwórz nowy projekt w Edge Impulse, otwórz terminal Linux i uruchom polecenie `edge-impulse-linux`.



Rysunek 1. Prototyp autorów



Rysunek 2. Uruchamianie Edge Impulse

Po wybraniu projektu urządzenie spróbuje połączyć się z Edge Impulse. Jeśli połączenie jest prawidłowe, zostanie wyświetlone wykrycie urządzenia kamery Edge Impulse Raspberry Pi i otrzymasz link. Otwórz łącze w przeglądarce internetowej i uzyskaj interfejs akwizycji danych, aby zrobić zdjęcie lub inne dane do trybu ML. Do modelu ML można robić zdjęcia przedmiotów, takich jak: butelka lub kubek, oraz dowolnych innych obiektów (w tym ludzi).

W celu pozyskania danych zrób co najmniej 100 zdjęć różnych obiektów, które chcesz sklasyfikować i użyć do trenowania oraz testowania modelu ML. Aby zrównoważyć

dane, stosunek podziału obiektów uczących do testowych powinien wynosić 70:30.

Przejdź do pulpitu nawigacyjnego, w którym jako „Labelling method” (metoda wykrywania) powinno być wybrane „Bounding boxes” (obwiednie). Oznacz wszystkie obiekty i zobacz je w „Labelling Queue” (kolejka oznaczania). Następnie przejdź do Impulse Design, gdzie szerokość i wysokość obrazu powinna wynosić 320×320 punktów. Możesz zmienić nazwę projektu wykrywania obiektów, a następnie kliknąć „Save Impulse”.

W sekcji „Image” (obraz) skonfiguruj blok przetwarzania i wybierz surowe dane u góry

ekranu. Teraz zapisz parametry w postaci RGB lub skali szarości.

Ze względu na różne wymiary obrazów następuje redukcja ich wymiarów.

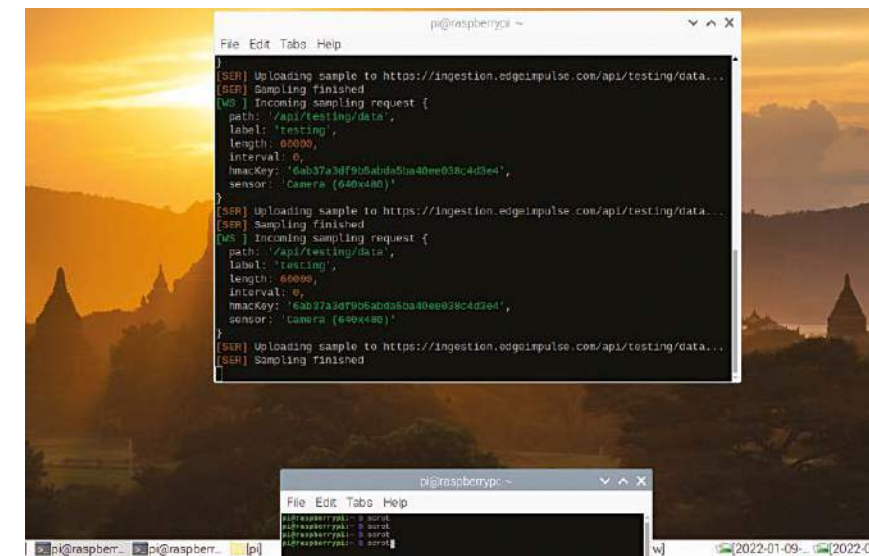
W sekcji Object Detection (wykrywanie obiektu) można ustawić liczbę cykli treningowych i szybkość uczenia się. W autorskim prototypie liczba cykli treningowych została ustalona na 25, a szybkość uczenia się na 0,015.

Możesz teraz rozpocząć trenowanie modelu. Po wytrenowaniu modelu możesz uzyskać stopień precyzji. Aby sprawdzić poprawność modelu, przejdź do „Model Testing” (testowanie modelu) i wybierz opcję „Classify All” (klasyfikuj wszystko). Następnie przejdź do klasyfikacji na żywo.

