



ELEKTRONIKA

dla wszystkich

nr 2/2025 (349) • luty • www.elportal.pl

DIY PLUS
tylko dla prenumeratorów

Synteзатор MIDI

PROJEKTY dla elektroników

- ▶ Regulowane „sztuczne obciążenie” z wbudowanym generatorem zegarowym do testowania zasilaczy, przetwornic napięcia i akumulatorów
- ▶ Mini sterownik LED

DIY dla wszystkich

- ▶ Zasilanie 5 V z pojedynczych ogniw baterii
- ▶ Bezpieczna przystawka żelazka do ochrony odzieży

TUTORIALE

- ▶ Migające diody LED i śliniacy się inżynierowie, część 17
- ▶ Audio OUT: Wokoder analogowy, część 5. Budowa filtrów
- ▶ Chirurgia obwodowa: Wzmacniacze operacyjne logarytmiczne i wykładnicze, część 2
- ▶ Co potrafi współczesny oscyloskop cyfrowy?



Programowalne obciążenie Arduino



Pomocna dłoń



automatykaB2B.pl

EP.com.pl

Największy portal dla elektroników konstruktorów

eprasa.pl/a64b3da2

FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

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

www.piekarz.pl





Najbardziej popularne kity AVT

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



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



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



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



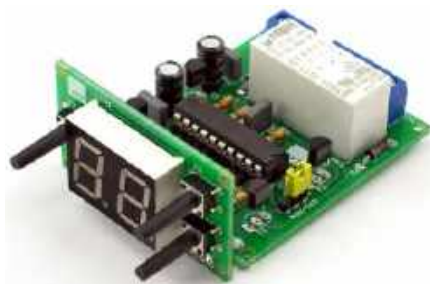
AVT5540 Radio FM z RDS
<https://sklep.avt.pl/avt5540.html>



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>



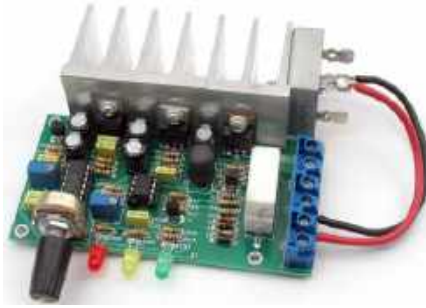
AVT990 Automatyczny włącznik światła
<https://sklep.avt.pl/avt990.html>



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



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



AVT3120 Automatyka ładowarka
akumulatorów ołowiowych
<https://sklep.avt.pl/avt3120.html>



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



Pełna oferta na: sklep.avt.pl

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

–20%
NA START
181,40 zł

–30%
po pierwszym roku
prenumeraty
158,80 zł

–40%
po drugim roku
prenumeraty
136,10 zł

–50%
po trzecim roku
nieprzerwanej prenumeraty
113,40 zł

Odkryj korzyści z **prenumeraty drukowanej** – większe oszczędności z każdym rokiem!

Rozpocznij swoją przygodę z *Elektroniką dla Wszystkich*. Decydując się teraz na roczną prenumeratę drukowaną, otrzymasz nie tylko dostęp do najnowszych wydań, ale i **znakomity start dzięki zniżce 20%** na pierwsze zamówienie!

Prenumerata to nie tylko wygoda dostępu do treści, ale także sposób na znaczące oszczędności. Dołącz do grona naszych stałych czytelników i ciesz się coraz lepszymi warunkami.

Im dłużej jesteś z nami, tym więcej oszczędzasz:

- po roku nieprzerwanej prenumeraty zapewnimy Ci **30% rabatu** na kolejny rok,
- po dwóch latach wierności zaoferujemy **40% rabatu**,
- po trzech latach lojalności osiągniesz **najwyższy poziom rabatu – 50%**!

Jak otrzymać rabat za lojalność?

Zaloguj się na swoje konto prenumeratora na www.UlubionyKiosk.pl i zamów prenumeratę, korzystając z przycisku PRZEDŁUŻ w zakładce „Prenumeraty”.

Przeglądaj wcześniej, płać mniej – postaw na **e-prenumeratę**!

Wybierz prenumeratę cyfrową PDF i ciesz się dostępem do czasopisma nawet 7 dni przed oficjalną premierą w kioskach. Oszczędzaj czas i pieniądze – skorzystaj z **rabatu 30%** na roczną e-prenumeratę w cenie 126,90 zł.

Dodatkowa oferta dla prenumeratorów wersji drukowanej: jeśli już subskrybujesz wersję papierową, możesz dokupić równoległe e-wydania w cenie 36,20 zł/rok – z **niesamowitym rabatem 80%**.

Zyskaj nieograniczony dostęp do zasobów dla pasjonatów elektroniki!

Tylko prenumeratorzy mają pełny dostęp do:

- cyfrowego archiwum *Elektroniki dla Wszystkich* na www.elportal.pl/archiwum
- projektów DIY+ na www.elportal.pl/diy

Zamów prenumeratę drukowaną lub e-prenumeratę na www.UlubionyKiosk.pl lub przez przelew na konto Wydawnictwa AVT, a po zaksięgowaniu wpłaty wyślemy Ci mailowo kod dostępu do portalu.

ARCHIWUM



Zacznij korzystać z pełnych zasobów już dziś!

8



Projekty dla elektroników:

Syntezator MIDI	8
Regulowane „sztuczne obciążenie” z wbudowanym generatorem zegarowym do testowania zasilaczy, przetwornic napięcia i akumulatorów	22
Mini sterownik LED	28
Programowalne obciążenie bazujące na Arduino	34

22



Tutoriale:

Ekscytacje Maxa:

• Migające diody LED i śliniacy się inżynierowie, część 17	39
• Sprytne porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania ..	42
Audio OUT: Wokoder analogowy, część 5. Konstrukcja filtrów	46
Chirurgia obwodowa:	
Wzmacniacze operacyjne logarytmiczne i wykładnicze, część 2	54
Edukacja w EdW dla szkół i uczelni.	
Wykład 27: Generatory wysokiego napięcia	58
Co potrafi współczesny oscyloskop cyfrowy?	69

28



DIY dla wszystkich:

Zasilanie 5 V z pojedynczych ogniw baterii	74
Bezpieczna przystawka żelazka do ochrony odzieży	77

Elektronika dla Wszystkich – Junior:

Ósme spotkanie z najmłodszymi pasjonatami elektroniki	81
---	----

Na zdjęciu na okładce Krystian, Młodzi Entuzjaści Elektroniki, Szkoła Podstawowa nr 86, Wrocław

34



DIY PLUS tylko dla prenumeratorów zamawiających prenumeratę na www.UlubionyKiosk.pl

8-kanałowy sterownik serwomechanizmu RC przez łącze RF z modułem RF NRF24L01 – kompatybilny z Arduino	91
Nadajnik zdalnego sterowania RF z dwoma joystickami i modułem RF NRF24L01 – sterowanie 2 joystickami	91

77



Rubryki stałe:

Prenumerata	3
Od redakcji	5
Pocztka	6

A za miesiąc w marcowym EdW

✱ Soundbar

Jak brzmią głośniki wbudowane we współczesne, coraz bardziej smukłe telewizory, wszyscy dobrze wiemy. Dlatego soundbar to dzisiaj podstawa każdego centrum domowej rozrywki. Z uwagi na cenę takiego sprzętu, samodzielna budowa podobnego urządzenia może okazać się ekonomicznie uzasadniona, przede wszystkim da jednak dużo frajdy amatorom samodzielnych konstrukcji audio, dla których średniej jakości półśrodek w postaci gotowego urządzenia ze sklepu często nie jest optymalnym wyborem.

✱ Tester i kalibrator multimetrów

Jako hobbysta elektroniki z pewnością posiadasz co najmniej kilka multimetrów cyfrowych różnej klasy, których używasz niemal na co dzień. Dzięki testerowi multimetrów wreszcie będziesz mógł zweryfikować dokładność pomiarową posiadanych egzemplarzy. W przypadku droższych modeli urządzenie umożliwi również ich kalibrację poprzez korektę odchyłek.

✱ Tłumik RF 0–110 dB do generatorów sygnałów

Ten tłumik stanowi doskonałe uzupełnienie opublikowanego w EdW 1/2025 precyzyjnego generatora sygnału AM-FM DDS. Może być również używany w pracy z dowolnym generatorem sygnałów, umożliwiając wygodną regulację poziomu wyjściowego w szerokim zakresie tłumienia.

✱ Szerokokokresowy omomierz

Mierzenie bardzo małych rezystancji przy użyciu jedynie podstawowego sprzętu pomiarowego, może być nie lada wyzwaniem. To urządzenie pozwoli Ci zmierzyć praktycznie każdą rezystancję – od kilku miliomów do wielu megaomów! Warto więc takowe zmontować. Z pewnością przyda się do pracy z cewkami w projektach audio, ale nie tylko.

✱ Wartościowe Tutoriale

✱ Projekty DIY

✱ Juniorzy EdW złożą kolejny zestaw z serii AVTEDU

**W kioskach
od 25 lutego**

Życie wyborami ustane

Cennym jest, gdy w życiu wszystko pięknie „gra i buczy”, toteż na okładkę numeru trafił tym razem zaawansowany syntezator MIDI. Jest to sprzęt zrealizowany w oparciu o mikrokontroler PIC18LF25K50 oraz procesory dźwięku dsPIC33EP512MC502-I/SP. Temu niebanalnemu cacku poświęciliśmy aż 14 stron bieżącego numeru EdW. Uznaliśmy, że naprawdę warto!

Z uwagi na porę roku, potyczki dnia z nocą wygrywa wciąż ów drugi zawodnik. W styczniu stanęliśmy z nim w szranki proponując Czytelnikom montaż sterownika paneli LED o dużej mocy. W tym numerze tematu nie odpuszczamy, ale do walki z mrokiem proponujemy zaprzęgnąć tym razem nieco mniejsze działko – mini przetwornicę do zasilania diod LED, czyli niewielki i niedrogi moduł pozwalający zasilić stosunkowo duże białe diody LED i niewielkie panele LED natywnie zasilane napięciem 12 V ze źródła zasilania 5 V, na przykład z power banku z portem USB.

Proponujemy również montaż dwóch projektów sztucznych obciążeń, niewątpliwie bardzo przydatnych przyrządów warsztatowych do testowania m.in.: zasilaczy, przetwornic napięcia i akumulatorów. Pierwsze potrafi rozproszyć, w zależności od przyjętego wariantu montażowego, do 18 W lub do 50 W mocy w zakresie napięć 3,3 V...30 V przy regulowanym, stałym prądzie w zakresie 30 mA...1900 mA. Co ważne, przyrząd posiada zabezpieczenia przed przegrzaniem i odwrotną polaryzacją zasilania. Druga konstrukcja oparta jest na Arduino i rezystorach dużej mocy. Może stanowić obciążenie do 70 W mocy ciągłej, przy napięciu do 15 V i natężeniu 4,7 A.

Jak każdego miesiąca Czytelnik znajdzie w numerze sporą garść wiedzy teoretycznej, m.in. w drugiej części artykułu o wzmacniaczach operacyjnych logarytmicznych i wykładniczych oraz w wykładzie na temat generatorów wysokiego napięcia. Ciekawą pozycją jest także artykuł poświęcony możliwościom współczesnych oscyloskopów cyfrowych.

Dla fanów audio z pewnością obowiązkową pozycją będzie materiał z cyklu Audio Out. Tym razem będzie to piąta już część serii poświęconej wokoderowi analogowemu, w której omówiona zostanie budowa filtrów. Z kolei pasjonaci programowania i malowania animacji za pomocą diod RGB będą mogli oddać się lekturze siedemnastej już części cyklu Maxa Maxfielda o NeoPixelach.

Zaciekawić mają prawo również treści z sekcji DIY. W bieżącym numerze poruszone zostaną tematy takie jak zasilanie 5 V uzyskiwane z pojedynczych ogniw baterii, czy rozbudowa klasycznego żelazka, mająca służyć ochronie ubrań podczas ich prasowania.

W ramach ósmego spotkania Juniorów EdW zbudowany zostanie pomysłowy gadżet o nazwie „wspomagacz wyboru”, który niezdecydowanemu Juniorowi w sposób efektywny i całkiem przyjemny pomoże podejmować, nie lada „trudne” skądinąd decyzje. Decyzje w których urządzenie faktycznie się sprawdzi mogą dotyczyć wyboru potrawy do spożycia w trakcie śniadania, rodzaju filmu do oglądnięcia na spotkaniu z przyjaciółmi, czy też kolejności prac porządkowych. Może nawet stać się kompanem do gry w „kamień, papier, nożyce”. Dodatkowo Junior będzie miał możliwość przygotowania własnych plansz wyboru, co doda urządzeniu jeszcze większej funkcjonalności.

„Wspomagacz wyboru” staje się również okazją i pretekstem do głębszej refleksji nad podejmowaniem decyzji i ich konsekwencjami. O ile zbudowana zabawka pozwoli sprowadzić drobne, choć dla dzieci i nastolatków czasem też „trudne” wybory do poziomu zabawy, to każdy Junior prędzej czy później stanie przed decyzjami niosącymi poważniejsze konsekwencje, których nie będzie można pozostawić ślepego losowi.

Junior, stając przed różnorodnymi wyborami, podmie niejedną decyzję – nie wszystkie będą trafne. W końcu nadejdzie moment, gdy trzeba będzie pożegnać się z juniorskim przydomkiem i wkroczyć w dorosłe życie. Nietrafiona decyzja, w zależności od jej wagi, może kosztować naprawdę wiele. Czasami nasze wybory wpływają na losy innych, a odpowiedzialność za nie wymaga trwania w podjętej decyzji przez dłuższy czas, nierzadko za cenę własnych wyrzeczeń. Jednak warto czasem zacisnąć zęby, by mogło się zrealizować coś ważnego i pięknego.

Wszystkim Czytelnikom życzymy udanych decyzji, tych codziennych, małych, i tych strategicznych, wielkich. Indywidualnie bądź wzajemnie satysfakcjonujących. Tych o życiu w pojedynkę i tych o życiu we dwoje.

Mariusz Ciszewski



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

Artykuły Czytelników

Droga Redakcjo,

czy nadal można publikować w EdW artykuły własnego autorstwa?

Rafał

Red. Jak najbardziej istnieje taka możliwość. Po zapoznaniu się z materiałem skontaktujemy się z autorem w celu ustalenia dalszych szczegółów w zakresie ewentualnej publikacji.

Zachęcamy do przysyłania nam Waszych autorskich materiałów, jednocześnie, z uwagi na spory ruch na redakcyjnej skrzynce pocztowej, z góry prosimy uzbroić się w cierpliwość w oczekiwaniu na naszą odpowiedź.

Sportowe rywalizacje

Szanowna Redakcjo,

rzadko zabieram głos w publicznych dyskusjach (lub wysyłam e-maile gdziekolwiek), ale postanowiłem podzielić się swoimi przemyśleniami na temat listu jednego z Czytelników „Jak to jest z błędami?” z ubiegłego miesiąca. Zgadzam się, że artykuły w sekcji DIY mają swój specyficzny charakter – jednych przyciągają, innych mogą drażnić. Osobiście, gdy zacząłem śledzić te teksty, poczułem pewnego rodzaju rywalizację z ekspertem, który w trakcie dodaje swoje uwagi, a na końcu szczegółowo analizuje oryginalny tekst autora. Traktując to jak sport, zaczynam samodzielnie wychwytywać ewentualne błędy, by potem sprawdzić, ile z nich udało mi się dostrzec, a co przeoczyłem. Okazało się, że stało się to dla mnie swoistą zabawą, wciągnęło mnie, podobnie jak wciągnąć potrafi człowieka gra w „znajdź pięć różnic” na dwóch obrazkach. Może nie jestem jedynym, kto w ten sposób podchodzi do tych artykułów, ale pomyślałem, że warto podzielić się tym z Redakcją i Czytelnikami. Jeśli chodzi o merytorykę, artykuły DIY są zazwyczaj łatwe do wykonania (nawet, jeśli trzeba wprowadzić sugerowane poprawki), a ich realizacja nie sprawia większych trudności. Czasami realizuję je po prostu z ciekawości lub buduję podobne projekty, albo też staram zasymulować, bo w zabieganym świecie rzadko znajduję czas na realizację bardziej zaawansowanych konstrukcji jak te w sekcji „Dla Elektroników”. Co nie zmienia faktu, że wiele z nich naprawdę miałbym chęć zrobić. Jeśli kiedyś będę miał więcej czasu, z pewnością wiem, do czego chcę wrócić.

Pozdrawiam Redakcję i Czytelników

Filip

Red. Dziękujemy za ciekawy głos i kibicujemy wspomnianym rywalizacjom. W przypadku wygranej w postaci większej liczby odnalezionych błędów, proszę się nimi z nami podzielić, z pewnością zainteresują Redakcję i Czytelników. Pozdrawiamy oczekując na dalsze wieści.

Inżynieria wsteczna

Szanowna Redakcjo,

pisze do Was te słowa pan już niemłody, taki, co to pamięta czasy, gdy tranzystor był nowinką, a lutownicę grzało się w piecu. Oczywiście to pierwsze to rzeczywistość, a drugie to taka satyra na upływający czas, ale do rzeczy. Od lat jestem wiernym czytelnikiem Waszych miesięczników – od „Elektronika dla Wszystkich” po „Elektronikę Praktyczną” – i nie tylko. Całkiem dobrze pamiętam lata dziewięćdziesiąte gdy Wasze czasopisma pojawiły się w kioskach Ruchu. Budowałem wtedy wiele układów, do których płytki wraz z instrukcjami nabywałem w lokalnym sklepie elektronicznym.

Ach, ileż to było radości z takich wynalazków jak prosta syrena elektroniczna, kit 2010 co wyla jak ta z Kojaka! Albo nieco trudniejszy zamek szyfrowy, kit AVT-144, podobny do tych, jakie widywało się tylko na filmach jak ten z udziałem pewnego sprytnego wynalazcy, co to zawsze miał spinacz w kieszeni. Dziś próżno szukać tych zestawów w Waszej ofercie (choć pewnie zamiast nich są teraz lepsze). A jednak mam gdzieś jeszcze w swoich gratach te płytki, czekające na lepsze czasy. Powiem szczerze, żal było je składać, bo to przecież ostatnie płytki, których nie będzie już potem gdzie kupić.

Choć lata swoje mam, to z komputerem jakoś sobie radzę, ale za tym całym nowym światem, co to niby jest cyfrowy i w chmurach zawieszony, trudno mi nadążyć. Więc pytam Was, Młodzi i Mądrzy – czy istnieje dziś taki cudowny program komputerowy, który weźmie skan płytki drukowanej i wypluje gotową dokumentację produkcyjną? Coś, co pozwoliłoby zlecić u lokalnego dostawcy kilka nowych płytek, zanim te moje resztki starych zestawów całkiem rozsypią się w proch?

Z wyrazami szacunku i ukłonami

**Wasz wierny czytelnik,
co się z lutownicą w rękę urodził**

Red. Dziękujemy za piękny list i sentymentalne nawiązanie do lat dziewięćdziesiątych ubiegłego wieku. Nam również zrobiło się bardzo nostalgicznie, a jednocześnie miło, gdyż poczuliśmy satysfakcję z dobrze wykonywanej pracy. Rzeczywiście, wiele materiałów, zwłaszcza sprzed niemal trzech dekad, nie przetrwało próby czasu. Aby wskrzesić je na nowo, należałoby w wielu przypadkach stworzyć dokumentację od podstaw, co choć nie jest niemożliwe, bywa czasochłonne.

Z tego, co nam wiadomo, obecnie nie istnieje ogólnodostępne i szeroko rozwijane narzędzie dedykowane inżynierii wstecznej. Ciekawą alternatywą może być program KrzysioPCB, będący nakładką na Eagle'a – popularne oprogramowanie do projektowania płytek drukowanych. Niestety, jego użytkowanie wymaga dobrej znajomości Eagle'a, a sam proces nie należy do szczególnie intuicyjnych, co może być dalekie od oczekiwań Autora listu.

Warto jednak zwrócić uwagę na KiCad, który w swoich najnowszych wersjach oferuje funkcje przydatne przy inżynierii wstecznej. Zarówno w edytorze schematów, jak i w edytorze PCB można wklejać obrazy jako tło obszaru roboczego. W praktyce oznacza to możliwość dodania zrzutu ekranu lub skanu schematu do edytora schematów oraz skanu czy zdjęcia płytki drukowanej do edytora PCB. Następnie, po odpowiednim skalowaniu obrazu, można sukcesywnie odtwarzać projekt – dodawać elementy, rysować ścieżki i jednocześnie uzupełniać schemat. Na końcu tego procesu otrzymujemy schemat i projekt PCB, z którego możemy wygenerować pliki Gerber i zlecić wykonanie płytek w fabryce.

Niestety, oba opisane rozwiązania są czasochłonne i wymagają sporo zaangażowania. W erze technologii chmurowych i rosnącej roli sztucznej inteligencji można by się spodziewać programu, który automatycznie analizuje skany płytek drukowanych, rozpoznaje kształt PCB, ścieżki miedziane, grafiki na warstwie silkscreen, otwory, przelotki i pady, a następnie generuje gotowy projekt do ewentualnych poprawek i eksportu w formacie Gerber. Na dzień dzisiejszy jednak takie rozwiązanie nie jest nam znane.

Subscribe to Elektor's newsletter and get the chance to

WIN

a Raspberry Pi Pico W board



www.elektor.com/eda



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



Be one of the 10 fortunate winners!



elektor
design > share > earn



Syntezytor MIDI mieści się w obudowie o wymiarach 150 mm × 100 mm × 40 mm. Można użyć innej obudowy, o ile jest większa, a jej wysokość zależy od zastosowanego radiatora

Syntezytor MIDI



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

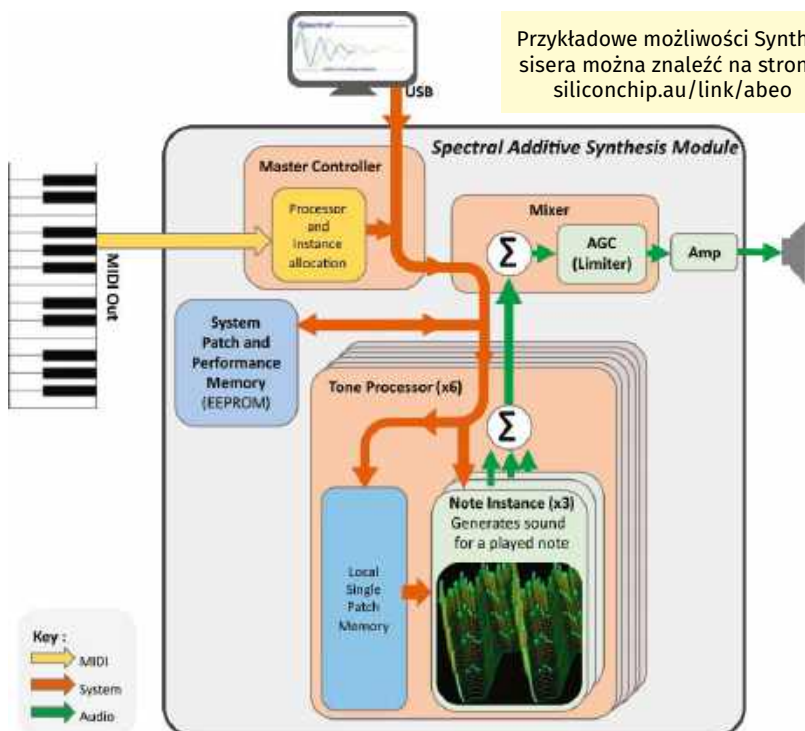
Przedstawiony w artykule zaawansowany syntezytor MIDI jest łatwy w budowie i można go podłączyć do dowolnego urządzenia kompatybilnego z MIDI. Urządzenie pozwala na eksplorację szerokiej gamy elementów akustycznych, które oddają charakterystykę prawdziwych i wymyślonych instrumentów. To więcej niż eksperyment – to w pełni rozwinięty instrument zdolny do tworzenia bogatych, szczegółowych dźwięków przy użyciu mnóstwa ustawień, obwiedni i przebiegów.

W syntezytorze Spectral Sound Synthesiser zostało zastosowanych siedem układów dsPIC działających z szybkością 70 lub 40 MIPS każdy. Ich zadaniem jest cyfrowe generowanie dźwięków. Daje to 18-nutową polifonię (tj. możliwość grania 18 różnych nut jednocześnie) i możliwość złożonego tworzenia dźwięku, z „morfingiem barwy” będącym kluczową cechą modułu.

Urządzenie charakteryzuje się również małymi opóźnieniami, co jest ważne, ponieważ nie chcesz, aby słyszalne było opóźnienie między naciśnięciem klawisza na klawiaturze a generowanym dźwiękiem.

Syntezytor jest samodzielnym modulem dźwiękowym, ma kuszącą możliwość wbudowania w niestandardowe instrumenty muzyczne DIY bez potrzeby korzystania z komputera.

Moduł ten to przygoda z syntezą dźwięku w czasie rzeczywistym, eksplorująca szeroką gamę elementów akustycznych oddających charakterystykę instrumentów prawdziwych i wymyślonych. Dzięki niemu docenimy różne rodzaje dźwięku, poznamy zasadę działania instrumentów muzycznych i dowiemy się dlaczego niektóre instrumenty są w wielu przypadkach trudne do naśladowania.



Przykładowe możliwości Synthesiseru można znaleźć na stronie siliconchip.au/link/abeo

Rysunek 1. Na schemacie blokowym zostało pokazane, jak działa syntezytor Spectral Sound. Główny kontroler odbiera komunikaty MIDI z portu MIDI In i dane patchy z komputera przez USB. Nakazuje on sekcji procesorom tonowym generowanie dźwięków w oparciu o zapisane patche i ewentualnie zapisane dane wydajności. Dźwięki te są przesyłane do miksera, a następnie do analogowego wyjścia audio

Jest to wprawdzie w pełni działające urządzenie, ale może również być przydatne podczas eksperymentowania z syntezą dźwięku. Jest to zabawne i stymulujące zajęcie, z ujmującą interakcją między cyfrowym generowaniem przebiegów a ludzką percepcją.

Konstrukcja modułu opiera się na prawdziwym przetwarzaniu równoległym, dzieląc obciążenie obliczeniowe na sześć procesorów tonów.

Dostępny jest cały kod źródłowy, zarówno dla oprogramowania układowego jak i towarzyszącego mu oprogramowania komputerowego dla systemu Windows. Dostępne jest również oprogramowanie w wersji wynikowej, włącznie z wszelkimi aktualizacjami wersji.

Jest to mocno zaawansowany projekt, więc dla tych, którzy chcą zgłębić jego tajniki, przygotowano stosowny podręcznik techniczny.

Cały temat syntezy dźwięku, przeplatający się z historią muzyki, jest bogatą i niezwykle interesującą ewolucyjną podróżą, pobudzającą naszym instynktownym pragnieniem zrozumienia i tworzenia dźwięku. Spectral Sound Synthesiser odpowiada na to pragnienie.

Przegląd systemu

Na **rysunku 1** pokazano, jak działa cały system. Posiada on wejście MIDI do odbierania nut oraz komunikatów sterujących MIDI (np. z klawiatury), wejście USB do konfiguracji przez oprogramowanie Windows oraz stereofoniczne wyjście analogowe audio do podłączenia wzmacniacza.

Możesz tworzyć patche (presety) za pomocą oprogramowania Windows, a pewna liczba tych presetów jest ładowana do modułu i przechowywana wewnętrznie. Wszelkie poprawki ustawień presetów można natychmiast wysłać do modułu i usłyszeć rezultat.

„Kontroler główny” to 8-bitowy mikrokontroler PIC18LF25K50 z przydatnym interfejsem USB. Układ ten jest powszechnie stosowany w systemach wbudowanych wymagających USB. Działa jako koncentrator, przetwarzając przychodzące komunikaty USB i MIDI. Przydziela również procesory do zadań w pozostałej części systemu.

Sześć procesorów tonów to 16-bitowe układy dsPIC33EP512MC502-I/SP z identycznym kodem do generowania cyfrowych próbek dźwięku. Każdy z nich oblicza do trzech wystąpień nut granych jednocześnie, więc system ma maksymalną polifonię $6 \times 3 = 18$ nut. Każdy procesor tonów przechowuje pojedynczy preset, ale różne układy scalone mogą mieć różne presety, dzięki czemu instrument ten zdolny jest do jednoczesnego odtwarzania kilku różnych barw dźwięku na różnych kanałach MIDI.

Pojedynczy układ „miksera”, 16-bitowy dsPIC33FJ128GP802-I/SP, miksuje próbki ze wszystkich procesorów tonów, ogranicza generowany poziom dźwięku za pomocą automatycznej regulacji wzmacnienia (AGC – automatic gain control), a następnie przekazuje dźwięk przez wbudowany stereofoniczny przetwornik cyfrowo-analogowy do wzmacniacza operacyjnego MCP6022.

Wyjście jest „pseudo stereo”. Wykorzystywana jest tu dobrze znana sztuczka audio zwana efektem Haasa, w której opóźnienie kopii sygnału z jednego ucha do drugiego daje bardzo przekonujące wrażenie pola stereo!

Moduł może pomieścić kilka presetów i „występów” w układzie scalonym EEPROM 24LC512.

Czym jest synteza addytywna?

Trwające badania nad słuchem i ludzką percepcją ujawniają, że wciąż jesteśmy daleko od całkowitego zrozumienia, w jaki sposób nasze mózgi przetwarzają i identyfikują dźwięk. Kluczowym elementem jest barwa, która jest związana z widmem częstotliwości dźwięku i tym, jak zmienia się ono w czasie.

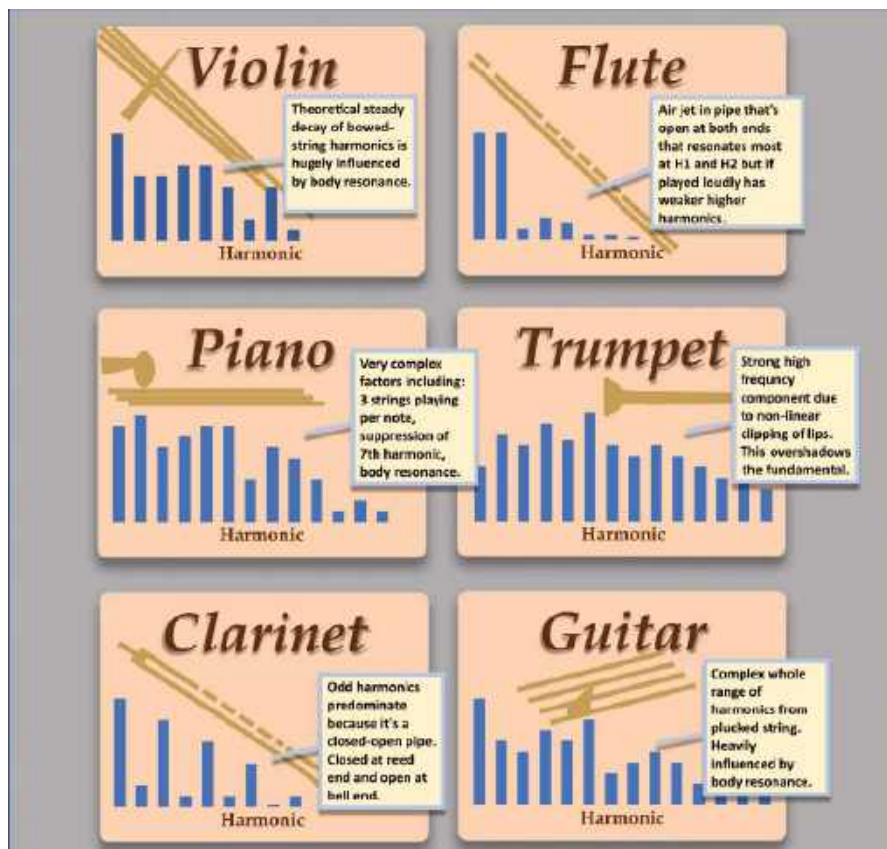
Synteza addytywna to metoda tworzenia i modulowania barw w oparciu o fakt,

że dowolną funkcję okresową można wyrazić jako sumę wielu fal sinusoidalnych, co nazywamy szeregiem Fouriera, opisanym przez Josepha Fouriera w 1822 roku. Używał jej do rozwiązywania funkcji przenoszenia ciepła, ale pomysł ten szybko stał się powszechny, od przewidywania pływów do ruchu planet, a znacznie później do syntezy dźwięku.

Prostota tego pomysłu jest atrakcyjna, ponieważ oznacza, że złożone barwy można konstruować po prostu dodając sinusoidy o odpowiednich wagach i fazach.

Okazuje się, że faza nie jest generalnie ważna, ponieważ nasz słuch ją lekceważy. Ma to sens, gdy zastanawiamy się nad falami dźwiękowymi odbijającymi się w pomieszczeniu. Pomimo tego, że fazy różnych częstotliwości są pomieszane, generalnie nie dostrzegamy żadnej różnicy w barwie.

Kolejną zaletą jest to, że dźwięki w naturze opierają się na wibracjach, w których sinusoidy tonalne mają częstotliwości będące całkowitymi wielokrotnościami lub harmonicznymi częstotliwości podstawowej („fundamentalnej”). Ten dobrze zdefiniowany związek nadaje się do obliczeń. Instrumenty muzyczne można rozpoznać po ich charakterystycznych poziomach



Rysunek 2. Przybliżenia struktury harmonicznej różnych instrumentów. Od lewej do prawej, słupki reprezentują kolejne harmoniczne powyżej podstawowej. Struktura harmoniczna jest tym, co definiuje barwę instrumentu, podczas gdy częstotliwość podstawowa jest określana przez wysokość granej nuty

harmonicznych, a kilka przykładów pokazano na **rysunku 2**.

Duże ewolucyjne drzewo genealogiczne syntezatorów elektronicznych zawiera wybitne przykłady syntezy addytywnej. Na przykład, ukochane organy Hammonda z 1935 roku łączą w sobie dźwięki generowane przez przetworniki umieszczone blisko obracających się mechanicznych „kół tonowych”.

Również wczesny syntezator Fairlight Quasar z lat 80. był syntezatorem addytywnym, podobnie jak Synclavier i kilka klawiatur Kawai. Loom, nowoczesny instrument VST, jest również typem addytywnym.

Przy wystarczającej mocy obliczeniowej, synteza addytywna umożliwia „morfinę” barw poprzez zmianę zestawu fal sinusoidalnych sumowanych w czasie – podobnie do tego, co dzieje się wokół nas z naturalnymi dźwiękami.

Synteza addytywna ma również tę wielką zaletę, że działa w dziedzinie częstotliwości, a nie czasu. Sprawia to, że filtrowanie jest prostą koncepcją, w której kontur filtra po prostu skaluje poziomy podstawowych fal sinusoidalnych. Filtrowanie typu „cegłana ściana” (brick-wall), to nic innego jak włączanie lub wykluczanie określonych przebiegów sinusoidalnych.

Taka metoda syntezy może tworzyć bogate, stymulujące i urzekające dźwięki. Gdy emuluje prawdziwe instrumenty, ma jednak pewne ograniczenia w porównaniu do syntezy opartej na samplach.

Problem polega na tym, że naturalny dźwięk to coś więcej niż tylko harmoniczne. W widmie występują częstotliwości „nieharmoniczne”, szczególnie w przypadku dźwięków perkusyjnych. Istnieją dźwięki spowodowane dmuchaniem, skrobaniem i drapaniem. Wytwarzane przez nie harmoniczne nie zawsze są dokładnymi wielokrotnościami liczb całkowitych dźwięku podstawowego wytwarzanego przez instrument.

Jako narzędzie dźwiękowe, synteza addytywna jest więc świetna. Nie zawsze jednak może z łatwością naśladować naturalny dźwięk.

Syntezator Fairlight CMI z lat 80. (który wziął swoją nazwę od promu Sydney Harbour) był przełomem w produkcji dźwięku poprzez samplowanie. Zrewolucjonizował on muzykę pop, wprowadzając prawdziwie nowe dźwięki.

Ironia polega na tym, że wynalazcy zaczęli od syntezy addytywnej, jak twierdzi współzałożyciel Kim Ryrie (co ciekawe, również założyciel magazynu ETI): „Traktowaliśmy używanie nagranych

rzeczywistych dźwięków jako kompromis – jako oszukiwanie – i nie czuliśmy się z tego szczególnie dumni”.

Ten model Fairlight był „samplerem” z możliwością nagrywania dźwięku, a wkrótce po nim pojawiły się tańsze „romplery” z nagraniem dźwiękiem zapisanym w pamięci ROM. W dzisiejszych czasach możemy mieć gigabajty próbek na półprzewodnikowych dyskach twardych.

Nie można zaprzeczyć, że syntezatory oparte na samplach mogą dawać niesamowite rezultaty, zwłaszcza z wieloma specyficznymi „warstwami” parametrów, takich jak *note velocity*, czyli prędkość uderzenia nuty – parametr MIDI, który określa siłę, z jaką klawisz lub inny kontroler jest naciśnięty. Wykorzystują one jednak masę pamięci zamiast modelować cokolwiek w sposób fizyczny. Mimo to, od wczesnych samplowanych pętli taśmowych Mellotronu, używanych w klasykach, takich jak organy piszczalkowe w „Strawberry fields” Beatlesów, jasne jest, że sample tu pozostaną.

Wraz ze stałym wzrostem mocy obliczeniowej w ostatnich latach, zaobserwowaliśmy rozwój fizycznie modelowanego dźwięku, takiego jak popularne pianina VST „Pianotech” oparte na fizyce prawdziwych instrumentów, dawnych i współczesnych. Syntezę addytywną można również w pewnym stopniu sklasyfikować jako modelowanie fizyczne ze względu na jej podejście oparte na barwie i dynamicznej naturze.

Harmonie i skala równomiernie temperowana

Prawdziwe dźwięki instrumentów są generowane poprzez vibracje, takie jak ruch powietrza wydobywanego przez flet, vibracje struny gitary lub drgania skóry bębna. Vibracje tworzą fale stojące, mające stałe węzły i ruchome anty-węzły. Węzły dzielą długość na równe części, prowadząc do integralnych harmonicznych widocznych w spektrum częstotliwości wielu instrumentów.

Na **rysunku 2** pokazano jak charakterystyczne poziomy harmonicznych rozkładają się w różnych instrumentach, choć zagadnienie jest przedstawione bardzo ogólnie.

Instrument gra dźwięk o wysokości, którą rozpoznajemy jako częstotliwość podstawową, ale ton ma „kolor” podyktowany względną zawartością harmonicznych. Częstotliwość podstawowa jest określana jako pierwsza harmoniczna. Druga harmoniczna ma podwójną częstotliwość podstawową (oktawę wyższą), trzecia harmoniczna ma trzykrotność częstotliwości podstawowej i tak dalej.

Rzadko jednak zdajemy sobie sprawę, że niektóre częstotliwości harmoniczne się tylko w przybliżeniu odpowiadają wysokościami, które rozpoznajemy w chromatycznej (12-dźwiękowej) skali muzycznej! Nasze mózgi słyszały wysokości dźwięków od naszych najwcześniejszych wspomnień. Jednak skala muzyczna, której używamy dzisiaj, jest stosunkowo nowa, a ludzie próbowali kilkunastu alternatyw, sięgając aż do czasów Pitagorasa.

Skala, której używamy dzisiaj, nazywana jest skalą „równomiernie temperowaną”, gdzie „równomiernie” odnosi się do stałego stosunku częstotliwości między dowolną nutą a jej sąsiadem. Aby obliczyć częstotliwość nuty o pół tonu wyższej, możemy pomnożyć ją przez ten stały współczynnik. Ponieważ nuty oddalone od siebie o jedną oktawę mają stosunek 2:1, więc jeśli każdy półton ma stały stosunek geometryczny, to musi on wynosić pierwiastek dwunastego stopnia z dwóch (około 1,0595:1).

Należy pamiętać, że jest to ludzki wynalazek mający na celu uzyskanie systemu równych proporcji, aby muzyka mogła być transponowana bez zmiany jej brzmienia. Pomimo tego, że poprzednie dziwaczne skale przyczyniły się do ubogacenia różnorodności muzyki, a niektórzy narzekają na ich upadek, skala równomiernie temperowana ma pewien sens.

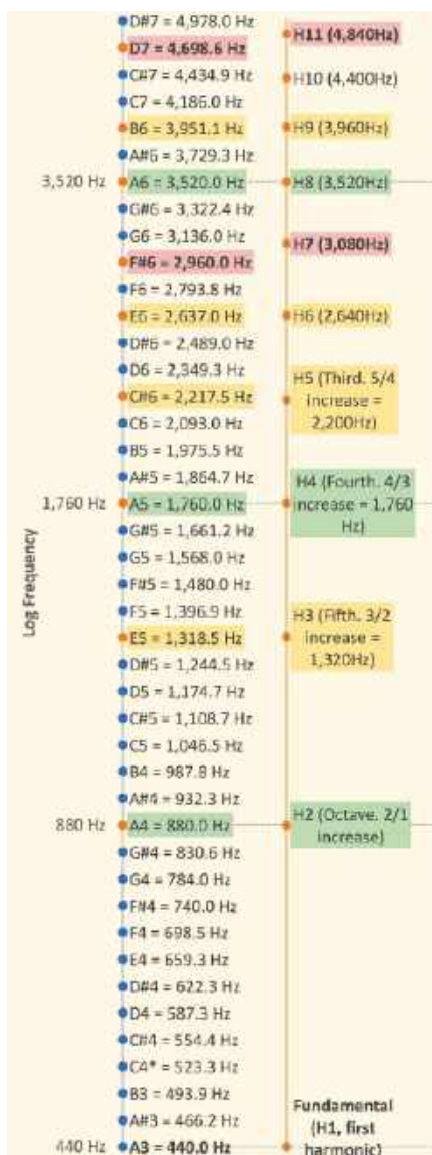
Na **rysunku 3** została zamieszczona szczegółowa analiza nuty A3 (440 Hz), pokazująca, że jej harmoniczne nie zawsze dokładnie pokrywają się z wysokościami skali równomiernie temperowanej. Potęgi dwóch harmonicznych mają dokładne dopasowanie w skali, ale inne nie (choć często są dość blisko).

Ujawnia to pewien stopień „nieharmoniczności” w harmonicznych (jak na ironię, biorąc pod uwagę nazwę). Analiza dotyczy wszystkich nut. 7. i 11. harmoniczna najbardziej odbiegają od rozpoznawalnych tonów, a jeśli są słyszalne, brzmią „płasko” i dysonansowo.

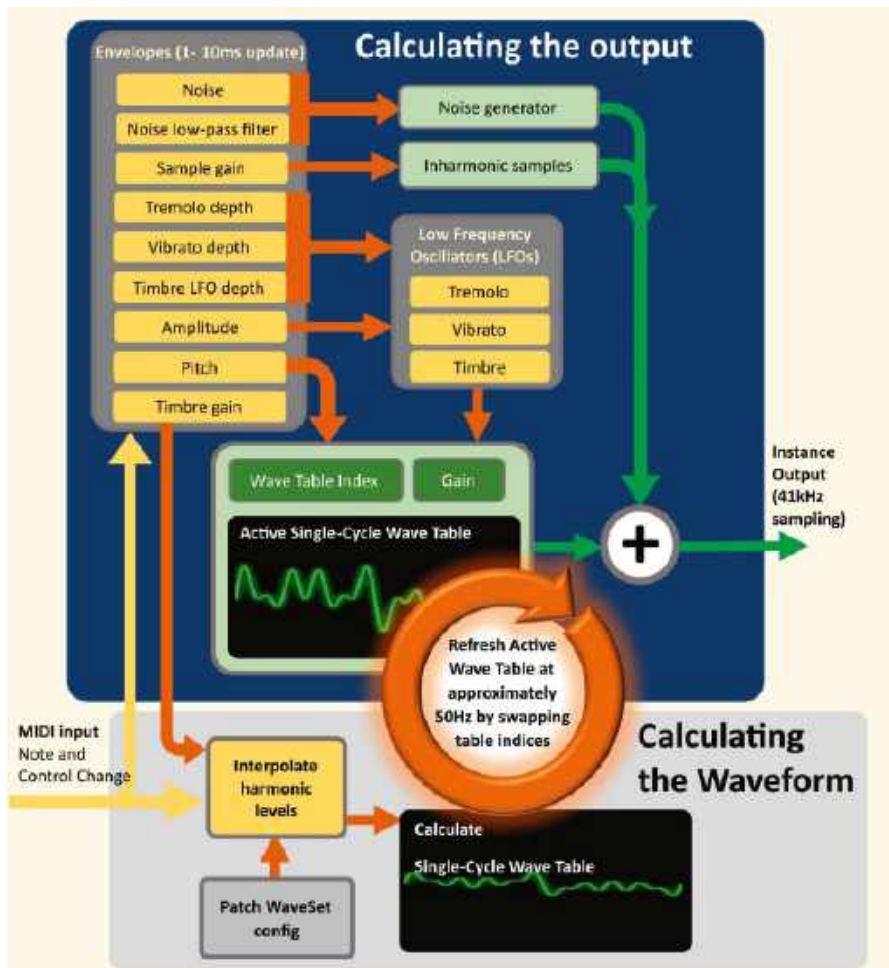
Te niedoskonałości są dobrze znane producentom instrumentów. Na przykład fortepiany są projektowane tak, aby młotek uderzał w struny w siódmym węźle vibracyjnym, aby stłumić tę „brzydką” siódmą harmoniczną! Jedenasta harmoniczna jest mniej zauważalna i często jest naturalnie cichsza.

Procesory dźwięku

Na **rysunku 4** przedstawiono serce jednego z naszych „instancji nut” – procesora tonów. Widać na nim, jak jest generowany dźwięk. Wysokość dźwięku reprezentuje nutę, którą gramy na klawiaturze podłączonej



◀ Rysunek 3. Skala równomiernie temperowana ma tę zaletę, że muzykę można grać w dowolnej tonacji bez konieczności przestrojenia instrumentu. Jednak niektóre harmoniczne nie pasują dokładnie do żadnej nuty w skali, a najgorsze z nich to 7. i 11. harmoniczna. Zazwyczaj jednak takie wysokie harmoniczne nie są szczególnie głośne, więc nie ma to znaczenia



Rysunek 4. Główne zadania i obliczenia, które są stale przetwarzane przez sześć układów procesora tonów wykonujących większość pracy związanej z syntezą

do MIDI. Sterownik główny zapewnia, że grane nuty są równomiernie rozłożone na dostępne procesory tonów.

Gdy nuta o danej wysokości jest uruchamiana w układzie procesora tonów, pierwszą rzeczą, która się dzieje, jest obliczenie kształtu przebiegu, który ma być użyty w oparciu o „statyczne” ustawienia patcha oraz inne „dynamiczne” czynniki, takie jak prędkość nuty (note velocity). Obejmuje to zaopatrzenie wszystkich aktywnych harmonicznych i dodanie odpowiednich fal sinusoidalnych.

Ponieważ harmoniczne są dokładnie całkowitymi wielokrotnościami częstotliwości podstawowej, wavetable przechowuje dokładnie jeden cykl zsumowanego przebiegu okresowego. Po jego obliczeniu, ta fizyczna tabela w pamięci staje się „aktywna” dla danej wysokości nuty. Aby to zrobić, wskaźniki tabeli są zamieniane, dzięki czemu nigdy nie

musimy powoli i nieefektywnie kopiować danych z jednej tabeli do drugiej.

Wavetable dla danej wysokości nuty jest regularnie odświeżany podczas jej aktywnego cyklu życia. Częstotliwość odświeżania nie jest stała, ale wynosi około 50 Hz. Na marginesie, fascynujące jest czytanie badań, w których odkryto, że próg, w którym ludzie mogą wykryć zmianę barwy dźwięku, jest często niższy. Wykrywanie barwy dźwięku jest wyraźnie wymagającym, abstrakcyjnym procesem rozpoznawania zachodzącym w mózgu.

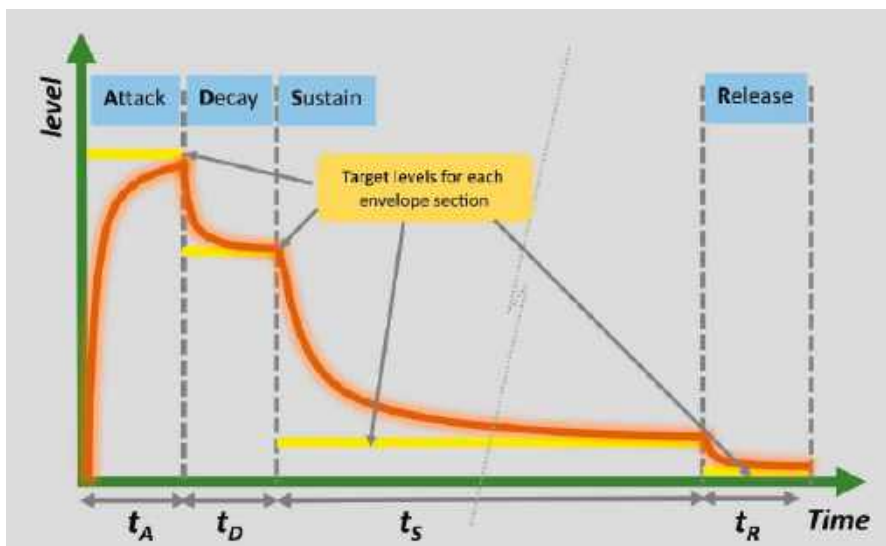
Podczas generowania nut można zastosować kilka „obwiedni wzmocnienia” do aspektów dźwięku. Jest to obwiednia amplitudy dźwięku. Na **rysunku 5** pokazano, w jaki sposób system oblicza obwiednie. Można tworzyć zarówno obwiednie liniowe, jak i wykładnicze, a każda sekcja obwiedni ADSR (attack, decay, sustain, release) ma wartość

„docelową”. Podczas każdej sekcji bieżąca wartość obwiedni przesuwa się w kierunku wartości docelowej.