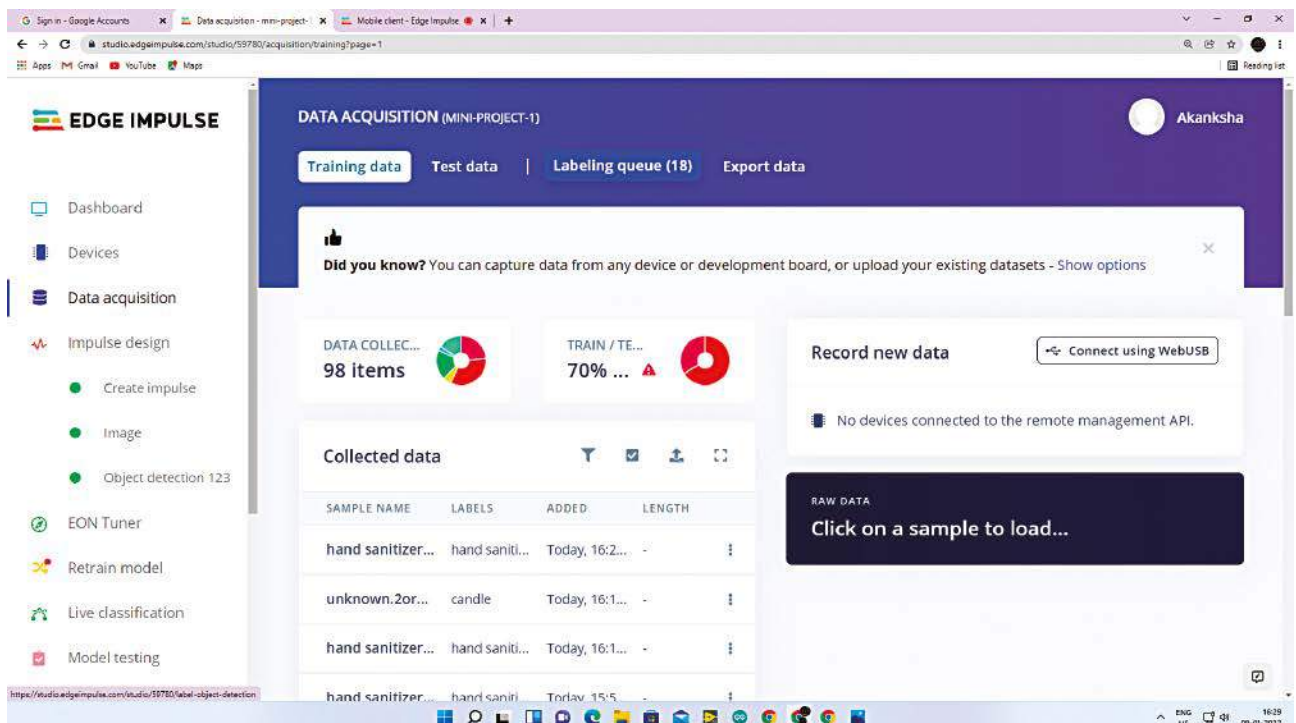
Umieść obiekt przed kamerą, a obraz obiektu będzie można zobaczyć w sekcji „Live Classification” w Edge Impulse bez żadnej etykiety sklasyfikowanego i wykrytego obiektu, takiego jak butelka, kubek lub osoba. Jeśli chcesz zobaczyć go z adresem IP, uruchom polecenie *edge-impulse-linux-runner*.