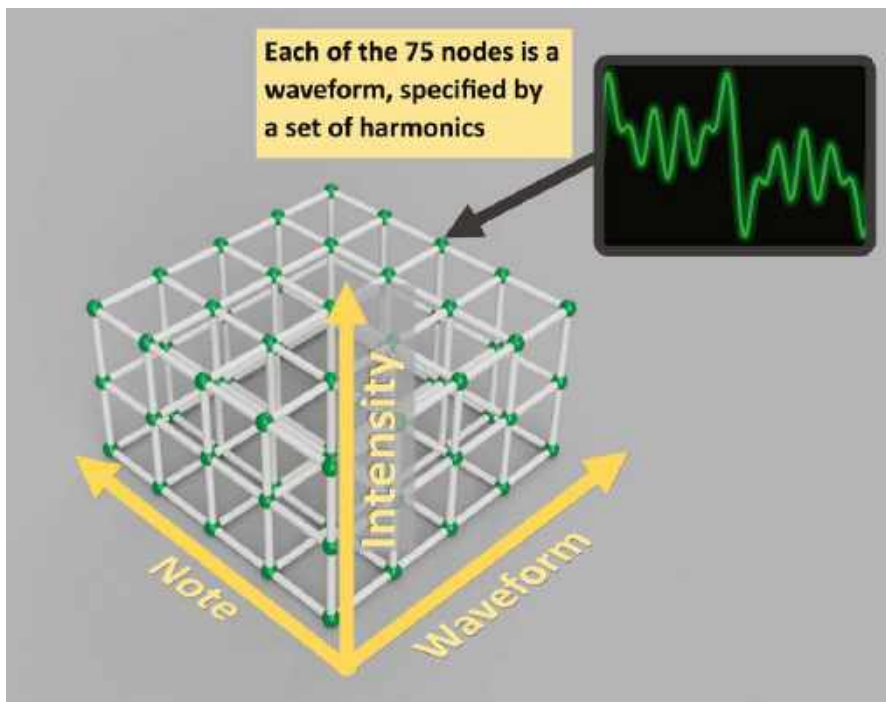
Krzywe wykładnicze występują wszędzie w przyrodzie, gdzie mamy do czynienia z rozpadem energii. Jest to więc naturalny wybór, nie tylko dla amplitudy. Na przykład, podczas szarpania struny, zawartość wysokiej częstotliwości zanika jako pierwsza. W wyniku drgań o wysokiej częstotliwości atomy zużywają więcej energii.

Dźwięk nuty zawiera również trzy „oscylatory niskiej częstotliwości” (LFO): Vibrato zmienia wysokość dźwięku, Tremolo amplitudę, a Timbre poziomy harmonicznych. Modulacja LFO może dodać wiele charakteru do dźwięku. Istnieją również obwiednie dla głębokości tej modulacji, pozwalające na przykład na stopniowe rozpoczęcie vibrato.

Inną interesującą kwestią jest to, że podczas gdy wiele syntezatorów używa Tremolo



Rysunek 5. Obwiednie to oparte na czasie profile, które można zastosować do różnych parametrów syntezy. Najłatwiejsza do zrozumienia jest obwiednia głośności, która zmienia głośność nuty od momentu jej wyzwolenia do momentu, gdy przestaje być słyszalna



Rysunek 6. „3D timbre morphing” jest rozwiązaniem problemu, polegającego na tym, że harmoniczne różnych instrumentów mogą się różnić w zależności od tego, która nuta jest grana, jak mocno jest grana itp. Jest to szczególnie oczywiste w przypadku instrumentów takich jak pianina, gdzie każdy klawisz może mieć unikatowy dźwięk, a głośniejsze dźwięki mogą wywoływać różne rezonanse

w całym dźwięku wyjściowym syntezatora, ten moduł moduluje każdą nutę, dzięki czemu ogólny dźwięk jest bardziej złożony i interesujący.

Zmiana barwy dźwięku (timbre morphing)

Określenie poziomów harmonicznych do wykorzystania podczas konstruowania kształtu przebiegu jest skomplikowanym procesem. Na **rysunku 6** pokazano, że konfiguracja przechowuje dane harmoniczne dla 75 przebiegów, a przebieg do zsyntezowania

zależy od trzech parametrów. Parametr „Note” to pozycja nuty na klawiaturze. Parametr „Intensity” najczęściej oznacza „velocity” (prędkość) nuty. Parametr „Waveform” wybiera jeden z pięciu przebiegów.

Każdy punkt w tej koncepcyjnej siatce przestrzeni 3D ma zestaw poziomów harmonicznych definiujących falę dźwiękową. Bieżące wartości parametrów definiują wymagany punkt wewnątrz tej przestrzeni, a poziomy harmoniczne, które należy zastosować

są interpolowane z najbliższych zdefiniowanych punktów siatki w tej przestrzeni.

Idąc dalej, modulowanie tej przestrzeni za pomocą Timbre LFO (modulacja LFO barwy) może dać imponujący realizm w porównaniu do zwykłego vibrato. Typowe vibrato czysto moduluje wysokość całego kształtu przebiegu, podczas gdy modulacja barwy jest oparta na harmonicznym, a zatem jest modulacją znacznie bardziej złożoną, którą nasz mózg może zauważyć.

To ekscytujące usłyszeć tę różnicę i zdać sobie sprawę, że nasze mózgi czerpią przyjemność i delectują się dźwiękiem. Być może, biorąc pod uwagę niewiarygodnie sprytnie przetwarzanie, które wykonują nasze mózgi, ta świadomość nie powinna być zbyt zaskakująca.

W kierunku większego realizmu

Wspomnieliśmy już, że naturalny dźwięk zawiera elementy nieharmoniczne, które nie pasują do czystych całkowitych wielokrotności częstotliwości harmonicznych. Aby temu zaradzić, nasz syntezator zawiera kilka dodatkowych funkcji wykraczających poza czysto addytywne podejście do syntezy.

Po pierwsze, dostępna jest obwiednia szumu, która pomaga symulować „dmuchane” instrumenty. Jest to generator białego szumu z regulowanym filtrem dolnoprzepustowym.

Po drugie, można dodawać krótkie próbki harmoniczne. Są to zakodowane na stałe klipy dźwięków stuknięć, zadrapań, kliknięć i uderzeń. Jednak funkcja próbek obejmuje również implementację dobrze znanej techniki linii opóźniającej „Karplus Strong” syntezy szarpanych strun. Technika ta została po raz pierwszy opisana w 1983 roku w artykule Kevina Karplusa i Alexa Stronga zatytułowanym „Digital Synthesis of Plucked-String and Drum Timbres”. Jest to prosta obliczeniowo, ale skuteczna metoda generowania realistycznych, zanikających dźwięków smyczków, które zaczynają życie w buforze opóźniającym jako szum. Dodaje to potężne narzędzie do tego modułu, mimo że nie ma nic wspólnego z syntezą addytywną!

Istnieją również ustawienia losowego lub systematycznego dostrajania częstotliwości granych nut, co pozwala wprowadzić swoiste odchyłki i nieczystości dźwięku, charakterystyczne dla prawdziwych instrumentów. Dodatkowo, istnieje opcja użycia dwóch oscylatorów wavetable na dźwięk nuty, odstrojonych o określoną wartość, dając efekt podobny do chóru, który próbuje uwzględnić fakt, że instrumenty takie jak fortepian używają wielu odstrojonych strun na nutę.

Ostatnią funkcją jest „Body Resonance Filter”. Próbuje on naśladować rezonans

korpusu prawdziwego instrumentu poprzez filtrowanie ogólnego dźwięku systemu. Po określeniu konturu filtra w aplikacji, skaluje on poziomy harmonicznych. Chociaż metoda ta jest bardzo atrakcyjna teoretycznie, ma ona ograniczenia matematyczne.

Pomimo tych wszystkich dodatkowych funkcji, nadal istnieją rzeczywiste złożoności, z którymi moduł po prostu nie może sobie poradzić. Na przykład, fortepian ma swoje osobliwości wynikające ze sztywności strun, gdzie wyższe harmoniczne stają się ostrzejsze w porównaniu do oczekiwanych harmonicznych struny. Dzieje się tak, ponieważ sztywność skutecznie skraca strunę dla wyższych częstotliwości, podnosząc częstotliwość oscylacji harmonicznych. Efekt ten ma zastosowanie do wszystkich instrumentów strunowych i jest czymś, czego nasz system nie może rozwiązać, ponieważ model opiera się całkowicie na integralnych harmonicznych.

System czasu rzeczywistego

Cały moduł jest przykładem „twardego” systemu czasu rzeczywistego, w którym momenty wytwarzania i przetwarzania próbek są niezmiennie. Stanowi to poważne wyzwanie, a rozwój modułu był powolną ewolucją kodowania, pomiarów, udoskonalania, a czasem przeprojektowywania.

Wszystkie procesory tonowe działają całkowicie równolegle i są odpytywane przez układ miksera w celu dostarczenia próbek z częstotliwością próbkowania dźwięku 41,7 kHz. Wewnątrz każdego procesora tonowego obowiązuje hierarchia przerwań, jak pokazano na **rysunku 7**, możliwa dzięki zdolności dsPIC do przypisywania priorytetów.

Główna procedura procesora tonowego, centralny element całego systemu, to tylko prosta pętla, która ponownie oblicza tabele przebiegów. To zadanie „w tle” ma nieprzewidywalny czas trwania, nie jest terminowe i może się różnić w zależności od aktywności przerwania i złożoności budowania kształtu przebiegu.

Oznacza to, że w pewnych okolicznościach częstotliwość odświeżania barwy dźwięku może ulec spowolnieniu – choć w praktyce wydajność jest bardzo akceptowalna, a zmiany barwy dźwięku są postrzegane jako płynne i gładkie.

„Warstwy” przetwarzania powyżej tej podstawowej pętli głównej zajmują się obliczaniem kroków obwiedni, obliczaniem wyjścia próbki i przetwarzaniem odebranych danych.

Sztuczka stosowana w celu poprawy przepustowości na magistrali SPI między procesorami tonów a mikserem polega na wysyłaniu tylko zmian wartości próbek. Sumowanie

sinusów w procesorze tonów może skutkować całkowitą wartością przekraczającą 16 bitów. Suma w procesorze tonów jest 32-bitową liczbą całkowitą ze znakiem, ale procesor tonów wysyła do miksera tylko zmianę tej sumy, ograniczoną do 16 bitów.

Może to spowodować zniekształcenie sygnału, ale statystycznie zdarza się to rzadko. Przewaga wydajności tej metody znacznie przewyższa rzadkie anomalie, które prawdopodobnie nie są nawet zauważalne.

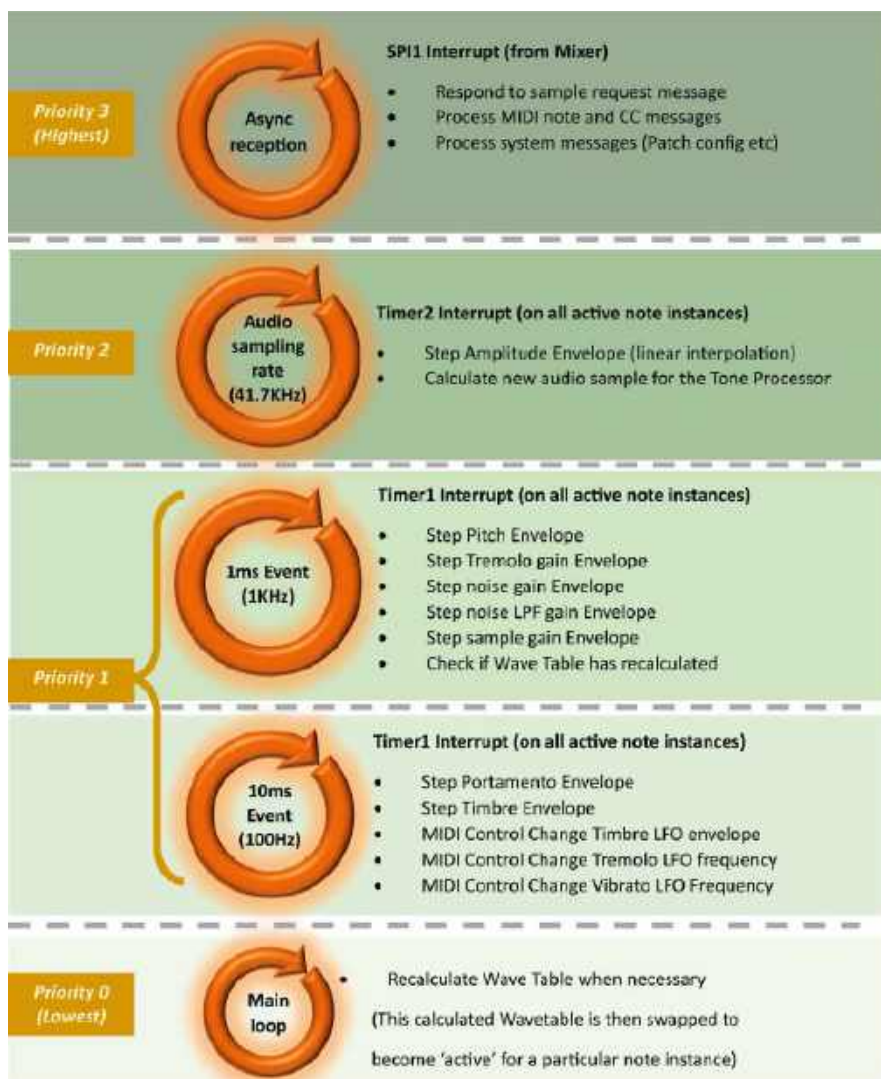
Oprogramowanie i sprzęt zostały zaprojektowane z myślą o szybkości. Wszystkie układy dsPIC pracują na najwyższych obrotach. Wszystkie obliczenia są oparte na liczbach całkowitych, w połączeniu z szerokim wykorzystaniem sprzętowego mnożnika w układzie poprzez polecenia „__builtin” kompilatora. Kod w znacznym stopniu wykorzystuje operacje przesunięcia do szybkiego mnożenia i dzielenia, a także liczne

tabele wyszukiwania (lookup tables), w tym szczegółową tablicę wartości sinus.

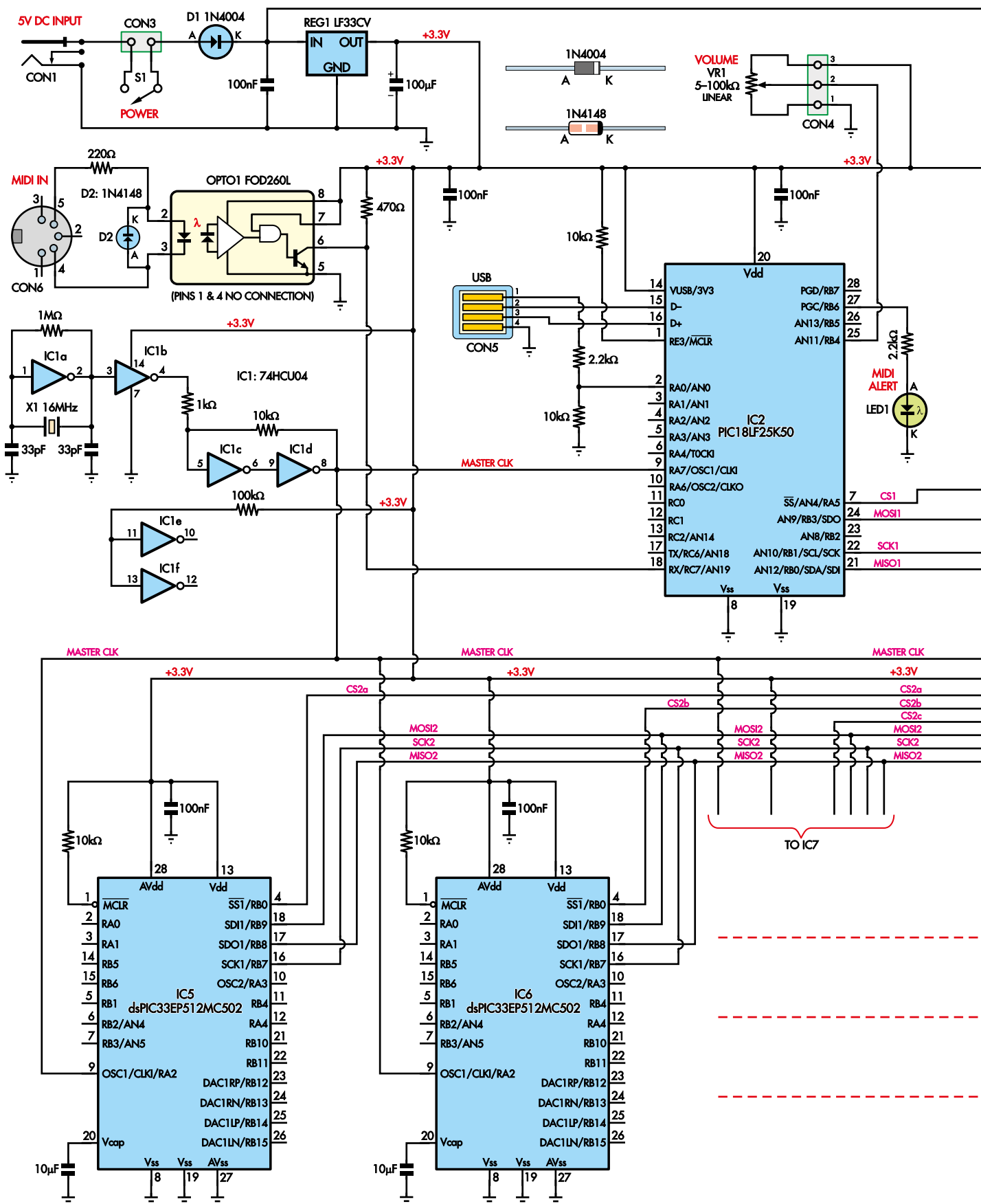
Szczegóły schematu

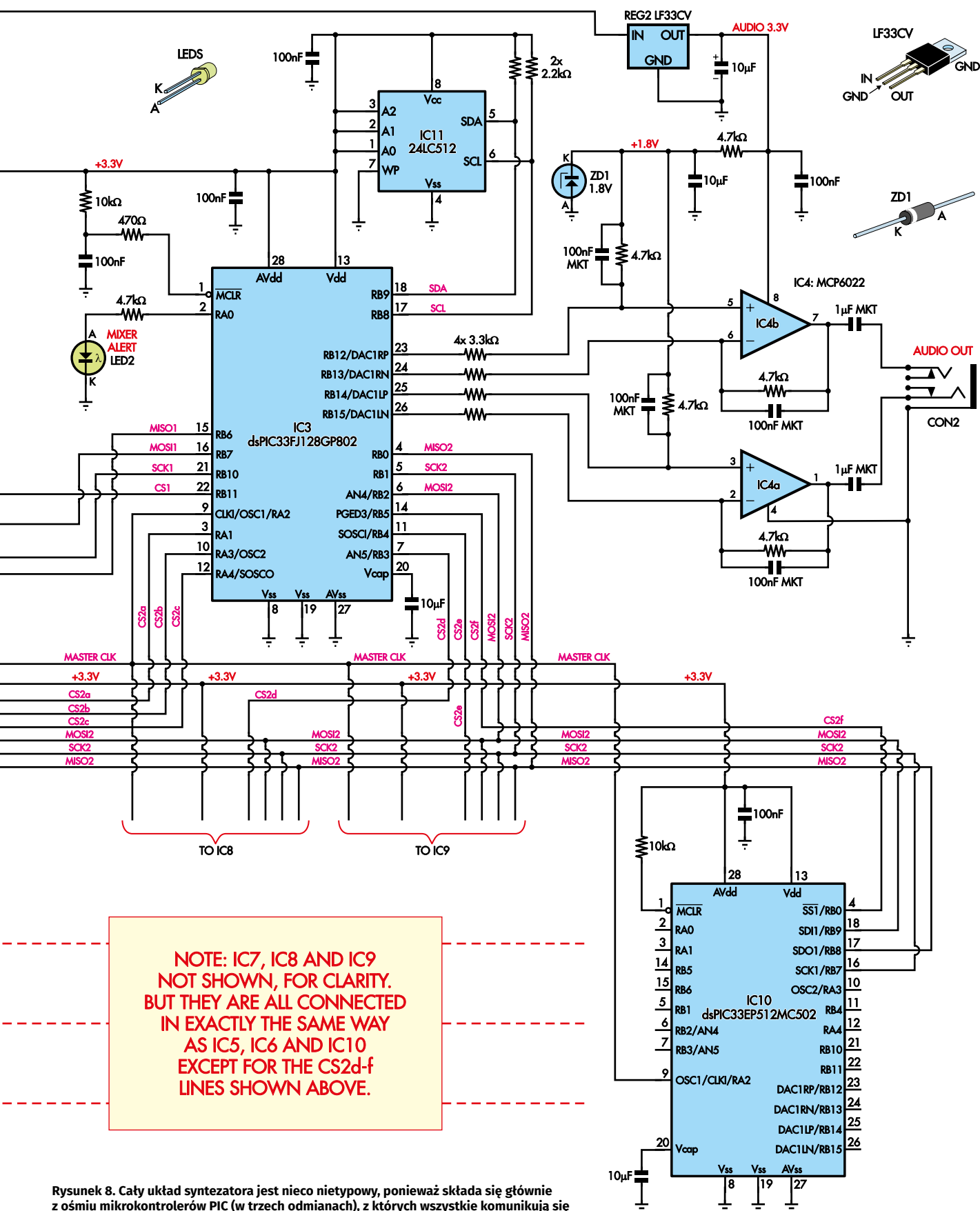
Pełny schemat pokazano na **rysunku 8**. Pierwszą rzeczą, na którą należy zwrócić uwagę, jest to, że sześć procesorów tonów (IC5...IC10) to identyczne układy dsPIC skonfigurowane w ten sam sposób, każdy z zaledwie kilkoma elementami pomocniczymi: jednym kondensatorem filtrującym VDD, jednym kondensatorem Vcap (wymagającym dla wewnętrznego stabilizatora układu) i jednym rezystorem 10 kΩ podciągającym do zasilania linię /MCLR a przez to zapobiegającym fałszywym sygnałom zerowania.

Poza zasilaniem, jedynymi połączeniami z tymi układami jest wspólna magistrala SPI, ponieważ są to typowe jednostki obliczeniowe, a wszystkie polecenia i dane są wysyłane przez tę magistralę. Jedyne

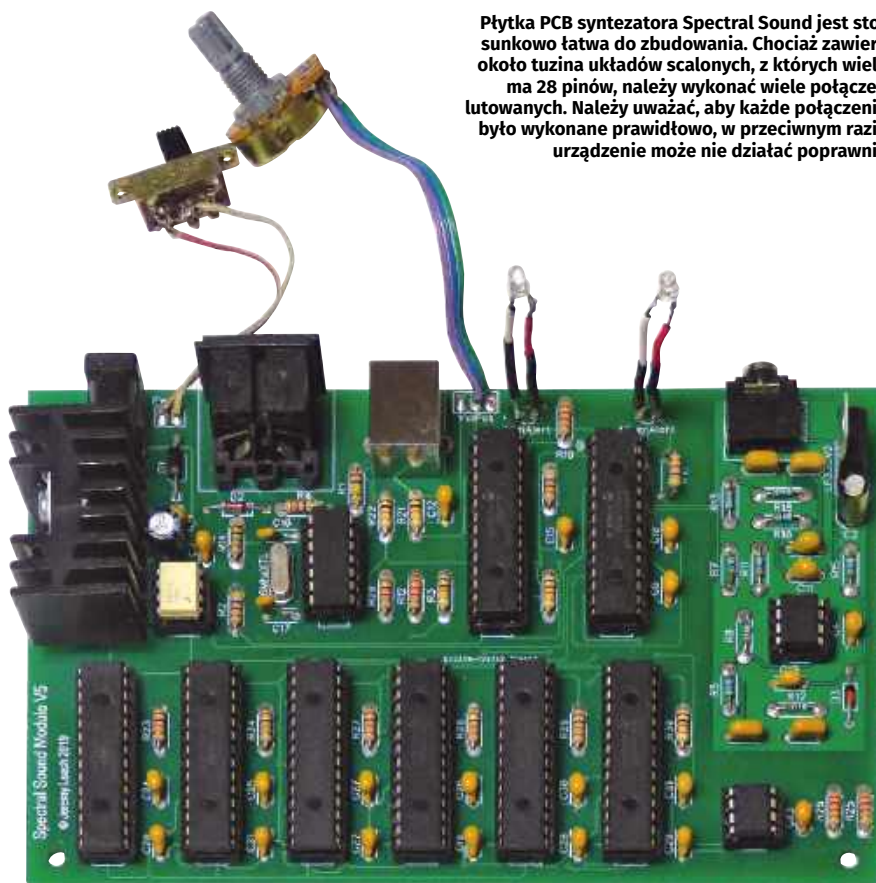


Rysunek 7. Przegląd wszystkich zadań, które musi obsłużyć procesor tonów, w kolejności według priorytetu. Zdarzenia o najwyższym priorytecie to te, które mogą spowodować przerwanie dźwięku lub inne nieoczekiwane rezultaty, jeśli zostaną opóźnione





Rysunek 8. Cały układ syntezy jest nieco nietypowy, ponieważ składa się głównie z ośmiu mikrokontrolerów PIC (w trzech odmianach), z których wszystkie komunikują się za pośrednictwem dwóch oddzielnych magistral szeregowych SPI. Pozostała część układu obejmuje pamięć EEPROM (IC11) używaną do przechowywania patchy i wydajności, zasilania, wejścia MIDI i wyjścia audio



Płytkę PCB syntezatora Spectral Sound jest stosunkowo łatwa do zbudowania. Choć zawiera około tuzina układów scalonych, z których wiele ma 28 pinów, należy wykonać wiele połączeń lutowanych. Należy uważać, aby każde połączenie było wykonane prawidłowo, w przeciwnym razie urządzenie może nie działać poprawnie

różnice między połączeniami z tymi układami polegają na tym, że wejście CS2a... CS2f każdego procesora tonowego (pin 4) łączy się z inną linią wyboru w mikserze IC3.

Mikser jest innym (ale pokrewnym) typem procesora dsPIC. Oprócz tego, że jest podłączony do magistrali SPI i sześciu linii wyboru układu procesora tonów, ma również dwa różnicowe wyjścia analogowe z wewnętrznego stereofonicznego przetwornika cyfrowo-analogowego.

Sygnały te są podawane do wzmacniacza operacyjnego IC4, który konwertuje sygnały różnicowe na sygnały audio single-ended odpowiednie do podania do gniazda wyjściowego CON2. Jednocześnie układ ten odfiltruje artefakty przetwornika cyfrowo-analogowego za pomocą filtrów dolno-przepustowych zbudowanych z dodanych kondensatorów i istniejących rezystorów ustawiających wzmacnienie.

Wirtualna szyna „połowy zasilania” jest generowana za pomocą diody Zenera ZD1 polaryzowanej z szyny +3,3 V, dzięki czemu sygnały audio z IC3 pozostają w obrębie szyn zasilania wzmacniacza operacyjnego.

Mikser IC3 łączy się również z pamięcią EEPROM 24LC512 (IC11) za pomocą dwuprzewodowego interfejsu szeregowego I²C (SDA i SCL). Na zasilaniu tego układu obecny

jest kondensator filtrujący a linie interfejsu szeregowego zostały podciągnięte za pomocą rezystorów do VCC.

Ostatnim zadaniem układu IC3 jest występowanie diody LED alarmu miksera, LED2, z wyjścia RA0 (pin 2).

Wejście MIDI, USB i inne zadania kontrolne przypadają układowi PIC18LF25K50, IC2. Monitoruje on obecność napięcia USB 5 V na wejściu cyfrowym RA0 (pin 2) za pomocą „dzielnika” 2,2 k Ω /10 k Ω , który służy głównie do ograniczania prądu do tego pinu i zapewnienia, że jest on ściągany do 0 V, gdy do wejścia USB nie zostało podłączone żadne urządzenie.

IC2 i IC3 komunikują się za pośrednictwem drugiej oddzielnej magistrali SPI, z przydzieloną linią wyboru układu, od pinu 7 IC2 do pinu 22 IC3. IC2 steruje również diodą LED1, diodą LED alarmu MIDI, z wyjścia cyfrowego RB6 (pin 27).

Zewnętrzny potencjometr VR1 (regulacja głośności) łączy się z CON4, równolegle do zasilania 3,3 V. Jego suwak trafia do wejścia analogowego AN11 układu scalonego IC2 (pin 25). Układ IC2 odczytuje napięcie na suwaku za pomocą przetwornika analogowo-cyfrowego (ADC) i przekazuje wartość cyfrową do układu IC3, który następnie skaluje swoje wyjście zapewniając żądany poziom głośności.

Rezystancja i typ potencjometru nie są krytyczne, ale 100 k Ω jest rozsądne. Skalowanie wartości próbek audio wchodzących do przetwornika cyfrowo-analogowego, zamiast bezpośredniej regulacji wzmacnienia wzmacniacza operacyjnego, upraszcza płytkę drukowaną kosztem zmniejszenia rozdzielczości bitowej dźwięku przy zmniejszonej głośności. W praktyce trudno jest usłyszeć tę degradację.

Pozostaje tylko opisać wejście MIDI, dystrybucję sygnału zegara i zasilanie.

Sygnał MIDI jest podawany do gniazda CON6 i zasila diodę LED IR w transopto-rze OPTO1 FOD260L. Rezystor 220 Ω zapewnia ograniczenie prądu, natomiast dioda D2 zapobiega odwrotnej polaryzacji diody LED. Konieczne jest użycie transoptora FOD260L, ponieważ jest on przystosowany do pracy z napięciem 3,3 V – inne typy mogą nie działać.

Tranzystor wyjściowy w OPTO1 pracuje w układzie ze wspólnym emiterem z rezystorem podciągającym 470 Ω . Wynikowy sygnał trafia do wejścia RX (pin 18) układu IC2.

Ze względu na to, że mamy osiem mikrokontrolerów, z których każdy potrzebuje źródła taktowania, zastosowany jest pojedynczy zewnętrzny oscylator. Jest on zbudowany przy użyciu rezonatora kwarcowego X1, dwóch kondensatorów 33 pF i niebuforowanego inwertera IC1a. Wynikowy sygnał 16 MHz jest odwracany przez IC1b i buforowany przez IC1c i IC1d, a następnie podawany na wejściowe piny zegarowe wszystkich mikrokontrolerów.

Nie zalecamy używania buforowanego inwertera zamiast IC1, takiego jak bardziej popularny 74HC04, ponieważ może on nie generować poprawnie przebiegów.

Zasilanie jest proste. Urządzenie jest zasilane napięciem 5 V DC z gniazda CON1. Napięcie to jest dalej kierowane przez diodę zabezpieczającą przed odwrotną polaryzacją do wejść stabilizatorów liniowych REG1 i REG2. REG1 zasila wszystkie układy cyfrowe, podczas gdy REG2 zasila obwody analogowe. Są to w zasadzie tylko wzmacniacze operacyjne IC4 i polaryzacja diody Zenera ZD1.

Zwiększenie stosunku sygnału do szumu (SNR)

Wyzwaniem dla każdego systemu zawierającego mieszane obwody cyfrowe i analogowe jest powstrzymanie przenikania szumów cyfrowych do wyjścia audio. Elementy audio i cyfrowe zostały na płycie modułu oddzielone w jak największym stopniu, z rozdzielczymi stabilizatorami, z odpowiednim wykorzystaniem płaszczyzny masy.

Wykaz elementów:

1 dwustronne PCB kod 01106221, 145 mm × 94 mm
1 futerał na instrument [Takachi YM-150; RS Cat 373-2255]
4 przyklejane gumowe nóżki
1 etykieta na panel przedni, 145 mm × 37 mm (rysunek 10)
1 etykieta na pokrywę, 141 mm × 85 mm (rysunek 11)
1 5...6 V DC 1 A regulowany zasilacz wtyczkowy *
1 rezonator kwarcowy 16 MHz, HC-49 (zwykły lub niskoprofilowy) (X1)
1 gniazdo barytkowe DC montowane na PCB, o średnicy wewnętrznej 2,1 mm lub 2,5 mm, pasujące do zestawu wtyczek (CON1)
1 gniazdo jack 3,5 mm stereo DPST (CON2) [Altronics P0094, RS Cat 913-1021 lub CUI SJ1-3555NG].
12-stykowe spolaryzowany gniazdo z pasującą wtyczką (CON3, do przełącznika zasilania)
13-stykowe spolaryzowane złącze z pasującą wtyczką (CON4, do regulacji głośności)
1 przełotowe pełnowymiarowe gniazdo USB typu B (CON5) [Jaycar PS0920, Altronics P1304A/P1304B].
15-stykowe gniazdo DIN 180°, montaż na płytce drukowanej pod kątem prostym (CON6) [Jaycar PS0350, Altronics P1188B lub RS Cat 491-087].
1 przełącznik suwakowy SPST lub SPDT do montażu panelowego (S1, zasilanie)
1 montowany na panelu potencjometr liniowy 100 kΩ i pokrętko (VR1, regulacja głośności)
8 28-pinowych wąskich podstawek DIL IC (opcjonalnie; dla IC2, IC3 i IC5...IC10)
3 8-pinowa podstawa DIL IC (opcjonalnie; dla IC4, IC11 i OPT01)
1 radiator TO-220 (REG1) [maks. szerokość 40 mm, głębokość 13 mm od zakładki, <18°C/W; RS Cat 263-251 używany w prototypie]
4 tuleje dystansowe 6,3 mm z gwintem M3
1 śruba z łbem walcowym M3 × 10 mm, podkładka sprężysta i nakrętka sześciokątna
8 śrub z łbem walcowym M3 × 5 mm
2 śruby z łbem stożkowym M2 × 10 mm i nakrętki (do montażu przełącznika suwakowego)
1 kabel taśmowy o długości 100 mm (do podłączenia do S1 i VR1)
1 mała tubka pasty termoprzewodzącej
* Można użyć napięcia do 9 V, ale 5...6 V zapewni bardziej rozsądne rozpraszanie energii.

Półprzewodniki:

1 74HCU04 niebuforowany inwerter hex, DIP-14 (IC1)
1 PIC18F25K50-I/SP 8-bitowy mikrokontroler zaprogramowany kodem 0110622A.HEX (IC2)
1 16-bitowy mikrokontroler dsPIC33F128GP802-I/SP zaprogramowany kodem 0110622B.HEX (IC3)
1 MCP6022-I/P wzmacniacz operacyjny typu rail-to-rail, DIP-8 (IC4)
6 16-bitowych mikrokontrolerów dsPIC33EP512MC502-I/SP zaprogramowanych kodem 0110622C.HEX (IC5-IC10)
1 24LC512-I/P 64 kbyte I²C EEPROM, DIP-8 (IC11)
1 FOD260L transceptor, DIP-8 (OPT01)
2 stabilizatory liniowe LF33CV 3,3 V o niskim spadku napięcia (REG1, REG2)
2 zielone diody LED 3 mm o wysokiej jasności (LED1, LED2)
1 1,8 V 250 mW dioda Zenera (ZD1) [np. 1N4614]
1 dioda 1N4004 400 V 1 A (D1)
1 1N4148 75 V 150 mA dioda sygnalizacyjna (D2)

Kondensatory:

1 elektrolityczny 100 µF 6,3 V
1 elektrolityczny 10 µF 16 V
8 10 µF 16 V X7R ceramiczny
2 1 µF 63 V MKT
4 100 nF 63 V MKT
13 100 nF 50 V X7R ceramiczny
2 33 pF 50 V ceramiczne

Rezystory: (wszystkie o mocy 1/4 W, 1%, metalizowane, osiowe)

11 MΩ 6,47 kΩ 11 kΩ 1100 kΩ 4,33 kΩ 2 470 Ω 10 10 kΩ 4 2,2 kΩ 1 220 Ω

Podjęto jednak dodatkowe środki w celu za-
pewnienia wolnego od szumów i przyzwoi-
tego systemu audio.

Jednym z takich środków jest ogranicznik
audio w kodzie audio miksera wykorzystujący
zaawansowane sterowanie AGC. Ogranicznik
nieznacznie zmniejsza zakres dynamiki
poprzez tłumienie szczytów, tym samym
skutecznie wzmacniając cichsze dźwięki
i obniżając poziom szumów. Nie wydaje się
to naturalne, gdy próbuje się naśladować
instrumenty polifoniczne, jednak limityery
i kompresory są powszechne w reprodukcji
dźwięku i mają znaczny korzystny wpływ
na ten system.

Używamy również techniki zwanej pre-
-emfazą i de-emfazą. Generowany dźwięk
cyfrowy ma zastosowane wzmocnienie wy-
sokich częstotliwości, a analogowy obwód
przetwarzania sygnału ma dopasowane tłum-
nienie wysokich częstotliwości zastosowane
przez filtr dolnoprzepustowy na wzmacnia-
czach operacyjnych. W ten sposób tłumiony

jest wyższy, bardziej zauważalny element
szumu obwodu.

Moduł faktycznie „wzmacnia” wyższe po-
ziomy harmonicznych poprzez ostrożne tłum-
nienie niższych poziomów harmonicznych.
Na szczęście nie są do tego potrzebne skom-
plikowane filtry cyfrowe.

Wreszcie, aplikacja Patch Editor automa-
tycznie zwiększa poziomy harmonicznych
do maksimum, zapewniając, że zsumowany
kształt przebiegu wavetable obejmuje cały
16-bitowy zakres, maksymalizując SNR.

Budowa

Syntezator Spectral Sound jest stosunkowo
prosty w budowie, ponieważ zastosowanie
wielu mikrokontrolerów minimalizuje liczbę
wymaganych oddzielnych komponentów.
Większość z nich montuje się na dwu-
stronnej płytce drukowanej oznaczonej
kodem 01106221 o wymiarach 145 mm ×
94 mm. Schemat montażowy płytki PCB
pokazano na **rysunku 9**.

W konstrukcji urządzenia nie ma nic nad-
zwyczajnego, potrzebne są jedynie dobre
umiejętności lutownicze przydatne podczas
dokładnego lutowania ponad 200 pinów!
Zalecamy użycie podstawek pod układy scalone,
także dla transoptora. Choć podstawki w dłuż-
szej perspektywie mogą powodować problemy
z niezawodnością z powodu utleniania styków,
w układzie brak jest możliwości programowa-
nia w systemie. Ponieważ jednak większość
konstruktorów będzie używać wstępnie zapro-
gramowanych mikrokontrolerów (lub progra-
mować je przed montażem), można rozważyć
przylutowanie ich bezpośrednio do płytki, o ile
masz pewność, że zostały poprawnie zapro-
gramowane. Zauważ, że nie przewidzieliśmy
podstawki dla IC1, ponieważ nie ma powodu,
aby jej tam użyć.

Zacznij od rezystorów, sprawdzając omo-
mierzem wartości każdego elementu przed
wzlutowaniem na miejsce. Następnie należy
wzlutować trzy diody. Każda z nich jest in-
nego typu, więc nie pomyśl ich i upewnij się,
że są zamontowane paskami katody skiero-
wanymi tak, jak pokazano na rysunku.

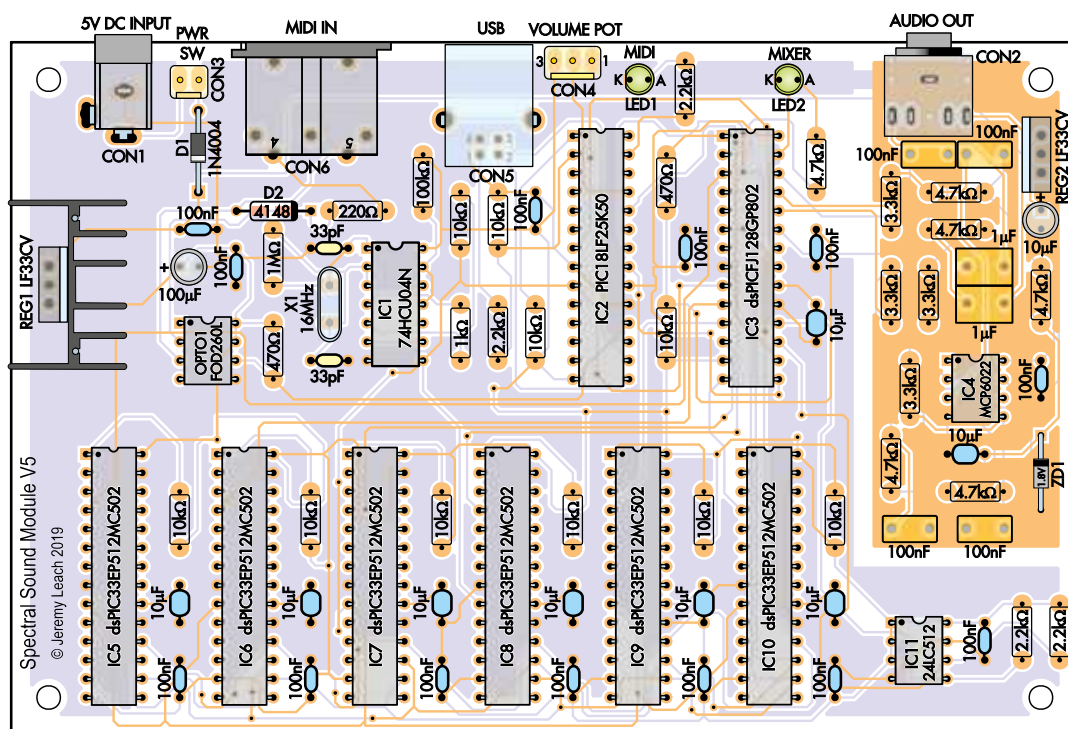
Następnie zamontuj podstawki pod układy
scalone i transoptor (lub układy scalone i tran-
soptor). Upewnij się, że wszystkie mają pin 1
skierowany w stronę górnej części płytki,
a jeśli lutujesz układy scalone do płytki, bądź
bardzo ostrożny, aby nie pomylić różnych
28-pinowych typów.

Następnie zamontuj wszystkie niespolary-
zowane kondensatory ceramiczne i MKT.
Są to kondensatory ceramiczne i MKT o po-
jemności 100 nF, więc upewnij się, że kon-
densatory MKT znajdują się w miejscach
pokazanych na rysunku 9.

Teraz zamontuj kondensatory elektroli-
tyczne z dłuższymi dodatnimi wyprowadze-
niami do pól oznaczonych + na rysunku 9,
a następnie spolaryzowane złącza pinowe
i gniazdo Jack CON2.

Następnie przylutuj diody LED dłuższymi
wyprowadzeniami do boku oznaczonego li-
terą A. Dopasuj długość ich wyprowadzeń
tak, aby sięgały do górnego panelu obudowy
po zainstalowaniu płytki drukowanej (zobacz
sekcję „Okablowanie” poniżej). Następnie
zamontuj dwa stabilizatory, najpierw mocując
radiator do REG1 za pomocą śruby, nakrętki
i podkładki.

Pozostaje tylko rezonator X1, gniazdo zasil-
ania CON1, gniazdo USB CON5 i gniazdo MIDI
CON6 na płytce drukowanej. Zamontuj je
w kolejności rosnącej wysokości. Wreszcie,
jeśli przylutowałeś do płytki wszystkie pod-
stawki pod układy scalone, umieść w nich
wszystkie układy scalone i transoptor, zwr-
cając szczególną uwagę na kierunek montażu



► Rysunek 9. Podobnie jak na schemacie, na PCB dominuje osiem układów PIC, wszystkie w 28-pinowych obudowach DIL. Upewnij się, że są one prawidłowo zorientowane i nie pomieszczone

► Rysunek 11. Grafika na panelu pokrywy (pokazana w około 85% rzeczywistego rozmiaru) stanowi eleganckie wykończenie projektu. Jest ona przeznaczona do wydrukowania na przezroczystym nośniku. Należy pamiętać, że dwie pozycje diod LED mogą się nieco różnić, zwłaszcza jeśli używasz innej obudowy; możesz po prostu odciąć tę część naklejki i umieścić ją osobno

pinów 1 i nie mieszając różnych 28-pinowych i 8-pinowych układów scalonych.

Wybór obudowy

Płytkę PCB została zaprojektowana tak, aby pasowała do obudowy określonej w wykazie elementów, a etykieta panelu przedniego (rysunek 10), grafika pokrywy (rysunek 11) i szablon wiercenia (rysunek 12) pasują do tej obudowy. Można je również pobrać w formacie PDF i PNG ze strony siliconchip.com.au/Shop/11/6416.

Możesz użyć innej obudowy, o ile jest wystarczająco duża, aby pomieścić płytkę drukowaną, ponieważ wszystkie złącza i elementy sterujące (oprócz dwóch, które należy zamontować na panelu) znajdują się wzdłuż jednej krawędzi płytki drukowanej.

Przy płytce drukowanej o wymiarach 145 mm × 94 mm, większość obudów o wymiarach co najmniej 165 mm × 100 mm powinna być odpowiednia. Wymagana wysokość zależy od radiatora używanego dla REG1. Podany radiator ma tylko 20 mm wysokości, więc obudowy o wysokości co najmniej

35 mm powinny być dobre. Jeśli używasz wyższego radiatora, dodaj 10...15 mm do jego wysokości, aby dowiedzieć się, jakie obudowy będą odpowiednie.

Możliwe alternatywne obudowy instrumentów to Altronics Cat H0374 lub Cat H0378 (z krótkim radiatorem), Jaycar Cat HB5912 lub Hammond RM2055M, który jest dostępny w Digi-Key i Mouser. Urządzenie powinno również zmieścić się w obudowie UB2 Jiffy, takiej jak Jaycar Cat HB6012, Altronics Cat H0152 lub H0202, ale nie wyglądają one tak dobrze jak skrzynki na instrumenty i nieco trudniej będzie w nich zmieścić płytkę.

Zamontuj płytkę w obudowie za pomocą wkrętów i gwintowanych tulei dystansowych, tak aby jej górna krawędź przylegała do jednej strony obudowy (w przypadku obudowy instrumentu powinien to być panel przedni). Następnie należy zaznaczyć i wywiercić/wyciąć otwory w sąsiednim panelu dla wtyczki zasilania DC, gniazda wejściowego MIDI, gniazda USB i gniazda wyjściowego audio.

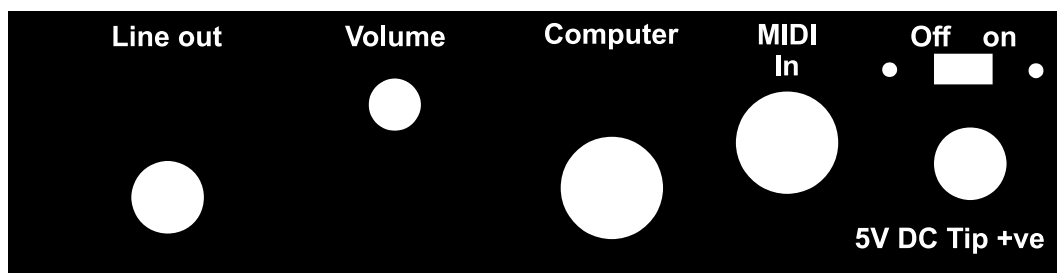
Jeśli używasz określonej obudowy, mocny będzie schemat wiercenia (rysunek 12).

Można go również użyć w innych obudowach, ale trzeba będzie dostosować położenie na panelu, aby pasowało do miejsca montażu płytki drukowanej.

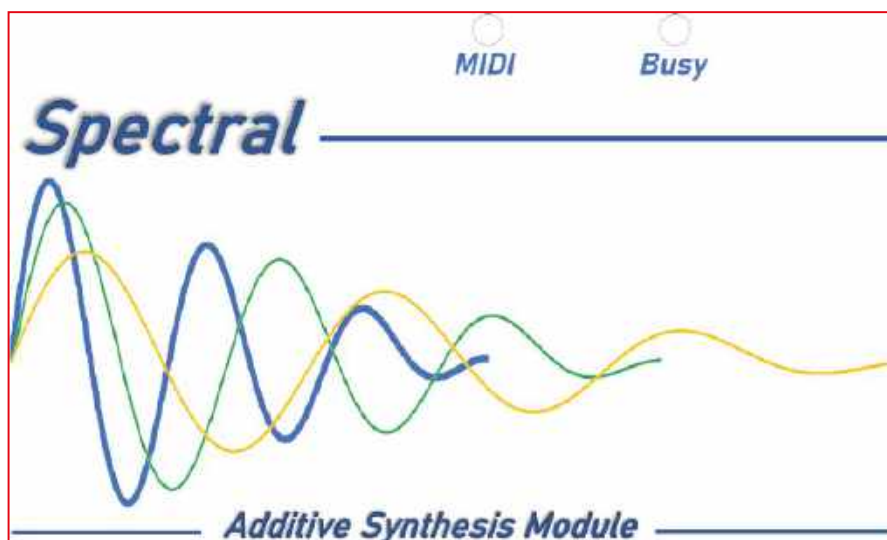
Okablowanie

Po upewnieniu się, że wszystkie te elementy są dostępne przez panel, jeśli jeszcze tego nie zrobiłeś, wywierć otwory na potencjometr głośności i przełącznik zasilania w dogodnych miejscach. Następnie przylutuj do tych elementów odpowiednie odcinki tasiemki i zaciśnij/lutuj piny na pozostałych końcach, które następnie wepchnij do plastikowych spolaryzowanych gniazd (lub przylutuj bezpośrednio do płytki drukowanej).

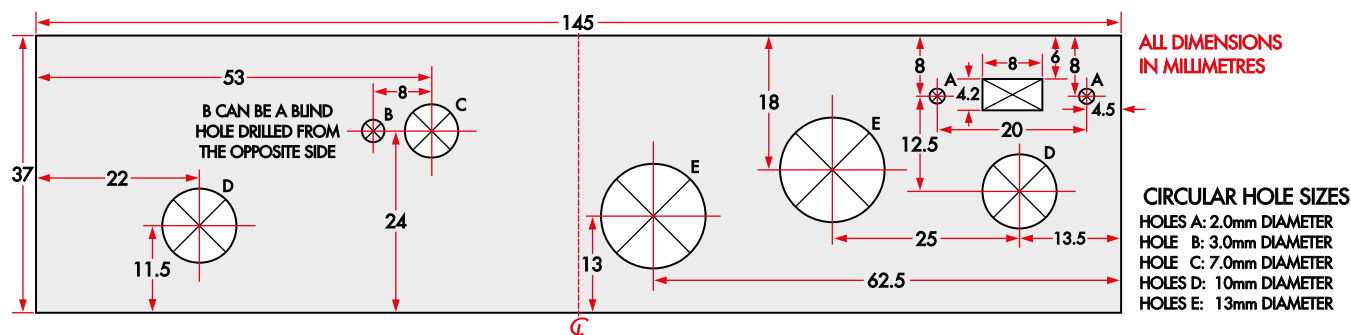
Konieczne będzie również zaznaczenie i wywiercenie dwóch otworów o średnicy 5 mm w pokrywie, przez które będą wystawać diody LED. W zależności od długości przewodów, możesz mieć trochę swobody w umieszczeniu tych diod, ponieważ możesz lekko zgiąć przewody. Pamiętaj, że jeśli naklejasz etykietę na panel pokrywy, będzie ona musiała pokrywać się z tymi otworami.



Rysunek 10. Grafikę panelu przedniego można pobrać, wydrukować, zaalaminować (lub zabezpieczyć w inny sposób), a następnie przymocować do wywierconego panelu



Gotowy syntezator posiada dwie diody LED w górnej części panelu, które wskazują, kiedy odbiera komunikaty MIDI i kiedy komunikuje się z komputerem (np. ładuje patche)



Rysunek 12. Pozycje otworów do wywiercenia/wycięcia w panelu przednim. Potencjometr głośności i włącznik/wyłącznik są zamontowane na panelu, więc można je przesunąć, ale te pozycje są zaprojektowane tak, aby usunąć inne pobliskie części. Można tego użyć w obudowach innych niż zalecana, ale do prawidłowego ustawienia potrzebny będzie co najmniej jeden punkt odniesienia

Teraz jest dobry moment, aby przykleić etykiety na panel przedni i/lub pokrywę (patrz dalszy opis) i wyciąć otwory za pomocą ostrego noża hobbystycznego.

Sprawdź, czy S1 i VR1 są prawidłowo podłączone, zamontuj je na panelu przednim, a następnie podłącz je do złączy CON3 i CON4. Następnie można „zapiąć” płytkę wewnątrz obudowy, włączyć ją i sprawdzić, czy działa.

Aby to zrobić, należy podłączyć go do komputera z systemem Windows, pobrać i zainstalować oprogramowanie opisane w poniższej sekcji oraz sprawdzić, czy może połączyć się z syntezatorem Spectral Sound Synthesiser.

Tworzenie etykiet

Przekonałem się, że drukowanie za pomocą drukarki atramentowej, a następnie spryskiwanie sprayem „utrwalającym” działa dobrze.

W przypadku etykiety na pokrywę użyłem specjalnych arkuszy naklejek wodnych do drukarki atramentowej z serwisu eBay. Mają one papierowy podkład. Proces wygląda następująco:

1. Drukuj tak, jak na dowolnej kartce papieru (ale błyszczącą stroną do góry).
2. Spryskaj lakierem/utrwalaczem.
3. Zanurz w wodzie na 30 sekund do minuty.
4. Delikatnie zsuń naklejkę (bardzo cienka naklejka odkleja się od papierowego podkładu, gdy jest mokra) i naklej ją na obudowę.
5. Dla dodatkowej ochrony warto ponownie polakierować wyschniętą naklejkę.

Oprogramowanie „Patch Editor”

Moduł posiada powiązany, potężny program Windows o nazwie „Patch Editor”, napisany w języku C# Winforms. Zrzut ekranu tego oprogramowania pokazano na **ekranach 1 i 2**. Jest to aplikacja „Click-Once”. NET, którą udostępniłem online pod adresem <https://collectany.blob.core.windows.net/ssm/SpectralSoundModuleApp1/setup.exe>. Można ją znaleźć w magazynie „blob”



Rysunek 13. Syntezator MIDI został połączony ze standardową klawiaturą kontrolera MIDI, wzmacniaczem i głośnikami, tworząc ten elektroniczny klawikord

Microsoft Azure, co oznacza, że użytkownik jest powiadamiany o ulepszeniach wersji, jeśli zostanie zainstalowany z tej lokalizacji online. Dostępna jest też obszerna instrukcja obsługi tego oprogramowania.

Aplikacja zawiera narzędzia pomagające w kształtowaniu „krajobrazu” barwy, obwiedni, filtrów i nie tylko. Obejmuje „wizualizatory” do przeglądania barw zarówno w dziedzinie czasu, jak i częstotliwości, a nawet zawiera analizator harmoniczny, w którym można pobrać zawartość harmoniczną z dźwięku!

Aplikacja posiada również własny unikatowy język programowania o nazwie „Spectral Definition Language” (SDL) [nie mylić z Simple DirectMedia Layer]. Możesz napisać kod SDL, aby precyzyjnie dostroić definicje patchy i łatwo ponownie wykorzystać fragmenty kodu. Pomysł polega na „abstrakcyjnym” projektowaniu dźwięku na wyższym poziomie, upraszczając wszystkie zawiłości szczegółowej konfiguracji.

W tym celu można przechowywać własne fragmenty kodu i uruchamiać je w razie potrzeby za pośrednictwem menu aplikacji

Programowanie mikrokontrolerów

W projekcie zostało zastosowanych osiem mikrokontrolerów trzech różnych typów. Wszystkie są produktami Microchip (jeden PIC i siedem dsPIC), więc można je zaprogramować za pomocą programatora PICkit 3, PICkit 4, Snap lub podobnego. Jeżeli budujesz urządzenie z zakupionego u nas zestawu nie musisz dodatkowo nic programować, ponieważ wszystkie mikrokontrolery zostały u nas wcześniej zaprogramowane.

Każdy typ mikrokontrolera ma własne oprogramowanie. Innymi słowy, istnieją trzy zestawy oprogramowania układowego. Kody są podane w wykazie elementów, a pakiet do pobrania na naszej stronie internetowej zawiera kod źródłowy dla wszystkich trzech oraz trzy pliki HEX potrzebne do ich zaprogramowania.

Jeśli chcesz przebudować kod źródłowy, aby utworzyć nowe pliki HEX (np. chcesz wprowadzić zmiany w sposobie działania), będziesz potrzebować kompilatora Microchip XC16 Pro (który można zamówić na stronie internetowej Microchip; dostępne są również bezpłatne wersje próbne). W przeciwnym razie poziom optymalizacji 3 nie będzie działał, a wynikowe oprogramowanie układowe nie będzie wystarczająco szybkie.

Kod PIC18 jest mniej krytyczny, więc prawdopodobnie można użyć darmowego kompilatora XC8 do zbudowania tego pliku HEX.

– potężna koncepcja, z gotowymi przykładami ustawień domyślnych takimi, jak np. obwiednia „uderzonej struny”!

Przemyślenia końcowe

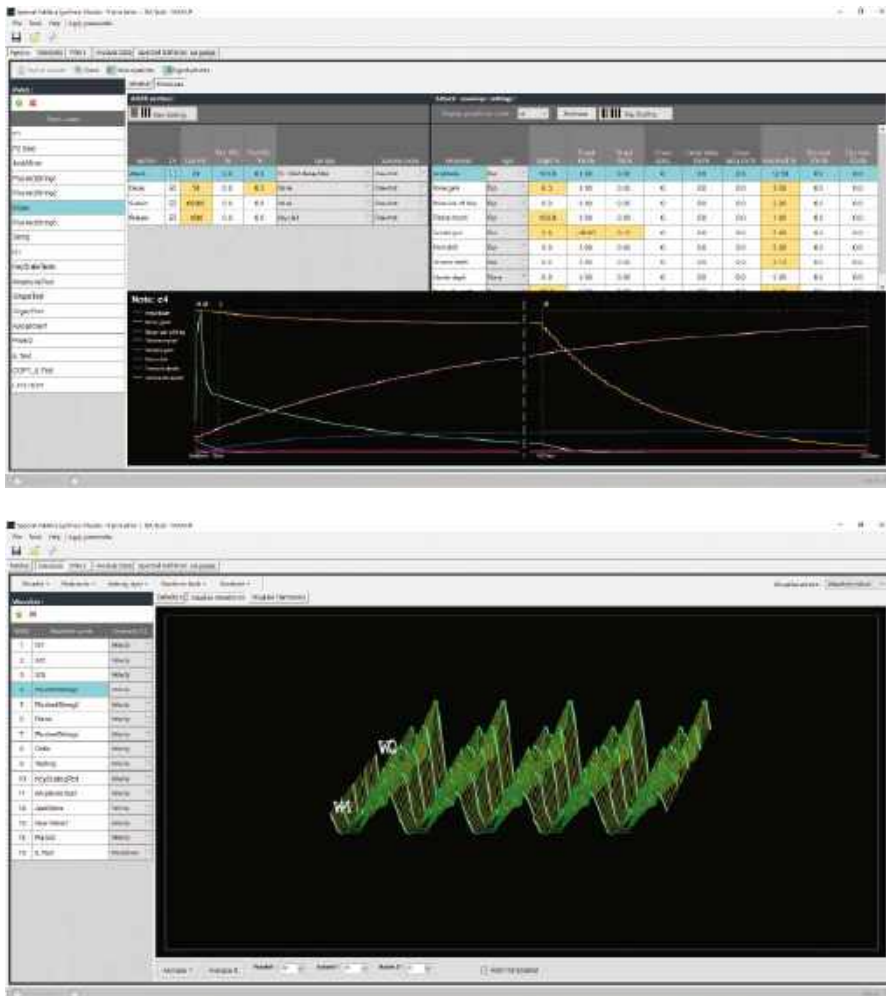
Przedstawiony projekt był bardzo intensywną, ale satysfakcjonującą podróżą, często sprawiającą wrażenie „strzelania do księżycy”. Pokazuje on, że synteza dźwięku jest wciąż żyzną glebą dla eksperymentów i inwencji. Na **rysunku 13** został przedstawiony mój DIY „Elektroniczny klawikord” zawierający standardową klawiaturę kontrolera MIDI połączoną z tym modulem oraz niewielkim wzmacniaczem i głośnikiem.

Jednym z kuszących pomysłów na przyszłość może być podejście do problemu systemu opartego na barwie dźwięku z bardziej holistycznego punktu widzenia. Zamiast do każdej zagranej nuty przypisywać własny wavetable, można by pomyśleć o wymaganych harmonicznch ze wszystkich granych nut jak o jednej gigantycznej puli oscylatorów. Moglibyśmy wtedy wykorzystać zjawisko „maskowania psychoakustycznego”, aby znacząco „przyciąć” rzeczywiste harmoniczne, które wymagają obliczeń.

Wymagałoby to zdolności do ustalania priorytetów ważności harmonicznch i ignorowania tych o najmniejszym znaczeniu. Interesującym aspektem tego podejścia byłoby to, że próg mógłby opierać się na wydajności systemu, zawsze przetwarzając maksymalną możliwą liczbę harmonicznch, ale w razie potrzeby pogarszając jakość dźwięku w kontrolowany sposób. Podejście to może również odbiegać od wymagań harmonicznch opartych na liczbach całkowitych, oferując większy realizm.

Innym pomysłem jest powrót do podejścia opartego w większym stopniu na próbkach, ale zamiast przechowywać próbki w konwencjonalnym formacie PCM, przechowuj je jako barwy lub nawet jako zmiany barwy. Może to przynieść znaczne oszczędności.

Inne pomysły zaczynają kwestionować podstawy precyzji wymaganej dla poziomów harmonicznch. Ponieważ ludzie postrzegają dźwięk logarytmicznie, odpowiednie skalowanie poziomów może wynikać z prostego przesunięcia bitów. Czy naprawdę jesteśmy w stanie dostrzec różnicę



Ekran 1 i 2. Przykładowe zrzuty ekranu z potężnego, oprogramowania Patch Editor dla systemu Windows, zaprojektowanego do współpracy z syntezatorem Spectral Sound Synthesiser. Jego kod źródłowy załączono w udostępnionej do pobrania paczce

Przydatne linki

- Biologiczne podstawy percepcji barwy dźwięku: siliconchip.com.au/link/abdc
- Synteza szarpanych strun: siliconchip.com.au/link/abdd
- Synteza instrumentów dętych: siliconchip.com.au/link/abde
- Jakość lub barwa dźwięku: siliconchip.com.au/link/abdf
- Szczegóły dotyczące barwy dźwięku: www.dspguide.com/ch22/2.htm

w poziomach harmonicznch w takim stopniu, że uzasadnia to cokolwiek lepszego?

Co więcej, musimy więcej myśleć o tym, jak nasze mózgi postrzegają dźwięk, a mniej o czyistości obliczeń matematycznych. Nasze mózgi pracują nad wrażeniami i rozpoznawaniem ogólnych cech, więc być może istnieje potencjał w skupieniu się na technikach, które przynoszą ogromne oszczędności obliczeniowe,

pomijając rzeczy, które po prostu nie mają znaczenia dla percepcji. Wygląda na to, że jest jeszcze wiele do przemyślenia w kwestii syntezy dźwięku! ■

Jeremy Leach's

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

REKLAMA

numery archiwalne • prenumerata • książki
www.UlubionyKiosk.pl



Regulowane „sztuczne obciążenie” z wbudowanym generatorem zegarowym do testowania zasilaczy, przetwornic napięcia i akumulatorów

Źródło prądu obciążenia, zwane również „sztucznym obciążeniem”, zwykle nie stanowi standardowego wyposażenia laboratorium elektronicznego, mimo że daje wiele korzyści w procesie testowania wszelkiego rodzaju źródeł zasilania. Komercyjne urządzenia są niestety dość drogie. Warto więc wziąć do ręki lutownicę, cząłki i zbudować sobie takie obciążenie elektroniczne.

Jak się testuje zasilacze, przetwornice DC/DC lub akumulatory? Większość elektroników ma jakiś zestaw rezystorów mocy. Każdy z nich można przyłożyć jako obciążenie do testowanego urządzenia i zmierzyć prąd oraz napięcie. Następnie można wykonać serię pomiarów w celu określenia statycznej charakterystyki pętli sterowania przy różnych obciążeniach i napięciach wejściowych. Na podstawie tych pomiarów można też obliczyć rezystancję wewnętrzną źródła zasilania. Procedura ta jest poprawna, ale skomplikowana i czasochłonna.

O wiele lepiej i łatwiej byłoby użyć regulowanego źródła prądu obciążenia, który pobiera stały prąd niezależnie od przyłożonego napięcia. Takie regulowane źródło nie umożliwia jednakże uchwycenia dynamicznego zachowania zasilacza, tj. tego, jak reaguje on na szybkie zmiany obciążenia. A dla pełnej oceny zasilacza ma to zasadnicze znaczenie. Zasilacze często reagują na nagle zmiany obciążenia silnymi przekroczeniami napięcia

wyjściowego, co może powodować nieprawidłowe działanie lub nawet uszkodzenie podłączonego do nich sprzętu. Co więcej, obwody sterowania zasilacza mogą być niestabilne i mieć tendencję do generowania oscylacji o wysokiej częstotliwości, szczególnie przy szybkich zmianach obciążenia. Oscylacje takie mają nie mniej negatywne skutki jak przeregulowania. W praktyce oznacza to, że źródło prądu obciążenia, oprócz działania statycznego, musi być również w stanie realizować szybkie zmiany prądu.

Takie zadanie spełnia sztuczne obciążenie tutaj opisywane. Dwie jego wersje pokazano na **rysunku 1**. Konstrukcja układu jest prosta i nie zawiera żadnych egzotycznych elementów, pozwala też uniknąć konieczności stosowania oddzielnego zasilacza. W zależności od tego, jakich użyjemy radiatorów, źródło może obsługiwać moc do 18 W (do 50 W z układem wspomagającym) i nadaje się dla maksymalnego napięcia wejściowego 30 V.

Chociaż maksymalna moc strat opisywanego układu może wydawać się niewysoka, można dzięki niemu oszacować charakterystykę dynamiczną nawet całkiem mocnych zasilaczy. Możliwości urządzenia są potwierdzone przez specyfikacje wymienione w ramce Parametry.

Jak to działa

Niemal każde źródło prądu jest układem liniowym, w którym na pewnej stałej rezystancji jest utrzymywane stałe napięcie. Zgodnie z prawem Ohma stałe napięcie na stałej rezystancji skutkuje przepływającym przez nią stałym prądem. Prąd ten – jeśli pominiemy niewielki prąd pobierany przez układ elektroniczny – przepływa przez obwód mierzony i w rezultacie stanowi prąd obciążenia testowanego zasilacza. Stały prąd występuje nawet wtedy, gdy zmienia się napięcie na wejściu układu. Mamy więc obciążenie, które pobiera stały prąd w każdych warunkach.

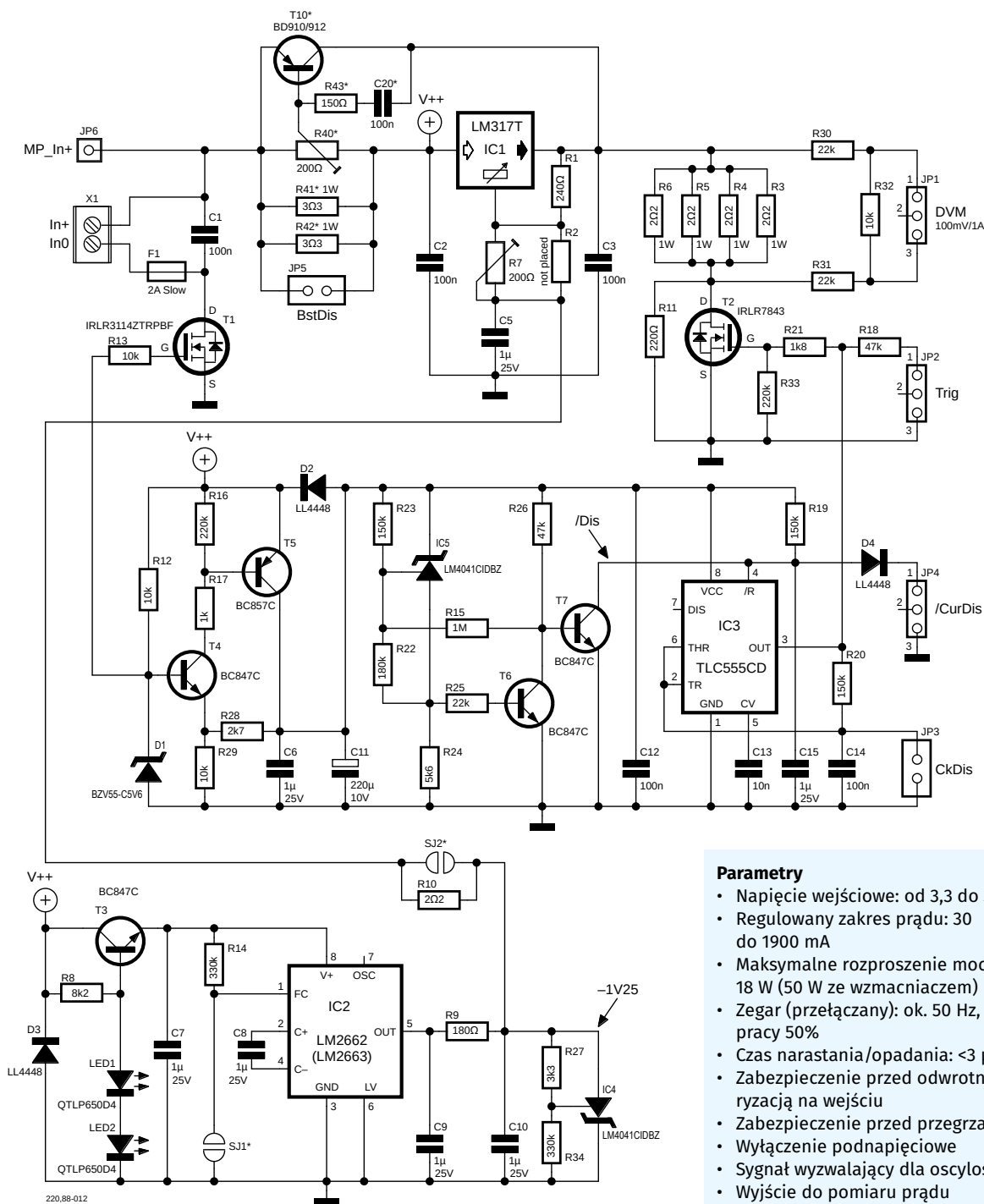
Prąd obciążenia może być ustawiany poprzez regulację napięcia na stałej rezystancji obciążenia. Naturalnym rozwiązaniem będzie użycie do tego celu scalonego stabilizatora napięcia. Oprócz prostoty ma to tę zaletę, że stabilizatory scalone mają wbudowane funkcje zabezpieczające.

Najniższym możliwym napięciem wyjściowym większości stabilizatorów jest wewnętrzne napięcie odniesienia, wynoszące zwykle 1,25 V. Aby umożliwić ustawienie prądu na zero, konieczne są dodatkowe środki. Napięcie wyjściowe będzie można regulować od 0 V, jeżeli potencjał odniesienia stabilizatora wyniesie $-1,25$ V. Wtedy zakres regulacji prądu rozciągnie się w dół do zera. Wymagane napięcie ujemne



Rysunek 1. Po lewej: sztuczne obciążenie w wersji 18 W, z dużą płytką radiatora. Przykręcenie LM317 plastikową śrubą było kiepskim pomysłem. Po prawej: płytka w pełni zmontowana z dodatkowymi elementami dla wersji do 50 W





Rysunek 2. Schemat układu sztucznego obciążenia

można stosunkowo łatwo uzyskać za pomocą pompy ładunkowej i prostego stopnia stabilizującego.

Najprostszym sposobem taktowania prądu obciążenia jest użycie tranzystora MOSFET, który kluczuje rezystor obciążenia na wyjściu stabilizatora napięcia. Częstotliwość taktowania powinna być zbliżona do częstotliwości sieci (50 Hz). W ten sposób testowany zasilacz będzie obciążany przez pełną połowę okresu napięcia sieciowego,

co pozwoli na sprawdzenie, czy kondensatory filtrujące zasilacza mają właściwą pojemność.

Układ

Układ sztucznego obciążenia pokazano na **rysunku 2**. Elementy oznaczone gwiazdką (*) są wymagane tylko dla układu w wersji umożliwiającej rozpraszanie większych obciążeń. Jako stabilizator napięcia IC1 został użyty sprawdzony typ LM317. Stabilizator

Parametry

- Napięcie wejściowe: od 3,3 do 30,0 V
- Regulowany zakres prądu: 30 do 1900 mA
- Maksymalne rozproszenie mocy: 18 W (50 W ze wzmacniaczem)
- Zegar (przetaczany): ok. 50 Hz, cykl pracy 50%
- Czas narastania/opadania: <3 μs
- Zabezpieczenie przed odwrotną polaryzacją na wejściu
- Zabezpieczenie przed przegrzaniem
- Wyłączenie podnapięciowe
- Sygnał wyzwalający dla oscyloskopu
- Wyjście do pomiaru prądu
- Zasilanie z testowanego urządzenia

ten ma optymalne właściwości termiczne, dopuszcza moc traconą do 20 W w rozsądnych temperaturach otoczenia, ma maksymalne napięcie wejściowe 37 V oraz zabezpieczenie przed przegrzaniem i przeciążeniem prądowym.

Układ IC2 to inwerter napięcia typu LM2662 oparty na pompie ładunkowej. Wytwarza on ujemne napięcie, które jest stabilizowane na poziomie -1,25 V przez układ IC4 – precyzyjne źródło napięcia odniesienia

LM4041. Ujemne napięcie trafia do potencjometru nastawnego R7, którym ustawiamy potencjał odniesienia w obwodzie stabilizatora napięcia IC1. Rezystory R1 i R7 tworzą dzielnik napięcia ze stałym spadkiem napięcia +1,25 V na R1 i napięciem -1,25 V na dolnym końcu R7. Potencjał odniesienia IC1 można regulować w zakresie od -1,25 V do -0,2 V. W rezultacie napięcie wyjściowe stabilizatora waha się od 0 V do około +1 V. Napięcie to jest przykładane do grupy połączonych równolegle rezystorów obciążających (R3...R6), określających prąd obciążenia. Szeregowo z rezystorami jest włączony tranzystor przełączający T2, który realizuje taktowanie prądu.

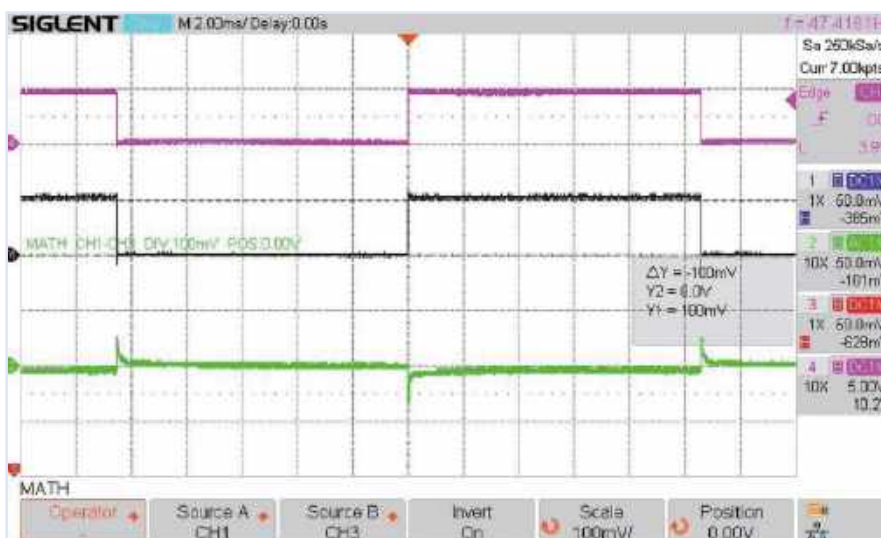
Aby utrzymać niską moc traconą w IC4, pompa ładunkowa IC2 działa przy napięciu zasilania zaledwie 3 V. Napięcie to pochodzi z wtórника emiterowego T3, zawierającego w obwodzie bazy dwie połączone szeregowo zielone diody LED. Zapewniają one bardziej stromą charakterystykę napięciowo-prądową (ostrzejsze „kolano”) niż dioda Zenera o tym samym napięciu.

Reszta obwodu jest zasilana przez prosty układ stabilizatora napięcia składający się z T4, T5 i diody Zenera D1, która zapewnia napięcie odniesienia. Układ ten ma zaletę w postaci niskiego spadku napięcia (LDO), dzięki czemu może zapewnić wymagane napięcie bramki dla T1 i T2 nawet przy niskim napięciu wejściowym.

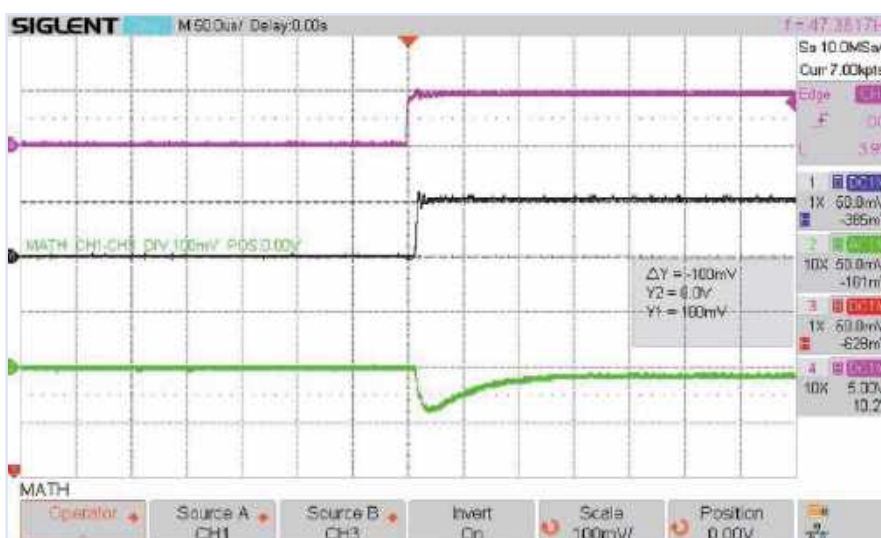
Sygnał sterujący dla tranzystora T2 pochodzi od kolejnego dobrego znajomego – klasycznego timera 555 (IC3). Gdy kondensator C14 (odpowiedzialny za wytwarzanie impulsów w timerze) jest zwarty zworką JP3 (ClkDis), wyjście IC3 jest w stanie wysokim, a T2 jest stale włączony. Odpowiada to statycznemu trybowi pracy. Gdy zworka zostanie zdjęta, prąd obciążenia będzie taktowany.

Po włączeniu zasilania, jeśli napięcie wejściowe jest niższe niż minimalne napięcie robocze, układ IC3 jest resetowany z wejścia Reset. W tej sytuacji wyjście timera jest w stanie niskim i prąd obciążenia nie może płynąć. Załączenie wyjścia jest opóźnione przez obwód czasowy R19/C15. Sztuczne obciążenie można wyłączyć stanem niskim zewnętrznego sygnału /CurDis na złączu JP4. Pozwala to między innymi na rozładowywanie akumulatora w określonym czasie z użyciem płytki Arduino.

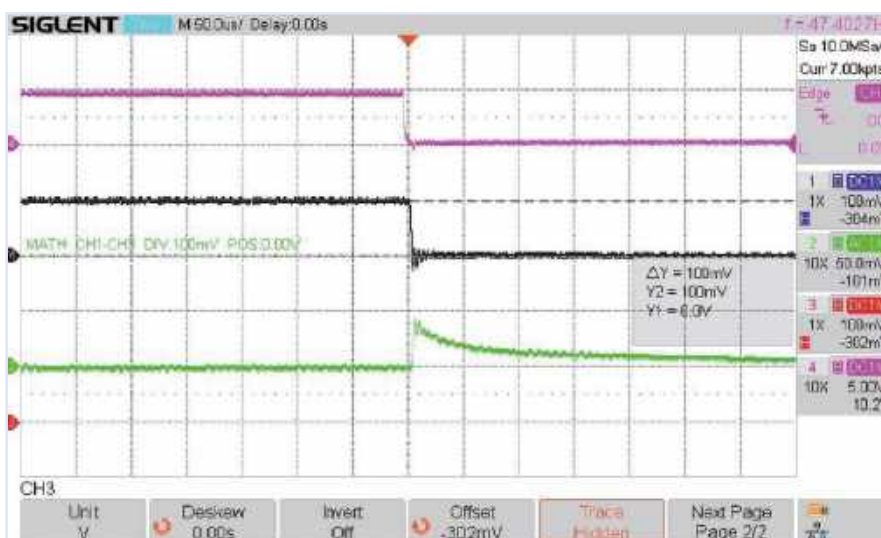
Wyłączenie podnapięciowe jest realizowane za pomocą stabilizatora równoległego IC5 (LM4041), który jest tak podłączony, że przy napięciu niższym niż 3 V nie przepływa przez niego prąd, co skutkuje stałym resetowaniem timera 555 tranzystorami T6 i T7. Gdy napięcie zasilania wzrośnie



Rysunek 3. Dynamiczne zachowanie liniowego zasilacza laboratoryjnego przy 12 V/1 A



Rysunek 4. Odpowiedź skokowa zasilacza laboratoryjnego z rysunku 3 na załączenie obciążenia



Rysunek 5. Odpowiedź skokowa zasilacza laboratoryjnego z rysunku 3 na wyłączenie obciążenia

powyżej około 3,1 V, IC5 zaczyna przewodzić, powodując włączenie T6 i poprzez T7 wycofanie resetu IC3.

Na JP1 jest dostępny sygnał napięciowy, proporcjonalny do prądu obciążenia, który można mierzyć woltomierzem. Współczynnik proporcjonalności wynosi 0,1 V/A. Na JP2 jest wyprowadzony sygnał o wartości szczytowej około 6 V. W trybie taktowania może on służyć do wyzwalania oscyloskopu.

Wejście układu jest chronione przed odwrotną polaryzacją przez tranzystor MOSFET T1. Gdy polaryzacja napięcia wejściowego jest prawidłowa, prąd przepływa najpierw przez wewnętrzną diodę T1, co powoduje uruchomienie wewnętrznego zasilacza i w konsekwencji pełne załączenie T1 napięciem bramki. W przypadku zwarcia w układzie sztucznego obciążenia przepali się bezpiecznik F1, zapobiegając bardziej dramatycznym skutkom.

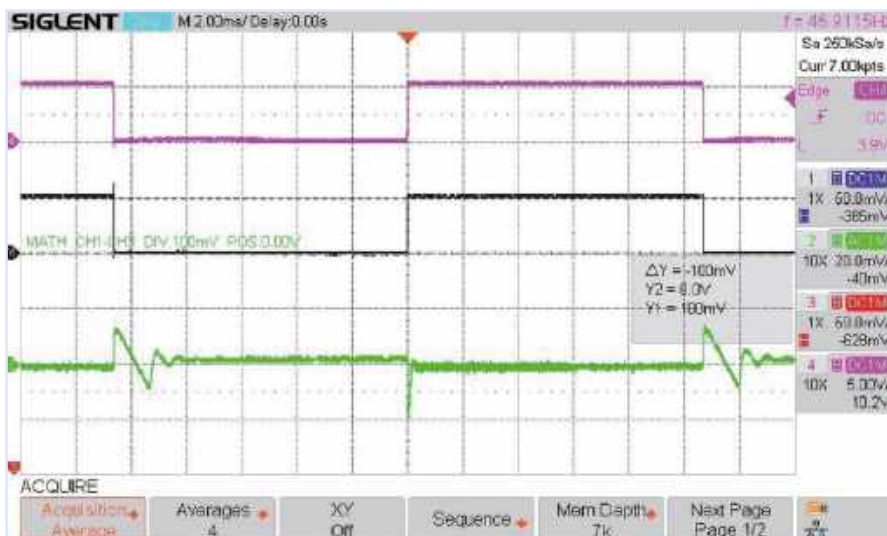
Wyższa moc

Jeśli chcemy zwiększyć moc maksymalną obciążenia, możemy przed stabilizatorem napięcia podłączyć układ wzmacniający, odciażający kostkę LM317. Zwiększy to dopuszczalne rozpraszanie mocy nawet do 50 W, w zależności od radiatora. Maksymalny prąd pozostaje jednak taki sam. Niestety, układ wzmacniający nie jest chroniony przez funkcje zabezpieczające LM317, więc ważne jest, aby radiator miał odpowiednie wymiary. Na układzie wzmacniacza odkłada się około 1,2 V napięcia, więc minimalne napięcie wyjściowe (MP_In+) musi wzrosnąć do 5 V. Dla napięć poniżej 8 V układ wzmacniający powinien być zawsze wyłączony poprzez założenie zworki na JP5 (BstDis).

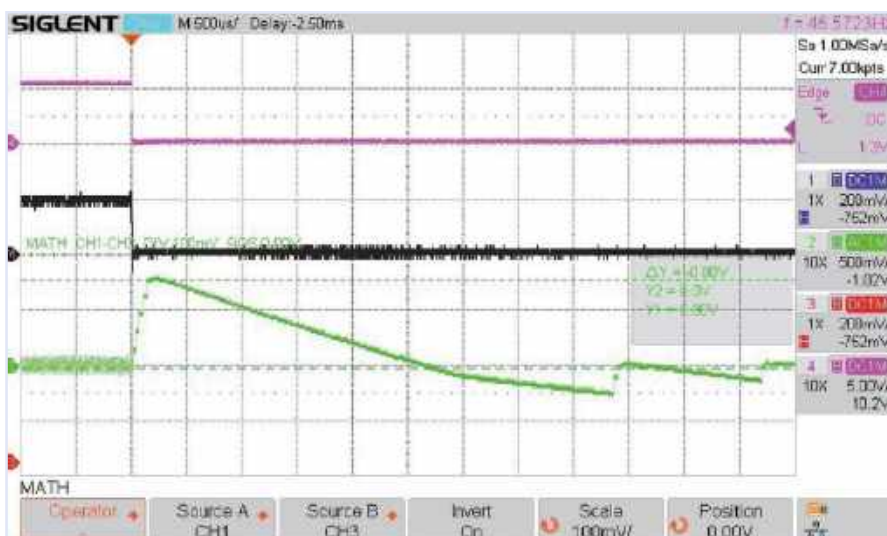
Zasada działania układu wzmacniającego jest dość prosta. Praktycznie cały prąd wejściowy przepływa przez rezystory szeregowo R41 i R42. Potencjometr R40 podaje część spadku napięcia na tych rezystorach na złącze baza-emiter tranzystora PNP (T10). R40 powinien być ustawiony tak, żeby T10 zaczął przewodzić przy prądzie wejściowym około 0,6 A. W rezultacie LM317 reguluje napięcie w zwykły sposób aż do prądu 0,6 A, a każda nadwyżka prądu powyżej tej wartości jest dostarczana przez T10 bezpośrednio do rezystorów obciążenia R3...R6. Obwód RC R43/C20 zapobiega oscylacjom.

Procedura testowa i interpretacja wyników

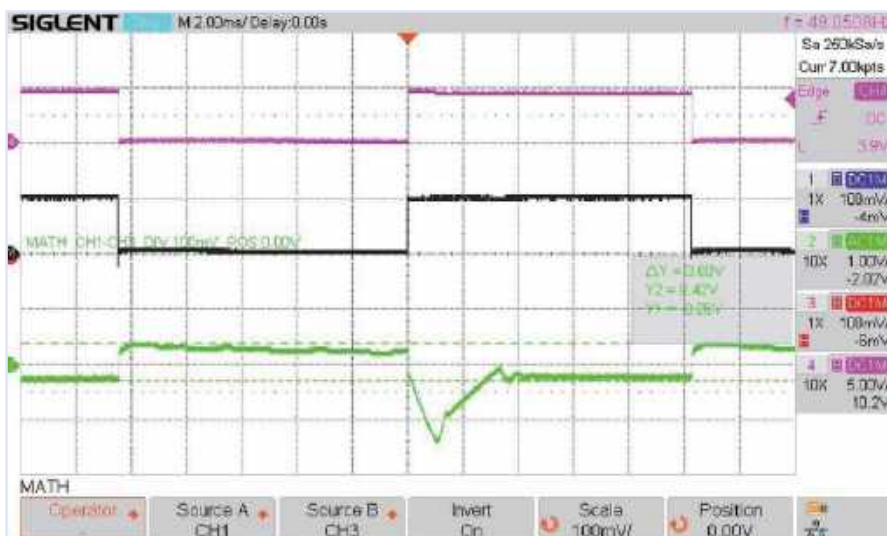
Aby mierzyć parametry statyczne źródła zasilania należy wyłączyć taktowanie, wkładając zworkę do JP3. Napięcie wyjściowe mierzymy woltomierzem podłączonym



Rysunek 6. Dynamiczne zachowanie zasilacza impulsowego przy 12 V/1 A



Rysunek 7. Odpowiedź skokowa przetwornicy napięcia typu flyback na wyłączenie obciążenia przy 13 V/1 A

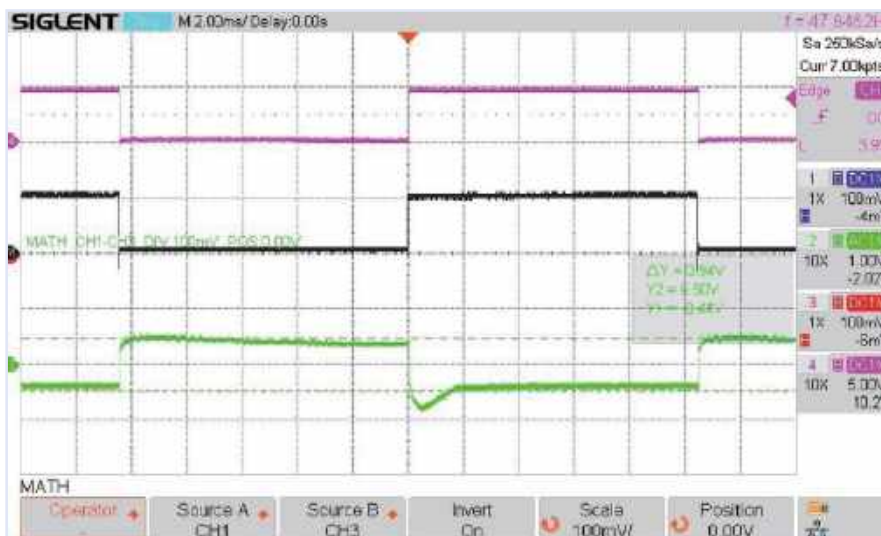


Rysunek 8. Dynamiczne zachowanie ładowarki USB przy prądzie 1 A

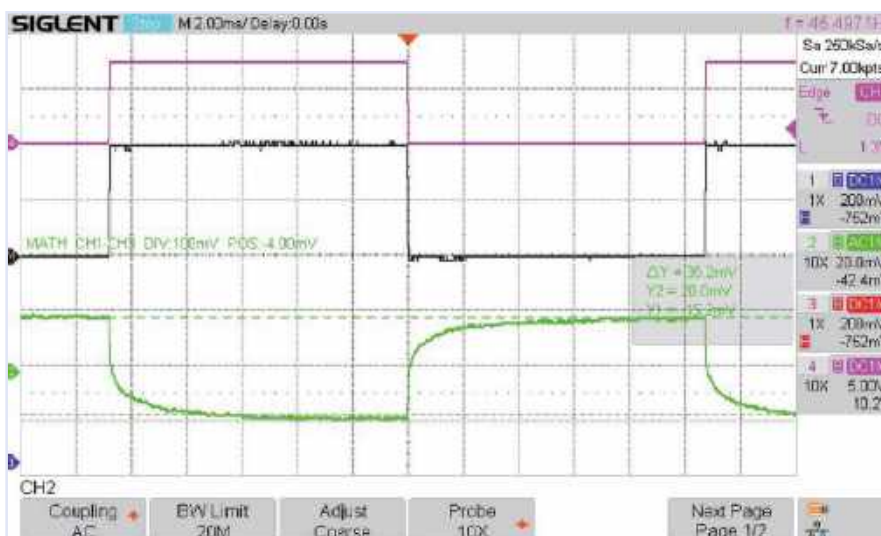
do wyjścia testowanego urządzenia. Pomiar bezpośrednio na sztucznym obciążeniu prowadziłby do nieprawidłowego wyniku z powodu spadku napięcia na przewodach łączących. Równolegle z woltomierzem podłączamy również oscyloskop z wejściem ustawionym w tryb AC. Następnie, mierząc napięcie źródła, stopniowo zwiększamy natężenie prądu. Pozwoli to uzyskać serię pomiarów pokazujących błąd regulacji napięcia testowanego urządzenia. Oscyloskop wskaże wszelkie tendencje do oscylacji w badanym urządzeniu. Podczas zwiększania natężenia prądu należy zawsze zwracać uwagę na maksymalną moc traconą w sztucznym obciążeniu. Z dobrym radiatorem na LM317 jest możliwe trzymanie do 18 W. Uwaga – przy wyższych napięciach wejściowych taka moc zostanie osiągnięta już przy niewysokim prądzie. Na szczęście przed uszkodzeniem ochronią nas wewnętrzne funkcje zabezpieczające LM317 (bezpieczny obszar działania; SOA).

Następnym krokiem jest pomiar charakterystyki dynamicznej testowanego urządzenia przy pomocy oscyloskopu. W tym celu należy włączyć taktowanie. Ważne jest, aby oscyloskop był podłączony bezpośrednio do źródła zasilania, ponieważ w przeciwnym razie wyniki będą zafałszowane przez skoki napięcia spowodowane indukcyjnością przewodów łączących. Podstawowa procedura pomiarowa jest taka sama jak w przypadku pomiarów statycznych. Moc tracona jest mniejsza o połowę dzięki taktowaniu z 50% wypełnieniem, dzięki czemu można teraz używać wyższych napięć i/lub prądów. Każdy z oscylogramów przedstawionych w tym artykule pokazuje sygnał wyzwalający na górze (CH4, fioletowy), prąd w środku (czarny, mierzony różnicowo na JP1 z czułością 0,1 V/A) i napięcie wyjściowe na dole (CH2, zielony). Należy zwrócić uwagę na różne rozdzielczości poziome i pionowe oscylogramów.

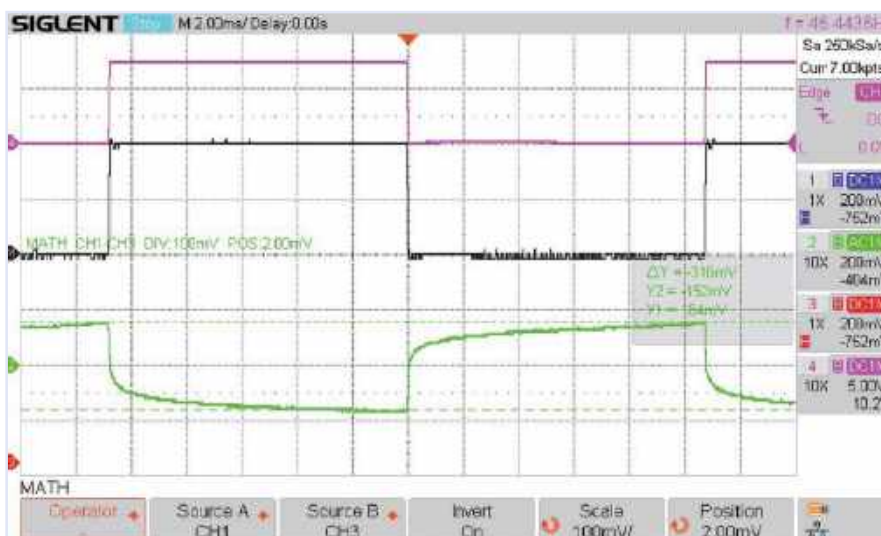
Rysunek 3 przedstawia zachowanie dynamiczne liniowego zasilacza laboratoryjnego. **Rysunki 4 i 5** pokazują reakcję na zmiany skokowe, gdy obciążenie jest odpowiednio załączane i wyłączane. Najbardziej interesuje nas tutaj amplituda i czas trwania przeregulowań. **Rysunek 6** jest odpowiednikiem rysunku 3, ale zmierzonym przy przełączaniu zasilacza laboratoryjnego. **Rysunek 7** przedstawia skokową reakcję przetwornicy typu flyback na zmiany obciążenia. **Rysunki 8 i 9** przedstawiają dynamiczne zachowanie dwóch ładowarek USB; tu zwracamy uwagę na stały spadek napięcia wyjściowego. Spadek napięcia pod obciążeniem i związany z nim prąd odpowiadają wewnętrznej rezystancji źródła zasilania.



Rysunek 9. Dynamiczne zachowanie innej ładowarki USB przy prądzie 1 A



Rysunek 10. Pomiar rezystancji wewnętrznej nowego trzyogniowego akumulatora LiPo przy prądzie 2 A



Rysunek 11. Pomiar rezystancji wewnętrznej używanego trzyogniowego akumulatora LiPo przy prądzie 2 A

Można również sprawdzić, czy wewnętrzne filtrowanie i tłumienie zakłóceń jest dobrane prawidłowo. Gdy źródło zasilania ma tendencję do oscylacji o wysokiej częstotliwości, pomiar niezawodnie to ujawni. Jeśli projektujemy samodzielnie zasilacz, prezentowane sztuczne obciążenie będzie nieocenionym narzędziem do optymalizacji parametrów czasowych i zapewnienia stabilności pętli sterowania.

Pomiar dynamiczny akumulatora dostarcza informacji na temat jego rezystancji wewnętrznej, która jest miarodajnym wskaźnikiem jego ogólnego stanu. Przy tym pomiarze prąd nie powinien przekraczać współczynnika 1...2 pojemności nominalnej. Rezystancja wewnętrzna ogniwa akumulatora litowo-jonowego o pojemności 1500 mAh wynosi zwykle poniżej 30 mΩ, a dla większych pojemności jest zwykle jeszcze mniejsza – poniżej 10 mΩ. Dobrze jest zawsze zmierzyć rezystancję wewnętrzną nowej baterii, aby móc później wykorzystać tę wartość do celów porównawczych.

Na oscylogramie na **rysunku 10** widać, że akumulator o rezystancji wewnętrznej wynoszącej 5,9 mΩ na ogniwo jest w dobrym stanie. Natomiast patrząc na **rysunek 11** widać, że akumulator o rezystancji wewnętrznej 53 mΩ na ogniwo szybko zbliża się do końca okresu użytkowania.

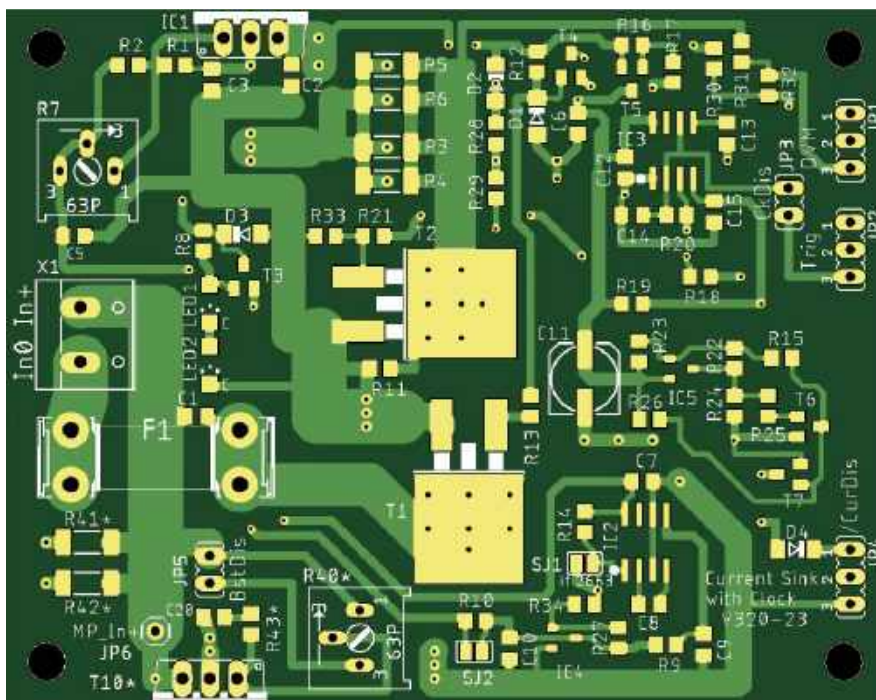
Konstrukcja

Projekt płytki uwzględnia opcjonalne dodanie układu wzmacniającego. Projekt z **rysunku 12**, a także schemat ideowy, można pobrać w formacie PDF z [1]. Do pobrania są również pliki Eagle.

Miejsca na płytce dla T1 i T2 pozwalają na umieszczenie zarówno elementów DPAK jak i D2PAK. W przypadku IC2 można użyć LM2662 lub LM2663, więc nie powinno być żadnych problemów z zakupem elementów. Jeśli nie potrzebujemy wzmacniacza, możemy pominąć elementy oznaczone gwiazdką (*). Należy zwrócić szczególną uwagę na radiatory dla IC1 i T10, ponieważ niezawodność sztucznego obciążenia zależy zasadniczo od nich. Rezystancja termiczna tych radiatorów musi być mniejsza niż 2 K/W, a elementy półprzewodnikowe muszą być zamontowane przy użyciu wysokiej jakości podkładek termicznych i metalowych śrub z odpowiednimi podkładkami, zapewniającymi izolację elektryczną.

Rozpoczęcie pracy ze sztucznym obciążeniem

Jeśli jako IC2 używasz LM2663, mostek lutowiczy SJ1 musi być zwarty. Załóż



Rysunek 12. Projekt płytki drukowanej sztucznego obciążenia z układem wzmacniacza

zworkę na JP5 (BstDis). Potencjometrem R7 wyreguluj układ na minimalny prąd, gdy jest on podłączony do regulowanego zasilacza z napięciem wyjściowym 3 V i ograniczeniem prądu 100 mA. Obciążenie powinno pobierać prąd spoczynkowy o wartości około 10 mA. Następnie zwiększ napięcie zasilania do 3,3 V. Potencjometrem R7 ustaw prąd na około 80 mA. Teraz zmniejsz napięcie zasilania. Przy około 3,1 V prąd powinien spaść do poziomu spoczynkowego. Będzie to oznaczało, że wyłącznik podnapięciowy działa prawidłowo. Przy napięciu wejściowym 5 V i potencjometrze R7 ustawionym na minimalną wartość, prąd powinien być mniejszy niż 30 mA. W przeciwnym razie można go zmniejszyć, zwierając mostek lutowiczy SJ2.