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

Raspberry Pi 3 B-1
USB camera/Raspberry Pi camera-1
Keyboard-1
Monitor-1
Mouse-1
SD Adaptor (32 GB)
HDMI to VGA cable
5 V power adaptor with USB Type-C connector
SD card reader



Rysunek 3. Terminal



Rysunek 4. Akwizycja danych



Akanksha Gupta i Sagar Raj

Sagar Raj jest założycielem i dyrektorem Shoolin Labs, Jaipur (Radžastan). Jest założycielem i dyrektorem Lifegraph Biomedical Instrumentation, Buxar oraz trenerem IoT w Lifegraph Academy, Incubation Centre, IIT Patna

Ten artykuł był wcześniej opublikowany na łamach „EFY”, lipiec 2022 (efymag.com)

Świat projektantów i programistów dla elektroniki w nowej
odśłonie.
Odwiedź wiecznie młody

ELPORTAL.pl

Przedstawiamy początkowe fragmenty dwóch projektów ze zbioru kilkudziesięciu projektów dostępnych wyłącznie dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl. W rubryce **DIY PLUS** zamieszczamy aktualnie najciekawsze projekty publikowane w Internecie w formacie open source. Prenumeratorów EdW zapraszamy do zapoznania się na www.elportal.pl z niezwykle inspirującymi zasobami rubryki **DIY PLUS**.

Wysokiej klasy przedwzmacniacz mikrofonowy ze zmienną kompresją, regulowanym wzmacnieniem oraz funkcją redukcji szumów

Układ pokazany w bieżącym odcinku „Electronics-Lab” bazuje na układzie scalonym SSM2166. To kompletny system kondycjonowania sygnału mikrofonowego. Zaprojektowany jest on dla zastosowań audio, przede wszystkim w paśmie „wokalnym”. Układ scalony realizuje funkcje wzmacnienia sygnału, rozpoznaje amplitudę (RMS) sygnału oraz dokonuje kondycjonowania polegającą na zmiennej kompresji, ekspansji w zakresie słabych sygnałów oraz ograniczenia poziomu sygnałów silnych. Kluczową częścią układu jest wzmacniacz VCA o zmiennym, regulowanym napięciem wzmacnieniu do 60 dB w paśmie do 30 kHz. Dodatkowe wzmacnienie sygnału fonicznego realizuje wzmacniacz wejściowy.

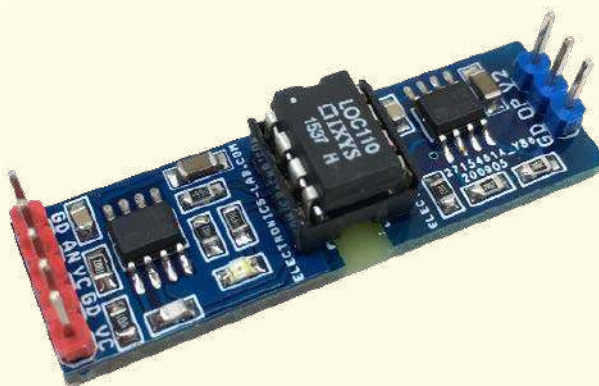


Dokończenie artykułu na stronie:
<https://bit.ly/3HFloaa>

Optycznie izolowane wejście analogowe dla Arduino

Bieżący projekt to niewielki moduł niezmiernie użyteczny, gdy zachodzi potrzeba podłączenia analogowego sygnału do systemu mikroprocesorowego, a równocześnie istnieje problem potencjałów mas. Moduł ten zapewnia optyczną izolację galwaniczną przenosząc wiernie sygnał analogowy. Warto go stosować jako układ pośredniczący w torze analogowym, co pozwala na bezpieczną współpracę systemu uP z różnego rodzaju czujnikami w automatyce, w przemyśle, a szczególnie w pracach „w terenie”.

Dokończenie artykułu na stronie:
<https://bit.ly/3HH0o0e>



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



Ulubiony Kiosk - Twoje internetowe centrum prasy specjalistycznej i hobbystycznej!

Poznaj najważniejsze na rynku tytuły z segmentów

ZDROWIE I RODZINA
DOM, OGRÓD I WNĘTRZA
FOTOGRAFIA, EDUKACJA I HI-TECH
ELEKTRONIKA I AUTOMATYKA
MUZYKA I DŹWIĘK

Sprawdź na **UlubionyKiosk.pl**