Aby uruchomić wzmacniacz, obróć potencjometr R40 w kierunku przeciwnym do ruchu wskazówek zegara, zapobiegając przepływowi prądu bazy T10, a następnie

Pytania lub komentarze?

W przypadku pytań technicznych lub komentarzy dotyczących tego artykułu prosimy o kontakt z redakcją „Elektor” pod adresem editor@elektor.com lub z redakcją „Elektroniki dla Wszystkich” – kontakt@elportal.pl.

usuń zworkę z JP5. Przy napięciu wejściowym 15 V i prądzie 1 A, wyreguluj R40 tak, aby prąd wejściowy stabilizatora napięcia wynosił 0,6 A. Prąd wejściowy stabilizatora można zmierzyć na podstawie spadku napięcia na rezystorach R41 i R42, mierzonym na JP5. 1,0 V odpowiada tam 0,6 A. Następnie zwiększ napięcie wejściowe do 25 V, a prąd do 2 A. W razie potrzeby można ponownie regulację po nagraniu się układu sztucznego obciążenia. ■

Roland Stiglmayr
(Niemcy)

O autorze

Roland Stiglmayr studiował informatykę w latach 70. i ma 40-letnie doświadczenie w pracach badawczo-rozwojowych. Zajmował się głównie opracowywaniem komputerów mainframe, światłowodowych systemów transmisji danych, zdalnych głowic radiowych do sieci bezprzewodowych oraz bezstykowych systemów transmisji energii. Obecnie pracuje jako konsultant, a jego szczególną pasją jest przekazywanie wiedzy.

Przydatne linki

[1] Materiały dodatkowe na stronie Elektor Labs: www.elektormagazine.com/labs/adjustable-current-sink-with-integrated-clock-generator

Mini sterownik LED

Opisany w artykule niewielki, tani moduł może sterować stosunkowo dużymi białymi diodami LED o napięciu 12 V ze źródła zasilania USB lub 5 V DC. W wielu sytuacjach do oświetlenia jakiejś powierzchni nie jest potrzebny reflektor, wystarczy niewielka ilość światła dostarczanego na przykład przez Mini LED Driver.

W czerwcowym numerze Silicon Chip z 2022 r. (oraz w EdW 1/2025) opisywaliśmy układ do zasilania paneli LED o mocy 70 W. Charakteryzują się one niezwykle dużą jasnością, gdy pracują z maksymalną mocą (około 6 A przy 12 V). Jednak mogą być nadal przydatne również przy niższych prądach. Wytwarzają sporo światła nawet przy prądzie 1 A (12 W). Ponadto istnieje wiele innych białych diod LED, które są zaprojektowane do pracy z mocą około 10 W. Prezentowany mini sterownik (zasilacz) LED jest dla nich idealny.

Główną motywacją stojącą za tym rozwiązaniem jest bezpieczne zasilanie paneli LED 12 V ze źródła 5 V DC. Jak wielu użytkowników różnych urządzeń elektronicznych masz prawdopodobnie wiele zapasowych zasilaczy USB lub banków energii, które można wykorzystać jako zasilacze 5 V.

Opisany sterownik może dostarczyć wystarczający prąd do zasilania większości białych diod LED, zapewniając całkiem zadawalający poziom oświetlenia. Jeśli są to duże panele, takie jak te o mocy 70 W, ich żywotność zostanie znacznie wydłużona dzięki zmniejszonemu wytwarzaniu ciepła.

Mini LED Driver bazuje na powszechnie dostępnych, tanich modułach boost, w których jest zastosowany układ scalony XL6009, ale dodaje im kilka dodatkowych funkcji. Moduły te nie mają wbudowanego ograniczenia prądu, z wyjątkiem ochrony przed zwarceniem. Nasze dodatkowe usprawnienia zapewniają regulowane ograniczenie prądu.

W ubiegłym miesiącu, w artykule na temat sterowników LED (*Buck-Boost LED Driver*) wyjaśniliśmy, dlaczego lepiej jest zasilac diody LED ze źródła z ograniczeniem prądu.

Krótko mówiąc, dostarczanie stałego napięcia do diod LED nie zapewni stałego strumienia świetlnego. Niewielkie wahania napięcia mogą powodować nieproporcjonalnie duże zmiany natężenia prądu, być może nawet na tyle duże, by uszkodzić diody LED.

Dodana przez nas funkcja ograniczania prądu będzie również chronić zasilanie wejściowe, szczególnie wtedy, gdy do zasilania diod LED, które pobierałyby zbyt duży prąd, aby mogły pracować z pełną jasnością używasz małego zasilacza USB.

Inną funkcją dodaną przez sterownik jest wyłącznik niskiego napięcia wejściowego. Pozwala on uniknąć sytuacji, w której moduł boost nie działa poprawnie przy niskim napięciu wejściowym. Ponadto, gdy urządzenie zasilane jest z akumulatora, rozwiązanie to chroni go przed nadmiernym rozładowaniem, które mogłoby prowadzić do jego uszkodzenia.

Moduł boost XL6009

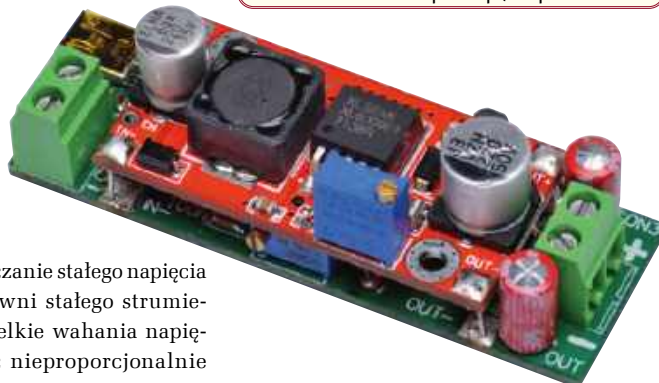
Liczne moduły przetwornic DC/DC są dostępne zarówno online, jak i w sklepach takich jak Jaycar i Altronics. Występują w dwóch głównych typach, boost i buck, choć niektóre łączą obie możliwości.

Typy buck redukują napięcie wejściowe do niższego poziomu. W przeciwieństwie do nich, konstrukcje typu buck/boost, takie jak Altronics Z6337 (patrz zdjęcie obok), zawierają dwa układy scalone sterownika (jak również powielają wiele innych elementów) i mogą zarówno zmniejszać, jak i zwiększać napięcie wejściowe.

Tego typu moduły są w rzeczywistości połączeniem modułów boost i buck. W tym projekcie celowo używamy jednak



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/9v1dr-20>
Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania



modułów typu boost specjalnie przeznaczonych do takich celów.

Gdyby czytelnik miał kłopot z nabyciem odpowiedniego modułu, w Silicon Chip Shop będzie dostępny właściwy moduł boost, który przetestowaliśmy pod kątem działania. Ten sam moduł zawarty jest w naszym zestawie. Jest to szczególnie ważne, biorąc pod uwagę, że istnieje wiele różnych konstrukcji modułów XL6009, z których nie wszystkie działają tak samo.

Moduły te mają niewielką płytkę drukowaną, która zawiera przetwornicę boost, minimalną liczbę elementów pasywnych oraz potencjometr montażowy do ustawiania napięcia wyjściowego. Połączenia wejściowe i wyjściowe są wykonane jako pola lutownicze.

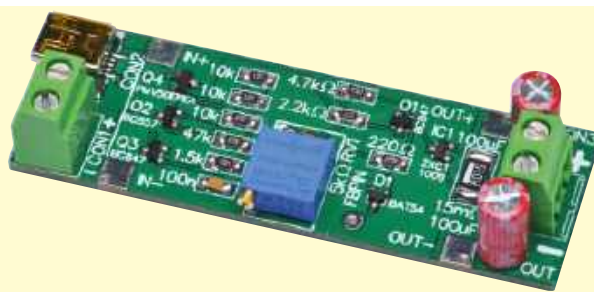
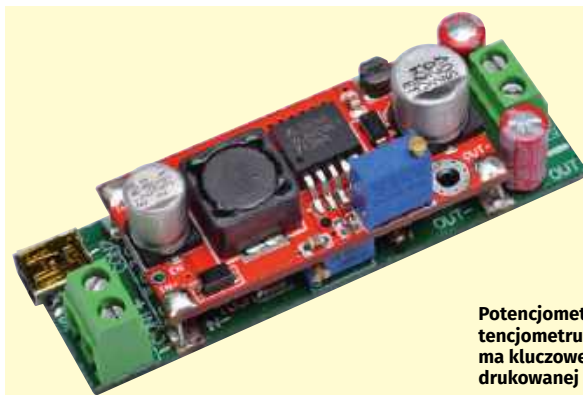
W niektórych poprzednich projektach używaliśmy modułów opartych na układzie scalonym MT3608. W takich przypadkach moduł jest lutowany bezpośrednio do innej płytki drukowanej i traktowany tak, jakby był tylko kolejnym elementem, podobnie jak Mini LED Driver.

Na przykład, miernik poziomu wody w zbiorniku z Wi-Fi z lutego 2018 r. (siliconchip.au/Article/10963) wykorzystywał taki moduł do dostarczania napięcia 24 V DC do czujnika głębokości wody z nominalnego zasilania 5 V. Nawiasem mówiąc, to zasilanie 5 V było dostarczane przez inny moduł, który zarządzał zasilaniem z ogniwa słonecznego i akumulatora.

Programator do dekodów DCC oparty na Arduino (wydanie z października 2018 r. siliconchip.au/Article/11261) podobnie wykorzystywał taki moduł do uzyskiwania zasilania 12 V z zasilacza USB 5 V. W tym przypadku napięcie 12 V było potrzebne do prawidłowego zasilania i programowania dekodów DCC.

Cechy i specyfikacja

- Może zasilac diody LED 12 V lub panele LED z zasilacza 5 V DC (np. ładowarka USB lub power bank)
- Regulowany prąd i napięcie wyjściowe, do 1 A/20 V
- Niewielkie rozmiary i niskie koszty
- Wejście do 4 A/20 V, w zależności od możliwości modułu boost



Potencjometr montażowy modułu boost służy do zmiany napięcia, natomiast śruba regulacyjna potencjometru prądowego jest widoczna poniżej. Połączenie widoczne poniżej górnego potencjometru ma kluczowe znaczenie dla działania układu Mini LED Driver. Prowadzi ono do punktu na płytce drukowanej modułu boost, który łączy się z pinem sprzężenia zwrotnego układu scalonego XL6009

bramkę Q4 w stanie wysokim, gdy Q3 jest wyłączony, więc jest on również wyłączany w razie potrzeby.

Zastosowano dwa rezystory, ponieważ bramka Q4 musi być podciągnięta o ponad 1 V poniżej napięcia zasilania, aby ją włączyć, podczas gdy złącze baza-emiter Q2 ogranicza jego napięcie bazowe do około 0,6 V poniżej przychodzącego zasilania.

Rezystor 47 kΩ pomiędzy kolektorem Q2 i bazą Q3 zapewnia pewną histerezę dla tego komparatora napięcia. Gdy Q3 włącza się, Q2 dostarcza niewielki dodatkowy prąd polaryzujący do punktu połączenia rezystorów dzielnika napięcia 10 kΩ/1,5 kΩ. Oznacza to, że napięcie wejściowe musi spaść do około 3,9 V przed wyłączeniem Q2, Q3 i Q4. Zmniejsza to ryzyko oscylacji wyłącznika niskiego napięcia, gdy napięcie wejściowe jest bliskie punktu odcięcia. Kondensator 100 nF równolegle z rezystorem 1,5 kΩ również zapobiega oscylacjom poprzez dalsze spowolnienie reakcji.

Domyślnie rezystory zostały dobrane tak, aby zapewnić prawidłowe działanie z nominalnym zasilaniem USB 5 V i chronić przed takimi zjawiskami, jak spadek napięcia zasilania na USB.

Chociaż Mini LED Driver nie został wyraźnie zaprojektowany do tego celu, może pracować z wyższymi napięciami. O niektórych

zastrzeżeniach i ograniczeniach wspomnimy później.

Nawiasem mówiąc, maksymalny limit 20 V w tym projekcie wynika z maksymalnego napięcia bramki-źródła MOSFET-a Q4, ale Q4 ogranicza również prąd podawany do modułu boost do 4 A, ponieważ jego limit prądu drenu wynosi 4,2 A. Mimo to, moduł XL6009 i tak osiąga swoje maksimum przy prądzie około 4 A, więc użycie mocniejszego MOSFET-a niewiele by dało.

Nie dodaliśmy żadnego ograniczenia prądu wejściowego, ponieważ większość zasilaczy USB zablokuje się przed dostarczeniem prądu 4 A.

Ograniczenie prądu

Układ scalony XL6009 zastosowany w module boost steruje napięciem wyjściowym, porównując wewnętrzne napięcie odniesienia z częścią napięcia wyjściowego i dostosowując swoje działanie tak, aby spróbować utrzymać je na tym samym poziomie. Potencjometr montażowy na module wzmacniającym jest częścią rezystancyjnego dzielnika napięcia używanego do próbkowania odpowiedniej części napięcia wyjściowego. Potencjometrem tym można więc ustawić napięcie wyjściowe.

Zapewniamy ograniczenie prądu poprzez wstrzykiwanie prądu do tego dzielnika napięcia, sprawiając, że układ przetwornicy traktuje napięcie wyjściowe jako wyższe niż jest w rzeczywistości, a to powoduje zmniejszenie jego mocy wyjściowej.

Rezystor bocznikowy 15 mΩ jest dołączony między wyjściem modułu boost (OUT+) a gniazdem wyjściowym CON3. Napięcie na tym rezystorze jest proporcjonalne do prądu pobieranego przez obciążenie z łączówki CON3. Układ scalony monitora bocznika ZXCT1009 (IC1) wzmacnia tę różnicę napięć i przekształca ją w prąd, który płynie z jego wyjścia przez nóżkę 3. Prąd ten wynosi 10 mA na każdy 1 V na boczniku.

Należy zauważyć, że rezystor bocznikowy 15 mΩ zmniejsza napięcie przyłożone

do obciążenia, ale ponieważ jego wartość jest niska, różnica wynosi tylko kilka miliwoltów (15 mV dla 1 A), więc jest to wartość pomijalna.

Ponieważ prąd obciążenia 1 A będzie wytwarzał 15-miliwoltowy spadek napięcia na rezystorze bocznikującym 15 mΩ, spowoduje to przepływ prądu 150 μA z nóżki 3 IC1 ($10 \text{ mA} \times 15 \text{ mV} \div 1 \text{ V}$). W rezultacie układ IC1 tworzy prąd o natężeniu 1/6667 (lub, jeśli wolisz, 3/20000) prądu wyjściowego.

Prąd ten jest doprowadzany do pinu FB (sprzężenia zwrotnego) dołączonego modułu boost poprzez rezystor 4,7 kΩ, potencjometr montażowy VR1 i diodę Schottky'ego D1. Prąd ten będzie miał tendencję do zmniejszania napięcia wyjściowego proporcjonalnie do jego natężenia, ale nie jest to główny czynnik w układzie ograniczającym prąd. Należy również wziąć pod uwagę tranzystor NPN Q1.

Baza i emiter tranzystora Q1 (z rezystorem emiterowym 220 Ω zmniejszającym jego wzmocnienie) są połączone przez rezystor 4,7 kΩ i VR1. Jeśli na tych dwóch elementach wystąpi napięcie wyższe niż 0,6 V, tranzystor Q1 zacznie przewodzić.

Działanie to stanowi większość funkcji ograniczania prądu, przy czym dodatkowy prąd jest doprowadzany do pinu FB przez kolektor i emiter Q1. Rezystor kolektorowy 2,2 kΩ ogranicza maksymalny prąd, który może zostać wstrzyknięty, pomagając utrzymać stabilność układu.

Ponieważ napięcie między bazą a emitrem Q1 zależy zarówno od prądu obciążenia, jak i ustawienia potencjometru VR1, to zgodnie z prawem Ohma oznacza to, że potencjometrem VR1 można ustawiać prąd obciążenia, przy którym Q1 zacznie przewodzić. Będzie to maksymalny prąd, jaki może dostarczyć całe urządzenie.

Zauważ, że jeśli użyjesz napięcia zasilania innego niż 5 V, ograniczenie prądu zmieni się ze względu na rezystor kolektora Q1 łączący się z napięciem zasilającym. Jednak większość źródeł 5 V DC jest regulowana,



Moduły z chipem XL6009 występują w kilku różnych wersjach. Widoczny tu, naszym zdaniem działa najlepiej i jest całkiem niedrogi. Będzie on również dostarczany jako część kompletnego zestawu dla płytki sterownika

więc generalnie nie ma to znaczenia. Należy o tym pamiętać, jeśli zamierzasz zasilac ten obwód bezpośrednio z akumulatora.

Na końcu znajdują się dwa kondensatory podłączone do wyjścia. Użyliśmy tutaj dwóch mniejszych elementów, ponieważ lepiej pasują do obrysu mini sterownika LED. Wygładzają one napięcie na rezystorze bocznikowym, które bez tych kondensatorów byłoby dość niestabilne ze względu na kondensatory poprzedzające w module boost.

Z tego powodu Mini LED Driver nie nadaje się dobrze jako regulowane prądowo źródło dla dynamicznych obciążeń, ponieważ kondensatory umożliwiają tylko na powolną reakcję. Gdyby rezystancja obciążenia nagle się zmieniła, kondensatory musiałyby się naładować lub rozładować, zanim system mógłby osiągnąć nowy stan ustalony. W tym czasie prąd płynący przez bocznik nie odzwierciedlałby tego, co dzieje się za łączówką CON3.

Na szczęście diody LED stanowią wolno zmieniające się obciążenie. Mini sterownik LED musi po prostu poradzić sobie ze zmianami, które występują, gdy napięcie przewodzenia LED zmienia się wraz z powoli zmieniającymi się zmiennymi, takimi jak temperatura.

Należy pamiętać, że przedstawiony schemat ograniczania prądu nie jest skuteczny jako zabezpieczenie przed zwarcie, ponieważ moduł boost nie może obniżyć swojego napięcia wyjściowego poniżej napięcia wejściowego (z wyjątkiem niewielkiego spadku spowodowanego wbudowaną diodą).

Zasadniczo, Mini LED Driver nie może ograniczyć swojego napięcia wyjściowego do poziomu znacznie poniżej napięcia wejściowego, a już na pewno nie do poziomu bliskiego zeru.

Regulacja prądu

Potencjometr VR1 jest podłączony w taki sposób, że w pozycji maksymalnego obrócenia ośki w prawo, między dwoma podłączonymi zaciskami występuje rezystancja 0 Ω . Pozycja zgodna z ruchem wskazówek zegara ustawia więc rezystancję 4,7 k Ω między wyjściem I_{out} IC1 a diodą, natomiast pozycja całkowicie przeciwna do ruchu wskazówek zegara ustawia rezystancję 9,7 Ω .

Zakładając próg około 0,6 V dla krzemowego złącza baza-emiter, Q1 zacznie przewodzić prąd o natężeniu 127 μ A z IC1, gdy VR1 jest skrócony całkowicie w prawo, i 62 μ A z IC1, gdy jest skrócony całkowicie w lewo.

Oznacza to, że użyteczny zakres ustawień prądu wyjściowego wynosi nominalnie od 0,85 A do 0,41 A (przywołując

współczynnik 6667 z poprzedniej wersji), choć nie są to sztywne limity.

Podczas jednego z naszych testów zaczęliśmy od ustawienia napięcia Mini LED Driver na 12 V bez obciążenia i VR1 ustawionym na minimum. Następnie podłączyliśmy jeden z dużych paneli LED o mocy 70 W i zmierziliśmy prąd panelu 0,48 A przy napięciu 11,1 V.

Ustawienie ograniczenia prądu na maksimum dało 0,84 A przy 11,3 V, ale prąd można było zwiększyć do 1 A, zwiększając wartość zadaną napięcia (bez obciążenia) do około 12,6 V. Zmierzyliśmy blisko 3 A na wejściu 5 V, więc nie spodziewamy się, że wiele zasilaczy USB będzie działać z takimi prądami.

Faktem jest, że ograniczenie prądu włącza się stopniowo, co jest niezbędne do utrzymania stabilności sterownika. Oznacza to również, że punkt pracy diody LED można dostosować poprzez ostrożną regulację zarówno ustawień prądu, jak i napięcia.

Na **rysunku 2** przedstawiono wpływ zmiany obciążenia na Mini LED Driver. Wykonaliśmy te wykresy przy napięciu bez obciążenia ustawionym na 12 V i potencjometrze ograniczającym prąd ustawionym w najniższej i najwyższej pozycji, plus trzeci punkt w pobliżu środka.

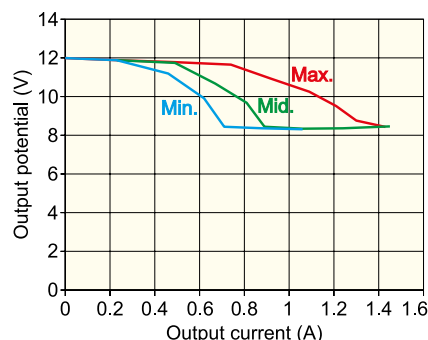
Istnieje granica tego, jak niskie napięcie może zostać osiągnięte przez obwód ograniczający prąd. W tym przypadku jest to około 8,3 V. Wynika to z faktu, że rezystor 2,2 k Ω ogranicza prąd wstrzykiwany do dzielnika napięcia.

Inne moduły boost, które używają różnych rezystorów dzielących do ustawienia napięcia, będą zachowywać się inaczej, ponieważ wstrzyknięty prąd zmieni wartość zadaną o inną wartość. Jest to jeden z powodów, dla których określamy i dostarczamy konkretny moduł, jak pokazano na zdjęciach obok. Jest to ten, który działał najlepiej w naszych testach.

Jeśli musisz wypróbować inny moduł boost, zalecamy dokładne przetestowanie kombinacji przed użyciem. Do większości naszych testów użyliśmy programowalnego obciążenia bazującego na Arduino z opisanego w numerze Silicon Chip z czerwca 2022 r. (siliconchip.au/Article/15341), oraz w niniejszym wydaniu EdW, w tym do wykreślenia rysunku 2.

Efektywność

Zmierzyliśmy również sprawność modułu i stwierdziliśmy, że nie osiągnęła wartości 96% deklarowanej przez dostawców wielu takich modułów. Zazwyczaj podają oni sprawność dla podnoszenia napięcia z 12 V do 20 V. Zwiększenie napięcia



Rysunek 2. Przedstawione krzywe pokazują zachowanie modułu Mini LED Driver przy nominalnym napięciu 12 V i trzech różnych ustawieniach ograniczenia prądu. Krzywe odpowiadają VR1 przy minimum (cyjan/niebieski), maksimum (czerwony) i mniej więcej w połowie ustawienia między dwoma pozycjami skrajnymi zielony)

z 5 V do 12 V to zarówno wyższy współczynnik, jak i rozpoczęcie od niższego napięcia, więc sprawność nie będzie najwyższa.

Przy regulowanym napięciu wejściowym 5 V i 12 V na wyjściu, pomocną zasadą jest pomnożenie prądu wyjściowego przez trzy. Operacja ta pozwala obliczyć teoretyczny prąd wejściowy. Odpowiada to przybliżonej sprawności na poziomie 80%.

Opcje

Możesz zdecydować się na pominięcie CON1 lub CON2, jeśli wiesz, że na pewno użyjesz tylko jednego z nich, ale wyjaśnimy procedurę budowy tak, jakbyś zamontował oba.

Należy pamiętać, że Mini LED Driver będzie pobierał znaczny prąd przy zasilaniu 5 V. Każdy znaczący spadek napięcia wejściowego może spowodować zadziałanie wyłącznika niskiego napięcia.

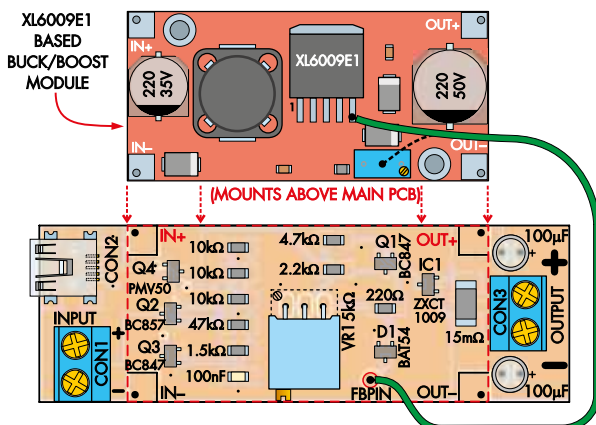
Złącze USB będzie miało zauważalnie wyższą rezystancję niż zaciski śrubowe. Dlatego zalecamy zamontowanie obu na wypadek, gdyby rezystancja okazała się zbyt wysoka i konieczne byłoby użycie zacisku śrubowego zamiast złącza USB.

Budowa

Płytkę modułu nie jest trudna w montażu, ale użyte są prawie wyłącznie elementy do montażu powierzchniowego. Upewnij się więc, że masz niezbędne narzędzia i materiały, w tym lutownicę, topnik, plecionkę lutowniczą, lutownicę z cienkim grotem, pęsetę, przyzwoite oświetlenie i lupe.

Więcej porad i wskazówek dotyczących lutowania SMD można znaleźć w naszym artykule na ten temat (Silicon Chip, grudzień 2021, siliconchip.au/Article/15138).

Płytkę drukowaną jest dobrze oznaczona, ale można również odnieść się do schematu montażowego (**rysunek 3**), który pozwoli sprawdzić dokładne lokalizacje poszczególnych



Rysunek 3. Krytyczną czynnością podczas montażu sterownika jest upewnienie się, że nie dojdzie do pomieszania różnych elementów w obudowach SOT-23. Przed ich przylutowaniem należy sprawdzić oznaczenia na płytce drukowanej. Moduł boost znajduje się nad górną częścią tej płytki drukowanej, co widać na innych zdjęciach. Chociaż sprzężenie zwrotne łączy się elektrycznie z pinem 5 układu scalonego XL6009, zwykle łatwiej jest dotaczyć się do wyprowadzenia potencjometru montażowego po sprawdzeniu niskiej rezystancji między nim a pinem sprzężenia zwrotnego układu scalonego



elementów. Płytkę jest oznaczona kodem 16106221 i ma wymiary 72 mm × 24 mm.

Zacznij od CON2, złącza mini-USB. Nałóż topnik na pady i umieść złącze na miejscu. Posiada ono piny pozycjonujące, więc powinno zablockować się we właściwej pozycji.

Wyczyść grot lutownicy i nałóż niewielką ilość świeżego lutu. Przyłóż lutownicę do dwóch przedłużonych końcówek w rzędzie pięciu – tylko te dwa są potrzebne do zasilania. Jeśli połączysz je z innymi pinami, przed dalszymi pracami użyj plecionki lutowniczej do usunięcia nadmiaru lutu.

Następnie nałóż dużą ilość lutu, aby zabezpieczyć cztery narożne wyprowadzenia na obudowie, co zapewni mechaniczne trzymanie złącza.

Teraz należy sprawdzić tranzystory, diody i układy scalone. Wszystkie są w identycznie wyglądających obudowach SOT-23, ale istnieje pięć różnych typów, więc należy uważać, aby ich nie pomylić. Płytkę drukowana jest oznaczona numerami części oraz oznaczeniami.

Sprawdź typy w odniesieniu do opisów na PCB, pracując jednocześnie tylko z jednym typem. Części SOT-23 są małe, ale wyprowadzenia mają dość duże odstępy, więc praca z nimi nie jest trudna.

Nałóż topnik na pady tych elementów, a następnie użyj pęsety do wstępnego umieszczenia każdego elementu po kolei. Przymocuj jedno wyprowadzenie i sprawdź, czy pozostałe nóżki znajdują się na swoich polach. Jeśli nie, dostosuj je w razie potrzeby za pomocą lutownicy i pęsety. Następnie przylutuj pozostałe wyprowadzenia.

Zrób to samo z ośmioma małymi (3,2 mm × 1,6 mm) rezystorami, sprawdzając ich wartości na sitodruku płytki. W przypadku tych elementów stosuje się tę samą technikę, jak dla półprzewodników.

Następnie zamontuj większy (6,3 mm × 3,2 mm) rezystor bocznikujący prąd. Jest on trudniejszy do pomylenia z innymi

elementami ze względu na jego odmienny rozmiar. Pojedynczy kondensator SMD znajduje się obok rezystora 1,5 kΩ i można go przylutować w podobny sposób.

Na tym kończy się montaż komponentów SMD. Przed przystąpieniem do dalszych czynności należy oczyścić płytkę drukowaną z wszelkich pozostałości topnika i pozwolić jej dokładnie wyschnąć.

Możesz przetestować funkcję odcięcia niskiego napięcia, jeśli możesz podłączyć regulowany zasilacz do wejść CON1 lub CON2. Najlepiej zrobić to teraz, przed podłączeniem modułu boost, ponieważ łatwiej jest naprawić wszelkie wykryte problemy. Nie przekraczaj napięcia 20 V i pamiętaj o polaryzacji połączeń z CON1.

Zwiększaj i zmniejszaj napięcie wejściowe. Sprawdź, czy między punktami IN+ i IN- występuje napięcie, gdy napięcie na CON1 lub CON2 jest powyżej górnego progu (około 4,6 V). Gdy napięcie wejściowe jest poniżej dolnego progu (około 3,7 V), powinno ono spaść.

Finalizacja montażu

Pozostałe elementy do zamontowania to CON1, CON3, potencjometr VR1, dwa kondensatory elektrolityczne i wreszcie moduł boost. Najpierw przylutuj złącza CON1 i CON3. Powinny znajdować się na tyle daleko od siebie, aby moduł boost mógł zmieścić się między nimi.

Następnie zamontuj potencjometr VR1. Chociaż można go przylutować w standardowej pozycji pionowej, moduł boost będzie znajdował się nad płytką drukowaną sterownika, blokując regulację. Zamiast tego zamontuj go z boku, jak pokazano na naszych zdjęciach. Upewnij się, że śruba regulacyjna jest ustawiona prawidłowo. W ramach przygotowań do testów należy również wyregulować potencjometr montażowy do minimum (całkowite skrócenie w lewo).

W pobliżu CON3 znajdują się dwa kondensatory elektrolityczne. Dłuższe wyprowadzenia

dotądnie wchodzą w pady oznaczone małymi symbolami +. Przed przylutowaniem i przecięciem wyprowadzeń kondensatory należy mocno docisnąć do płytki drukowanej.

Ostrzeżenie przed zamontowaniem modułu boost – widzieliśmy kilka modułów boost, które (myląc) zwiększają napięcie, gdy potencjometr jest regulowany w kierunku przeciwnym do ruchu wskazówek zegara.

Jeśli używasz innego modułu niż ten, który dostarczamy, sprawdź przed przylutowaniem go do płytki sterownika jego napięcie, zasilając go, a następnie mierząc jego napięcie wyjściowe multimetrem. W przeciwnym razie można ugotować dwa kondensatory wyjściowe przy pierwszym włączeniu zasilania.

Po upewnieniu się, że wszystko jest skonfigurowane poprawnie przylutuj krótki odcinek przewodu do pinu sprzężenia zwrotnego na odwrocie potencjometru regulacji napięcia na spodzie modułu boost. Ma on być podłączony do środkowego pinu, aby wyrównać z drugą płytką drukowaną, ale może się okazać, że dwa piny potencjometru modułu boost i tak są ze sobą połączone.

Możesz zobaczyć, jak wyglądają te połączenia na naszych zdjęciach. W modułach XL6009, których używamy, powinno to być zgodne bezpośrednio z padem FBPIN na płytce drukowanej sterownika, ale może być w innym miejscu, jeśli używasz innego modułu.

Może się okazać, że konieczne będzie odciecie wyprowadzenia elementu, ale jeśli nie możesz poprowadzić wyprowadzenia bezpośrednio, użyj zamiast tego krótkiego odcinka cienkiego izolowanego drutu (np. kynar).

Teraz przylutuj końcówki wyprowadzeń elementów lub krótkie odcinki sztywnego przewodu do czterech rogów modułu boost na padach IN+, IN-, OUT+ i OUT-. Wszystkie powinny być skierowane w dół w taki sam sposób, jak wyprowadzenie do FBPIN. Do chwycenia wyprowadzenia podczas lutowania przydała nam się

Wykaz elementów

- 1 dwustronna płytka drukowana o wymiarach 72 mm × 24 mm, kod 16106221
- 1 moduł DC-DC boost oparty na sterowniku XL6009 z czerwoną płytką drukowaną (MOD1, patrz tekst) [SC6546]
- 1 2-pinowe zaciski śrubowe 5,08 mm (CON1) I/LUB
- 1 gniazdo mini-USB (CON2)
- 1 2-pinowe zaciski śrubowe 5,08 mm (CON3)
- 5 odcinków drutu o średnicy 1 mm i długości 20 mm lub ścinki wyprowadzeń (patrz tekst)

Półprzewodniki:

- 1 ZXCT1009 monitor bocznika prądowego high-side, SOT-23 (IC1)
- 1 dioda Schottky'ego BAT54 (lub BAT54C lub BAT54S), SOT-23 (D1)
- 2 tranzystory bipolarnie BC847B NPN, SOT-23 (Q1, Q3)
- 1 tranzystor bipolarny PNP BC857B, SOT-23 (Q2)
- 1 tranzystor MOSFET z kanałem P PMV50EPEA lub AO3407, SOT-23 (Q4)

Kondensatory:

- 2 100 µF 25 V elektrolityczny
- 1 100 nF 50 V X7R M3216/1206 SMD ceramiczny

Rezystory: (wszystkie 1206/M3216 1/8 W, chyba że określono inaczej)

- 1 47 kΩ 3 10 kΩ 1 4,7 kΩ 1 2,2 Ω
- 1 1,5 Ω 1 220 Ω
- 1 5 kΩ górny pokrętko wieloobrotowe (VR1)
- 1 15 mΩ 2512/M6432 3 W rezystor bocznikujący [SC3943]

Zestaw (SC6405 – 25 USD): zawiera płytkę drukowaną i wszystkie niezbędne elementy, w tym moduł XL6009

pęseta lub szczypce. Dzięki temu uniknęliśmy poparzenia palców.

Teraz możesz połączyć dwie płytki razem z modulem boost nad płytką drukowaną sterownika, upewniając się, że etykiety padów pasują do siebie. Jeśli to możliwe, pozostaw trochę wolnej przestrzeni między dwiema płytkami drukowanymi i przylutuj jeden przewód na miejscu.

Dopasuj płytki, pilnując, by nic nie stykało się tam, gdzie nie powinno, sprawdź, czy elementy są prostopadłe i równoległe, a następnie przylutuj cztery narożne pady, a w dalszej kolejności wyprowadzenie do FBPIN. Przytnij wszystkie wyprowadzenia, które są dłuższe niż to konieczne.

Testowanie

Podczas testowania należy pamiętać, że Mini LED Driver nie jest odporny na zwarcia. Należy więc uważać na podłączone obciążenia, i dopilnować, by nie było szans na zwarcie lub bardzo niską rezystancję, która mogłaby przeciążyć MOSFET Q4.

Jak wspomnieliśmy wcześniej, programowalne obciążenie oparte na Arduino nada się do testów, ale można też użyć rezystora (np. 22 Ω 10 W lub dwóch rezystorów 10 Ω 5 W połączonych szeregowo) lub białej diody LED o dużej mocy. Poniższy schemat zakłada zasilanie 5 V i może nie działać, jeśli masz znacznie wyższe zasilanie.

Podłącz zasilanie bez obciążenia i wyreguluj napięcie wyjściowe na CON2 do 12 V za pomocą potencjometru montażowego w module



Wyraźnie widać przewód od nóżki FBPIN do potencjometru powyżej

XL6009. Jeśli nie możesz płynnie wyregulować napięcia na wyjściu, przed kontynuowaniem prac sprawdź zespół sterownika. Pamiętaj, aby nie ustawiać napięcia wyjściowego powyżej 20 V!

Przy VR1 w sterowniku ustawionym w pozycji minimalnej, obciążenie 20 Ω lub o niższej rezystancji powinno pobierać około 0,6 A i powodować ograniczenie prądu na wyjściu. Odnosząc się do rysunku 2, sprawdź, czy Twoje urządzenie reaguje podobnie do naszego prototypu.

Jeśli napięcie lub prąd wyjściowy wydaje się spadać bardziej, sprawdź, czy zasilacz USB działa w granicach swoich ograniczeń. Może on mieć własne wewnętrzne ograniczenie prądu. Jeśli napięcie na CON1 odbiega znacznie od 5 V, jest to znak, że używany zasilacz nie radzi sobie z obciążeniem.

Dobrym pomysłem jest sprawdzenie napięcia zasilającego moduł boost na zaciskach IN+ i IN-. Jeśli jest ono znacznie niższe niż napięcie na CON1, działa wyłącznik niskiego napięcia. Może to być spowodowane spadkami napięcia w kablu lub spadkiem napięcia zasilania USB pod obciążeniem.

Wyreguluj VR1 i sprawdź, czy ograniczenie prądu uległo zmianie. Konieczne może być zwiększenie obciążenia (zmniejszenie rezystancji) poprzez dodanie dodatkowych rezystorów równoległych. Ustaw prąd na żadaną wartość i podłącz żądane obciążenie. Następnie potwierdź, że układ działa zgodnie z oczekiwaniami.

Korzystanie z innych modułów boost

Nie zalecamy tego, chyba że masz doświadczenie. Znalazienie pinu sprzężenia zwrotnego (FB) może być trudne, jeśli moduł boost nie jest oznaczony. Dobrym miejscem do rozpoczęcia jest środkowy pin potencjometru regulacyjnego, chociaż widzieliśmy kilka modułów, które nie podążają za tym trendem.

Pin FB jest wyprowadzony na układ scalony XL6009, a większość sterowników układów boost powinna mieć zewnętrzny pin sprzężenia zwrotnego, więc warto zacząć od niego. W XL6009 jest to najbardziej

wysunięty na prawo z małych pinów (pin 5). W modułach, które wypróbowaliśmy, jest to mniejszy pin najbliższy potencjometrowi regulacji napięcia.

Można przylutować przewód bezpośrednio do tego pinu, choć nie będzie to tak estetyczne, jak podłączenie do zacisku potencjometru regulacyjnego. Zamiast tego można użyć multimetru ustawionego na tryb pomiaru ciągłości połączenia i znaleźć inne bardziej dostępne (np. jakąś przelotkę) połączenie lutowane o rezystancji bliskiej zera z pinem sprzężenia zwrotnego. W razie wątpliwości należy zajrzeć do noty katalogowej układu sterownika przetwornicy w zastosowanym module.

Zasady korzystania z Mini sterownika LED

Wypróbowaliśmy Mini LED Driver z jednym z dużych 70 W paneli LED, których używaliśmy w czerwcu z Buck-Boost LED Driver (dostępny w sklepie internetowym Silicon Chip, nr kat. SC6307 lub SC6308). Podłączyliśmy panel LED po ustawieniu napięcia wyjściowego na 12 V bez obciążenia i ustawieniu ograniczenia prądu na minimum.

Zasiłał on panel prądem 480 mA, a napięcie wyjściowe wynosiło 11,1 V. Powolne zwiększanie limitu prądu zwiększyło prąd i jasność panelu. Aby potwierdzić, że prąd jest odpowiednio regulowany, odłącz panel LED i sprawdź, czy napięcie wyjściowe wzrasta o co najmniej pół wolta. Oznacza to, że Mini LED Driver ma rezerwę regulacji prądu.

Zauważyliśmy, że panel ściemniał się i czasami migotał po ustawieniu prądu powyżej pewnego punktu, co oznaczało, że zasilacz USB osiągnął swój limit.

Innym objawem przeciążenia jest wysoki dźwięk z modułu boost pod obciążeniem. W takim przypadku, aby zapobiec uszkodzeniu zasilacza USB należy zmniejszyć ograniczenie prądu. ■

Tim Blythman

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

Programowalne obciążenie bazujące na Arduino

Do testowania urządzeń, takich jak zasilacze, obwody sterowników i źródła prądu, często potrzebna jest stała lub zmienna rezystancja obciążenia, na której będzie mogła wydzielić się określona moc. Opisanie w artykule Programowalne Obciążenie jest oparte na łatwym w zrozumieniu i obsłudze module Arduino. Może być sterowane ręcznie lub automatycznie w sposób właściwy dla danego zastosowania.

Podczas projektowania i testowania naszego sterownika LED High Power Buck-Boost chcieliśmy sprawdzić, jak radzi sobie z różnymi obciążeniami, aby przetestować poprawność i wszechstronność projektu. W tym celu opracowaliśmy opisanie dalej urządzenie, które okazało się tak przydatne, że powstał samodzielny projekt.

W przeciwieństwie do innego projektu – 50 W DC Electronic Load (wrzesień 2002; siliconchip.com.au/Article/4029) – prezentowane tutaj Programowalne Obciążenie nie jest regulowane bezstopniowo i nie jest przeznaczone do pobierania stałego prądu. Zamiast tego wykorzystuje przełączane elementy rezystancyjne, które powodują, że użytkownik może korzystać z dyskretnych stopni obciążenia.

Podłączenie do mikrokontrolera Arduino oznacza jednak, że możliwe jest dodanie kilku inteligentnych funkcji. Układ zawiera również elementy umożliwiające pomiar przyłożonego napięcia i przepływającego przez Obciążenie prądu. Oznacza to, że może również obliczyć moc rozpraszaną w Obciążeniu ($P=U \cdot I$).

W ten sposób można zaprogramować Obciążenie tak, aby zachowywało się różnie, w zależności od zastosowania. Jego funkcje obejmują tryby stałej rezystancji lub śledzenia prądu. Można je nawet zaprogramować tak, aby zapewniało obciążenie dynamiczne, dzięki czemu można testować sprzęt w zmieniających się warunkach.

Typowym testem zasilacza lub stabilizatora jest sprawdzenie, jak reaguje on na nagłe zmiany rezystancji obciążenia.

Nasz przykładowy kod zapewnia tylko podstawowe funkcje, w tym ręczne tryby śledzenia rezystancji i prądu, ale łatwo jest zmodyfikować kod tak, aby dodać niestandardowe funkcje. Nasz przykładowy kod wyświetla również wszystkie zebrane dane.



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/kwxnmz84>
Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania

Szczegóły schematu

W Obciążeniu Elektronicznym DC o mocy 50 W z 2002 roku był zastosowany pojedynczy MOSFET przykręcony do dużego radiatora jako element obciążenia. Wymaga to starannego zaprojektowania układu, który powinien zapewniać reakcję obciążenia na dynamiczne warunki.

Z drugiej strony, nasze Programowalne Obciążenie składa się z 15 rezystorów o dużej mocy, które nie mają problemów z radzeniem sobie z szybko zmieniającymi się warunkami. Co najważniejsze, o ile układ działa w zakresie napięcia roboczego nie ma szans na zwarcie.

Koncepcja jest prosta. Istnieją cztery grupy rezystorów o mocy 5 W i rezystancji 47 Ω. Grupy składają się odpowiednio z jednego, dwóch, czterech i ośmiu rezystorów,

które mogą być przełączane w dowolną kombinację od braku do 15 rezystorów równolegle.

Obciążenie jest zoptymalizowane do użytku ze źródłami napięcia do 12 V nominalnie. Pamiętaliśmy jednak, że mogą występować pewne wahania napięcia. Na przykład na akumulatorze 12 V podczas ładowania może występować napięcie do 14,4 V, a dioda LED 12 V może do wydzielenia pełnej mocy wymagać napięcia 13 V lub wyższego. Dlatego wybraliśmy elementy, które będą obsługiwać do 15 V w sposób ciągły (więcej w sposób pulsacyjny lub przerywany).

47 Ω to najniższa wartość rezystora z serii E24, na której moc wydzielana przy napięciu 15 V wynosi mniej niż 5 W. Stąd nasze zastosowanie rezystorów 47 Ω.

Na **rysunku 1** przedstawiono opracowany przez nas schemat. Cztery N-kanałowe



Cechy i specyfikacja



- Obsługuje do 70 W mocy ciągłej, przy napięciu do 15 V i natężeniu 4,7 A
- Prezentuje rezystancję obciążenia w zakresie od 3,1 Ω do 47 Ω w 15 krokach lub 43 kΩ, gdy jest „wyłączony”
- Pobiera od 255 mA do 3,83 A w krokach co 255 mA z idealnie stabilizowanego źródła 12 V
- Ręczne sterowanie obciążeniem lub rezystancją
- Oprogramowanie zapewnia tryb w przybliżeniu stałoprądowy
- Mierzy napięcie do 20 V
- Mierzy prąd do 6,5 A
- Oblicza moc do 130 W

MOSFET-y, Q1...Q4, podłączają lub odłączają do badanego źródła zasilania podległe im rezystory mocy. Źródła MOSFET-ów są podłączone do masy obwodu, a dreny trafiają odpowiednio do grup jednego, dwóch, czterech lub ośmiu rezystorów. Ich bramki są utrzymywane w stanie niskim przez rezystory 10 k Ω , więc zwykle są wyłączone.

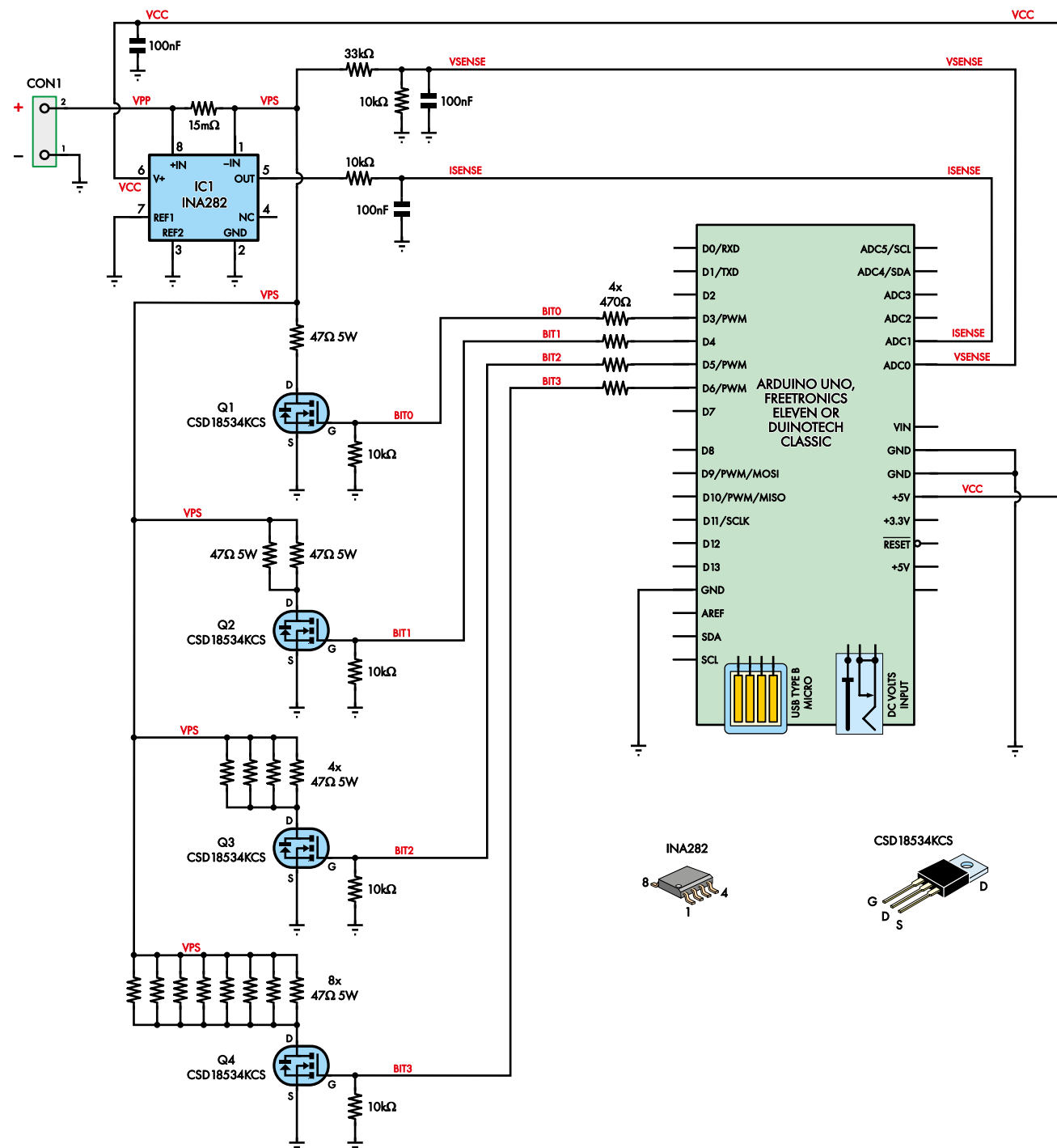
Bramki łączą się również z czterema cyfrowymi pinami I/O (D3, D4, D5 i D6)

podłączonej płytki Arduino poprzez rezystory 470 Ω . Rezystory zapewniają pewien stopień ochrony w przypadku katastrofalnej awarii. Obwody są oddzielone, ale łączy je wspólna masa.

Drugi koniec rezystorów Obciążenia łączy się z bocznikiem pomiarowym prądu 15 m Ω , a następnie z dodatnim zaciskiem Obciążenia. Do połączeń z układami zewnętrznymi zastosowano zaciski śrubowe CON1.

Do górnej części rezystorów obciążenia podłączony jest również dzielnik 33 k Ω /10 k Ω z kondensatorem 100 nF na dolnym rezystorze. Dzięki niemu Arduino może mierzyć napięcia do 21,5 V. Jest tak przy założeniu, że napięcie referencyjne przetwornika analogowo-cyfrowego (ADC) na płytce Arduino jest równe 5 V.

Podzielone i wygładzone napięcie jest podawane na wejście analogowe A0 dołączonej



Rysunek 1. Do przełączania max. piętnastu rezystorów 47 Ω używane są cztery MOSFET-y, Q1...Q4, uzyskując zmienną rezystancję obciążenia na CON1. IC1 i bocznik 15 m Ω umożliwiają pomiar prądu obciążenia, natomiast dzielnik 33 k Ω /10 k Ω mierzy napięcie, umożliwiając obliczenie mocy wydzielanej w obciążeniu

płytki Arduino. Dzielnik ten oznacza, że Programowalne Obciążenie zawsze prezentuje minimalne obciążenie 43 kΩ.

Napięcie na boczniku jest mierzone przez układ IC1. Jest to monitor bocznika prądowego INA282 o wzmacnieniu 50. Prąd o natężeniu 1 A powoduje 15-milivoltowy spadek napięcia na tym rezystorze 15 mΩ. Tym samym na pinie 5 układu IC1 występuje napięcie 750 mV. Maksymalny mierzalny prąd przy napięciu referencyjnym 5 V wynosi zatem 6,67 A.

Napięcie to trafia przez rezystor 10 kΩ do innego kanału ADC dostępnego na wejściu A1 Arduino i jest filtrowane przez kondensator 100 nF. Napięcie wyjściowe układu IC1 jest ustawione jako odniesienie do masy obwodu poprzez podłączenie jego pinów 3 i 7 do masy.

Układ IC1 jest zasilany napięciem 5 V z dołączonej płytki Arduino doprowadzonym do pinu 6 i filtrowanym za pomocą kondensatora 100 nF. Masa zasilania tego układu to nóżka 2 dołączona do GND.

Zmieniając stany pinów wyjściowych Arduino D3...D6 można zmieniać rezystancję obciążenia między wartościami odpowiadającymi 1, 2, 3, ..., 15 równolegle połączonych rezystancji 47 Ω. Można też całkowicie je odłączyć. Arduino monitoruje napięcie oraz prąd i zgłasza je wraz z obliczoną mocą wydzielaną na Obciążeniu.

W zależności od zaprogramowanego trybu, Obciążenie może zapewniać stałą rezystancję lub próbować naśladować stały prąd, a nawet dynamicznie zmieniające się obciążenie w celu sprawdzenia reakcji zasilania.

Wybór płytki Arduino

W wykazie elementów podaliśmy Arduino Uno, ale każda płytki Arduino 5 V, w tym inne płytki kompatybilne z R3 shield oparte na AVR, takie jak Leonardo lub Mega, powinny działać poprawnie. Przykładowy kod nie wykorzystuje żadnych urządzeń peryferyjnych specyficznych dla pinów, więc nie jest powiązany z konkretną płytką. Jednak, aby zapewnić pełne włączenie MOSFET-ów, niezbędne są poziomy 5 V cyfrowych wejść/wyjść.

Jeśli naprawdę chcesz użyć płytki 3,3 V, możesz to zrobić z pewnymi zmianami. Pamiętaj, że wiele z nich nie jest kompatybilnych z modułem R3 (zwykle używają zamiast tego modułu MKR). Jednym z wyjątków jest Due. Nie testowaliśmy tego projektu z płytką Arduino 3,3 V, ale uważamy, że będzie działać z następującymi zmianami.

Po pierwsze, upewnij się, że używasz MOSFET-ów IPP80N06S4L-07 lub podobnych, ponieważ CSD18534KCS nie nadają się do sterowania bramki 3,3 V.

Po drugie, zmień rezystor 33 kΩ na 56 kΩ i zmień bocznik 15 mΩ na 10 mΩ. Ma to na celu uniknięcie przeciążenia pinów ADC napięciami powyżej 3,3 V i zakłada domyślne napięcie odniesienia ADC 3,3 V (zgodnie z Due).

W szkicu zmień definicję V_CONST na 0.0212695, a definicję I_CONST na 0.0064453. Pozwoli to uwzględnić różne wartości elementów.

Budowa

Obciążenie jest prezentowane jako goła płytki drukowana z zewnętrznymi zaciskami śrubowymi. Oczekuje się, że będzie ono używane podobnie jak zasilacz Arduino (luty 2021; siliconchip.com.au/Article/14741), jako nakładka dla płytki prototypowej kompatybilnej z Arduino.

Brak obudowy w rzeczywistości nieco nam pomaga. Przy wydzielaniu mocy do 70 W, niezbędna jest duża przestrzeń dla swobodnej konwekcji powietrza pozwalająca uniknąć przegrzania. Idealnie byłoby, gdyby na moduł był skierowany wentylator, gdy jest on używany z maksymalną mocą znamionową lub do niej się zbliża.

Obciążenie jest zmontowane na dwustronnej płytce drukowanej z kodem 04105221 o wymiarach 89 mm × 54 mm. Na **rysunku 2** zostało przedstawione rozmieszczenie elementów na PCB.

Zacznij od montażu małych elementów. IC1 to układ SMD w obudowie SOIC-8. Najlepiej lutować go używając pasty topnikowej i pęsety, chociaż można sobie poradzić bez nich.

Nałóż topnik na pady i przyłutuj jedną nóżkę czystym grotiem lutownicy, upewniając się, że pin 1 znajduje się w tej samej pozycji, co kropka na płytce drukowanej. Jeśli układ jest prawidłowo wyśrodkowany względem padów, przyłutuj pozostałe piny. W przeciwnym razie skoryguj jego położenie za pomocą pęsety. W podobny sposób może być montowany również rezystor bocznikowy

15 mΩ sąsiadujący z CON1, chociaż lutuje się go znacznie prościej.

Oczyść teraz nadmiar topnika, ponieważ pozostałe elementy są przewlekane. Należy pamiętać, że płytki PCB w roli bocznika pozwala użyć zarówno rezystor SMD jak i przewlekany, możesz więc zastosować taki, który bardziej ci odpowiada. Jeśli jednak zmienisz jego wartość, będziesz musiał dostosować kalibrację w oprogramowaniu.

Następnie zamontuj pozostałe małe rezystory osiowe, zgodnie z oznaczeniami na sitodruku płytki drukowanej. Sprawdź rezystory multimetrem, jeśli nie masz pewności co do ich wartości.

Postępuj zgodnie z trzema kondensatorami 100 nF, z których wszystkie znajdują się w pobliżu IC1. Nie są one spolaryzowane. Przytnij wszystkie wyprowadzenia od spodu płytki drukowanej. Następnie można przyłutować złącze śrubowe CON1. Upewnij się, że wyprowadzenia są skierowane na zewnątrz płytki.

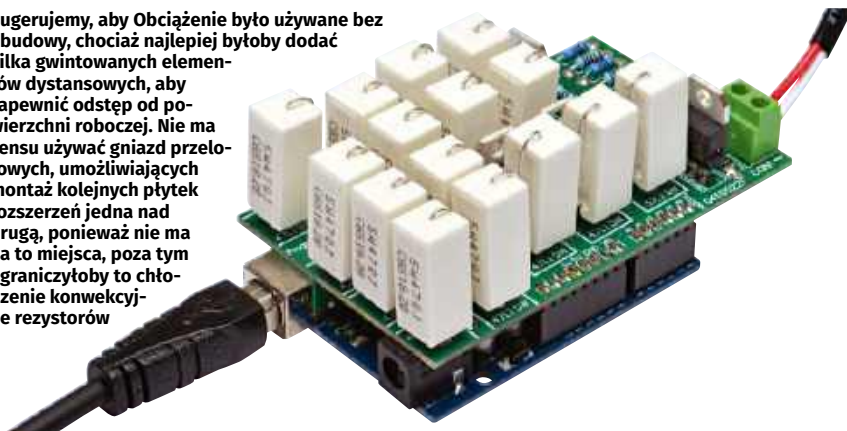
Kolejnymi najwyższymi elementami są MOSFET-y Q1...Q4. Wszystkie są tego samego typu. Upewnij się, że są one prawidłowo zorientowane, a wypustki są wyrównane z oznaczeniami na sitodruku. W celu potwierdzenia orientacji MOSFET-ów możesz również odnieść się do zdjęć i rysunku 2.

MOSFET-y należy zamontować wolnostojąco i pionowo. Nie powodują zbyt dużego spadku napięcia, gdy są włączone i nie przepływa przez nie zbyt duży prąd w stosunku do wartości znamionowych, więc nie potrzebują radiatora.

Przygotuj rezystory ceramiczne 5 W, zginając jedno wyprowadzenie o 180° w dół z jednej strony, tak aby można je było umieścić pionowo na płytce drukowanej. Zginanie wyprowadzenia w dół po stronie przeciwnej do oznaczeń daje najlepszy rezultat.

Podczas montażu rezystorów 5 W pomocne będzie również umieszczenie ich nieco powyżej płytki drukowanej. Zapewni to więcej miejsca na cyrkulację powietrza. Można to zobaczyć

Sugerujemy, aby Obciążenie było używane bez obudowy, chociaż najlepiej byłoby dodać kilka gwintowanych elementów dystansowych, aby zapewnić odstęp od powierzchni roboczej. Nie ma sensu używać gniazd przelotowych, umożliwiających montaż kolejnych płytek rozszerzeń jedna nad drugą, ponieważ nie ma na to miejsca, poza tym ograniczyłoby to chłodzenie konwekcyjne rezystorów



na naszych zdjęciach. Zrobiliśmy odstęp 3 mm, chociaż długość ich wyprowadzeń może stanowić pewne ograniczenie.

Zacznij od rezystorów w pobliżu środka płytki drukowanej i pracuj w kierunku do zewnątrz, starając się utrzymać wierzchołki na równym poziomie w celu zapewnienia jednolitości i wyśrodkowania elementów w ich padach. Należy pamiętać, że niektóre elementy nie mają wyprowadzeń mieszczących się w rastrze, aby zapewnić odstęp od gniazda DC i gniazda USB.

Po stronie lutowania starannie i równo przytnij wszystkie wystające wyprowadzenia elementów przewlekanych.

Pozostały jeszcze do zmontowania złącza. Najpierw podłącz je do płytki Arduino tak, aby były prawidłowo wyrównane, a następnie umieść na nich nakładkę. Przed lutowaniem należy upewnić się czy nie ma kontaktu między wyprowadzeniami złącza programowania szeregowego (ICSP) na płytce Uno z płytką Obciążenia.

Upewnij się również, że płytka PCB jest mocno dociśnięta do złącza, a następnie przylutuj złącza.

Programowanie

Nasz podstawowy szkic (program) dla Sztucznego Obciążenia jest dla uproszczenia sterowany przez Arduino Serial Monitor. Poprzez monitor raportowane jest napięcie, prąd i moc. Na **ekranie 1** przedstawiono typowy wygląd ekranu monitora szeregowego Arduino podczas użytkowania.

Jeśli nie posiadasz Arduino IDE (zintegrowanego środowiska programistycznego), zacznij od pobrania go ze strony siliconchip.com.au/link/aaq, a następnie zainstaluj.

Teraz otwórz plik szkicu (pobierz z siliconchip.com.au/Shop/6/6330) i wybierz

swoją płytkę (np. Uno, Leonardo lub Mega) oraz port szeregowy z menu *Narzędzia*. Prześlij szkic, a następnie otwórz Monitor szeregowy z menu *Narzędzia*. Ustaw szybkość transmisji na 115200.

Powinieneś zacząć widzieć komunikaty podobne jak te widoczne na ekranie 1, z aktualizacjami występującymi kilka razy na sekundę. Należy pamiętać, że napięcie jest mierzone na samym obciążeniu, więc pokazana moc jest tą, która jest rozpraszana w Obciążeniu.

Testowanie i użytkowanie

Dobrym sposobem na przetestowanie Obciążenia jest podłączenie multimetru do CON1 w celu zmierzenia rezystancji między jego zaciskami. Dodatni przewód multimetru powinien być podłączony do zacisku „+”, a ujemny do „-”. Należy pamiętać, że jeśli zostanie przyłożony prąd wsteczny, będzie on przewodzony przez wewnętrzne diody MOSFET-ów (a tym samym wszystkie rezystory) i pojawi się jako obciążenie 3 Ω.

Nasze oprogramowanie może działać w trzech trybach. Pierwszym z nich jest tryb ręczny, wybierany poprzez wpisanie litery „m” w Serial Monitor, a następnie liczby od 0 do 15. Jest to po prostu liczba rezystorów, które zostaną połączone równolegle i przedstawione jako obciążenie.

Przykładowo, po wprowadzeniu „1”, włączony jest Q1, wpisanie „2” oznacza, że włączony jest Q2, „3” powoduje, że zarówno Q1, jak i Q2 są włączone itd. Trwa to aż do wpisania „15” oznaczającej, że wszystkie MOSFET-y są włączone.

Wpisanie „m1” i naciśnięcie przycisku Enter (upewniając się, że wybrano zakończenie linii „CR”) spowoduje wystąpienie obciążenia 47 Ω na CON1. Wpisanie „m2” spowoduje wybranie obciążenia 23,5 Ω. Można to sprawdzić

Wykaz elementów

1 dwustronna płytka drukowana kodowana, kod 04105221, 89 mm × 54 mm
 1 płytka kompatybilna z Arduino 5 V (np. Uno, Leonardo lub Mega)
 1 10-pinowe złącze szpilkowe o rastrze 2,54 mm
 2 8-pinowe złącze szpilkowe o rastrze 2,54 mm
 1 6-pinowe złącze szpilkowe o rastrze 2,54 mm
 1 2-pinowe złącze śrubowe o rastrze 5/5,08 mm (CON1)

Półprzewodniki:

1 Monitor bocznika prądowego INA282, SOIC-8 (IC1)
 4 CSD18534KCS, IPP80N06S4L-07 lub podobne N-kanalowe MOSFET-y sterowane poziomami logicznymi, TO-220 (Q1...Q4) [2× Cat. SC4177 lub 4× Cat. SC6184]

Kondensatory:

3 kondensatory MKT 100 nF

Rezystory: (wszystkie 1% 1/4 W ośowe, chyba że podano inaczej)

1 33 kΩ 6 10 kΩ 4 470 Ω
 1 15 mΩ 1...3 W M6332/2512 SMD [Cat SC3943]
 15 47 Ω 5 W 10% drutowy

Q1...Q4 mogą być dowolnymi N-kanalowymi MOSFET-ami sterowanymi poziomem logicznym (tj. odpowiednimi do zasilania napięciem 5 V) w obudowach TO-220 o wystarczającym napięciu i prądzie znamionowym

za pomocą multimetru, chociaż ze względu na rezystancję przewodów można zobaczyć nieco wyższe wartości niż oczekiwano.

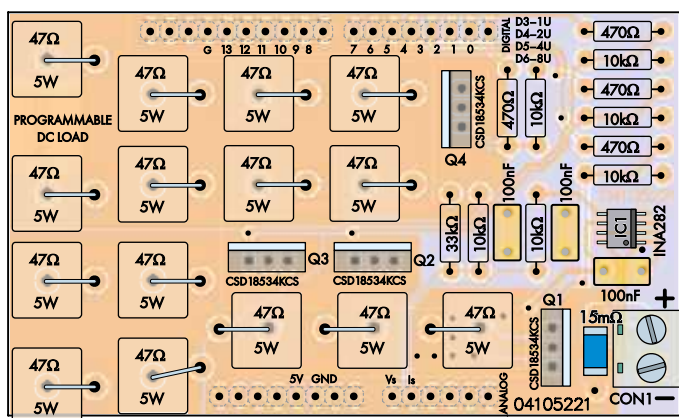
W dowolnym momencie polecenie „m0” odłączy wszystkie rezystory, więc warto o tym pamiętać, jeśli coś pójdzie nie tak.

Drugi tryb polega na wprowadzeniu rezystancji za pomocą polecenia „r”. Oprogramowanie znajduje najbliższą możliwą wartość rezystancji do wprowadzonej wartości. Oczywiście jest tylko 15 dyskretnych kroków, więc prawie nigdy nie będzie to dokładne. Jest to jednak dobry sposób na przybliżenie obciążenia rezystancyjnych o znanej wartości.

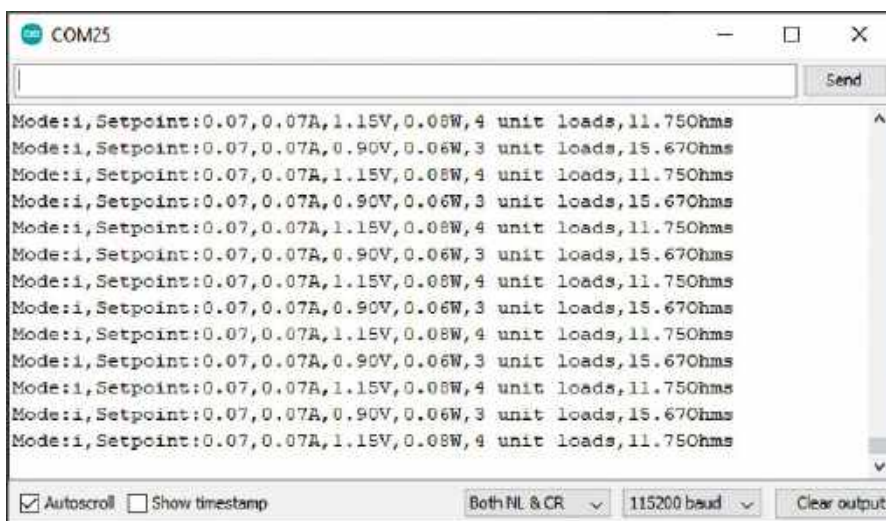
Emulowany tryb stałoprądowy jest uruchamiany poleceniem „i”. W tym trybie podejmowana jest próba dopasowania zmierzzonego prądu do wartości zadanej poprzez zwiększanie i zmniejszanie liczby rezystorów Obciążenia.

Przy ograniczonej liczbie kroków, w trybie stałoprądowym może być również jedynie przybliżony ustawiony prąd i nie będzie reakcji na szybko zmieniające się warunki. Praktycznie we wszystkich przypadkach będziemy obserwować przesłakiwanie między dwoma sąsiednimi poziomami obciążenia, a prąd będzie zygzakiem podążał wokół wartości zadanej. Na ekranie 1 widoczna jest sytuacja, w której Obciążenie przełącza się między 3 i 4 rezystorami, aby utrzymać prąd w pobliżu 70 mA. Przypadek taki występuje po wydaniu polecenia „i0.07”.

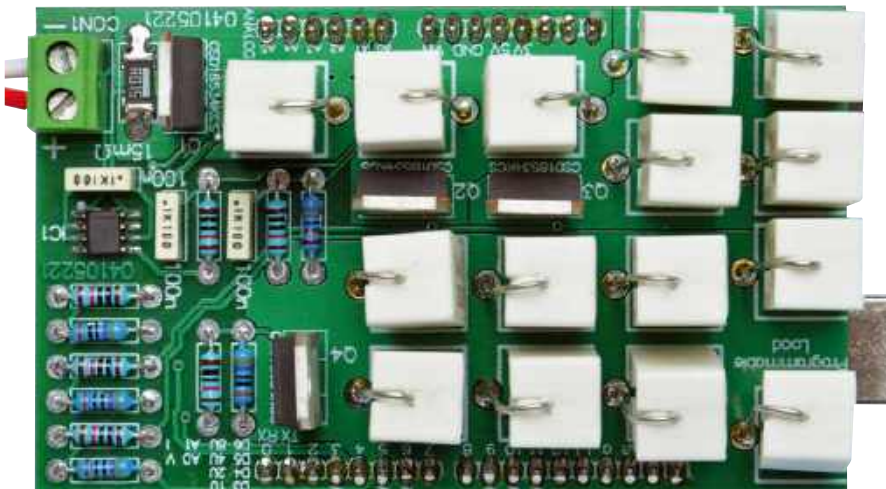
Jeśli napięcie wzrośnie powyżej 15 V lub moc przekroczy 70 W przez dłuższy czas, wyłącz Obciążenie za pomocą polecenia „m0”, aby uniknąć uszkodzenia rezystorów. MOSFET-y nie powinny ulec uszkodzeniu,



Rysunek 2. Płytkę jest łatwa w montażu, ale najlepiej jest zachować ostrożność, aby starannie ułożyć rezystory 5 W. W przeciwnym razie urządzenie będzie wyglądać niechlujnie. Zwróć uwagę na kierunek montażu MOSFET-ów oraz układu IC1. Sprawdź również spód płytki drukowanej, gdy jest ona przymocowana do płytki Arduino, i dopilnuj, by żaden z elementów Obciążenia nie był zwarty z elementami na płytce Arduino. Bocznik 15 mΩ może być zamontowany jako rezystor SMD lub przewlekany



Ekran 1. Do sterowania urządzeniem i wyświetlania jego stanu służy Serial Monitor (lub inny wybrany program terminala szeregowego). Umożliwia on odczyty prądu, napięcia i mocy, a zastosowane Obciążenie jest wyświetlane zarówno jako omy, jak i liczba jednostek 47 Ω . W używanym tutaj trybie „statego prądu” w celu utrzymania jego wartości w pobliżu wartości zadanej sterowana jest rezystancja obciążenia



dopóki napięcie pozostaje poniżej znamionowego napięcia dren-źródło MOSFET-ów wynoszącego dla zalecanych typów 60 V.

Pamiętaj, że wyświetlane napięcie nie może przekroczyć 21,5 V, więc może być znacznie wyższe niż pokazane, jeśli przekracza 20 V.

Więcej wskazówek dotyczących użytkowania

Podłącz ujemny zacisk Programowalnego Obciążenia Arduino do masy układu (pamiętaj, że jest on również wspólny z komputerem sterującym), a zacisk „+” do wyjścia dodatniego.

Na przykład, zasilacz powinien być po prostu podłączony „+” do „+” i „-” do „-”. Jeśli inne obciążenia muszą być podłączone szeregowo, należy je podłączyć między zasilaczem „+” a obciążeniem „+”, aby zadbać o to, by obciążenie „-” pozostało na potencjale masy.

Obciążenie świetnie nadaje się do testowania paneli słonecznych, z zastrzeżeniem, że jest dopilnowana poprawna wartość napięcia dren-źródło MOSFET-ów, szczególnie w warunkach otwartego obwodu, gdy panele wytwarzają najwyższe napięcia. Ogranicza to przede wszystkim panele słoneczne o nominalnym napięciu

wyjściowym 24 V. Mogą one wytwarzać do 44 V w warunkach otwartego obwodu.

Ręczne skanowanie szesnastu różnych poziomów obciążenia utworzy szesnaście punktów danych, które można wykreślić na krzywej I/U lub P/U.

Wprowadzanie modyfikacji

Oprogramowanie jest napisane tak, że większości parametrów jest ustawionych instrukcją #define na początku.

Jeśli chcesz zmodyfikować rezystory obciążenia, wszystkie muszą zachować tę samą rezystancję (chyba że wprowadzisz znaczące zmiany w oprogramowaniu). Jednostkowa rezystancja obciążenia jest określona przez wartość R_CONST.

Wyższe napięcie testowe może przy zmianie zakresu wymagać innego dzielnika (choć trzeba będzie sprawdzić, czy MOSFET-y mogą również obsługiwać wyższe napięcie). Inny dzielnik będzie oznaczał konieczność zmiany mnożnika V_CONST.

Aby obliczyć nową wartość V_CONST, należy obliczyć, jakie napięcie należy przyłożyć, aby na pinie A0 Arduino wystąpiło napięcie 5 V, a następnie podzielić to wyższe napięcie przez 1024. Domyślna wartość 0,0209961 to po prostu 21,5 V podzielone przez 1024.

Użyliśmy (w miarę możliwości) pinów obsługujących PWM, aby możliwe było emulowanie pośrednich wartości rezystancji poprzez zastosowanie sygnałów PWM do MOSFET-ów. Nie wypróbowaliśmy tej techniki, ale można z nią poeksperymentować, jeśli potrzebne jest dokładniejsze sterowanie rezystancji niż przedstawione tutaj wyłącznie dyskretne kroki.

Należy pamiętać, że spowoduje to impulsowe obciążenie źródła prądu/napięcia, co może spowodować nieoczekiwane reakcje. ■

Tim Blythman

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

REKLAMA





Migające diody LED i śliniacy się inżynierowie (17)

Oooh! Mam coś tak kusząco smacznego do omówienia, że aż wiercę się w fotelu z niecierpliwości, ale najpierw musimy przeprowadzić ostatni eksperyment, przejrzeć i zastanowić się nad moim zestawem piłeczek pingpongowych 12×12.

Życie, co można zrobić, co?

Każda z piłeczek pingpongowych w moim układzie jest wyposażona w trójkolorową diodę LED. Kilka felietonów temu (EdW 12/2024) użyliśmy tablicy do realizacji pierwszego przejścia implementacji gry „Gra w życie” Hortona Conway’a (GOL). W poprzednim odcinku (EdW 1/2025) ulepszyliśmy nasz oryginalny program: po pierwsze, aby dodać płynne zanikanie między przejściami, a następnie użyć dodatkowych kolorów, aby odzwierciedlić wszelkie stany międzypokoleniowe (rysunek 1).

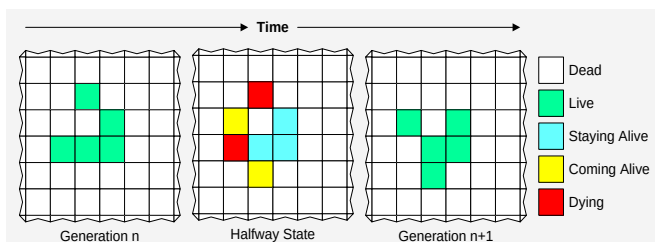
Zanim przejdziemy dalej, dobrym pomysłem może być zapoznanie się z najnowszą wersją tego programu, aby przypomnieć sobie, jak działała jego magia (plik **CB-June21-02.txt**, który jest dostępny na stronie PE z czerwca 2021 r. pod adresem <https://bit.ly/3oouhbl>). Ponadto, by sprawić Wam nieco przyjemności, umieściłem na YouTube film pokazujący to wszystko w akcji (<https://bit.ly/3sXsHPr>).

Sposób, w jaki zostawiliśmy rzeczy, polegał na posiadaniu dwóch wzorów komórek, zwanych „szybowcami”, które przekładają się na siatkę w ciągu wielu pokoleń. Chociaż było to dość skuteczne, chciałem zobaczyć coś nieco bardziej interesującego, więc zdecydowałem się losowo zasiać początkową populację tablicy, pozwolić ewoluować, aż nic się nie zmieni, a następnie rozpocząć wszystko od nowa z nową losowo wygenerowaną populacją.

Należy tu rozważyć kilka kwestii. Pierwszym z nich był przybliżony procent żywych komórek, którymi chcemy wypełnić tablicę. Zacząłem od dwóch definicji: **RAND_MAX**, którą ustawiłem na 100, i **RAND_CUT**, której używam jako wartości odcięcia. Następnie utworzyłem funkcję **InitializeUniverse()**, która „przechodzi” przez każdy wiersz (indeksowany przez **yInd**) i każdą kolumnę (indeksowaną przez **xInd**) w tablicy, wykonując następujące czynności:

```
int tmpRand = random(0, RAND_MAX);

if (tmpRand > RAND_CUT)
    Cells[yInd][xInd].nGenState = ALIVE;
else
    Cells[yInd][xInd].nGenState = DEAD;
```



Rysunek 1. Reprezentacja stanów międzypokoleniowych w Grze Życia Conwaya

Po niewielkich eksperymentach odkryłem, że ustawienie **RAND_CUT** na zbyt małą wartość (powiedzmy 15), co skutkowało ~85% żywych komórek w populacji zalążkowej, powodowało, że wszystkie komórki umierały w następnym pokoleniu z powodu przeludnienia. Analogicznie, zbyt duże **RAND_CUT** (powiedzmy 95), które powodowało, że tylko ~5% komórek było żywych w populacji nasiennej, zwykle powodowało, że wszystkie komórki umierały w ciągu kilku pokoleń z powodu niedostatecznej populacji. Ostatecznie ustaliłem wartość **RAND_CUT** na 85, co spowodowało, że ~15% komórek zaczynało żyć w populacji zalążkowej. Cały program można zobaczyć, pobierając plik **CB-Jul21-01.txt**, który jest dostępny na stronie PE z lipca 2021 r.

Zwróćmy teraz uwagę na fakt, że moja funkcja **InitializeUniverse()** w pierwszym przebiegu ustawia wartości bieżącej („n” dla „teraz”) generacji (**nGenState**) na **ALIVE** lub **DEAD**. Chociaż początkowo wydawało się, że jest to dość intuicyjny sposób robienia rzeczy, pojawiły się pewne niezamierzone konsekwencje, w tym konieczność użycia naszej starej funkcji **DisplayCurrentGeneration()**, której jedynym celem było wyświetlenie generacji nasion, co wydaje się nieco marnotrawstwem, jeśli o mnie chodzi.

Po krótkim zastanowieniu zmodyfikowałem rdzeń funkcji **InitializeUniverse()**, aby ustawić wartości następnej (**n + 1**) generacji (**xGenState**) w przeciwieństwie do bieżącej generacji w następujący sposób (patrz także plik **CB-Jul21-02.txt**):

```
tmpRand = random(0, RAND_MAX);

if (tmpRand > RAND_CUT)
    Cells[yInd][xInd].xGenState = COMING_ALIVE;
else
    Cells[yInd][xInd].xGenState = STAYING_DEAD;
```

Oprócz umożliwienia nam zrzucenia funkcji **DisplayCurrentGeneration()**, pozwala nam to również sprawić, by wszystko wyglądało nieco bardziej efektownie, ponieważ zaczynamy od wszystkich martwych komórek (czarnych), a następnie nowa generacja nasion najpierw zanika z czarnego na cyjan, a następnie z cyjanu na zielony, po czym pozwalamy grze rozpocząć ewolucję. Wreszcie, pod koniec każdego przebiegu, gdy wszystkie komórki umrą, tablica powraca do wszystkich pikseli w kolorze czarnym.

Nagrałem film pokazujący to wszystko w akcji (<https://youtu.be/znysyVUITA>). Jak widzimy, pierwszy losowo zasiany wszechświat ewoluuje przez zaledwie kilka pokoleń, zanim wszystkie komórki opuszczą życie na ziemskim (bądź innym) łóżku padole. Podobnie jest w przypadku drugiego losowo zasianego wszechświata. Jednak trzeci losowy siew okazał się znacznie bardziej interesujący, skutkując wzorcami komórek, które ewoluowały przez wiele pokoleń, zanim ostatecznie ustabilizowały się w prostym oscylatorze „mrugającym” z okresem dwóch pokoleń.

Szczerze mówiąc, nie zastanawiałem się, co by się stało, gdyby system ewoluował i zawierał jeden lub więcej oscylatorów. Mój obecny kod czeka, aż rzeczy przestaną się zmieniać przed ponowną inicjalizacją wszechświata, ale posiadanie wzorca oscylatora oznacza, że rzeczy nigdy nie przestają się zmieniać. Aby skomplikować sprawę, różne oscylatory mogą mieć różne okresy, więc co możemy zrobić, aby określić, że ewoluowaliśmy tak daleko, jak to możliwe i przejść do nowego losowego materiału siewnego? Jeśli masz jakieś sugestie, byłby to idealny moment, aby wyjaśnić, wyjaśnić i wyjaśnić.



Rysunek 2. Klasyczny 7-segmentowy wyświetlacz LED

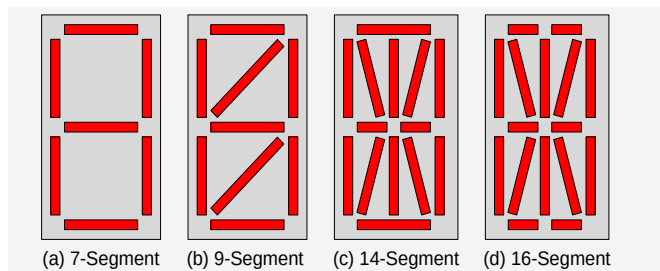
Wystawne segmenty

Do niedawna, gdybym usłyszał termin „wyświetlacz siedmiosegmentowy”, moją odruchową reakcją byłoby pomyślenie o urządzeniach opartych na diodach LED, które mogą być używane do wyświetlania cyfr dziesiętnych od 0 do 9 (**rysunek 2** i **rysunek 3a**). Te maleństwa powstały w latach 70. i są z nami do dziś. Kiedy owe maleństwa po raz pierwszy pojawiły się na scenie, można było mieć dowolny kolor, o ile był to kolor czerwony. Szybko zaczęły pojawiać się w różnego rodzaju urządzeniach, które w tamtych czasach były uważane za niezwykle fajne, takich jak 4-funkcyjne kalkulatory elektroniczne i zegarki cyfrowe. Niestety, nie można było pozostawić włączonego wyświetlacza, ponieważ wyczerpywałoby to baterię, więc trzeba było nacisnąć przycisk z boku zegarka, gdy chciałeś sprawdzić godzinę lub pochwalić się chronometrem na nadgarstku.

Czy wspominałem już, jak fajne były te wyświetlacze w tamtych czasach? Czuję się jak honorowy członek skeczu Czterech Yorkshiremenów Monty Pythona (<https://bit.ly/3xFUYOb>), ale to prawda, że możesz powiedzieć dzisiejszym młodym ludziom, że te wyświetlacze były „bee’s knees” (czyli mega fajne), a oni pomyślą, że przesadzasz... i nie uwierzą ci!

Oprócz cyfr dziesiętnych, wyświetlacze 7-segmentowe mogą być również używane do reprezentowania wartości szesnastkowych, chociaż – w tym przypadku – znaki alfa muszą być prezentowane jako mieszanka wielkich i małych liter: A, b, C, d, E i F.

Oczywiście dodanie większej liczby segmentów pozwala nam reprezentować więcej znaków ze zwiększoną wiernością, więc nie minęło wiele czasu, zanim na scenie zaczęły pojawiać się wyświetlacze z 9, 14 i 16 segmentami (odpowiednio **rysunek 3b**, **rysunek 3c** i **rysunek 3d**), przy czym te ostatnie mogą być używane do wyświetlania wszystkich cyfr arabskich („0” do „9”), liter łacińskich („A” do „Z”) oraz dużej liczby znaków interpunkcyjnych i innych symboli.



Rysunek 3. Coraz większa liczba segmentów wyświetlacza zwiększa zakres znaków, które mogą być wyświetlane

Obławy na Wiktorianów

Niedawno widziałem kreskówkę, w której osoba przeprowadzająca rozmowę kwalifikacyjną pyta: „Dlaczego chcesz zostać redaktorem?”. Rozmówca odpowiada: „Cóż, żeby skrócić długą historię” (pomyśl o tym). Więc skróć długą historię. Dwóch facetów o nazwiskach Chris Barron i John Smout założyło (i obecnie działa jako współmoderatorzy) grupę o nazwie Smartsockets na Groups.io (<https://bit.ly/2PPnVGI>).

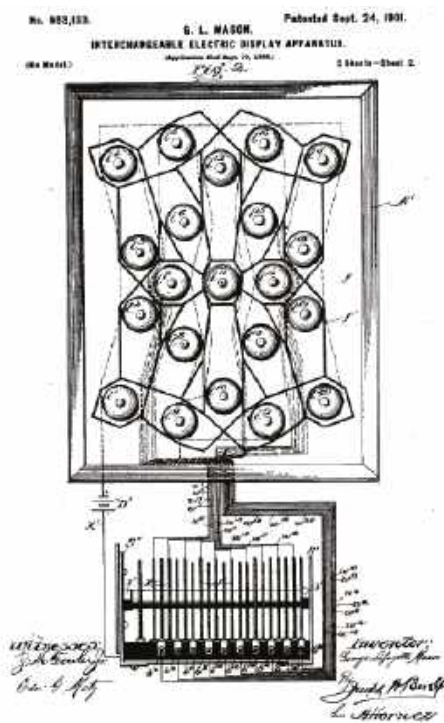
Koncepcja Smartsocket dotyczy systemu oprogramowania i sprzętu do sterowania wielosegmentowymi wyświetlaczami alfanumerycznymi. Każda cyfra jest wyposażona we własny mikrokontroler PIC. Za pomocą prostych instrukcji typu ASCII i protokołów zgodnych ze standardami branżowymi – w połączeniu z wbudowanymi czcionkami i efektami przejścia – możliwe jest wykorzystanie Smartsocket do tworzenia tablic zawierających wiele wyświetlaczy podobnych i różnych typów.

Jakiś czas temu John odkrył, że niejaki George Lafayette Mason złożył patent na 21-segmentowe wyświetlacze w 1898 roku. Patent ten został ostatecznie przyznany w 1901 roku (<https://bit.ly/3gWH17s>). Oryginalne wersje tych pięknych wyświetlaczy składały się z 21 małych żarówek (po jednej na segment), które były kontrolowane przez skomplikowany przełącznik elektromechaniczny, który mógł aktywować różne grupy segmentów, aby reprezentować różne symbole w zależności od potrzeb (**rysunek 4**).

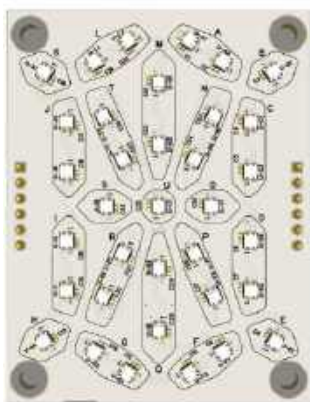
Po tym, jak John dowiedział się o tej 21-segmentowej koncepcji, zdecydował się dodać trójkolorową wersję LED do portfolio Smartsockets, a prace te są kontynuowane.

Wielkie umysły

Chociaż uwielbiam koncepcję Smartsocket, sam nigdy nie byłem zwolennikiem PIC. Ponadto, zamiast mieć oddzielny mikrokontroler dla każdej cyfry, mam tendencję do używania pojedynczego mikrokontrolera – takiego jak Arduino – do sterowania wieloma wyświetlaczami. Mój kumpel Steve Manley jest



Rysunek 4. Kopia oryginalnego patentu na 21-segmentowy wyświetlacz z 1901 roku



Rysunek 5. Płytkę LED i wydrukowaną na drukarce 3D obudowę 21-segmentowego wyświetlacza: a) płytkę LED (po lewej); b) wydrukowaną na drukarce 3D obudowę (po prawej)

podobnego zdania. „Wielkie umysły myślą podobnie”, jak to mówią (oczywiście oni – kimkolwiek są – mówią też, że „Głupcy rzadko się różnią”, ale jestem pewien, że nie mówili o nas).

Aby skrócić kolejną długą historię, Steve stworzył własną wersję 21-segmentowej wiktoriańskiej tablicy wyświetlacza (**rysunek 5a**) z 35 trójkolorowymi diodami LED WS2812B-2020 (po jednej na siedem mniejszych segmentów i po dwie na 14 większych segmentów). Płytki te mają 50 mm szerokości i 64 mm wysokości. Ponadto Steve zaprojektował obudowę do druku w 3D, która jest montowana na powierzchni płytki, aby zapewnić 10 mm odstępu między diodami LED a dyfuzorem przyciemnianym do przedniej części obudowy.

Steve podzielił się swoimi sprytnymi kreacjami i obaj jesteśmy w trakcie tworzenia 10-znakowych wyświetlaczy, które – trzeba przyznać – wyglądają dość niesamowicie, z czym na pewno się zgodzisz, gdy zobaczysz je w przyszłych odcinkach.

Przejęcie kontroli

Steve i twój skromny narrator w przeszłości współzawodniczyli w tworzeniu komplementarnych projektów opartych na diodach LED WS2812. Ogólnie rzecz biorąc, podstawowe zasady dla tych projektów były takie same (np. liczba, lokalizacja i orientacja diod LED), ale poszliśmy różnymi drogami poczynając od użytego sprzętu (w postaci platform programistycznych mikrokontrolerów, których używaliśmy), przez oprogramowanie (w postaci bibliotek LED, których używaliśmy) po wszelkie dodatkowe rozszerzenia, takich jak strategie przetwarzania sygnałów i tym podobne. W rezultacie, co nie jest zaskoczeniem, dzielenie się kodem i szczegółowymi pomysłami projektowymi okazało się trudne, jeśli nie niemożliwe.

W przypadku naszych wiktoriańskich wyświetlaczy zdecydowaliśmy się użyć tego samego sprzętu, co pozwoliło nam wspólnie pracować nad oprogramowaniem (myślę, że „wspólnie” w tym zdaniu było zbędne, ale nie obchodzi mnie to). W związku z tym postanowiliśmy stworzyć płytkę testową, która nie tylko zaspokoi bieżące potrzeby naszych 10-znakowych wyświetlaczy, ale która będzie również miała zastosowanie do szerokiej gamy innych projektów w przyszłości.

Jeśli chodzi o mikrokontrolery, Steve preferuje obecnie Teensy 3.2 od PJRC (<https://bit.ly/3h5tifW>). Posiada on 32-bitowy mikrokontroler Arm Cortex-M4 pracujący z częstotliwością 72 MHz (może być podkreślony do 96 MHz), 256 kB pamięci Flash, 64 kB pamięci SRAM i 2 kB pamięci EEPROM. Osobiście wolę Teensy 3.6, który ma więcej pinów I/O i jest wyposażony w 32-bitowy mikrokontroler Arm Cortex-M4F (tj. ma jednostkę zmiennoprzecinkową (FPU)) pracujący z częstotliwością 180 MHz (można go podkreślić do 240 MHz) z 1 MB pamięci



Rysunek 6. Korzystanie z taniego kontrolera IR znacząco upraszcza życie

Flash, 256 kB pamięci SRAM i 4 kB pamięci EEPROM. W ramach kompromisu zdecydowaliśmy, że nasza płytkę kontrolna będzie obsługiwać oba układy, choć oczywiście tylko jeden na raz.

Następnie przeprowadziliśmy burzę mózgów dotyczącą wszelkich istotnych cech i funkcjonalności, które w różny sposób wykorzystywaliśmy w poprzednich projektach. Obejmują one użycie bardzo dokładnego układu zegara czasu rzeczywistego (RTC) do śledzenia daty i godziny po odłączeniu zasilania od systemu. Obejmują również użycie rezystora zależnego od światła (LDR), czyli fotorezystora, który może być wykorzystany do zmiany jasności wyświetlacza w zależności od poziomu oświetlenia otoczenia.

Zależy nam również na dostępie do wejścia audio, abyśmy mogli używać dźwięku do kontrolowania wyświetlanych wzorów. W jednym z naszych poprzednich projektów użyłem wejścia liniowego audio w moim projekcie, podczas gdy Steve użył w swojej implementacji mikrofonu, więc zdecydowaliśmy się wesprzeć obie opcje na naszej nowej płytce. Ponadto obaj używaliśmy układów analizatora widma audio MSGEQ7 w poprzednich projektach (<https://bit.ly/3h23QYv>). Ostatnio Steve eksperymentował z zaawansowanym układem kodeka audio w połączeniu z biblioteką Teensy Audio Library (<https://bit.ly/3tpnrEA>) i zdecydowaliśmy, że podejście Steve’a oferuje najlepszą drogę na przyszłość.

Jeśli chodzi o sterowanie systemem za pomocą przycisków, jeden z moich wcześniejszych projektów wykorzystywał schemat 3-przyciskowy, podczas gdy inny wykorzystywał podejście 6-przyciskowe. Na potrzeby naszej nowej płytki sterującej zdecydowaliśmy się na 5-przyciskowe rozwiązanie, które Steve rozwinął w wielu swoich ostatnich projektach. Możemy używać przycisków w lewo, w prawo, w górę, w dół i „OK”, aby uzyskać dostęp do menu, wybierać tryby i efekty oraz wprowadzać wartości (np. datę i godzinę). Co więcej, jeden z ostatnich projektów Steve’a miał również możliwość obsługi taniego kontrolera podczerwieni (IR) (<https://amzn.to/3b6ij8x>). Oprócz



Rysunek 7. Płytkę sterującą, która będzie sterować naszymi 10-znakowymi, 21-segmentowymi wyświetlaczami

przycisków w lewo, w prawo, w górę, w dół i „OK” (**rysunek 6**), ma on również przyciski dla cyfr od 0 do 9, co znacznie przyspiesza i ułatwia wykonywanie zadań, takich jak wprowadzanie daty i godziny, dlatego zdecydowaliśmy się włączyć tę funkcję do naszej nowej płytki.

W ramach tego zdecydowaliśmy się dodać Seeedunio XIAO od Seeed Studio (<https://bit.ly/2QYem8v>) do obsługi funkcji sterowania IR (Steve wcześniej używał do tego celu Adafruit Trinket). Podczas gdy XIAO jest wielkości zwykłego znaczka pocztowego, posiada 32-bitowy mikrokontroler Arm Cortex-M0+ pracujący z częstotliwością 48 MHz, 256 kB pamięci Flash i 32 kB pamięci SRAM. W rzeczywistości prawdopodobnie moglibyśmy obsługiwać funkcje sterowania IR za pomocą głównego Teensy, ale czasami łatwiej jest zastosować podejście „dziel i zwyciężaj”. Ponadto, chociaż XIAO jest całkowicie przesadzony do użytku jako kontroler IR, jest niezwykle tani i – podobnie jak mikrokontrolery Teensy – można go programować za pomocą zintegrowanego środowiska programistycznego Arduino (IDE). Ponadto, zawsze warto mieć zapasową moc obliczeniową na wypadek, gdyby była potrzebna w przyszłości.

W oparciu o wszystkie powyższe, wraz z szeregiem innych rozważań, Steve zaprojektował niesamowicie fajną płytkę, która będzie nam dobrze służyć w wielu przyszłych projektach (**rysunek 7**).

W chwili pisania tego tekstu nie planujemy udostępniać tej tablicy kontrolnej jako produktu – nie zamierzamy też sprzedawać samych wiktoriańskich wyświetlaczy – wszystko to jest tylko projektem hobbystycznym. Z drugiej strony, będę dzielił się niektórymi szczegółami w moich artykułach, ponieważ wierzę, że mogą ci się one przydać w twoich projektach.

Niezwykle ekscytującą wiadomością jest również to, że właśnie zaprojektowaliśmy specjalną płytkę SMAD (Steve and Max's Awesome Display) z 45 trójkolorowymi diodami LED o niskim poborze prądu, które mogą być sterowane za pomocą zwykłego Arduino. Płytkę ta może być używana do prototypowania efektów specjalnych, których będziemy używać na naszych wiktoriańskich wyświetlaczach, i udostępnimy ją do zakupu za pośrednictwem sklepu PE (ale o tym w następnym odcinku).

Kilka najdrobniejszych szczegółów

Oszczędzając tekst, wybierzmy się na krótką przechadzkę po tablicy kontrolnej, aby zapoznać się i rozważyć interesujące fakty. Po lewej stronie widzimy Teensy 3.6. W razie potrzeby można go usunąć i zastąpić układem Teensy 3.2. Po prawej stronie widzimy XIAO, który jest również podłączony do niektórych pinów i który będzie używany do przetwarzania sygnałów przychodzących z kontrolera IR.

Tuż za Teensy znajdują się dwa potencjometry, które mogą być używane do regulacji jasności wyświetlacza i czułości dźwięku (lub

czegokolwiek innego, do czego je zaprogramujemy). Tuż przed Teensy znajduje się bateria pastylkowa, która podtrzymuje działanie zegara RTC po odłączeniu zasilania od systemu.

Jeśli chodzi o RTC, we wcześniejszych projektach korzystałem z ultraprecyzyjnej płytki ChronoDot (BOB) firmy Adafruit (<https://bit.ly/3tkCHml>). W międzyczasie Steve używał układu scalonego DS3231 na swoich niestandardowych płytkach. DS3231 to niedrogi, niezwykle dokładny RTC z interfejsem I²C oraz zintegrowanym oscylatorem kwarcowym z kompensacją temperatury (TCXO). Dziwnym zrządzeniem losu jest to ten sam układ, który jest używany w ChronoDot.

DS3231 i wszelkie inne układy scalone do montażu powierzchniowego znajdują się w dolnej części płytki i dlatego nie są widoczne na **rysunku 7**. Te układy scalone obejmują stereofoniczny kodek audio o niskim poborze mocy SGT5000 firmy NXP, układ scalony debounced LS119 firmy LogiSwitch oraz transceiver magistrali ósemkowej SN74HCT245 firmy Texas Instruments działający jako przesuwnik poziomu napięcia.

Zwróć uwagę na pięć przycisków sterujących, które znajdują się w prawej środkowej części płytki. Możliwe jest również równoległe podłączenie pięciu kolejnych przycisków zamontowanych na obudowie za pomocą zielonych złączy śrubowych. Powyżej środkowego przycisku znajduje się LDR i mikrofon elektretowy; poniżej środkowego przycisku widzimy detektor podczerwieni.

Sposób, w jaki płytka została zaprojektowana, większość jej cech i funkcji, takich jak sterowanie IR, RTC i kodek audio, jest opcjonalna. Ponadto, w przeciwieństwie do lutowania komponentów, takich jak mikrofon, LDR i detektor podczerwieni bezpośrednio na płycie, można je zamontować zdalnie i podłączyć za pomocą zielonych zacisków śrubowych.

Piny nagłówka po prawej stronie Teensy zapewniają dostęp do wszystkich niezaangażowanych pinów mikrokontrolera. Piny nagłówka w prawym dolnym rogu baterii pastylkowej umożliwiają podłączenie do płytki dodatkowych czujników z obsługą I²C i innych urządzeń. Na koniec należy zwrócić uwagę na klaster 8×3 pinów nagłówka w lewym dolnym rogu baterii pastylkowej. Są one przedstawione jako osiem grup, z których każda składa się z 5 V, 0 V i sygnału danych, który może być używany do sterowania trójkolorowymi diodami LED WS2812. Omówiony wcześniej układ SN74HCT245 służy do konwersji sygnałów wyjściowych 3,3 V z układu Teensy na sygnały 5 V wymagane do sterowania diodami LED. Powodem zastosowania tego bufora ósemkowego jest to, że Teensy jest w stanie sterować ośmioma pasmami diod LED jednocześnie. Zagłębimy się w tę możliwość w następnym odcinku. W międzyczasie, jak zawsze, czekam na wszelkie komentarze, pytania i sugestie. ■

Clive „Max” Maxfield

Sprytne porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania



Zanim wskoczmy do walki z zapalonym i porzuceniem (i oczywiście zuchwałością), rozmawiałem ostatnio z moim kumplem Stevem Manleyem, który zajmuje tak ważne miejsce w treści artykułu powyżej. Nie jestem pewien, jak doszliśmy do tego konkretnego tematu, ale Steve nauczył mnie bardzo sprytnej sztuczki kodowania, o czym poniżej.

Osiągnąłem swój limit!

W moich programach bardzo często zdarza się, że mam zmienną całkowitą, której używam jako wskaźnika (lub indeksu) do liczby elementów w tablicy lub podobnej jednostce. Na przykład, powiedzmy, że mamy dziesięć elementów ponumerowanych od 0

do 9. Dla przejrzystości naszego kodu i ułatwienia jego modyfikacji w przyszłości, możemy zdefiniować `NUM_E` („liczba elementów”) jako 10, a `MAX_E` („maksymalny element”) jako 9. (Możemy również zdefiniować `MIN_E` jako 0, ale zazwyczaj używamy po prostu 0).

Załóżmy, że nasz wskaźnik liczby całkowitej nazywa się `PtrP`. Czasami chcemy zwiększyć ten wskaźnik, dodając 1 do jego bieżącej wartości. Oczywiście, gdy jego aktualna wartość wynosi 9, chcemy, aby jego zwiększona wartość „zawinęła się” do 0. W tym przypadku, pamiętając, że operator `%` modulo zwraca resztę z dzielenia liczb całkowitych i opierając się na naszych dyskusjach z poprzedniego spotkania (EdW 10/2024), przed dyskusją ze Stevem, mógłbym użyć instrukcji jak poniżej:

```
// Zwiększenie wskaźnika
PtrP = (PtrP + 1) % NUM_E;
```

W innych przypadkach chcemy zmniejszyć nasz wskaźnik, odejmując 1 (lub dodając -1) do jego bieżącej wartości. W tym przypadku, gdy jego aktualna wartość wynosi 0, chcemy, aby jego zmniejszona wartość „zawinęła się” do 9. Ponownie, w oparciu o nasze wcześniejsze dyskusje, mógłbym użyć instrukcji jak poniżej:

```
// Zmniejszenie wskaźnika
PtrP = (PtrP + MAX_E) % NUM_E;
```

Cóż, Steve wskazał, że sposób, w jaki to robi, zakładając, że ma zmienną o nazwie `Delta`, której można przypisać wartość +1 lub -1, jeśli chcemy odpowiednio zwiększać lub zmniejszać, jest następujący:

```
// Zwiększanie/zmniejszanie wskaźnika
PtrP = (PtrP + NUM_E + Delta) % NUM_E;
```

Dobry Boże, panno Molly! To jest takie proste... takie zwięzłe... takie ostre. Uwielbiam to!

Dzięki za pamięć!

Nie wiem jak Ty, ale ja martwię się, że z wiekiem stracę pamięć. Zdarzyło się to siostrze mojej mamy – mojej ciotce Shirley – która obecnie mieszka w domu opieki, ponieważ nie może już o siebie zadbać i nie rozpoznaje już żadnych członków naszej rodziny, w tym własnych dzieci. Na szczęście moja 90-letnia mama wciąż ma umysł jak pułapka (od Red. EdW: wyrażenie „umysł jak pułapka” pochodzi od angielskiego idiomu „mind like a steel trap” i oznacza, że dana osoba ma niezwykle bystry, przenikliwy i szybki umysł, który potrafi natychmiast uchwycić szczegóły, zapamiętać informacje i wyciągać trafne wnioski). W rzeczywistości pamięć mojej mamy jest tak dobra, że czasami pamięta rzeczy, które nawet się jeszcze nie wydarzyły!

To meandryczne rozmyślanie zostało wywołane moimi przemyśleniami na temat pamięci komputerowej. W głównym wątku naszego spotkania zauważyliśmy, że Teensy – podobnie jak praktycznie wszystkie mikrokontrolery – zawiera trzy rodzaje pamięci. Po pierwsze, mamy pamięć Flash, która służy do przechowywania głównego programu (czyli szkicu). Następnie mamy pamięć SRAM (statyczną pamięć o dostępie swobodnym), w której główny program tworzy, przechowuje i manipuluje zmiennymi podczas działania. Wreszcie, mamy formę pamięci, o której większość początkujących nawet nie wie i rzadko z niej korzysta – niewielką ilość EEPROM (elektrycznie kasowalna programowalna pamięć tylko do odczytu) – w której programiści mogą przechowywać skromne ilości długo-terminowych informacji.

Problem z pamięcią SRAM polega na tym, że jest ona ulotna, co oznacza, że zapomina swoją zawartość po odłączeniu zasilania od systemu. Pamięć Flash jest nieulotna, co oznacza, że pamięta swoją zawartość po odłączeniu zasilania, ale to właśnie w niej przechowywany jest główny program. Załóżmy, że zdecydujemy się napisać program, który mierzy temperaturę otoczenia raz na godzinę i chcemy zapisać te wartości w taki sposób, abyśmy mogli je później odzyskać, nawet jeśli zasilanie mikrokontrolera ulegnie awarii w dowolnym momencie. W takim przypadku jedną z opcji byłoby użycie pamięci EEPROM.

Podobnie, w przypadku 10-znakowych 21-segmentowych wyświetlaczy wiktoriańskich, które Steve i ja konstruujemy, chcemy użyć bajtowych wartości całkowitych bez znaku, aby śledzić różne ustawienia użytkownika, takie jak preferowany format daty (np. 0 = RRRR/MM/DD, 1 = MM/DD/RRRR, 2 = DD/MM/RRRR) i format czasu (np. 0 = 12-godzinny, 1 = 24-godzinny) oraz lokalizację (np. 0 = Wielka Brytania, 1 = USA) i sposób obsługi czasu letniego (np. 0 = ręcznie, 1 = automatycznie) i tak dalej. (FYI, „czas letni” jest nazywany „czasem letnim” (DST) w USA i „brytyjskim czasem letnim” (BST) w Wielkiej Brytanii).

Oczywiście w naszym głównym programie ustalimy wartości domyślne dla tych ustawień. Chcemy jednak również umożliwić użytkownikowi zmianę tych ustawień podczas działania programu i chcemy, aby nasz wyświetlacz pamiętał te zdefiniowane przez użytkownika ustawienia po odłączeniu zasilania od systemu. Ponownie, jednym ze sposobów osiągnięcia tego celu jest użycie pamięci EEPROM.

Należy pamiętać, że termin „bajt” odnosi się do wielkości 8-bitowej. Jak omówiliśmy w głównej części artykułu, mikrokontrolery Teensy 3.2 i 3.6 mają odpowiednio 2 kB (2,048 bajtów ponumerowanych od 0 do 2,047) i 4 kB (4,096 bajtów ponumerowanych od 0 do 4,095) pamięci EEPROM.

Jeśli chcemy użyć tej pamięci EEPROM w programach, musimy najpierw dołączyć specjalną bibliotekę, która jest dostarczana jako część IDE Arduino:

```
#include <EEPROM.h>
```

Załóżmy teraz, że chcemy zapisać wartość 128 do pamięci EEPROM pod adresem 0. Możemy to zrobić za pomocą następującej instrukcji:

```
EEPROM.write(0, 128);
```

Alternatywnie, jeśli zadeklarujemy zmienną całkowitą o nazwie `Address`, której przypiszemy wartość 0, wraz ze zmienną o rozmiarze bajtu o nazwie `Data`, której przypiszemy wartość 128, możemy użyć następującej instrukcji:

```
EEPROM.write(Address, Data);
```

W związku z tym, jeśli na pewnym etapie chcemy odczytać bajt danych z adresu 0 pamięci EEPROM, możemy użyć jednej z poniższych instrukcji:

```
Data = EEPROM.read(0);
Data = EEPROM.read(Address);
```

Więcej o bibliotece EEPROM można przeczytać w podręczniku Arduino (<https://bit.ly/2QYwsHO>) i na stronie Teensy (<https://bit.ly/3y0L8H1>), ale sama znajomość funkcji `read()` i `write()` zapewni nam wystarczającą wiedzę, byśmy mogli zacząć się bać samych siebie.

Ile kopii?

W rzeczywistości będziemy mieć wiele różnych ustawień, które chcemy śledzić. Jednak wyłącznie na potrzeby tej dyskusji założymy, że mamy tylko cztery ustawienia wielkości bajtu, które omówiliśmy wcześniej: `vdDate`, `vdTime`, `vdLocation` i `vdSummer` (gdzie „vd” oznacza „Victorian Display”). W rzeczywistości będziemy chcieli zachować trzy kopie tych ustawień (w przeciwieństwie do „kopii”, możemy myśleć o nich jako o wersjach lub instancjach). Po pierwsze, będziemy potrzebować kopii naszych wartości domyślnych, które będziemy przechowywać w samym programie. Po drugie, będziemy potrzebować kopii wartości zdefiniowanych przez użytkownika, które będziemy przechowywać w pamięci EEPROM. Wreszcie, będziemy potrzebować kopii roboczej wartości, których faktycznie używamy.

Na początku może się to wydawać nieco zagmatwane, ale ma to sens, gdy się nad tym zastanowić. Założymy, że ładujemy nasz program na nowy mikrokontroler. W tym przypadku, gdy uruchomimy program po raz pierwszy, w pamięci EEPROM nie będzie zapisane nic użytecznego. Kiedy wykryjemy ten fakt (omówimy to dalej w następnym odcinku Sprytnie porady i sztuczki), załadujemy nasze wartości domyślne zarówno do pamięci EEPROM, jak i do naszych wartości roboczych.

Ponieważ jest to pierwszy raz, kiedy uruchamiamy program, prawdopodobnie skorzystamy z okazji, aby zmodyfikować różne ustawienia tak, aby były takie, jak chcemy. Oczywiście w przyszłości możemy wprowadzić dodatkowe zmiany. Chodzi o to, że dla każdego ustawienia, które zmienimy, zastąpimy odpowiednią wartość roboczą i wartość EEPROM nową wartością.

Czasami czujesz się jak tablica

Na przykład, gdy przychodzi do kopiowania ustawień do i z pamięci EEPROM, łatwiej jest myśleć o rzeczach jako o tablicach. Na przykład, zakładając, że zdefiniowaliśmy `NUM_SETTINGS` jako 4, możemy zdefiniować nasze ustawienia domyślne i nasze ustawienia robocze w postaci tablic w następujący sposób:

```
uint8_t DefSettings[NUM_SETTINGS];
uint8_t WrkSettings[NUM_SETTINGS];
```

Założymy, że używamy lokalizacji 0, 1, 2 i 3 w tych tablicach, aby śledzić odpowiednio nasze ustawienia `vdDate`, `vdTime`, `vdLocation` i `vdSummer`. Nie będziemy się tutaj martwić o sposób inicjalizacji, po prostu założymy, że nasza tablica `DefSettings[]` została załadowana odpowiednimi wartościami. Jeśli chcemy skopiować wartości z tablicy `DefSettings[]` do tablicy `WrkSettings[]`, możemy użyć:

```
for (i = 0; i < NUM_SETTINGS; i++)
    WrkSettings[i] = DefSettings[i];
```

Podobnie, jeśli chcemy skopiować wartości z tablicy `DefSettings[]` do pamięci EEPROM, możemy użyć:

```
for (i = 0; i < NUM_SETTINGS; i++)
    EEPROM.write(i, DefSettings[i]);
```

I oczywiście, jeśli chcemy skopiować wartości z EEPROM do naszej tablicy `WrkSettings[]`, możemy użyć:

```
for (i = 0; i < NUM_SETTINGS; i++)
    WrkSettings[i] = EEPROM.read(i);
```

Czasami czujesz się jak struktura

Problem z myśleniem o rzeczach jako tablicach polega na tym, że nie czyni to naszego kodu zbyt czytelnym później. Na przykład, co pomyślimy, jeśli czytamy główny program i widzimy coś takiego:

```
if (WrkSettings[2] == 0)...
```

Pamiętajmy, że w prawdziwej aplikacji możemy mieć dziesiątki lub setki takich ustawień. Dlatego znacznie ułatwiłoby nam życie, gdybyśmy mogli myśleć o naszych wartościach jako o polach i powiedzieć coś w stylu:

```
if (WrkSettings.vdLocation == UK)...
```

Jak być może pamiętasz, wprowadziliśmy pojęcia `typedef` (definiacje typów), `enum` (typy wyliczeniowe) i `struct` (struktury) w Sprytnie porady i sztuczki, EdW 7/2024. Na tej podstawie możemy zdecydować się na zdefiniowanie i zadeklarowanie niektórych struktur w następujący sposób:

```
typedef struct Settings
{
    uint8_t vdDate;
    uint8_t vdTime;
    uint8_t vdLocation;
    uint8_t vdSummer;
};
```

```
Settings DefSettings;
Settings WrkSettings;
```

Ponownie, nie będziemy się martwić o to, jak zainicjujemy rzeczy tutaj, po prostu założymy, że nasza struktura `DefSettings` została załadowana odpowiednimi wartościami. Gdy mamy już odpowiednie wartości w naszej strukturze `WrkSettings`, jesteśmy gotowi do pracy. Problem pojawia się, gdy chcemy załadować tę strukturę. Jeśli ładujemy ją z naszej struktury `DefSettings`, będziemy musieli użyć serii instrukcji takich jak:

```
WrkSettings.vdDate =
    DefSettings.vdDate;
WrkSettings.vdTime =
    DefSettings.vdTime;
:
itp.
```

Alternatywnie, jeśli ładujemy wartości do naszej struktury `WrkSettings` z pamięci EEPROM, będziemy musieli użyć serii instrukcji takich jak:

```
WrkSettings.vdDate =
    EEPROM.read(0);
WrkSettings.vdTime =
    EEPROM.read(1);
:
itp.
```

Nie trzeba długo czekać, aby zdać sobie sprawę, że jeśli mamy dziesiątki lub setki ustawień, szybko stanie się to uciążliwe, a – uwierz mi – jest to ostatnie miejsce, w którym chcemy mieć problem. Gdyby

tylko istniał jakiś sposób, w jaki moglibyśmy traktować nasze ustawienia zarówno jako tablicę, jak i strukturę...

Stwórzmy związek!

To prawie tak, jakby ludzie, którzy stworzyli język programowania C, czytali w naszych myślach, ponieważ stworzyli specjalny typ danych zwany unią, który pozwala nam przechowywać różne typy danych w tych samych lokalizacjach pamięci. Innym sposobem myślenia o tym jest to, że unia pozwala nam wyświetlać te same lokalizacje pamięci na różne sposoby. Na przykład, rozważmy następującą sytuację (pamiętajmy, że `NUM_SETTINGS` zostało zdefiniowane jako 4):

```
typedef union Settings
{
    struct
    {
        uint8_t vdDate;
        uint8_t vdTime;
        uint8_t vdLocation;
        uint8_t vdSummer;
    } vds;

    uint8_t vda[NUM_SETTINGS];
};
```

```
Settings DefSettings;
Settings WrkSettings;
```

Najpierw definiujemy nowy typ w postaci unii, którą nazwalimy `Settings`. Jak widzimy, ta unia oferuje dwa różne sposoby

postrzegania/myślenia/traktowania tych samych czterech bajtów pamięci. Pierwszą metodą jest myślenie o tych czterech bajtach jako o strukturze, którą nazwaliliśmy `vdS`; drugim podejściem jest myślenie o tych samych czterech bajtach jako o tablicy, którą nazwaliliśmy `vdA`.

Następnie deklarujemy dwie zmienne, `DefSettings` i `WrkSettings`, z których obie są typu `Settings`. Tak jak poprzednio, nie będziemy się tutaj martwić o sposób inicjalizacji. Wystarczy powiedzieć, że jeśli ustalimy, że pamięć EEPROM zawiera prawidłowy zestaw wartości ustawień, możemy załadować nasze ustawienia robocze w następujący sposób:

```
for (i = 0; i < NUM_SETTINGS; i++)
    WrkSettings.vdA[i] = EEPROM.read(i);
```

Później, w treści programu, możemy użyć instrukcji takich jak:

```
if (WrkSettings.vdS.vdLocation
    == UK)...
```

Dotknęliśmy tutaj tylko mocy typu `union`, ponieważ po prostu zdefiniowaliśmy dwa różne sposoby patrzenia na te same cztery bajty pamięci. W rzeczywistości unia może zapewnić trzy lub więcej sposobów patrzenia na ten sam obszar pamięci, gdzie jeden członek może myśleć o rzeczach jako 4-bajtowych liczbach całkowitych bez znaku, inny może myśleć o każdej z tych liczb całkowitych jako o czterech oddzielnych bajtach, a jeszcze inny może myśleć o rzeczach jako o indywidualnie nazwanych bitach... i wtedy sprawy zaczynają się komplikować, ale możemy to zostawić na inny dzień. Jak zawsze, czekam na komentarze, pytania i sugestie. ■

Clive „Max” Maxfield

REKLAMA

UWAGA! Tylko prenumeratorzy czasopism Elektronika dla Wszystkich, Elektronika Praktyczna, Świat Radio oraz Elektronik mogą korzystać z atrakcyjnych rabatów w Sklepie AVT:

- ✓ do 50% na wydania specjalne czasopism Wydawnictwa AVT
- ✓ 20% na kity w wersji A (płytki drukowane do projektów AVT)
- ✓ 10% na pozostałe wersje kitów: (A+, B, C, D)
- ✓ 10% na książki
- ✓ 5% na pozostałe produkty z oferty sklepu

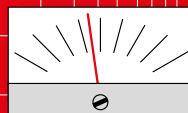
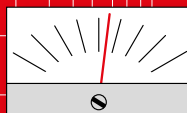
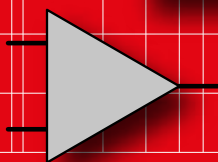
Ponadto każdy prenumerator ww. czasopism korzysta z rabatów od 30% do 50% na zakup czasopism z oferty www.UlubionyKiosk.pl

K L U B
AVT
ELEKTRONIKA

Jak uzyskać rabat? Podczas zamówienia powołaj się na swój numer prenumeraty – otrzymasz go mailowo po zakupie prenumeraty wraz z kartą członkowską Klubu AVT-Elektronika.

Regulamin Klubu AVT-Elektronika znajdziesz na stronie <https://sklep.avt.pl/klub-avt-elektronika>

AUDIO OUT



Wokoder analogowy, część 5

Konstrukcja filtrów

W tym odcinku przechodzimy niejako do sedna projektu – do filtrów kanałowych. Filtry te są sercem każdego wokodera. Upprzedzam, że zgromadzenie wszystkich elementów dla filtrów i ich zmontowanie zajmie dużo czasu!

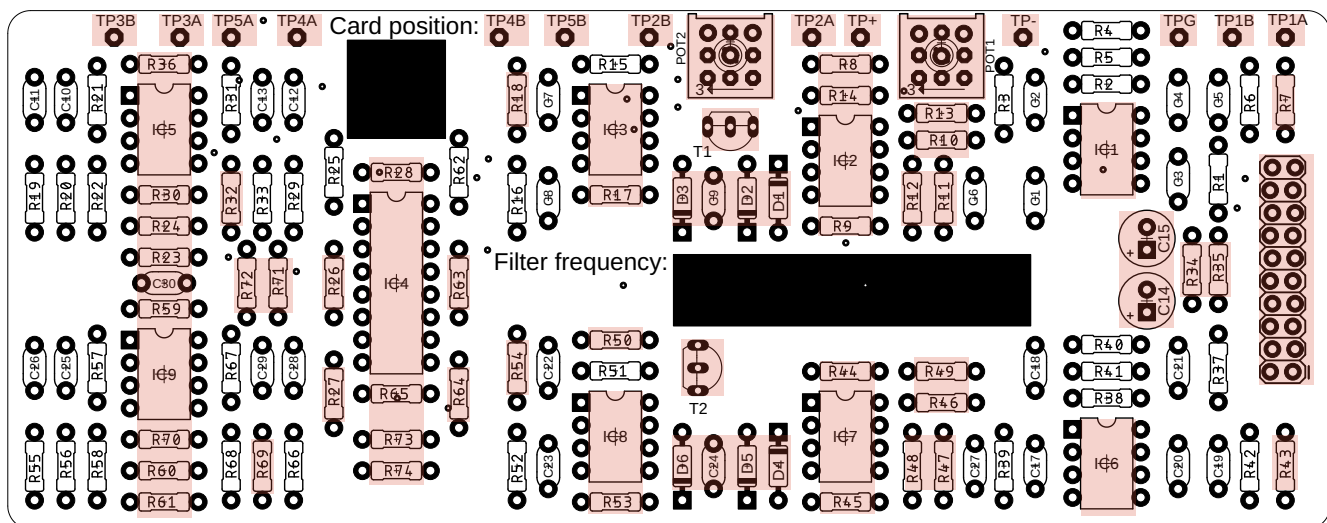
Jest sześć płytek pasmowoprzepustowych, z których każda zawiera po dwa filtry. Istnieje również jedna płytka zawierająca filtry górno- i dolnoprzepustowy, oznaczana skrótem „HP/LP”. Wszystkie te płytki są dołączone do magistrali z rozprowadzonymi wspólnymi liniami zasilania, dwoma sygnałami wejściowymi i dwoma liniami sumującymi. Płytki filtrów są specyficzne dla wokodera; nie są to moduły ogólnego zastosowania, takie jak płytki przedwzmacniacza mikrofonowego i sterownika przedstawione w poprzednim odcinku cyklu. Całą konstrukcję pokazano na rysunku 1.

Układ płytki

Kompletny schemat płytki pasmowoprzepustowej przedstawia rysunek 2. Przypominam, że na jednej płytce są dwa kanały wokodera. Na szczęście wszystkie elementy są „ze sklepu za rogiem” – nie ma tu komponentów egzotycznych ani kosztownych. Należy zauważyć,

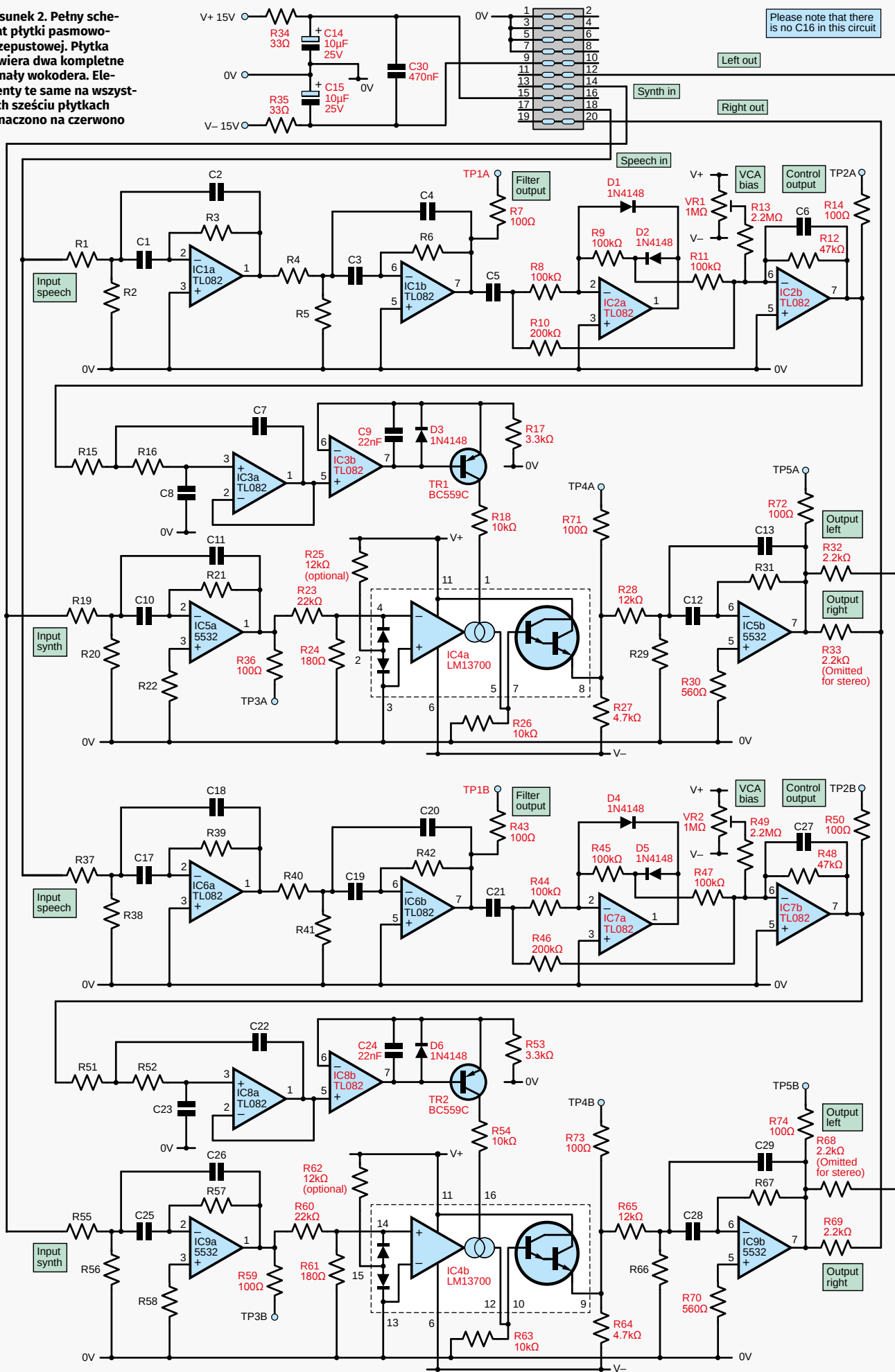


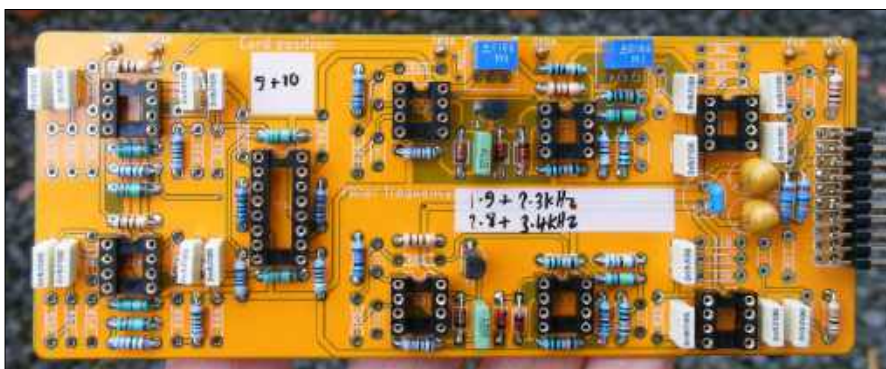
Rysunek 1. Serce wokodera – sześć płytek z filtrami pasmowoprzepustowymi i płytka górno-/ dolnoprzepustowa. Wszystko podłączone do magistrali



Rysunek 3. Schemat montażowy płytki pasmowoprzepustowej. Elementy wspólne dla wszystkich płytek oznaczono na czerwono. Płytkę zaprojektował Mike Grindle

Rysunek 2. Pełny schemat płytki pasmowo-przepustowej. Płytkę zawiera dwa kompletne kanały wokodera. Elementy te same na wszystkich sześciu płytkach oznaczono na czerwono





Rysunek 4. Zdjęcie płytki pasmowoprzepustowej ze zmontowanymi elementami wspólnymi oraz białymi kondensatorami 3,9 nF użytymi w filtrach. Zwracam uwagę na białe pola sitodrukowe na opisy

że jest kilka ulepszeń (a więc różnic!) w porównaniu ze schematem opublikowanym poprzednio w EdW 11/2024. Po pierwsze:

- R24 wynosi 180 Ω;
- kondensatory odsprężające C14, C15 i C16 zostały inaczej ponumerowane;
- ceramiczny kondensator odsprężający C16 nazywa się teraz C30,

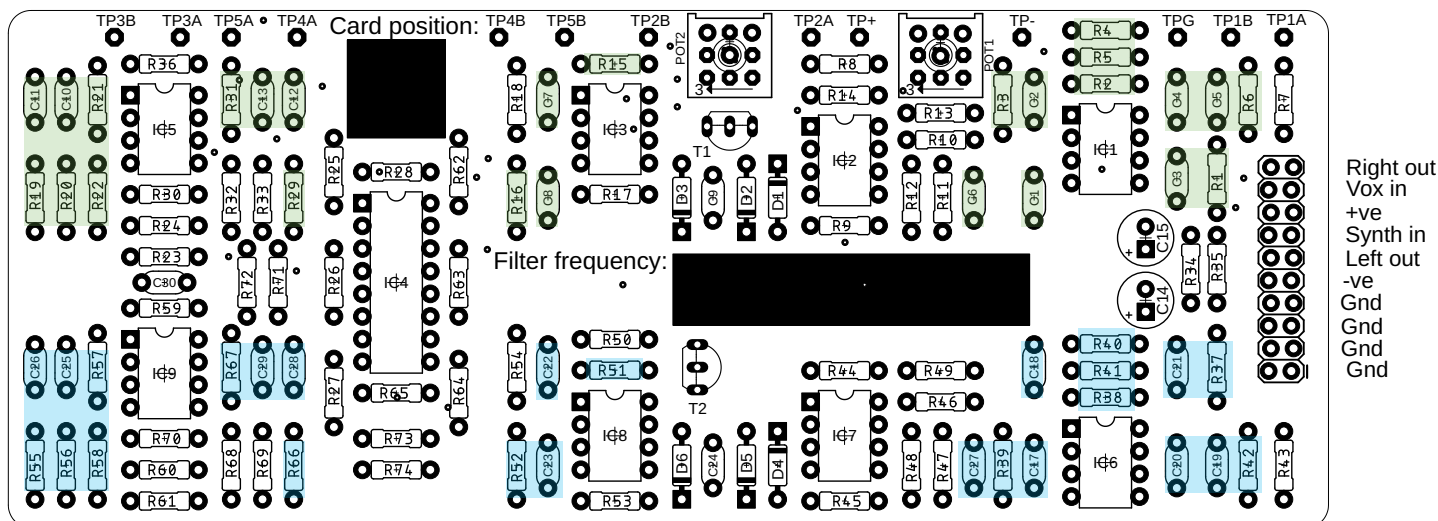
ponieważ musiałem przenieść go na drugi koniec płytki, tam gdzie znajdują się układy wyjściowe NE5532;

- było to konieczne, ponieważ układ 5532 jest bardziej wymagający w kwestii odsprężania zasilania niż np. TL082;
- jeszcze jeden drobiazg: na poprzednim schemacie układ IC3 miał pozamieniane wejścia; rysunek 2 przedstawia wersję poprawioną.

Drugą ważną zmianą jest to, że odwróciłem fazę co drugiego kanału poprzez zamianę wejść odwracających i nieodwracających w transkonduktancyjnym wzmacniaczu operacyjnym IC4b. Zapewnia to mniejsze znoszenie się fazy w obszarach nakładania się odpowiedzi częstotliwościowej filtrów. Po zmiksowaniu monofonicznym dźwięk jest wyraźniejszy. Do odwracania fazy

Tabela 1. Elementy specyficzne dla każdego filtra pasmowoprzepustowego, zależne od zakresu częstotliwości

Band-pass channel and frequency		R2, 5, 20, 29 R38, 41, 56, 66 (kΩ)	R1, 19 R37, 55 (kΩ)	R3, 21, 22 R39, 57, 58 (kΩ)	C1, 2, 3, 4, 10, 11, 12, 13 C17, 18, 19, 20, 25, 26, 28, 29 (nF)	R4 R40 (kΩ)	R6, 31 R42, R67 (kΩ)	C5 C21 (nF)	C6 C27 (nF)	C7 C22 (nF)	C8 C23 (nF)	R15, 16 R51, R52 (kΩ)
Ch	Freq (Hz)											
1	86/125	0,22	1,6	56	330	12	68	100	330	100	22	390
2	135/60	1,5	15	390	33	75	470	68	220	68	15	390
3	190/230	1,0	13	270	33	56	330	47	150	47	10	390
4	280/340	0,68	5,1	180	33	39	220	33	100	330	68	82
5	420/500	0,47	3,6	120	33	27	150	22	82	220	47	82
6	580/720	0,33	2,4	82	33	18	100	15	56	150	33	82
7	860/1250	0,27	1,6	56	33	12	68	10	39	100	22	82
8	1350/1600	0,15	1,1	39	33	8,2	47	6,8	22	68	15	82
9	1900/2300	1,0	7,5	270	3,9	56	330	4,7	18	47	10	82
10	2800/3400	0,68	5,1	180	3,9	39	220	3,3	12	33	6,8	82
11	4200/5000	0,47	3,6	120	3,9	27	150	2,2	8,2	22	4,7	82
12	5800/7200	0,33	2,4	82	4,7	18	100	1,5	5,6	15	3,3	82



Rysunek 5. Rysunek montażowy płytki pasmowoprzepustowej z wyróżnionymi elementami specyficznymi dla poszczególnych kanałów – kolorem zielonym dla nieparzystych, niebieskim dla parzystych. Projekt płytki: Mike Grindle

kanałów w wokoderach stosuje się zwykle dodatkowe wzmacniacze operacyjne w układzie wzmacniacza odwracającego – tak jak w konstrukcji Richarda Beckera (ETI, wrzesień 1980). Kiedy zdałem sobie sprawę, że mogę dokonać tej modyfikacji bez stosowania żadnych dodatkowych układów, to aż się prosiło, aby je dokonać.

Montaż

Na wszystkich siedmiu płytkach filtrów („HP/LP” i „bandpass”), elementy w układach VCA i prostownika dwupołkowego są we wszystkich kanałach identyczne, więc podczas montażu płytek najlepiej umieścić je w pierwszej kolejności. Schemat montażowy z zaznaczonymi elementami wspólnymi jest pokazany na **rysunku 3**, a zdjęcie płytki obsadzonej tylko wspólnymi elementami – na **rysunku 4**. W ten sposób wstępnie montujemy sześć płytek pasmowoprzepustowych oraz płytkę „HP/LP”, również zawierającą te wspólne elementy. Następny krok to zmontowanie na płytkach pasmowoprzepustowych ośmiu głównych kondensatorów filtrujących. Są tylko cztery wartości tych kondensatorów, pokrywające wszystkie częstotliwości. Większość z nich ma wartość 33 nF. Na koniec montowane są pozostałe elementy, specyficzne dla każdej częstotliwości filtra i dla płytki „HP/LP”. Wartości tych elementów podano w **tabeli 1**. Mike Grindle, mój „człowiek od płytek”, wpadł na świetny pomysł umieszczenia na płytce w warstwie sitodrukowej białych pól, na których można potem ręcznie zapisać częstotliwości filtrów. Opis płytki zanim zaczniesz je montować!

Elementy filtrów pasmowoprzepustowych

Poniżej podano listę elementów wspólnych dla wszystkich sześciu płytek pasmowoprzepustowych. Elementy zależne od częstotliwości filtrów, specyficzne dla poszczególnych płytek i kanałów, nie są wymienione tutaj, ale znajdują się w tabeli 1. Uwaga: elementy w tej tablicy to rezystory metalizowane 1% 0,25 W oraz kondensatory foliowe 5% (najlepiej) lub 10% o rozstawie wyprowadzeń 5 mm.

Podane ilości dotyczą jednej płytki pasmowoprzepustowej – należy je pomnożyć przez liczbę stosowanych płytek (w normalnym przypadku sześć).

Półprzewodniki

IC1...IC3, IC6...IC8: TL082 podwójny wzmacniacz operacyjny z wejściem JFET; ew. odpowiednik, np. LF353, TL072; IC4: LM13700 podwójny transkonduktancyjny wzmacniacz operacyjny;

IC5, IC9: NE5532 podwójny niskoszumny wzmacniacz operacyjny; ew. odpowiednik, np. LM833; TR1, TR2: BC559C małosygnałowy tranzystor PNP (może być dowolny ogólnego przeznaczenia, z wyprowadzeniem bazy pośrodku, np. BC212); D1...D6: 1N4148 dioda małosygnałowa

Rezystory

Wszystkie 1% 0,25 W, metalizowane
R8, R9, R11, R19: 100 kΩ
R10, R46: 200 kΩ
R7, R14, R36, R43, R50, R59, R71...
R74: 100 Ω
R13, R49: 2,2 MΩ
R12, R48: 47 kΩ
R17, R53: 3,3 kΩ
R18, R54: 10 kΩ
R23, R60: 22 kΩ
R24, R61: 180 Ω
R26, R63: 10 kΩ
R25, R28, R62, R65: 12 kΩ (R25 i R62 zwykle nie należy zamontować; jeśli są, to R24 i R61 muszą zostać zmniejszone w celu skompensowania
R30, R70: 560 Ω
R34, R35: 33 Ω
R32, R33, R68, R69: 2,2 kΩ patrz panoramowanie stereo opisane dalej; 4,7 kΩ dla kanałów mono
VR1, VR2: 1 MΩ potencjometr pionowy z regulacją z boku, raster 2,54 mm

Kondensatory

C9, C24: 22 nF 20%, rozstaw 5 mm, ceramiczny lub foliowy
C14, C15: 10 μF 25V, osiowy, tantalowy lub elektrolityczny
C30: 470 nF 20% X7R, ceramiczny

Pozostałe

Piny testowe – 13 szt.
Podstawki do układów scalonych DIP8 – 8 szt.

Wtyk „goldpin”, 2,54 mm, kątowy, 2×10 pinów
Wymaganych jest również 7 gniazd „goldpin” żeńskich, 2,54 mm, prostych, 2×10-pinowych do płyty głównej (magistrali).

Częstotliwości filtrów

Częstotliwości filtrów w każdym kanale są różne, więc większość powiązanych z nimi elementów ma różne wartości. Należy zachować ostrożność i uniknąć błędów montażowych, trudnych potem do wykrycia, chyba że mamy możliwość zmierzenia charakterystyki częstotliwościowej układu. Wokoder jest „procesorem równoległym” i może działać zwodniczo dobrze z kilkoma brakującymi lub źle nastrojonymi kanałami filtrów. Wykryć błędy samym słuchaniem może być trudno. Wyznaczenie wartości elementów dla filtrów stanowiło jedną z najbardziej żmudnych czynności, jakie kiedykolwiek wykonałem przy sporządzaniu listy elementów. Musiałem się uciec do pomocy arkusza kalkulacyjnego, czego efekt stanowi tabela 1. Arkusz można pobrać ze strony Practical Electronic, marzec 2022.

Kolejnym etapem budowy jest zamontowanie elementów filtra pokazanych na schemacie montażowym pokazanym na **rysunku 5**. Rysunek jest dość gęsty, ale można pobrać jego wersję elektroniczną wraz z rysunkiem 8 ze strony Practical Electronic, marzec 2022 r. i obejrzeć lub wydrukować ją w powiększeniu.

Rysunek 6 przedstawia płytkę kompletnie zmontowaną.

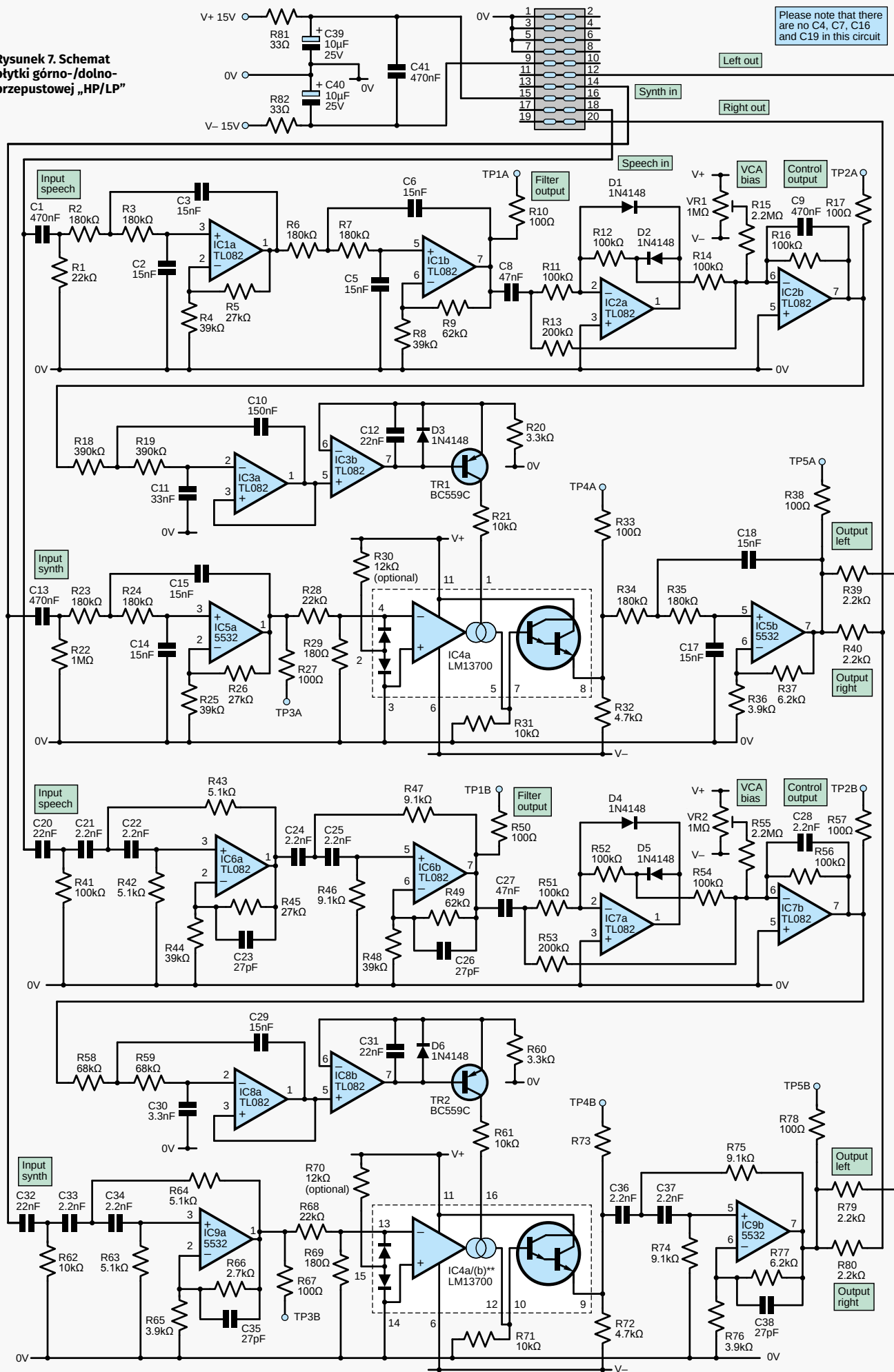
Płytki górno-/dolnoprzepustowa

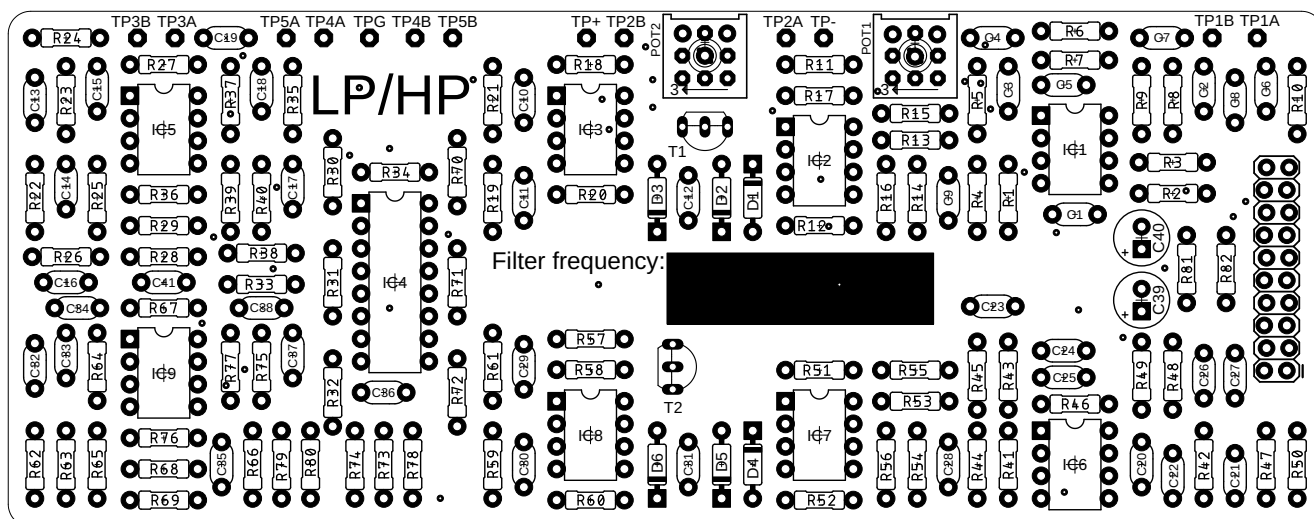
Wspólne elementy i podukłady są takie same jak w kanałach pasmowo-przepustowych, ale topologia filtrów górnoprzepustowych i dolnoprzepustowych jest inna. Filtry te nie składają się z układów pasmowo-przepustowych z wielokrotnym sprzężeniem zwrotnym, lecz ze standardowych kaskadowych



Rysunek 6. W pełni zmontowana płytka pasmowoprzepustowa. Zauważmy, że w stosunku do prototypu zmieniono położenie kilku kondensatorów

**Rysunek 7. Schemat
płytki górno-/dolno-
przepustowej „HP/LP”**





Rysunek 8. Rysunek montażowy płytki „HP/LP”. Projekt płytki: Mike Grindle

sekcji Sallena i Key'a drugiego rzędu. Schemat płytki „HP/LP” pokazano na **rysunku 7**, rysunek montażowy na **rysunku 8**, a zdjęcie zmontowanej płytki na **rysunku 9**.

Lista elementów płytki „HP/LP”

Półprzewodniki

IC1...IC3, IC6...IC8: TL082 podwójny wzmacniacz operacyjny z wejściem JFET; ew. odpowiednik, np. LF353, TL072;
IC4: LM13700 podwójny transkonduktancyjny wzmacniacz operacyjny;
IC5, IC9: NE5532 podwójny niskoszumny wzmacniacz operacyjny; ew. odpowiednik, np. LM833;
TR1, TR2: BC559C małosygnałowy tranzystor PNP (może być dowolny ogólnego przeznaczenia, z wyprowadzeniem bazy pośrodku, np. BC212);
D1...D3: 1N4148 dioda małosygnałowa

Kondensatory

Wszystkie foliowe, 5% (najlepiej) lub 10%, rozstaw wyprowadzeń 5 mm – chyba że oznaczono je symbolem *
C1, C8, C9, C13: 470 nF

C8, C27: 47 nF
C2, C3, C5, C6, C14, C15, C17, C18, C29: 15 nF
C10: 150 nF
C11: 33 nF
C12, C20, C31, C32: 22 nF 20%, ceramiczny lub foliowy *
C30: 3,3 nF
C21, C22, C24, C25, C28, C33, C34, C36, C37: 2,2 nF
C23, C26, C35, C38: 27 pF, ceramiczny *
C39, C40: 10 µF 25V, tantalowy lub elektrolityczny *
C41: 470 nF 20% X7R, ceramiczny *

Rezystory

Wszystkie 1% 0,25 W, metalizowane
R1, R22: 1 MΩ
R2, R3, R6, R7, R23, R24, R34, R35: 180 kΩ
R4, R8, R25, R44, R48: 39 kΩ
R5, R26, R45: 27 kΩ
R9, R49: 62 kΩ
R10, R17, R27, R33, R38, R50, R57, R67, R73, R78: 100 Ω
R11, R12, R14, R16, R51, R52, R54, R56, R41: 100 kΩ
R13, R53: 200 kΩ

R15, R55: 2,2 MΩ
R18, R19: 390 kΩ
R20, R60: 3,3 kΩ
R30, R70: 12 kΩ (zwykle nie należy montować)
R21, R31, R61, R62, R63, R71: 10 kΩ
R32, R72: 4,7 kΩ
R28, R68: 22 kΩ
R29, R69: 180 Ω
R36, R65, R76: 3,9 kΩ
R37, R77: 6,2 kΩ
R39, R40, R79, R80: 2,2 kΩ
R42, R43, R63, R64: 5,1 kΩ
R46, R47, R74, R75: 9,1 kΩ
R58, R59: 68 kΩ
R66: 2,7 kΩ
R81, R82: 33 Ω
VR1, VR2: 1 MΩ potencjometr dostrojczy pionowy z regulacją z boku, raster 2,54 mm

Nie ma dymu bez ognia

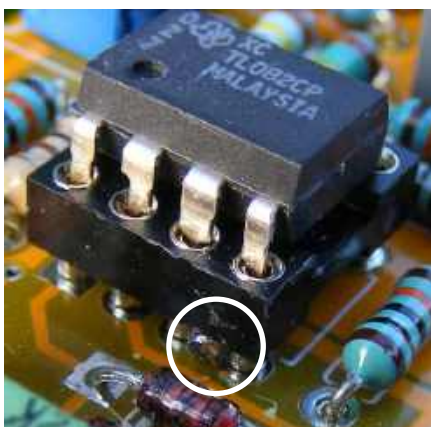
Rezystory doprowadzające napięcia zasilania (R34/R35, R81/R82) filtrują zakłócenia, ale również zapewniają ochronę w przypadku zwarcia na płycie. Niebezpieczeństwo zrobienia przypadkowego zwarcia podczas testowania zachodzi zawsze. W wyniku zwarcia



Rysunek 9. W pełni zmontowana płytka „HP/LP”. Jest to wersja wcześniejsza z dodanymi kilkoma elementami do celów eksperymentalnych. Wszystkie płytki sprzedawane przez Practical Electronics są dokładnie zgodne z rysunkami montażowymi



Rysunek 10. Rezystory R34 i R35 doprowadzające zasilanie powinny być montowane ponad płytką, bo w razie awarii mogą się mocno nagrzewać



Rysunek 11. Zwarcia mogą wystąpić w najdziwniejszych miejscach. Tu widzimy rekonstrukcję przypadku, jaki wystąpił na jednej z płytek. W rzeczywistości było gorzej – kulka cyny była ukryta za pinami, więc nie można jej było zobaczyć. Do jej znalezienia trzeba było użyć miliomierza Model 1000 Tracer

wspomniane rezystory mogą ulec spaleniu, więc powinny być one montowane kilka milimetrów ponad płytką (rysunek10). Spalone rezystory kosztują grosze, spalona płytka drukowana o wiele więcej!

Na jednej z płytek mała kulka lutownia utknęła pod podstawką układu scalonego IC2 (rysunek 11), zwierając ujemną linię zasilania do masy (piny 3 i 4). Znalezienie tego zwarcia zajęło wieki. Pierwszy raz widziałem coś takiego.

Panoramowanie stereo

Na płytkach pasmowoprzepustowych w każdym kanale znajdują się po dwa rezystory, wypuszczające sygnał do lewej (R32 i R68) i prawej (R33 i R69) szyny sumującej. Normalnie montowane są tylko

R32 i R69. Kanały są zatem doprowadzane naprzemiennie do lewej i prawej szyny sumującej, jak pokazano na **rysunku 12**. W przypadku kanałów monofonicznych (płytki „HP/LP” i oba kanały na płytce pasmowoprzepustowej o najniższej częstotliwości) montowane są wszystkie cztery rezystory (na płytce „HP/LP” rezystory te to R39, R40, R79 i R80). W kanałach mono, aby utrzymać te same poziomy względne co na płytce pasmowoprzepustowej, rezystory sumujące zostały zwiększone z 2,2 kΩ do 4,3 kΩ, jak pokazano na **rysunku 13**. Zauważmy, że w filtrach dolno- i górnoprzepustowym pozostawiono wartości rezystorów 2,2 kΩ, ponieważ filtry te mają niższe wzmocnienie niż filtry pasmowoprzepustowe.

Płytki magistrali

Wszystkie płytki filtrów są podłączone do płytki magistrali, która przypomina płytkę uniwersalną, tyle, że nie odpadają na niej ścieżki (**rysunek 14**). Zauważmy, że złącza są dwurzędowe, co zmniejsza rezystancję styku i zwiększa stabilność mechaniczną (**rysunek 15**). Na tej płytce są również umieszczone dodatkowe złącza, doprowadzające zasilanie i łączące płytkę ze wzmacniaczem sterującym i płytkami miksera opisanymi w zeszłym miesiącu. Przewidziano także dodatkowe wejścia do szyn sumujących. W razie potrzeby można przez nie dołączyć nieobrobione sygnały z kanału mikrofonowego i syntezatora.

Testowanie krok po kroku

Urządzenie testuj tak samo jak w przypadku każdej innej elektroniki – małymi

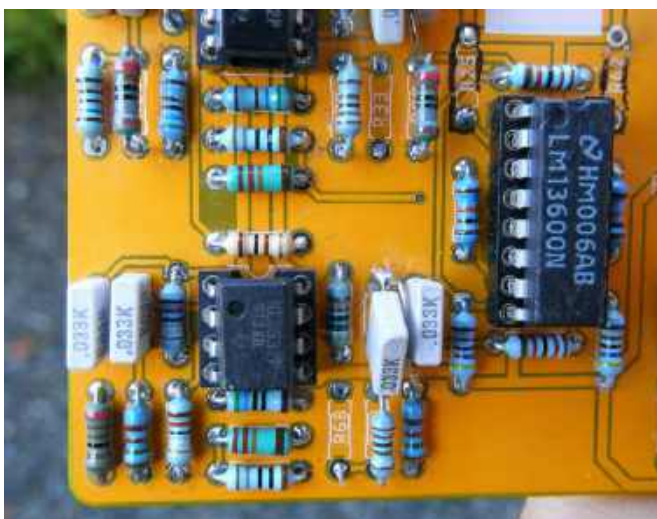
kroczkami, po jednej płytce na raz. Jeśli będziesz uruchamiał od razu całość, na pewno poniesiesz porażkę. Każdą płytkę należy najpierw sprawdzić wzrokowo, czy wszystkie elementy o istotnej kolejności wyprowadzeń, takie jak wzmacniacze operacyjne i kondensatory elektrolityczne, są właściwie zorientowane. Następnie podłączamy zasilanie i sprawdzamy, czy R34 i R35 nie nagrzewają się z powodu nadmiernego prądu. Dalej kontrolujemy napięcia stałe na wyjściach wzmacniaczy operacyjnych (piny testowe TP1A i TP1B przy górnej krawędzi płytki). W szereg z pinami są włączone rezystory 100 Ω, aby zapobiec oscylacjom po podłączeniu długich przewodów testowych.

Jeśli wszystko gra, to jesteśmy gotowi do testów sygnałowych. Opiszemy je w przyszłym miesiącu.

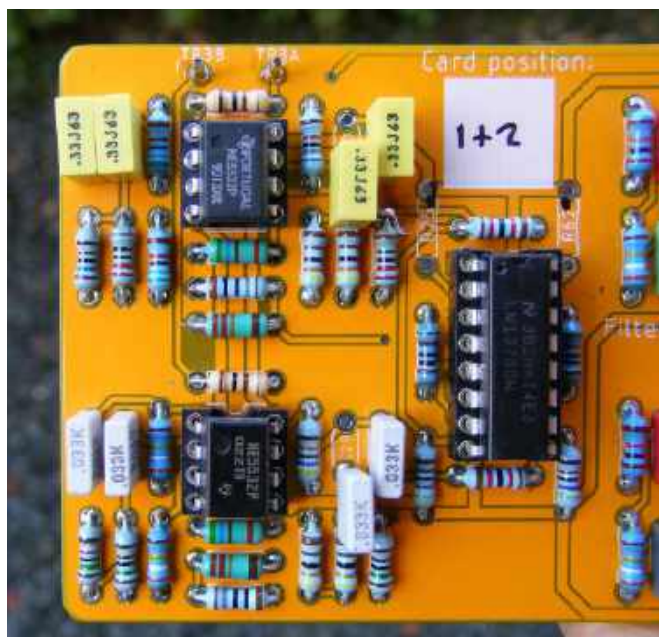
Strojenie

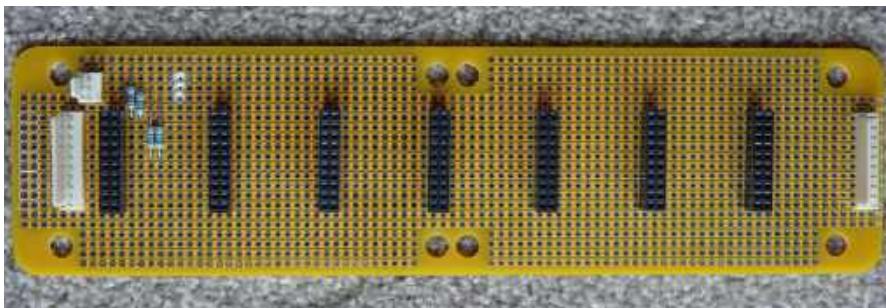
Wokodery droższe – takie jak wspomniany projekt Richarda Beckera – mogą mieć np. pięć elementów regulacyjnych na kanał. Ja ich ilość zredukowałem do absolutnego minimum, czyli do jednego – ustawiającego wstępne napięcie sterujące wzmacniacza VCA w danym kanale. VR1 i VR2 są nastawiane tak, żeby przy braku sygnału modulującego sygnał nośnej nie był słyszalny, ale żeby odpowiedni VCA znajdował się tuż przed otwarciem. Takie ustawienie zapewni optymalną liniość i najlepszą wyrazistość.

Wartość dobroci (Q) filtrów pasmowoprzepustowych może wykazywać znaczne odchyłki w zależności od tolerancji kondensatorów. Dla filtru z wielokrotnym sprzężeniem zwrotnym teoria zakłada, że oba kondensatory



Rysunek 12 (powyżej). Aby uzyskać efekt stereofoniczny, kanały są panoramowane naprzemiennie do lewej i prawej strony. Uzyskano to przez pominięcie rezystorów R33 i R68
Rysunek 13 (po prawej). W kanałach monofonicznych montowane są wszystkie cztery rezystory, a ich wartość jest zwiększona do 4,3 kΩ





Rysunek 14. Płytkę magistrali. Zauważmy, że do szyn sumujących dołączono kilka rezystorów, co umożliwia domiksowanie sygnałów z wejść dodatkowych

określające częstotliwość mają identyczne wartości. W przypadku użycia tanich kondensatorów o tolerancji 10% ich wartości mogą się jednak różnić między sobą nawet o 20%. W takim przypadku wzmocnienie danego kanału nadmiernie wzrośnie. Jeśli stwierdzimy, że w miksie wyjściowym jakieś częstotliwości dominują, powinniśmy zwiększyć wartości rezystorów wejściowych filtrów R1 i R19 lub R37 i R55. Poziomy sygnału można sprawdzać w punktach testowych TP1A i TP1B. Jeśli mielibyśmy w torze filtru dać dodatkowe elementy regulacyjne, to przede wszystkim właśnie do ustawiania poziomu. Innym środkiem jest po prostu użycie w filtrach kondensatorów polistyrenowych o wąskiej tolerancji. Są one często dostępne tylko w formie osiowej,

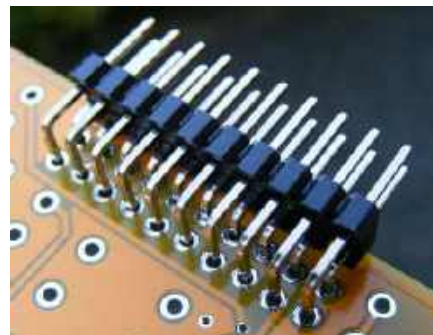
a to nieco komplikuje ich zamontowanie. Przykład widzimy na rysunku 16.

Za miesiąc

W następnym odcinku omówimy pozostałe elementy konstrukcji oraz testy. Przedstawię również zasilacz trójnapięciowy o bardzo niskim poziomie tętnień, nadający się nie tylko do wokodera, ale i innych układów audio, jak np. miksery, które zawierają wiele wzmacniaczy operacyjnych i/lub wymagają zasilania phantom. ■

Jake Rothman

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, marzec 2022 (www.epemag3.com)



Rysunek 15. Wtyki łączące się dwurzędowo. Zwiększa to sztywność mechaniczną. Masa jest łączona przez 8 pinów łączących, co skutkuje bardzo małą impedancją



Rysunek 16. Na płytkach filtrów jest wystarczająco dużo miejsca, aby dać kondensatory polistyrenowe o tolerancji 1% – o ile uda się je zdobyć. Do filtrów elektroakustycznych doskonale nadawały się kondensatory Philips z serii 424, ale zostały wycofane z produkcji, ponieważ ich dielektryk zawierał neurotoksyczny ołów



Generatory wysokiego napięcia

Rozwiązanie znajdziesz na www.elportal.pl/quizy

1. Dlaczego nie powinno się pracować z generatorami wysokiego napięcia w pobliżu sprzętu pomiarowego?

- ☐ wyładowania mogą uszkodzić instrumenty pomiarowe;
- ☐ wyładowania generują silne fale radiowe, które mogą skasować pamięci Flash i EEPROM;
- ☐ bo tworzywa odbarwiają się z powodu ozonu.

2. Przebicie powstaje wtedy, gdy:

- ☐ pod wpływem pola elektrycznego elektrony są rozpędzone tak bardzo, że „przebijają się” przez materiał;
- ☐ pod wpływem pola elektrycznego materiał zostaje zjonizowany i spontanicznie staje się przewodnikiem;
- ☐ pod wpływem pola elektrycznego elektrony w atomach materiału są „wyrwane” i tworzą kanał przepływu dla prądu.

3. Townsend wyjaśnił, jak:

- ☐ pojedyncze wolne elektrony w powietrzu są przyspieszane w polu elektrycznym;
- ☐ pojedyncze wolne elektrony w powietrzu są przyspieszane w polu elektrycznym, i zderzając się z neutralnymi atomami wybijają więcej wolnych elektronów;
- ☐ jak pod wpływem wyładowania powstaje ozon.

4. Wytwarzanie ładunków elektrostatycznych przez pocieranie nosi nazwę efektu:

- ☐ piezoelektrycznego;
- ☐ termoelektrycznego;
- ☐ tryboelektrycznego.

5. Teflon szczególnie dobrze nadaje się do budowy generatorów Van de Graaffa gdyż:

- ☐ atomy fluoru na jego powierzchni są wyjątkowo elektroujemne;

- ☐ bo ma niski współczynnik tarcia;
- ☐ bo elektrony do niego nie przywierają.

6. Obwód używający diod i kondensatorów do powielania napięcia to generator:

- ☐ Graetza;
- ☐ Cockcrofta-Waltona;
- ☐ Greinachera

7. Obwód używający rezystorów, kondensatorów i iskierników do generowania wysokich napięć to generator:

- ☐ Marxa;
- ☐ Marksa;
- ☐ Markowa;

8. Cewka Tesli zbudowana zgodnie z oryginalnym projektem Nicolii Tesli nosi skrócone oznaczenie:

- ☐ OLTC;
- ☐ VTTC;
- ☐ SGTC.

9. Cewka Tesli zbudowana na radzieckiej pentodzie mocy GU-81M nosiłaby oznaczenie:

- ☐ OLTC;
- ☐ VTTC;
- ☐ SGTC.

10. Układ cewki Tesli z jednym tranzystorem pokazany w artykule nosi nazwę:

- ☐ Slayer Exciter;
- ☐ Sodom Oscillator;
- ☐ Sabbath Resonator.



Chirurgia obwodowa

Wzmacniacze operacyjne logarytmiczne i wykładnicze, część 2

W zeszłym miesiącu zaczęliśmy się przyglądać wzmacniaczom logarytmicznym i wykładniczym (zwanym również antylogarytmicznymi), opartym na wzmacniaczach operacyjnych. Koncentrowaliśmy się na układach logarytmicznych. W tym miesiącu przyszła kolej na wzmacniacze wykładnicze (od Red EdW: Dla zachowania zgodności z ilustracjami zachowano anglosaską nomenklaturę oznaczania napięcia za pomocą litery V).

Potęgowanie

Funkcje wykładnicze są przeciwieństwem (funkcjami odwrotnymi) do logarytmów, stąd termin „antylogarytm”. Funkcja wykładnicza oznacza podnoszenie jednej liczby (podstawy) do potęgi drugiej liczby (wykładnika), na przykład a do potęgi y , co zapisujemy a^y . Jak wspomnieliśmy w zeszłym miesiącu: jeśli $y = \log_{10}(x)$ (logarytm o podstawie 10), to na podstawie wartości y możemy znaleźć wartość x ze wzoru $x = 10^y$. Często używane są tzw. logarytmy naturalne o podstawie $e \approx 2,71828$. Odpowiednia funkcja wykładnicza nazywana jest eksponentą i oznaczana e^y lub $\exp(y)$: jeśli $y = \ln(x)$ to $x = e^y$. Od Red. EdW: funkcja wykładnicza przy podstawie e występuje niezwykle często we wzorach opisujących różne zjawiska fizyczne; przypomnijmy chociażby wzór na napięcie na kondensatorze C rozładowywanym przez rezystor R :

$$u(t) = u(0) \cdot e^{-\frac{t}{RC}}$$

Zauważyliśmy również, że możemy łatwo zmieniać podstawę logarytmu poprzez współczynnik skalowania, na przykład:

$$\log_{10}(x) = \frac{\ln(x)}{\ln(10)} \approx \frac{\ln(x)}{2,303}$$

Wykorzystujemy ten fakt np. wtedy, gdy potrzebujemy układu elektronicznego realizującego funkcję 10^x albo $\log_{10}(x)$, podczas gdy wzmacniacze logarytmiczne i antylogarytmiczne są z natury oparte na zależności typu e^x między napięciem przewodzenia

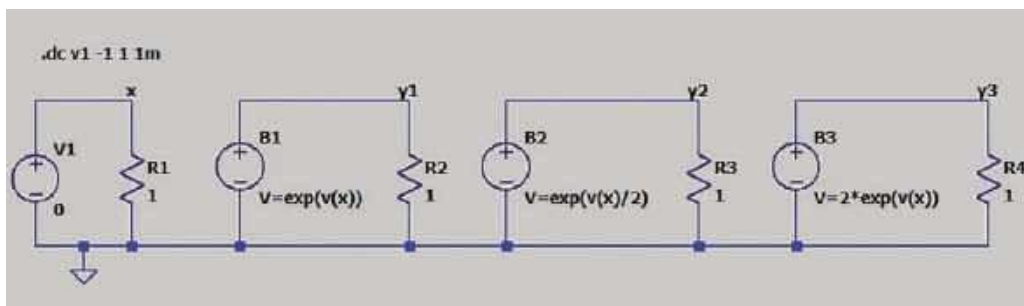


Rysunek 1. Idealne wykładnicze (antylogarytmiczne) zależności wejście-wyjście: funkcja nieskalowana (/zielony), wejście skalowane przez 0,5 (/żółty) i wyjście skalowane przez 2 (/czerwony)

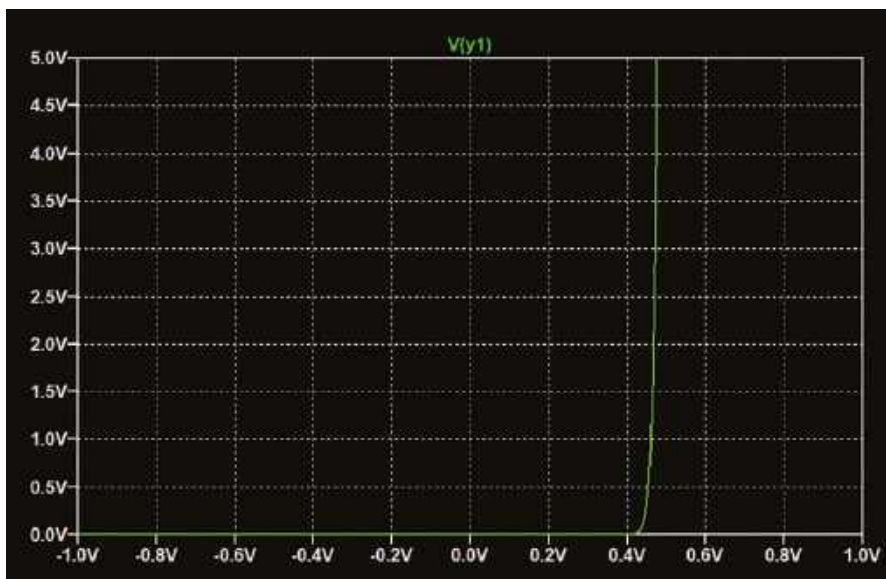
złącza półprzewodnikowego a przepływającym przez nie prądem. Korzystamy wtedy z zależności:

$$a^x = e^{x \cdot \ln(a)}$$

Możemy więc użyć układu opartego na funkcji wykładniczej e^x , żeby uzyskać



Rysunek 2. Schemat w LTspice do wykresu na rysunku 1



Rysunek 3. Funkcja wykładnicza może, w zależności od zakresu obserwacji, wykazywać zmiany stosunkowo powolne (rysunek 1) lub niezwykle szybkie (jak tutaj)

funkcję związaną z potęgowaniem dowolnej (w granicach rozsądku) podstawy (na przykład 2^x lub 10^x).

Synteзаторы analogowe

Jak wspominaliśmy w zeszłym miesiącu, wzmacniacze wykładnicze mogą być wykorzystywane w analogowych układach mnożących. Korzystamy tu z zależności

$$x \cdot y = e^{\ln(x) + \ln(y)}$$

czyli: logarytmujemy napięcia wejściowe x i y , dodajemy do siebie logarytmy, a następnie tworzymy funkcję wykładniczą tej sumy. Jest to jeden ze sposobów mnożenia napięć analogowych – choć niekoniecznie najlepszy.

Wzmacniacze wykładnicze mają ogólnie mniejszy zakres zastosowań niż wzmacniacze logarytmiczne. Jednym z interesujących przykładów ich zastosowań są analogowe synteзаторы muzyczne. Niniejszy cykl artykułów został w istocie zainspirowany wykorzystaniem wzmacniaczy wykładniczych w projekcie MIDI Ultimate Synthesiser, który był prezentowany na łamach „Practical Electronics” w 2019 roku.

W systemach takich jak MIDI Ultimate, do ustawienia częstotliwości odtwarzanej nuty służy napięcie sterujące zmieniające się liniowo. Ale częstotliwości nut w skali muzycznej tworzą postęp geometryczny, to znaczy częstotliwość każdej nuty jest równa częstotliwości nuty poprzedniej, pomnożonej przez pewien stały współczynnik a . Istnieje częstotliwość odniesienia f_0 , wynosząca najczęściej 440 Hz, odpowiadająca nucie A w oktawie razkreślnej. Kolejne nuty (A#, H, C...) mają częstotliwości $f_0 \cdot a$,

$f_0 \cdot a^2, f_0 \cdot a^3$ i tak dalej. Ogólnie: n -ta nuta w górę od nuty odniesienia ma częstotliwość $f_0 \cdot a^n$.

Napięcie sterujące w syntezorze jest liniowo zależne od numeru nuty, reprezentuje więc n . Chcemy, aby generator syntezatora wytworzył częstotliwość $f_0 \cdot a^n$. Generatory wytwarzają jednak częstotliwość zależną liniowo od napięcia sterującego. Musimy zatem przekonwertować im napięcie sterujące z proporcjonalnego do n na proporcjonalne do a^n .

Zachodnia skala muzyczna dzieli się na oktawy (stosunek częstotliwości nut odległych o oktawę wynosi 2), a każda oktawa składa się z 12 półtonów. Zatem współczynnik a wynosi

$$a = 2^{\frac{1}{12}} \approx 1,05946$$

Zależność wejście-wyjście, wymagana w konwerterze napięcia sterującego, uzyskamy, stosując wzmacniacz wykładniczy oparty na złączu półprzewodnikowym i używając odpowiedniego skalowania.

Funkcje wykładnicze

Tak jak w przypadku wzmacniacza logarytmicznego, zaczniemy od przyjrzenia się samej funkcji wykładniczej. Rysunek 1 przedstawia trzy odpowiedzi idealnych wzmacniaczy wykładniczych na napięcie wejściowe x , gdzie:

$$y1 = \exp(x)$$

$$y2 = \exp\left(\frac{x}{2}\right)$$

$$y3 = 2 \cdot \exp(x)$$

Wykresy te, podobnie jak w przypadku wyidealizowanych krzywych logarytmicznych z ubiegłego miesiąca, zostały utworzone w programie LTspice przy użyciu behawioralnych źródeł napięcia i symulacji przemiatania DC, jak pokazano na rysunku 2. Krzywe

te bardziej reprezentują idealne funkcje matematyczne niż prawdziwe odpowiedzi rzeczywistego układu elektronicznego. Kształt tych krzywych ilustruje ogólne zachowanie odpowiedzi wykładniczej. Wraz ze wzrostem napięcia wejściowego tempo wzrostu wyjścia rośnie. Wynika z tego, że krzywa wejście/wyjście jest płaska dla małych napięć wejściowych, a staje się stroma dla dużych. Od Red. EdW: stykamy się tu z niezwykle ważną cechą funkcji wykładniczej, szczególnie funkcji przy podstawie e (eksponens). Funkcja ta w każdym punkcie rośnie z prędkością dokładnie taką jak wartość funkcji w tym punkcie – co matematycznie zapisujemy jako

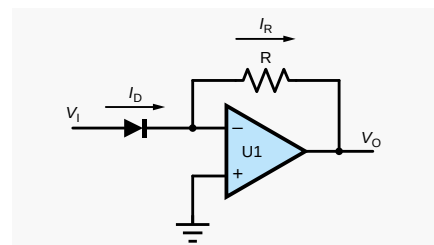
$$\frac{d}{dx} \cdot e^x = e^x$$

Ta prosta własność sprawia, że od funkcji wykładniczej aż roi się w przyrodzie...

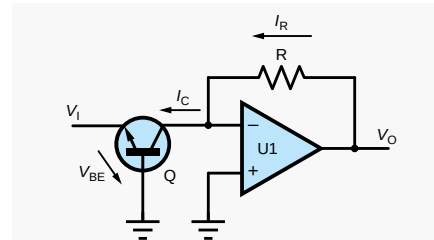
Jeśli nie skalujemy wyjścia (funkcje $y1$ i $y2$), to krzywe na rysunku 1 dla wejścia równego 0 przyjmują wartość 1. Wynika to stąd, że dla dowolnej liczby (nierównej zero) $n^0=1$, więc również $\exp(0)=e^0=1$. Skalowanie wyjścia przez K skutkuje tym, że dla wejścia równego 0 wyjście wynosi K (patrz rysunek 1, krzywa $y3$, gdzie $K=2$). Skalowanie wejścia x w dół sprawia, że tempo zmian jest wolniejsze (patrz rysunek 1, krzywa $y2$, w której x jest przeskalowane przez 0,5). I odwrotnie – skalowanie wejścia w górę daje szybsze tempo zmian. Skalowanie x ma efekt taki sam jak zmiana podstawy potęgowania.

Używanie i nadużywanie

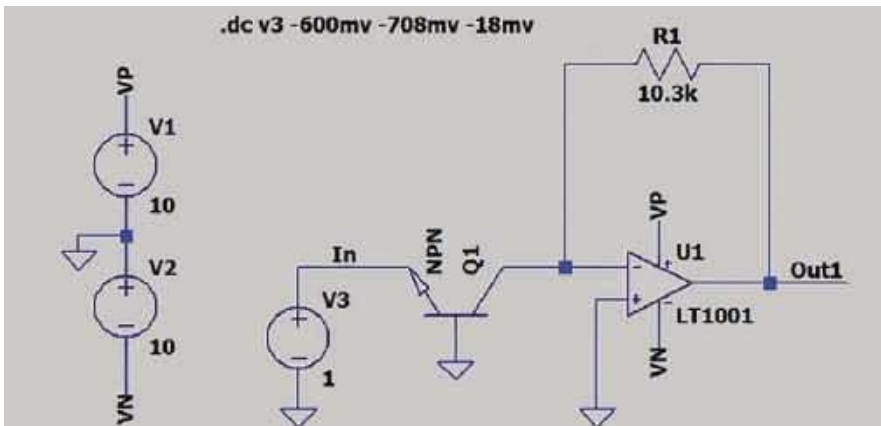
Określenie „wykładniczy” jest często nadużywane – interpretowane błędnie



Rysunek 4. Wykładniczy wzmacniacz napięciowy oparty na diodzie i wzmacniaczu operacyjnym



Rysunek 5. Wykładniczy wzmacniacz napięciowy oparty na tranzystorze bipolarnym i wzmacniaczu operacyjnym



Rysunek 6. Schemat układu z rysunku 5, symulowanego w LTspice w celu zilustrowania podwajania napięcia wyjściowego przy zmianie napięcia wejściowego o 18 mV

jako po prostu „szybko rosnący”, podczas gdy prawidłowo odnosi się ono do sytuacji, w której tempo zmian pewnej wielkości jest proporcjonalne do samej wielkości. Termin „wykładniczy” zaczął być chyba lepiej rozumiany dopiero podczas pandemii covid-19 za sprawą faktu, że infekcje wirusowe nasilają się wykładniczo.

Zależność wykładnicza dla małych argumentów zmienia się powoli, i czasem łatwo jest przeoczyć fakt, że tempo jej wzrostu znacznie rośnie dla argumentów większych. Czysta funkcja wykładnicza dla dostatecznie dużych argumentów daje do wolnie duże wyniki. Natomiast w systemach rzeczywistych „zjawisko wykładnicze” osiąga zawsze pewną fizyczną granicę, poza którą dalszy wzrost nie jest możliwy. W przypadku rozwoju wirusa, jeśli nie uwzględnimy zjawiska odporności, wzrost zatrzyma się, gdy zostanie zainfekowana cała populacja. W przypadku wzmacniaczy wykładniczych, maksymalną wartość wyjściową określa napięcie zasilania lub inne ograniczenie układu. Rysunek 3 przedstawia funkcję wykładniczą $y=10^{-23} \cdot \exp(100 \cdot x)$, wykreśloną w tej samej skali co rysunek 1. W odróżnieniu od przykładów z rysunku 1 funkcja ta przy zakresie wejściowym $-1 \text{ V} \dots +1 \text{ V}$ daje bardzo duży zakres wyjściowy. Dla $x=1$ wartość funkcji wynosi $y=2,7 \cdot 10^{23}$.

Równanie prądu diody

Jak już wspominaliśmy, wzmacniacze wykładnicze mogą wykorzystywać wykładniczą zależność prądowo-napięciową diody półprzewodnikowej. Zależność ta ma postać:

$$I_D = I_S \cdot \left[\exp\left(\frac{V_D}{V_T}\right) - 1 \right]$$

V_D to napięcie na diodzie, a I_D to prąd przez nią płynący. I_S to tzw. prąd nasycenia diody – parametr specyficzny dla konkretnej diody lub tranzystora. V_T to tzw. napięcie

termiczne (zdefiniowane w ubiegłym miesiącu). Równanie zawiera składnik -1 , który przy większych napięciach przewodzenia V_D (już od $0,2 \dots 0,3 \text{ V}$) może zostać pominięty, bo wtedy dominujący staje się składnik $\exp()$ (wynosi co najmniej kilka tysięcy). Zależność przybierze w takim przypadku uproszczoną postać:

$$I_D \approx I_S \cdot \exp\left(\frac{V_D}{V_T}\right)$$

Wersję pełną równania obowiązkowo stosujemy dla małych napięć przewodzenia, a także dla napięć wstecznych (ujemnych). Jak wiadomo, $\exp(0)=1$. Składnik -1 w pełnym wzorze zapewnia, że dla napięcia $V_D=0$ prąd diody I_D wyniesie również 0.

Wzmacniacz wykładniczy

Możemy uzyskać prąd, który jest wykładniczo związany z napięciem, przykładając po prostu napięcie do diody. Prąd ten można z kolei przekształcić w napięcie, podając go do wzmacniacza transimpedancyjnego (wzmacniacz z wejściem prądowym i wyjściem napięciowym). Taki układ pokazano na rysunku 4, na którym

wzmacniacz operacyjny i rezystor tworzą wzmacniacz transimpedancyjny o wzmocnieniu R (woltów na amper, V/A). W tym układzie wejście odwracające wzmacniacza operacyjnego działa jak „wirtualna masa” (podobnie jak w przypadku wzmacniacza logarytmicznego omówionego w zeszłym miesiącu), zatem napięcie na diodzie jest równe napięciu wyjściowemu, czyli $V_D=V_r$. Prąd I_D wynika z przytoczonego wcześniej równania prądu diody. Rezystor R jest podłączony między wirtualną masą a wyjściem, więc napięcie na nim jest równe napięciu wyjściowemu ze znakiem minus ($-V_o$), a z prawa Ohma prąd rezystora wynosi

$$I_R = -\frac{V_o}{R}$$

Zakładamy, że wzmacniacz operacyjny jest idealny, więc do wejść wzmacniacza operacyjnego nie wpływa żaden prąd. Oznacza to, że cały prąd diody musi przepływać przez rezystor, czyli $I_R=I_D$. Jeśli w tym równaniu zastąpimy I_D przez wzór na prąd diody, a I_R wyrazimy jako $-\frac{V_o}{R}$, to po prostych przekształceniach otrzymamy:

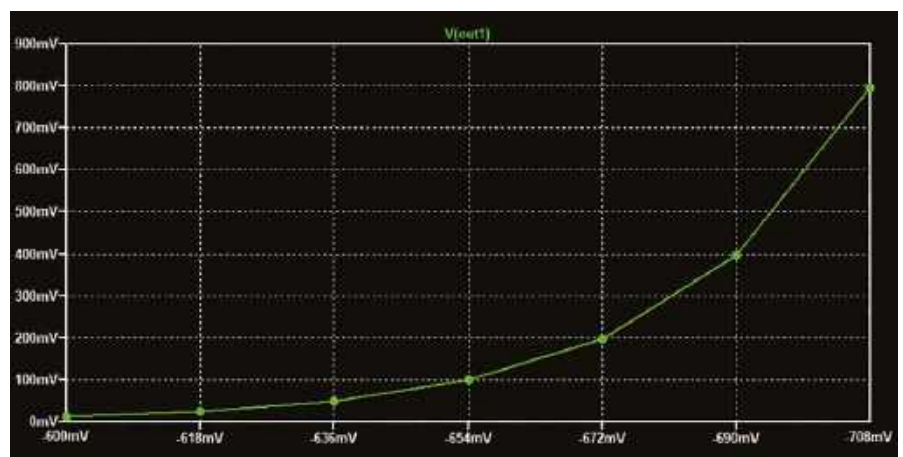
$$V_o = -I_S \cdot R \cdot \left[\exp\left(\frac{V_i}{V_T}\right) - 1 \right]$$

Niniejsze równanie opisuje wzmacniacz wykładniczy. Wejście jest skalowane przez $\frac{1}{V_T}$, a wyjście przez $-I_S \cdot R$. Charakterystyka odbiega od idealnej wykładniczej; napięcie wyjściowe zawiera stałe przesunięcie (offset) o wielkości $I_S \cdot R$, co sprawia, że gdy $V_i=0$, to $V_o=0$.

V_T i I_S zależą od temperatury, a więc zależą od niej również napięcie wyjściowe V_o wzmacniacza – tak jak w przypadku wzmacniacza logarytmicznego. Więcej na ten temat wkrótce.

Tranzystorowy wzmacniacz wykładniczy

Podobnie jak w przypadku wzmacniaczy logarytmicznych, w roli złącza półprzewodnikowego często stosuje się tranzystor



Rysunek 7. Wyniki symulacji LTspice układu z rysunku 5

zamiast diody. Podstawowy układ pokazano na **rysunku 5**. Wejściem układu jest złącze baza-emiter tranzystora, a do wzmacniacza transkonduktancyjnego wpływa prąd kolektora. Baza tranzystora znajduje się na potencjale masy, a wejście układu jest dołączone do emitera. Można sobie wyobrazić inną konfigurację, istotne jest tylko, by napięcie baza-emiter V_{BE} było sterowane przez napięcie wejściowe. Jeśli jest użyty tranzystor NPN, jak na rysunku, wówczas napięcie wejściowe V_i musi być ujemne; wtedy napięcie V_{BE} jest dodatnie i tranzystor przewodzi. Układ jest odwracający (podobnie jak ten z rysunku 4) – napięcie wejściowe jest ujemne, a napięcie wyjściowe dodatnie. Można użyć tranzystora PNP, wtedy sytuacja stanie się odwrotna: V_i będzie dodatnie, V_o a ujemne.

W zastosowaniach muzycznych musimy wiedzieć, jaka zmiana napięcia na wejściu układu wykładniczego spowoduje podwojenie napięcia wyjściowego (co w generatorze da zmianę częstotliwości o oktawę). Przyjmijmy, że dwukrotną zmianę napięcia na wyjściu powoduje zmiana napięcia wejściowego z V_{BE1} na V_{BE2} . Stosunek prądów płynących przez diodę w każdym z tych przypadków wynosi:

$$\frac{I_{D2}}{I_{D1}} = \frac{I_S \cdot \exp\left(\frac{V_{BE2}}{V_T}\right)}{I_S \cdot \exp\left(\frac{V_{BE1}}{V_T}\right)} = 2$$

Po skróceniu czynnika I_S i uporządkowaniu otrzymujemy:

$$\exp\left(\frac{V_{BE2}}{V_T} - \frac{V_{BE1}}{V_T}\right) = \exp\left(\frac{V_{BE2} - V_{BE1}}{V_T}\right) = 2$$

Biorąc logarytmy naturalne z obu stron i przekształcając dostajemy:

$$V_{BE2} - V_{BE1} = V_T \cdot \ln(2)$$

V_T w temperaturze $+27^\circ\text{C}$ wynosi około 26 mV, więc $V_T \cdot \ln(2)$ to około 18 mV. W celu uzyskania wykładniczego napięcia sterującego dla generatora, na złączu baza-emiter tranzystora potrzebne jest napięcie zmieniające się o 18 mV na oktawę. Aby otrzymać wygodniejszy współczynnik 1 V na oktawę (właśnie taki jest zwykle używany), należy w jakimś układzie przeskalować napięcie wejściowe.

Krok 18 mV, powodujący podwajanie się napięcia wyjściowego, można zweryfikować w symulacji LTSpice w układzie pokazanym na **rysunku 6**. Rezystor R1 jest dobrany tak, aby zapewnić wartości napięć wyjściowych dogodne do zaobserwowania, na przykład 100, 200, 400 i 800 mV. Symulacja zmienia napięcie wejściowe w krokach co 18 mV. Wynik, pokazany na **rysunku 7**, pokazuje, że przy każdej zmianie wejścia o 18 mV napięcie wyjściowe rzeczywiście się podwaja. Wykres jest linią łamaną ze względu na małą ilość punktów pomiarowych. Wykonanie symulacji

ze znacznie mniejszym krokiem wejściowym dałoby w wyniku krzywą gładką.

Wpływ temperatury i kompensacja termiczna

Zarówno tutaj jak i przy opisywaniu wzmacniacza logarytmicznego wspominaliśmy, że podstawowe wersje omawianych układów są bardzo wrażliwe na temperaturę. Stwierdziliśmy, że zmiana napięcia wejściowego, wymagana do dwukrotnego zwiększenia lub zmniejszenia napięcia wyjściowego, wynosi $V_T \cdot \ln(2)$ czyli 18 mV przy $+27^\circ\text{C}$. Zależy ona od V_T (na szczęście nie zależy od I_S). V_T zmienia się wraz z temperaturą liniowo, więc stosunkowo łatwo jest skompensować ten efekt. Wystarczy w wejściowym wzmacniaczu skalującym dać, jako rezystor określający wzmocnienie, termistor o odpowiednim współczynniku temperaturowym.

Główny problem polega jednak na tym, że **bezwzględny** prąd tranzystora zależy zarówno od I_S , jak i V_T , a tym samym jest zależny od temperatury w dość skomplikowany sposób. Jednym ze sposobów rozwiązania tego problemu jest wprowadzenie pewnego potencjału odniesienia, zależnego od temperatury, i takie podłączenie tranzystora, by przesunąć jego napięcie V_{BE} o ten potencjał. Na tym pomysły jest oparty układ z **rysunku 8**. Do realizacji funkcji wykładniczej służy tranzystor Q1. Jego napięcie baza-emiter równe jest napięciu baza-emiter tranzystora Q2 plus napięcie wejściowe V_{IN} . Jeśli napięcie V_{IN} wynosi 0 V, wówczas oba tranzystory będą miały takie samo V_{BE} , a przez tranzystor Q1 popłynie prąd taki sam, jaki płynie przez Q2. Podanie niezerowego napięcia wejściowego spowoduje wykładniczą zmianę V_o . Napięcie V_{IN} równe $+18$ mV zwiększy V_{BE} tranzystora Q1 o 18 mV, a V_o wzrośnie wtedy dwa razy. V_{IN} wynoszące -18 mV spowoduje dwukrotne zmniejszenie V_o . **Od Red. EdW: w oryginalnym tekście mylnie podano, że V_{IN} równe $+18$ mV spowoduje dwukrotne zmniejszenie, a -18 mV – dwukrotne zwiększenie napięcia V_o .**

Prąd tranzystora Q2 to „prąd odniesienia”. Jest on utrzymywany na stałym poziomie przez pętlę sprzężenia zwrotnego wzmacniacza operacyjnego U2. Sprzężenie zwrotne powoduje, że różnica napięć między wejściami wzmacniacza operacyjnego jest sprowadzana do zera. Wskutek tego kolektor Q2 i prawa

końcówka R2 mają potencjał 0 V. Sprawia to, że prąd płynący przez R2 jest stały i wynosi $\frac{V_P}{R2}$ (V_P to dodatnie napięcie zasilania). Zakładamy, że do wejścia wzmacniacza operacyjnego nie wpływa żaden prąd, więc cały prąd R2 popłynie przez Q2, który dzięki temu pracuje przy stałym prądzie kolektora. V_{BE} tego tranzystora będzie przez wzmacniacz operacyjny tak regulowane, aby prąd kolektora stale utrzymywał na poziomie $\frac{V_P}{R2}$. Jeśli zmieni się temperatura, a tym samym parametry tranzystora Q2, wówczas wzmacniacz skoryguje mu V_{BE} tak, by prąd kolektora nie uległ zmianie.

Przy zerowym napięciu V_{IN} , Q1 będzie miał zawsze takie samo V_{BE} jak Q2, a więc nawet w różnych temperaturach będzie przez niego płynął taki sam prąd. A każde niezerowe napięcie będzie niejako przesunąć prąd Q1 względem wartości prądu odniesienia $\frac{V_P}{R2}$ – niezależnej od temperatury.

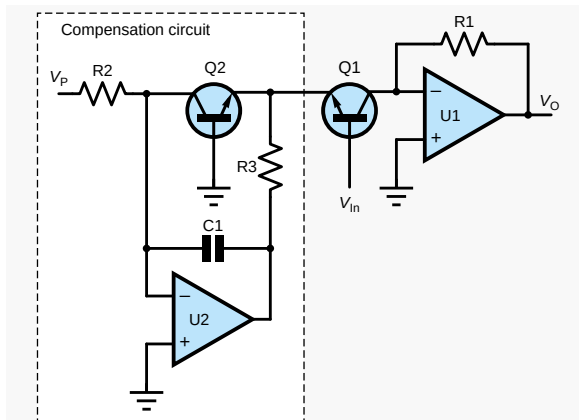
R3 ogranicza prąd wyjściowy wzmacniacza operacyjnego, zapobiegając ewentualnemu uszkodzeniu tranzystorów, a C1 poprawia stabilność, redukując wzmocnienie układu przy wysokich częstotliwościach.

Analogiczny układ jest używany do zapewnienia kompensacji temperatury dla wzmacniacza logarytmicznego.

Układ wymaga, aby oba tranzystory były parowane (miały zbliżone parametry). Powinny też znajdować się w tej samej temperaturze. Można to zapewnić przez sklejenie ich razem lub umieszczenie bardzo blisko siebie na wspólnym radiatorze. **Od Red. EdW: istnieją tranzystory podwójne, umieszczone na jednym kawałku krzemu, predestynowane do użycia właśnie w tego typu zastosowaniach. Przykład: BC846S. ■**

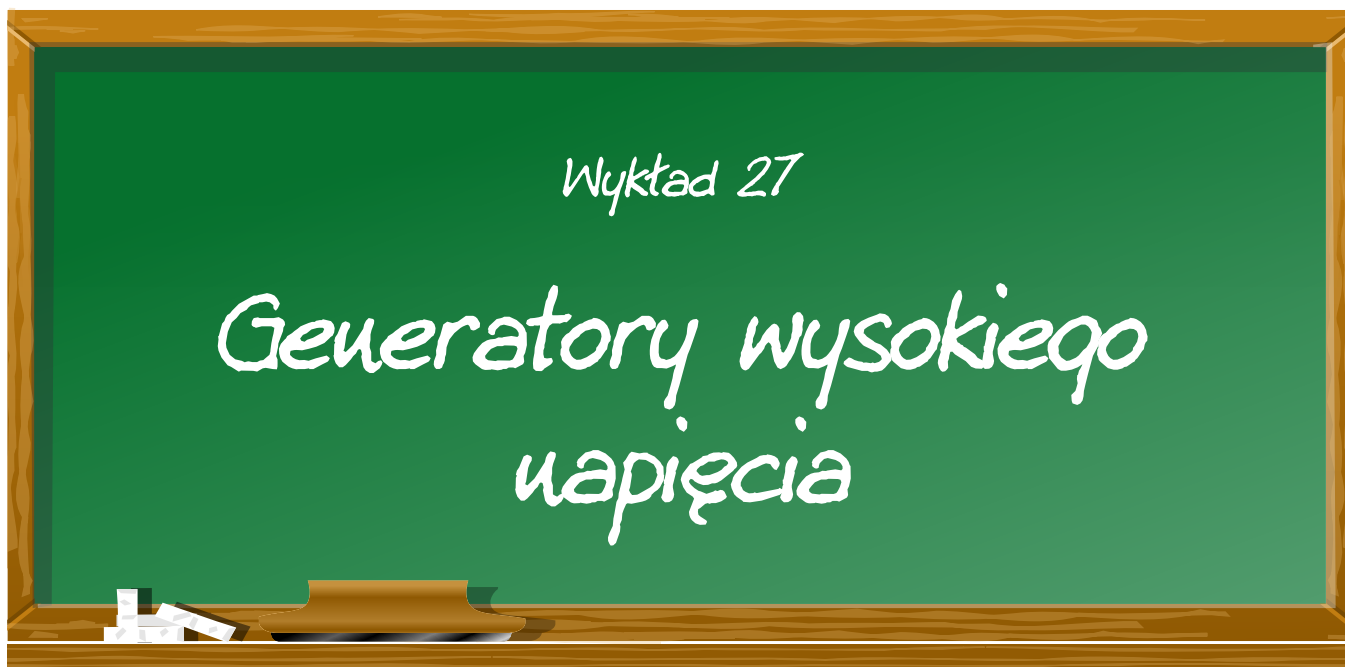
Ian Bell

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, styczeń 2022 (www.epemag3.com)



Rysunek 8. Wzmacniacz wykładniczy z kompensacją temperatury

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 wiadomości od uczniów, nauczycieli i studentów. Dla nich jest ta rubryka.



Wytwarzanie efektownych wyładowań elektrycznych cieszy się dużym zainteresowaniem wielu hobbystów. Jednakże by je pozyskać potrzebne jest wysokie napięcie. W tym artykule zajmiemy się generatorami napięć mierzonych w kilowoltach.

Podstawowe informacje na temat wyładowań

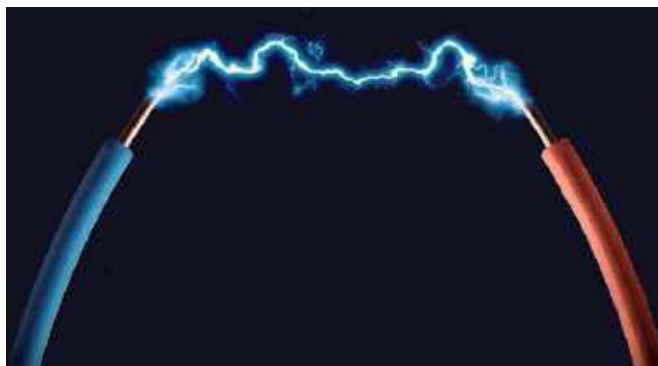
Przedstawione metody są niebezpieczne! Pracując z generatorami wysokich napięć należy zawsze pamiętać o niebezpieczeństwie, które im towarzyszy! Generatory te wytwarzają napięcia od kilkunastu tysięcy voltów wzwyż, co może stanowić zagrożenie dla życia i zdrowia. Nawet po wyłączeniu niektórych z przedstawionych generatorów, kondensatory w nich występujące mogą posiadać ładunek o napięciu setek lub tysięcy voltów nawet wiele minut po odłączeniu zasilania. Dlatego po wyłączeniu takiego generatora należy użyć dobrze zaizolowanego narzędzia lub przewodu by rozładować te kondensatory, albo odczekać 10...20 minut.

Wyładowania elektryczne, a sprzęt pomiarowy. Ponadto należy pamiętać, że duże wyładowania generują silne pola elektromagnetyczne między elektrodami wyładowczymi, a wrażliwy sprzęt pomiarowy może nie być w stanie tego wytrzymać. Nie należy zatem uruchamiać takich generatorów w pobliżu delikatnego sprzętu pomiarowego, jak multimetry czy oscyloskopy. Nie dość, iż zbyt mała odległość może uszkodzić delikatne obwody wejściowe tych przyrządów, to jeszcze, w większości przypadków, oscyloskop czy multimetr i tak jest bezużyteczny. Dlatego lepiej albo przenieść generator na inny blat, albo przenieść sprzęt pomiarowy.

Należy też pamiętać o tym, iż wyładowania elektryczne wytwarzają dość silne zakłócenia w paśmie radiowym. Sąsiedzi mogą nie być zachwyceni, gdy ich ulubiona audycja radiowa zostanie przerwana trzaskami i szumami wytwarzanymi przez taki generator. Uporczywi niszczyciele miru radiowego mogą spodziewać się odwiedzin pracowników Urzędu Komunikacji Elektronicznej i policji. Przyp. tłum.

Zależność między napięciem a długością łuku elektrycznego.

Eksperymentując z opisanymi obwodami i technikami, oczywiście chcemy wiedzieć, jakie napięcie jest generowane przez dany prototyp. Nie można zmierzyć tej wielkości za pomocą standardowego sprzętu pomiarowego. Bardzo niedokładną, ale praktycznie użyteczną metodą pomiaru jest zmierzenie długości łuku elektrycznego, który można wytworzyć. Ogólnie rzecz biorąc, przyjmuje się wartość orientacyjną 10 000 V (10 kV) na centymetr długości łuku między dwiema ostro zakończonymi elektrodami w suchym powietrzu. W przypadku pracy z elektrodami o kulistych końcach przyjmuje się napięcie 30 kV na centymetr długości łuku.



1. Typowa iskra między dwoma przewodnikami (© Adobe Free Stock)

Ta metoda pomiaru nie jest dokładna, ponieważ napięcie przebicia powietrza zależy od wielu czynników, takich jak:

- wilgotność,
- temperatura,
- ciśnienie powietrza,
- rodzaj napięcia,
- skład powietrza.

Niektórzy hobbyści budują samodzielnie dzielniki rezystorowe pozwalające na pomiar bardzo wysokich napięć. Dzielniki takie wykonywane są z wielu rezystorów o wartości 1 MΩ połączonych szeregowo i umieszczonych w pojemniku wypełnionym olejem syntetycznym. Dzielnik 1:100 wymaga setki takich rezystorów, a napięcie odkładające się na ostatnim z nich będzie równe 1/100 napięcia na całym dzielniku. Przyp. tłum.

Definicja iskry. Iskra to nagle wyładowanie elektryczne, które występuje, gdy pole elektryczne między dwiema elektrodami staje się tak silne, że w normalnie nieprzewodzącym powietrzu powstaje zjonizowany, przewodzący prąd kanał. Gwałtowne przejście ze stanu nieprzewodzącego do przewodzącego powoduje krótki błysk światła i głośny huk. Mówi się wtedy, że następuje „przebiecie”, a wielkość napięcia, przy którym to następuje, nazywa się „napięciem przebicia”.

Niebezpieczeństwa związane z wyładowaniami. Nawet niewielkie wyładowania mogą spowodować zapłon łatwopalnych materiałów, cieczy i gazów. Wszystkie wyładowania wytwarzają w powietrzu „tunel plazmowy”, przez który przepływają elektrony. Plazma ta ma temperaturę wyższą niż temperatura na powierzchni Słońca. Dlatego nawet najmniejsze iskry z zapalarki gazowej są w stanie zapalić gaz. Iskry mogą również łatwo spowodować niewielkie, miejscowe oparzenia.

Bardzo wysoka temperatura w tunelu plazmowym może spowodować uszkodzenie powierzchni metalowych przedmiotów, ponieważ część metalu odparowuje lokalnie. Zjawisko to jest wykorzystywane podczas elektrycznego wytrawiania tekstów lub piktogramów na gładkich powierzchniach metalowych. Dobrze znanym urządzeniem opracowanym do tego celu jest „Sparcatron”.

Nawet iskry o niskiej energii mogą przeciążać ścieżki przewodzenia ludzkiego układu nerwowego, powodując miejscowe skurcze mięśni lub zakłócając rytm serca. Wreszcie, wszystkie wyładowania wytwarzają ozon, który w odpowiednio wysokich stężeniach może powodować problemy z oddychaniem i jest szkodliwy dla niektórych tworzyw sztucznych.

Zjawisko przebicia elektrycznego. Prąd elektryczny powstaje, gdy duża ilość naładowanych elektrycznie cząstek zaczyna przepływać w jednym kierunku przez materiał. Cząstki te są indukowane do takiego zachowania po przyłożeniu napięcia elektrycznego do materiału, które wytwarza pole elektryczne. Naładowane cząstki tworzące prąd elektryczny nazywane są „nośnikami ładunku”. W metalach te nośniki ładunku składają się z elektronów obecnych na zewnętrznych powłokach atomów. Są one bardzo ruchliwe i mogą przeskakiwać z atomu na atom. Zjawisko to powoduje, że metale dobrze przewodzą prąd elektryczny. W przewodzących cieczach (elektrolitach) i plazmie nośnikami ładunku są naładowane elektrycznie atomy zwane „jonami”.

W materiałach, które nie przewodzą prądu elektrycznego, takich jak powietrze, nośniki ładunku są ściśle związane z atomami. Potrzeba bardzo silnego pola elektrycznego, tj. bardzo wysokiego napięcia, aby zerwać te wiązania i wywołać przepływ prądu elektrycznego. Przy określonym natężeniu pola liczba swobodnych nośników ładunku w materiale gwałtownie wzrasta, powodując, że materiał staje się przewodnikiem. Zjawisko to nazywane jest „przebieciem elektrycznym”.

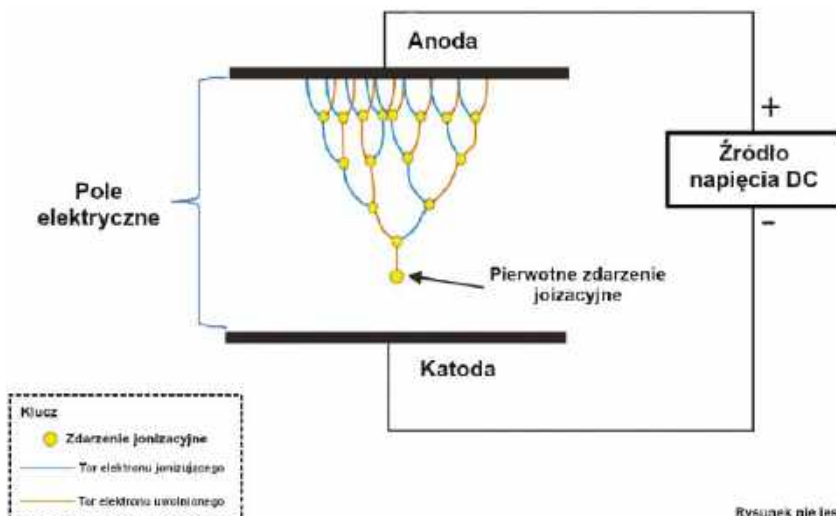
Wyładowanie Townsenda. W powietrzu, ze względu na procesy takie jak fotojonizacja i promieniowanie z rozpadu radioaktywnego, zawsze obecnych jest kilka wolnych nośników ładunku. Są one przyspieszane przez pole elektryczne wysokiego napięcia i uwalniają dodatkowe elektrony z atomów, gdy zderzają się z atomami znajdującymi się w powietrzu. Elektrony te są również przyspieszane i z kolei wyzwalały nowe wolne elektrony. W ten sposób powstaje reakcja łańcuchowa zwana „wyładowaniem Townsenda”. Takie wyładowanie Townsenda objawia się w postaci łuku elektrycznego (zwanego potocznie iskrą bądź iskrami).

Atomy w powietrzu, które utraciły jeden lub więcej elektronów w wyniku wyładowania Townsenda, nazywane są „jonami”, a proces, w którym to się dzieje, nazywany jest „jonizacją powietrza”.

Zjawisko świetlne isker. Dlaczego iskra to błysk światła? Wyładowanie Townsenda tworzy jony w powietrzu. Taki jon jest w stanie niestabilnym i będzie dążył do uzyskania stabilności poprzez rekombinację elektronów. Ten proces rekombinacji uwalnia energię w postaci cząstek światła lub „fotonów”. Fotony te mają różne poziomy energii, co skutkuje światłem o różnych kolorach i intensywności.

Jednak podczas tworzenia iskry, lokalnie generowane jest również dużo ciepła, które może spowodować, że cząsteczki w powietrzu będą przez chwilę świecić. To również przyczyni się do emisji światła.

Fakt, iż różne pierwiastki emitują światło o różnych długościach jest wykorzystywany



2. Wizualizacja zasady wyładowania Townsenda (© Wikipedia Commons)

od lat w systemach oświetleniowych opartych o jonizację gazów polem elektrycznym i wyładowania elektryczne. Znane z oświetlenia ulicznego niskoprężne lampy sodowe używają niewielkiej ilości rtęci do wstępnego zapalenia łuku w specjalnej rurce. Wysoka temperatura tego wstępnego wyładowania rozgrzewa atomy metalicznego sodu w rurce do temperatury parowania, a zjonizowany sód zaczyna emitować charakterystyczne dla siebie pomarańczowe światło. Z kolei w tradycyjnych lampach błyskowych i stroboskopach wypełniający rurkę palnika z kwarcowego szkła ksenon jest jonizowany za pomocą zewnętrznej elektrody napięciem kilku kV. Zjonizowany gaz zaczyna przewodzić prąd między dwoma elektrodami głównymi zwierając tym samym kondensator wyładowczy, co kończy się bardzo krótkim błyskiem intensywnego, białego światła. Przyp. tłum.

Zjawisko dźwiękowe isker. Podczas procesu iskrzenia wytwarzana jest znaczna ilość ciepła. Ciepło to zwiększa temperaturę otaczającego powietrza i powoduje jego rozszerzanie się. Ta gwałtowna ekspansja powietrza wokół iskry tworzy falę ciśnienia, która rozprzestrzenia się w powietrzu. Ta fala uderzeniowa jest odbierana jako huk.

Co więcej, w przypadku iskry generowana fala ciśnienia może poruszać się szybciej niż sam dźwięk, co skutkuje hukiem, który słychać nieco wcześniej niż dźwięk fali ciśnienia.

Prawo Paschena. Prawo to zostało nazwane na cześć niemieckiego fizyka Friedricha Paschena, który odkrył i opisał to zjawisko w 1889 roku. Prawo to matematycznie opisuje minimalne napięcie (U) wymagane do zjonizowania gazu pod określonym ciśnieniem (p) i rozpoczęcia wyładowania iskrowego.

Matematyczne wyrażenie prawa Paschena jest następujące:

$$U = \frac{B_{pd}}{\ln(A_{pd}) - \ln\left[\ln\left(1 + \frac{1}{\gamma_{se}}\right)\right]}$$

gdzie:

- U to napięcie przebicia w woltach,
- p to ciśnienie w paskalach,
- d to odległość między elektrodami w metrach,
- γ_{SE} to współczynnik emisji elektronów wtórnych,
- A to współczynnik jonizacji nasycenia w powietrzu,
- B to stała związana z energiami wzbudzenia i jonizacji powietrza.

Minimalne potrzebne wysokie napięcie. Aby wygenerować iskry, nawet o długości jednego centymetra, potrzebne jest napięcie rzędu dziesiątek kV. Nie można ich wygenerować za pomocą tradycyjnych obwodów, z którymi można się spotkać będąc hobbystą. Można to jednak zrobić za pomocą niecodziennych rozwiązań opisanych poniżej:

- generator kroplowy Kelvina,
- generator Van de Graaffa,
- maszyna elektrostatyczna Wimshursta,
- generator Cockcrofta-Waltona,
- generator Marksa,
- cewka Tesli.

Omówimy te sześć urządzeń w kolejnych rozdziałach.

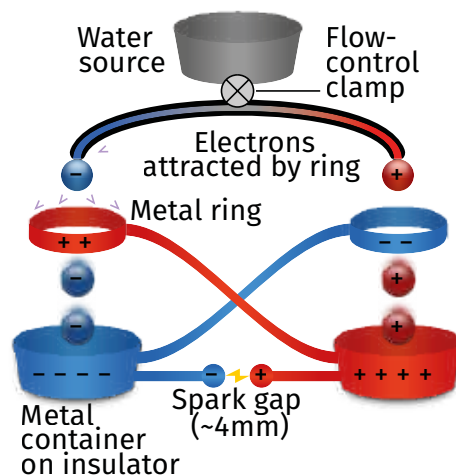
Generator kroplowy Kelvina

Wprowadzenie. Generator kroplowy Kelvina został wynaleziony w 1867 roku przez szkockiego naukowca Williama Thomsona, zwanego Lordem Kelvinem. Jest to generator elektrostatyczny, który wykorzystuje spadające krople wody do generowania różnic napięcia między dwoma kubkami.

Opis urządzenia. Budowę standardowego generatora kroplowego Kelvin przedstawiono na ilustracji. Z naczynia zawierającego wodę (szare) woda kapie przez dwie rurki do dwóch pojemników (niebieskiego i czerwonego). Są one całkowicie odizolowane elektrycznie od otoczenia. Jednak krople najpierw spadają przez metalowe pierścienie połączone elektrycznie z przeciwnymi pojemnikami.

Pomiędzy dwoma pojemnikami zamontowano iskrownik, składający się z dwóch spiczastych elektrod oddalonych od siebie o 4 mm.

Działanie urządzenia. Bez względu na warunkiem działania urządzenia jest to, że przed otwarciem kranu istnieje niewielka różnica ładunków między czerwonym i niebieskim pojemnikiem. Załóżmy, że czerwony pojemnik ma niewielki ładunek dodatni. Ładunek ten jest przenoszony przez przewód do lewego czerwonego pierścienia. Dodatni ładunek na tym pierścieniu wpłynie na równowagę ładunków kropelek wody spadających przez pierścień. Jest to konsekwencja przyciągania elektrostatycznego Coulomba. Niemniej jednak w neutralnych kropelkach wody obecna jest ograniczona liczba wolnych nośników ładunku ujemnego, tj. elektronów. Są one przyciągane przez dodatni ładunek na pierścieniu. Tak więc powierzchnia kropli wody spadającej przez ten pierścień będzie zawierać wolne elektrony. Gdy kropla spadnie do lewego pojemnika,



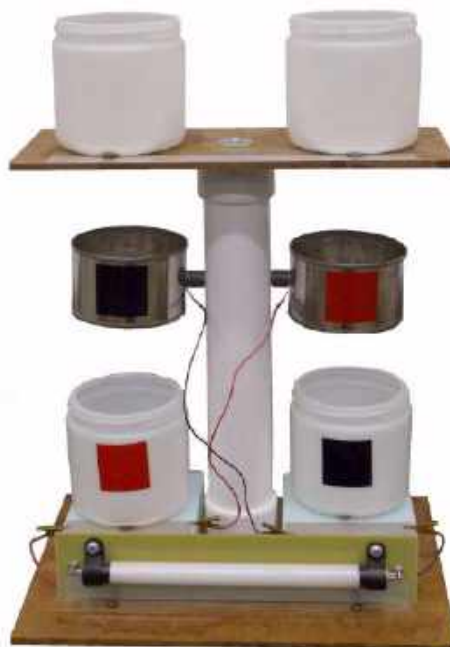
3. Budowa generatora kroplowego Kelvina
(© Wikimedia Commons)

ten ujemny ładunek rozproszy się w tym pojemniku. Pojemnik ten stanie się nieco bardziej naładowany ujemnie. Ładunek ten jest przenoszony do prawego pierścienia przez niebieski przewód.

Ujemny ładunek na tym pierścieniu przyciągnie wolne jony dodatnie w kropelkach wody. W ten sposób kropla spadająca przez prawy pierścień będzie miała powierzchnię zawierającą wiele jonów dodatnich. Jeśli kropla wpadnie do prawego pojemnika, stanie się bardziej naładowana dodatnio.

Oczywiście proces ten nasila się z każdą spadającą kroplą wody. Dzieje się tak, ponieważ ilość ładunków obu pojemników, a tym samym obu pierścieni, wzrasta. Różnica potencjałów między dwoma pojemnikami również wzrasta, aż napięcie stanie się tak wysokie, że między dwiema elektrodami iskiernika powstanie iskra.

Generator kropłowy Kelvina w praktyce. Poniższe zdjęcie przedstawia praktyczny przykład generatora kropłowego. W tej wersji jeden górny zbiornik na wodę został zastąpiony dwoma kubkami, każdy z zakraplaczem na dole. Na pierwszym planie widać metalowy pasek łączący oba zbiorniki, dzięki czemu pozostają one w identycznym stanie naładowania. Pierścienie to dwie puszki po konserwach ze zdjętymi górnymi i dolnymi denkami. Pojemniki stykają się z dwoma metalowymi paskami, które zapewniają transfer ładunku do pierścieni za pośrednictwem przewodów. W tym modelu iskiernik został zastąpiony małą świetlówką, która jest zapalana, gdy generowane napięcie wzrośnie do określonej wartości. Świetlówka wytwarza wówczas błysk światła.



4. Łatwy do samodzielnej budowy generator kropłowy Kelvina (© phys-office.phys.washington, edit 2024 Jos Verstraten)

Generator Van de Graaffa

Efekt tryboelektryczny. Powszechnie wiadomo, iż można wytworzyć małe iskry, zakładając wełniany sweter. Można usłyszeć, aż trzeszczy! Jest to wynikiem „efektu tryboelektrycznego”. Kiedy pocierasz o siebie dwa materiały, elektrony przeskakują z jednego materiału do drugiego. W rezultacie jeden materiał otrzymuje nadwyżkę elektronów i staje się naładowany ujemnie. Drugi materiał otrzymuje niedobór elektronów i staje się naładowany dodatnio. W ten sposób między dwoma materiałami powstaje różnica potencjałów, która może stać się wystarczająco duża, by między oboma materiałami zaczęły przeskakiwać iskry.

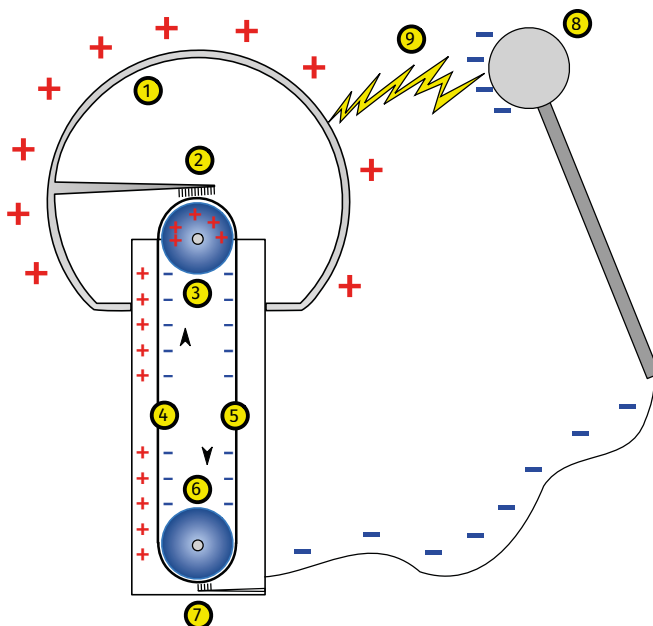
Naukowcy opracowali tak zwany „szereg tryboelektryczny”, który przedstawiono w poniższej tabeli. Im dalej od siebie znajdują się materiały w tym szeregu, tym wyższe jest napięcie elektrostatyczne między nimi, gdy wchodzi one ze sobą w kontakt przez tarcie. Jeśli więc pocieramy o siebie arkusz poliuretanu i arkusz teflonu, arkusz poliuretanu jest naładowany dodatnio, a arkusz teflonu ujemnie. Zjawisko to jest wykorzystywane w generatorze Van de Graaffa, w którym poruszający się pas przenosi ładunek elektryczny z jednego materiału na drugi.

Trochę historii. Idea przenoszenia ładunku przez poruszający się pas ma swoje korzenie w generatorze kropłowym Kelvina. Sam Kelvin jako pierwszy zaproponował wykorzystanie pasa zamiast kropielek wody do przenoszenia ładunku. Pierwsza maszyna elektrostatyczna, która wykorzystywała zapętlony pas do przenoszenia ładunku, została zbudowana w 1872 roku przez Augusto Righiego. W kolejnych dekadach opracowano kilka maszyn wykorzystujących tę zasadę:

- w 1890 r. John Gray,
- w 1903 r. Juan Burboa,
- w 1920 r. W. F. G. Swann.

Prawdziwy generator Van de Graaffa został opracowany w 1929 roku przez fizyka Roberta J. Van de Graaffa na Uniwersytecie Princeton, z pomocą Nicholasa Burke'a. Pierwszy model został zademonstrowany w październiku 1929 roku. W 1931 roku Van de Graaff zbudował urządzenie, które mogło generować napięcie 1,5 MV (1 500 000 V). Van de Graaff złożył wniosek o patent w grudniu 1931 roku, ale został on przyznany Massachusetts Institute of Technology.

Zasada działania generatora Van de Graaffa. Przekrój przez generator Van de Graaffa przedstawiono na poniższym rysunku. Urządzenie składa się z gumowego pasa [4, 5] poruszającego się na dwóch rolkach [3, 6] wykonanych z różnych materiałów. Górna rolka [3] jest otoczona dużą wydrążoną metalową kulą [1]. Na każdej rolce zamontowana jest



5. Budowa generatora Van de Graaffa (© 2016 Wikimedia Commons – Pilar Mareca)

Tabela 1. Seria tryboelektryczna (© dr W.B. Lee, AlphaLabs Inc.)

Nazwa izolatora	Powinowactwo ładunku nC/J (nano amper/wat tarcia)	Efekt metaliczny (W=słaby, N=normalny lub zgodny z powinowactwem)	Uwagi
Pianka poliuretanowa	+60	+N	Wszystkie materiały są dobrymi izolatorami (>1000 T Ω/cm), chyba że zaznaczono inaczej
Sorbotan	+58	-W	Lekko przewodzący (120 G Ω/cm)
Taśma pakowa do kartonów (BOPP)	+55	+W	Nieklejąca strona. Staje się bardziej negatywny po zeszlifowaniu do folii BOPP
Włosy, skóra tłusta	+45	+N	Skóra jest przewodząca. Nie można wytworzyć ładunku przez pocieranie metalu
Lity poliuretan, wypełniony	+40	+N	Lekko przewodzący (8 T Ω/cm)
Fluorek magnezu (MgF2)	+35	+N	Antyreflekcyjna powłoka optyczna
Nylon, sucha skóra	+30	+N	Skóra jest przewodząca. Nie można wytworzyć ładunku przez pocieranie metalu
Olej maszynowy	+29	+N	
Nylatron (nylon wypełniony MoS ₂)	+28	+N	
Szkło (sodowe)	+25	+N	Lekko przewodzący (zależy od wilgotności)
Papier (xero, niepowlekana)	+10	-W	Większość papierów i tektury ma podobne powinowactwo. Lekko przewodzący
Drewno (sosna)	+7	-W	
Silikon II marki GE (twardnieje na powietrzu)	+6	+N	Bardziej dodatni niż inne produkty silikonowe (patrz poniżej)
Bawełna	+5	+N	Lekko przewodząca (zależy od wilgotności)
Kauczuk nitrylowy	+3	-W	
Wełna	0	-W	
Poliwęglan	-5	-W	
ABS	-5	-N	
Akryl (polimetakrylan metylu) i klejąca strona przezroczystej taśmy samoprzylepnej	-10	-N	Kilka klejów stosowanych w bezbarwnych taśmach samoprzylepnych powinowactwo prawie identyczne z akrylem, mimo iż mają odmienne składniki
Żywica epoksydowa (płytki drukowane)	-32	-N	
Kauczuk styrenowo-butadienowy (SBR, Buna S)	-35	-N	Czasami niedokładnie nazywany „neoprenem” (patrz poniżej)
Farby w sprayu na bazie rozpuszczalników	-38	-N	Mogą się różnić między markami i kolorami
Tkanina PET (mylarowa)	-40	-W	
PET (mylar) stały	-40	+W	
Guma EVA do uszczelnień, wypełniona	-55	-N	Lekko przewodząca (10 T Ω/cm). Wypełniona guma zwykle przewodzi
Guma naturalna	-60	-N	Ledwo przewodząca (500 T Ω/cm)
Klej do klejenia na gorąco	-62	-N	
Polistyren	-70	-N	
Poliamid	-70	-N	
Silikony (utwardzany na powietrzu i termoutwardzalny, ale nie GE)	-72	-N	
Winył: elastyczny (przezroczyste rurki)	-75	-N	
Taśma do pakowania (BOPP), szlifowana	-85	-N	Surowa powierzchnia jest bardzo + (patrz wyżej), ale po przeszlifowaniu zbliża się do PP
Olefiny (alkeny): LDPE, HDPE, PP	-90	-N	UHMWPE znajduje się poniżej. W przypadku metali PP jest bardziej ujemny niż PE
Azotan celulozy	-93	-N	
Podkład taśmy biurowej (kopolimer winylowy?)	-95	-N	
UHMWPE	-95	-N	

Tabela 1. Seria tryboelektryczna (© dr W.B. Lee, AlphaLabs Inc.) – cd.

Nazwa izolatora	Powinowactwo ładunku nC/J (nano amper/wat tarcia)	Efekt metaliczny (W=stały, N=normalny lub zgodny z powinowactwem)	Uwagi
Neopren (polichloropren, nie SBR)	-98	-N	Lekko przewodzący, jeśli jest wypełniony (1,5 T Ω /cm)
PVC (sztywny winyl)	-100	-N	
Lateks (naturalny) kauczuk	-105	-N	
Viton, wypełniony	-117	-N	Lekko przewodzący (40 T Ω /cm)
Kauczuk epichlorohydrynowy, wypełniony	-118	-N	Lekko przewodzący (250 G Ω /cm)
Guma santoprenowa	-120	-N	
Guma hypalonowa, wypełniona	-130	-N	Lekko przewodząca (30 T Ω /cm)
Kauczuk butylowy, wypełniony	-135	-N	Przewodzący (900 M Ω /cm). Test został wykonany szybko
Guma EDPM, wypełniona	-140	-N	Lekko przewodząca (40 T Ω /cm)
Teflon	-190	-N	Powierzchnia to atomy fluoru – bardzo elektryzujące

metalowa elektroda w kształcie grzebienia z ostrymi punktami [2, 7]. Punkty te nie stykają się jednak z gumowym paskiem, ale znajdują się jak najbliżej pasa. Górny grzebień [2] jest elektrycznie połączony z kulą, a dolny grzebień [7] z uziemieniem i z drugą elektrodą [8]. Wyładowania [9] występują między kulą [1], a drugą elektrodą [8].

Działanie generatora. Gdy pasek jest napędzany, efekt tryboelektryczny powoduje przenoszenie nośników ładunku między różnymi materiałami paska i dwóch rolek. W przedstawionym przykładzie wewnętrzna guma paska jest naładowana ujemnie, a górna rolka jest naładowana dodatnio. Ładunek na dodatniej górnej rolce [3] generuje bardzo wysokie pole elektryczne w pobliżu końców metalowego grzbietu [2]. Pole to jest tak silne, że jest w stanie zjonizować cząsteczki powietrza wokół punktów grzbietu. Uwolnione elektrony z cząsteczek powietrza są przyciągane na zewnątrz pasma, podczas gdy jony dodatnie trafiają do grzebienia. Na grzebieniu są one neutralizowane przez elektrony z metalu kuli, pozostawiając grzebień i powierzchnię metalowej kuli bez elektronów. W rezultacie metalowa kula uzyskuje ładunek dodatni.

W miarę obracania się paska, coraz więcej ładunku jest transportowane z dolnej rolki do górnej, a metalowa kula nabiera coraz większego ładunku dodatniego. Rzeczywiście, z praw elektrostatyki wynika, że cały ładunek dodatni jest skoncentrowany na powierzchni metalowej kuli i żaden ładunek nie jest „odczuwalny” wewnątrz kuli. Maszyna może zatem nadal przenosić ładunek.

W pewnym momencie pojawia się jednak równowaga między dostarczaniem ładunku przez gumowy pas a odprowadzaniem ładunku przez upływ i wyładowania koronowe.

W ten sposób na powierzchni metalowej kuli powstaje bardzo wysoki potencjał. Wartość tego potencjału zależy od wielkości kuli. Im większa kula, tym wyższy potencjał.

Pozostaje pytanie, w jaki sposób dolna rolka może nadal dostarczać ładunek. Jest to możliwe dzięki dolnemu grzbietowi [7], który jest połączony z uziemieniem i może nadal dostarczać dolnej rolce nośniki ładunku z ziemi za pośrednictwem grzebienia [7].

Narysowany przykład zakłada wytworzenie dodatniego ładunku na kuli. Wybierając materiały dwóch rolek, można oczywiście wygenerować również ładunek ujemny na kuli.

Wielkość generowanego napięcia. Maksymalne napięcie, jakie można uzyskać na metalowej kuli, jest w przybliżeniu równe promieniowi R kuli pomnożonemu przez pole elektryczne Emax, przy którym w powietrzu zaczynają występować wyładowania koronowe, około 30 kV/cm. Generator Van de Graaffa z metalową kulą o średnicy 30 centymetrów może zatem generować maksymalne napięcie 450 kV.

Tani generator eksperymentów. Samodzielne wykonanie generatora Van de Graaffa jest nie lada wyzwaniem. Istnieją oczywiście generatory oferowane za pośrednictwem AliExpress. Jeden z najtańszych jest pokazany na fotografii i kosztuje trochę ponad 750 złotych.



6. Generator Van de Graaffa oferowany przez „Module Parts Factory Store” (© AliExpress)

Maszyna elektrostatyczna Wimshursta

Wprowadzenie. Maszyna elektrostatyczna Wimshursta to elektrostatyczny generator iskieł opracowany w 1883 roku przez brytyjskiego wynalazcę Jamesa Wimshursta. Urządzenie charakteryzuje się dwoma dużymi, dobrze izolowanymi plastikowymi (we współczesnych modelach – przyp. tłum.) tarczami obracającymi się w przeciwnych kierunkach wokół tej samej osi. Każda

tarcza wyposażona jest w dużą liczbę segmentów przewodzących. Segmenty obracają się pod dwoma prętami neutralizującymi i dwoma kolektorami ładunków. Ładunki zebrane z segmentów są gromadzone w dwóch butelkach lejdejskich. Są one połączone z iskiernikiem zbudowanym z dwóch metalowych kul.

Maszyna Wimshurst gromadzi ładunki elektryczne poprzez indukcję elektrostatyczną, a zatem nie działa na zasadzie tarcia, jak generator Van de Graaffa.

Typowa maszyna Wimshurst może wytwarzać iskry o długości około jednej trzeciej średnicy tarcz i może dostarczać kilkadziesiąt μA prądu.

Maszyna Wimshurst gromadzi ładunki elektryczne poprzez indukcję elektrostatyczną, a zatem nie działa na zasadzie tarcia, jak generator Van de Graaffa.

Typowa maszyna Wimshurst może wytwarzać iskry o długości około jednej trzeciej średnicy tarcz i może dostarczać kilkadziesiąt μA prądu.

Działanie maszyny elektrostatycznej Wimshurst. Działanie tego urządzenia jest dość skomplikowane i musimy przyznać, że nawet po długich poszukiwaniach w internecie nie udało nam się znaleźć opisu, który w pełni wyjaśniłby nam jego działanie. Możemy więc niestety opisać jego działanie jedynie w bardzo ogólnym zarysie.

Dwie przeciwbieżne tarcze, zwykle wykonane ze szkła, są wyposażone w dużą liczbę metalowych segmentów. Maszyna zaczyna działać w wyniku niewielkich ładunków na tych metalowych segmentach. Ładunki te są spowodowane przypadkowymi zjawiskami w pomieszczeniu, w którym maszyna jest ustawiona.

Gdy dwie tarcze zaczynają obracać się w przeciwnych kierunkach z dużą prędkością, dany segment otrzymuje przeciwny ładunek indukowany jako ładunek obecny w segmencie na drugiej tarczy obracającej się obok danego segmentu.

Pręty neutralizujące mają na końcach szczotki zbierające i łączą każdy segment z segmentem na przeciwległej stronie tej samej tarczy.

W rezultacie segment na jednym końcu pręta otrzymuje przeciwny ładunek od segmentu na drugim końcu. Następnie przeciwbieżnie obracające się tarcze przesuwają naładowane segmenty do kolektorów ładunku, gdzie segment dotykający kolektora zawsze znajduje się naprzeciwko identycznie naładowanego segmentu na drugiej tarczy. Ponieważ te same ładunki odpychają się wzajemnie, część ładunku jest usuwana przez kolektory ładunku i przechowywana w butelkach lejdejskich.

Butelki lejdejskie gromadzą ładunek do momentu, w którym napięcie na iskierniku staje się tak duże, że dochodzi do przebicia i przeskoku iskry, po czym proces się powtarza.

Tania maszyna elektrostatyczna Wimshurst. Maszyny Wimshurst są oferowane w sprzedaży w cenach zaczynających się od około 160 złotych na Aliexpress. Istnieje też tani zestaw maszyny Wimshurst kosztujący 50,00 euro, opracowany przez niemiecką firmę AstroMedia. Wyjątkowość tej konstrukcji polega na tym, że urządzenie składa się niemal w całości z kartonowych części, które trzeba skleić ze sobą. Ta maszyna Wimshurst może generować iskry o długości do 5 cm, jak twierdzi producent. Zrobiliśmy kiedyś taką i naprawdę działa, choć nie osiągnęliśmy 5 cm!

[W internecie można spotkać konstrukcje DIY wykonane z tarcz akrylowych lub poliwęglanowych oraz pasków folii aluminiowej lub taśmy miedzianej na nich naklejonych.](#) Przyp. tłum.



7. Typowy przykład maszyny elektrostatycznej Wimshurst (© AliExpress)



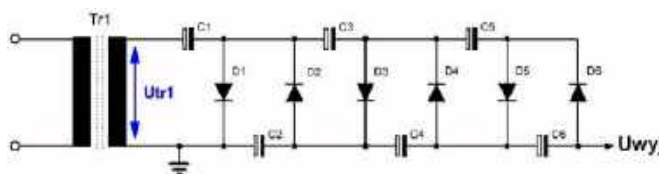
8. Kartonowy zestaw konstrukcyjny Astromedia (© Astromedia)

Generator Cockcrofta-Waltona

Wprowadzenie. Teraz powoli, ale zdecydowanie zmierzamy w kierunku „normalnej” elektroniki! Generator Cockcrofta-Waltona to obwód, który wytwarza wysokie napięcie stałe ze znacznie niższego napięcia przemiennego lub ze znacznie mniejszego impulsowego napięcia stałego. Nazwa obwodu pochodzi od brytyjskich fizyków Johna Douglasa Cockcrofta i Ernesta Thomasa Sintona Waltona, którzy po raz pierwszy użyli tego obwodu w 1932 roku do wygenerowania wysokiego napięcia przyspieszającego niezbędnego dla akceleratora cząstek.

W rzeczywistości fundamenty tego obwodu zostały położone już w 1919 roku przez Heinricha Greinachera, szwajcarskiego fizyka. Cockcroft i Walton rozbudowali prosty podwajacz napięcia Greinachera.

Działanie obwodu. Jest to bez wątpienia najprostszy obwód, który umożliwia uzyskanie wysokiego lub bardzo wysokiego napięcia stałego z niskiego napięcia przemiennego. Zasada działania została przedstawiona na poniższym rysunku. Obwód ten znajdował się kiedyś w każdym telewizorze kineskopowym, gdzie taka kaskada diod i kondensatorów zapewniała wysokie napięcie przyspieszające wynoszące około 20 000 V DC, które znajdowało się po wewnętrznej stronie ekranu kineskopu. Kaskady Cockcrofta-Waltona są również wykorzystywane, na przykład, w słynnych lampach plazmowych, które pozwalają za pomocą palca przyciągać wyładowania do szklanej bańki wypełnionej gazem szlachetnym.



9. Generator wysokiego napięcia Cockcroft-Walton (© 2024 Jos Verstraten)

Zarówno w telewizorach kineskopowych, jak i w kulach plazmowych napięcie jest wstępnie podnoszone do wartości kilku kilowoltów, by liczba potrzebnych segmentów tego powielacza napięcia nie była zbyt duża. Przyp. tłum.

Na parzystych kondensatorach występuje napięcie stałe w przybliżeniu równe dwukrotnej amplitudzie napięcia wtórnego transformatora Tr1. To się ładnie sumuje! W przypadku zastosowania transformatora separującego 230 V dla Tr1, tj. z napięciem wtórnym również 230 V, na kondensatorach występuje napięcie:

$$U_{dc} = 2 \cdot 1,41 \cdot 230V = 648,6V$$

W narysowanym przykładzie na wyjściu występowałoby napięcie stałe o wartości nie mniejszej niż ~1,945 kV względem masy! Odwracając wszystkie diody i kondensatory, można oczywiście wygenerować dodatnie napięcie wyjściowe.

Zalety i wady. Ogromną zaletą obwodu Cockcrofta-Waltona jest to, że na wszystkich komponentach występuje maksymalne napięcie o amplitudzie tylko dwukrotnie większej od amplitudy napięcia wejściowego. Można więc użyć tanich kondensatorów elektrolitycznych i zwykłych diod krzemowych.

Główną wadą jest to, że obwód działa jak prostownik jednopółkowy. Gdy obwód jest obciążony prądem wyjściowym, tętnienia napięcia na wyjściu gwałtownie rosną wraz ze wzrostem liczby stopni. Chociaż można temu zaradzić za pomocą filtra wyjściowego, wymaga to kondensatorów wyładowczych o bardzo wysokim napięciu roboczym i dlatego nie jest praktycznym rozwiązaniem. Wraz ze wzrostem liczby stopni, napięcia na wyższych kondensatorach zaczynają spadać z powodu impedancji kondensatorów w niższych stopniach.

Praca z prądem przemiennym o wysokiej częstotliwości. Niektóre z tych wad można wyeliminować, zasilając obwód Cockcrofta-Waltona nie transformatorem sieciowym, ale transformatorem HF. Niezbędna wartość kondensatorów jest wtedy znacznie mniejsza, a ponadto tętnienia stają się znacznie mniej uciążliwe i łatwiejsze do usunięcia. Podczas pracy z napięciem sieciowym 50 Hz, bezwzględnie trzeba używać kondensatorów elektrolitycznych, które są drogie i duże, jeśli napięcie robocze ma wynosić setki woltów. W przypadku korzystania z transformatorów HF kondensatory mogą mieć znacznie niższą wartość, a technologicznie znacznie łatwiej jest wykonać takie kondensatory o napięciu roboczym kilku kV.

Generator Cockcrofta-Waltona w formie zestawu do samodzielnego montażu.

Prostota obwodu i fakt, że zabawa z obwodami, które mogą generować wyładowania, jest bardzo popularna, sprawiły, że chińscy dostawcy elektroniki rzucili się na ten obwód. Za niewielkie pieniądze można kupić płytki drukowane i zestawy, które pozwalają na wykonanie kaskady HF Cockcroft-Walton generującej dziesiątki kV.

Poniższy przykład przedstawia taką płytkę PCB oferowaną na AliExpress za 80...110 PLN. Ten obwód od Qinda Electronics, kod AC-Z107, ma osiem kondensatorów i diod, które mogą wytrzymać napięcie robocze 10 kV. Wygląda na to, że dzięki temu można łatwo generować napięcia do 60 kV!



10. Generator Cockcrofta-Waltona AC-Z107 (© Alibaba)

Druga dość tania kaskada Cockcrofta-Waltona została przedstawiona na poniższym rysunku. Kosztuje ona 63 złote na AliExpress i zawiera 24 kondensatory i diody. Moduł ten należy zasilать napięciem przemiennym o częstotliwości co najmniej 10 kHz, a zalecane jest 50 kHz. Obwód zawiera kondensatory o napięciu roboczym 6 kV i wartości 1000 pF. Pozwala to również na uzyskanie napięcia wyjściowego do 60 kV przy napięciu wejściowym 2,5 kV.



11. 24-stopniowy generator Cockcrofta-Waltona (© AliExpress)

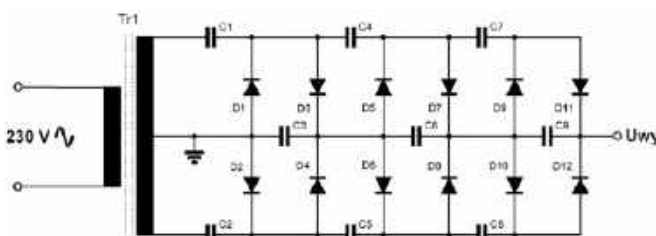
Generator wysokiej częstotliwości jako obwód wejściowy. Oczywiście potrzebny jest obwód generujący niezbędne napięcie przemiennie rzędu kilku kV o wysokiej częstotliwości. Można poeksperymentować z układem NE555 ze stopniem wyjściowym zasilającym transformator wysokiego napięcia ze starego telewizora (metoda odradzana ponieważ istnieje wiele dedykowanych układów do pracy ze starymi transformatorami WN od telewizorów kineskopowych, które dużo lepiej sprawdzają się w tej roli – przyp. tłum.). Ale oczywiście można też kupić gotowe płytki drukowane. Poniższy obrazek przedstawia płytkę drukowaną, która jest zasilana napięciem stałym od 5 V do 15 V i generuje impuls napięcia przemiennego od 5 kV do 20 kV na wyjściu dwadzieścia razy na sekundę. Ten mały obwód można kupić za około 20...40 złotych na AliExpress.



12. Generator 20 kV do zasilania generatora Cockcrofta-Waltona (© AliExpress)

Obwód Cockcrofta-Waltona z prostowaniem pełnookresowym.

Opisany obwód ma tę wadę, że działa przy użyciu prostowania jednopółkowego, a zatem cierpi na duże tętnienia napięcia wyjściowego. Opracowano alternatywny obwód Cockcrofta-Waltona, który wykorzystuje prostowanie dwupółkowe. Schemat jest narysowany na poniższym rysunku i jest zasadniczo łączonymi ze sobą.



13. Generator Cockcrofta-Waltona z prostowaniem dwupółkowym (© 2021 Jos Verstraten)

Generatory Cockcrofta-Waltona i rozwój nauki. Istnieje wiele obszarów badań naukowych, w których bardzo wysokie napięcia DC są niezbędne. Wystarczy pomyśleć o akceleratorach cząstek! W tym celu nadal budowane są bardzo duże generatory Cockcrofta-Waltona. Na poniższym zdjęciu znajduje taki generator zbudowany przez firmę Philips w 1937 roku, który był wykorzystywany w amerykańskim Projekcie Manhattan w celu opracowania technologii rozszczepienia jądowego. Ten generator Cockcrofta-Waltona znajduje się obecnie w Narodowym Muzeum Nauki w Londynie.



14. Generator Cockcrofta-Waltona zbudowany przez firmę Philips w 1937 roku (© 2012 Wikimedia Commons – Geni)

Generator Marxa

Wprowadzenie. Generator Marxa został po raz pierwszy opisany przez Erwina Otto Marxa w 1924 roku. Obwód ten jest w stanie wygenerować krótki impuls wysokiego napięcia ze znacznie niższego napięcia stałego. Generatory Marxa są nadal wykorzystywane w nauce i inżynierii, na przykład w eksperymentach fizyki wysokoenergetycznej i do symulacji skutków uderzeń piorunów w linie energetyczne i sprzęt lotniczy.

Zasada działania generatora Marx. Zasada działania obwodu, przedstawiona na poniższym rysunku, jest zarówno niezwykle oczywista, jak i niezwykle pomysłowa. Pomysł polega na tym, że identyczne kondensatory wysokonapięciowe są najpierw połączone równolegle do napięcia wejściowego i ładowane do tego napięcia za pomocą szeregowych rezystorów. Następnie kondensatory są łączone szeregowo, dzięki czemu wszystkie napięcia kondensatorów są sumowane na wyjściu obwodu.

Na poniższym rysunku koncepcja ta została opracowana dla czterostopniowego generatora Marxa. W pewnym momencie napięcie na pierwszym kondensatorze C1 staje się tak wysokie, że iskiernik SG1 umieszczony na tym kondensatorze doznaje przebicia i wytwarza iskrę. Iskra ta jonizuje atomy między dwoma stykami iskrownika, powodując, że rezystancja między tymi stykami staje się bardzo niska. W rezultacie kondensatory C1 i C2 zostają połączone szeregowo, a napięcie sumaryczne, równe dwukrotności napięcia wejściowego, powoduje przebicie iskiernika SG2. W ten sposób, w ciągu ułamka sekundy, wszystkie iskierniki zadziałają jeden po drugim, powodując szeregowe połączenie wszystkich kondensatorów.

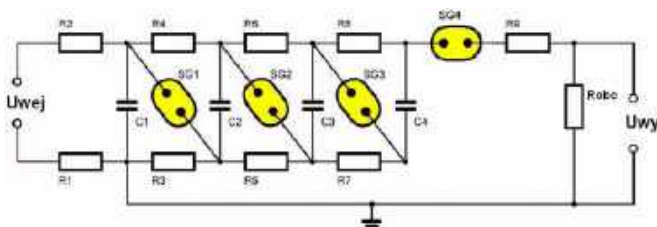
Chociaż w tym momencie nie wszystkie kondensatory są naładowane do napięcia wejściowego, między masą a wyjściem generowane jest bardzo wysokie napięcie impulsowe. Jest ono tak wysokie, że można wytworzyć iskry o długości kilkudziesięciu centymetrów. Ostatni iskiernik SG4 izoluje obwód od obciążenia podczas ładowania kondensatorów. Bez niego obciążenie Robc uniemożliwiłoby ładowanie kondensatorów. W momencie przebicia iskierników, napięcie na SG4 naturalnie staje się tak duże, że również iskrzy i oferuje napięcie na kondensatorach podłączonych szeregowo do obciążenia. Na wyjściu pojawia się krótki impuls napięcia idealnie równy $[n \cdot U_{in}]$, gdzie n to liczba stopni kaskady. Jednak w praktyce napięcie wyjściowe będzie znacznie niższe, ponieważ, jak już napisano, nie wszystkie kondensatory są naładowane do maksymalnej wartości.

Prąd rozładowania powoduje gwałtowny spadek napięcia na kondensatorach. Napięcia na iskiernikach również zaczynają spadać, w wyniku czego gasną one jeden po drugim. Gdy ostatni z nich również zgaśnie, proces rozpoczyna się od nowa, a kondensatory zaczynają się ponownie ładować.

Optymalizacja obwodu Marxa. Istnieją źródła, które twierdzą, że transfer napięcia z poszczególnych kondensatorów do wyjścia może być zoptymalizowany, jeśli różne iskierniki widzą się nawzajem. Wynikałoby to z faktu, iż zapłon iskiernika jest szybszy, jeśli ten jest napromieniowany światłem ultrafioletowym. Ten rodzaj promieniowania znajduje się w spektrum elektromagnetycznym uwalnianym przez wyładowanie elektryczne. Światło ultrafioletowe emitowane przez pierwszy iskiernik spowodowałoby szybszy zapłon drugiego iskiernika, itd.

Zgodnie z tą samą teorią, generator Marxa można zainicjować na polecenie, naświetlając pierwszy iskiernik światłem diody LED UV. Następnie należy przy tym upewnić się, iż napięcie wejściowe jest nieco za niskie na spontaniczny zapłon pierwszego iskiernika.

Generator Marxa jako zestaw do samodzielnego montażu. Kilka chińskich firm oferuje tanie generatory Marxa, które pozwalają wytwarzać spektakularne wyładowania. Poniższy zestaw, na przykład, jest dostarczany przez Banggood w cenie około 50 euro. Zestaw ten zawiera obwód Marxa, który należy zasilić napięciem stałym o wartości 20 kV i generuje napięcie impulsowe o wartości około 200 kV!



15. Podstawowy schemat obwodu Marxa (© 2024 Jos Verstraten)



16. Zestaw generatora Marxa (© Banggood)

Cewka Tesli

Wprowadzenie. Cewka Tesli to obwód zaprojektowany przez Nikołę Teslę w 1891 roku, który działa jak transformator wprowadzony w rezonans wysokiej częstotliwości. Służy do wytwarzania wysokiego napięcia o wysokiej częstotliwości. Cewki Tesli były wykorzystywane komercyjnie do lat dwudziestych XX wieku w nadajnikach iskrowych do telegrafii bezprzewodowej. Obecnie są one używane głównie w celach rozrywkowych i edukacyjnych.

Charakterystyka cewki Tesli. Cewka Tesli to nic innego jak specjalny transformator, charakteryzujący się następującymi cechami:

Uzwojenie pierwotne zawiera tylko kilka zwojów.

Uzwojenie wtórne zawiera od setek do tysięcy zwojów.

Jeden z końców uzwojenia wtórnego kończy się bardzo wąską elektrodą w kształcie igły, albo dużą elektrodą kulistą bądź w kształcie torusa.

Pomiędzy uzwojeniem pierwotnym i wtórnym znajduje się tylko powietrze.

Oba uzwojenia są tak dobrane, aby wchodziły razem w rezonans.

Powstałe pole magnetyczne o wysokiej częstotliwości wytworzone wokół uzwojenia pierwotnego generuje bardzo wysokie napięcie wtórne w uzwojeniu wtórnym.

Poniższe zdjęcie przedstawia typową cewkę Tesli, którą można kupić w Chinach. Tutaj kilka zwojów uzwojenia pierwotnego jest wytrawionych na kawałku płytki drukowanej zamontowanej wokół uzwojenia wtórnego.

Typy cewek Tesli. Istnieje kilka podstawowych typów cewek Tesli różniących się konstrukcją i sposobem generowania napięcia po stronie pierwotnej. Podstawowe typy oznaczane są skrótami:

SGTC, cewka Tesli z iskiernikiem. Jest to oryginalny projekt Tesli. Kondensator i uzwojenie pierwotne tworzą w nim równoległy obwód rezonansowy. Między nimi jednak znajduje się iskiernik. Uzwojenie wtórne i duży torus na jednym jego końcu tworzą szeregowy obwód LC. Gdy kondensator, ładowany z wydajnego źródła wysokiego napięcia osiągnie wystarczający stopień naładowania, iskra przeskakuje przez iskiernik, a oba uzwojenia rezonują wytwarzając bardzo wysokie napięcie.

VTTC, cewka Tesli z lampą próżniową. W tym typie elementem wykonawczym jest lampa elektronowa dużej mocy, zazwyczaj tetroda lub pentoda nadawcza produkcji sowieckiej bądź amerykańskiej. Charakterystyczną cechą tych cewek jest obecność dodatkowego uzwojenia odpowiedzialnego za sprzężenie zwrotne. Wyładowania mają też charakterystyczną formę wynikającą z faktu, iż cewka pracuje na częstotliwościach kilku do kilkunastu MHz.

SSTC, półprzewodnikowa cewka Tesli. Istnieje szereg różnych projektów tego typu cewek, ale w każdym elementem wykonawczym jest przynajmniej jeden tranzystor. W wersji „tradycyjnej” pracują tranzystory MOSFET. Sterownik bywa dość rozbudowany i używa anteny jako sprzężenia zwrotnego.

DRSSTC, podtyp półprzewodnikowej cewki Tesli, zwykle oparty o moduły IGBT, z rozbudowanym układem sterowania, i podwójnym rezonansem stron pierwotnej i wtórnej. Te cewki wytwarzają naprawdę imponujące wyładowania.

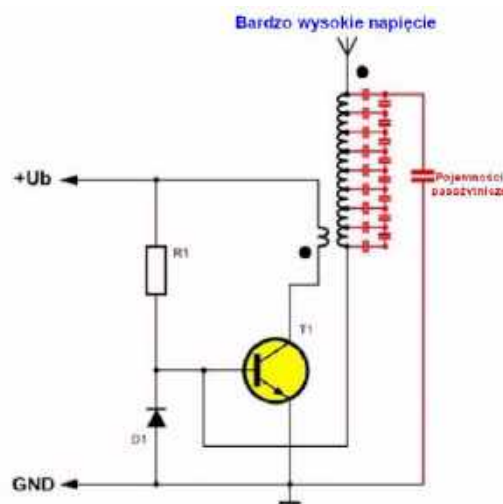
Class E TC, drugi podtyp cewki Tesli, gdzie sterownik oparty jest o wzmacniacz klasy E, zwykle zbudowany na jednym tranzystorze MOSFET. Pracują z częstotliwością kilku MHz, ale wyładowania nie są zbyt spektakularne.

OLTC, cewki Tesli zasilane bezpośrednio z napięcia sieciowego. Konstrukcyjnie łatwiejsze w realizacji ze względu na tańsze komponenty wykonawcze.

Każda półprzewodnikowa cewka Tesli może być zmodyfikowana tak, by można było modulować wyładowania dźwiękiem z zewnętrznego źródła. Tak powstają grające cewki Tesli. Ten krótki opis typów nie wyczerpuje



17. Typowy wygląd cewki Tesli (© AliExpress)



18. Najprostszy schemat cewki Tesli, „Slayer Exciter” (© 2024 Jos Verstraten)



19. Electrum, największa na świecie cewka Tesli w Nowej Zelandii (© 2015 Wikimedia Commons – Joe Decker)

oczywiście bogatego świata amatorskich konstrukcji cewek Tesli. Tłumacz zachęca Czytelników do samodzielnej eksploracji. Przyp. tłum.

Nie każda cewka jest użyteczna. Często przedstawia się to tak, jakby wystarczyło zastosować kilkaset zwojów wokół dowolnego rdzenia powietrznego, a następnie kilka kolejnych zwojów znacznie grubszego drutu. Jasne, z wtórnego wyjdzie napięcie, ale taka konstrukcja nie spełnia podstawowego warunku cewki Tesli: rezonansu. Jeśli chcemy stworzyć obwód, który naprawdę zasługuje na miano cewki Tesli, musimy zadbać o to, by uzwojenia pierwotne i wtórne miały identyczne częstotliwości rezonansowe. Częstotliwość ta zależy od typu cewki Tesli, może się mieścić w zakresie od kilkudziesięciu kHz do kilkunastu MHz.

Częstotliwość rezonansową cewki określa wyrażenie:

$$f = \frac{1}{2 \cdot \pi \cdot \sqrt{L \cdot C}}$$

gdzie L i C to odpowiednio indukcyjność i pojemność własna uzwojenia.

Od razu będzie jasne, że warunek rezonansu jest spełniony, jeśli:

$$L_p \cdot C_p = L_s \cdot C_s$$

gdzie „p” oznacza oczywiście uzwojenie pierwotne, a „s” uzwojenie wtórne.

Warto dodatkowo wspomnieć o takich kwestiach jak współczynnik sprzężenia między uzwojeniami, pojemność kuli lub torusa na szczycie uzwojenia wtórnego czy dodatkowe kondensatory w obwodzie pierwotnym, zależnie od typu konstrukcji. W Internecie dostępnych jest również sporo kalkulatorów do obliczania i projektowania cewek Tesli. Przyp. tłum.

Działanie cewki Tesli. Poniższy schemat przedstawia najprostszy obwód cewki Tesli. Obwód ten nosi nazwę „Slayer Exciter”. Jak to może działać? Można to zrozumieć tylko wtedy, gdy weźmie się pod uwagę, że każdy obiekt ma pewną niewielką pojemność pasywną w stosunku do otoczenia. Tak więc uzwojenie wtórne cewki Tesli również ją ma i to właśnie ta pojemność pasywna zapewnia sprzężenie zwrotne z wyjścia do wejścia, które jest absolutnie niezbędne dla oscylatora.

W momencie włączenia napięcia zasilania, prąd przepływa przez rezystor R1 do bazy tranzystora T1. Ten półprzewodnik zaczyna przewodzić i w rezultacie przez kilka zwojów uzwojenia pierwotnego przepływa bardzo duży prąd. Silne pole magnetyczne wokół tego uzwojenia generuje wysokie napięcie w uzwojeniu wtórnym. Dolne połączenie uzwojenia wtórnego biegnie do bazy tranzystora. Pętla sprzężenia zwrotnego jest zamykana przez pasywną pojemność uzwojenia wtórnego. Na bazie indukowane jest ujemne napięcie, które jest ograniczane do bezpiecznej wartości przez diodę D1. Oczywiście jest to możliwe tylko wtedy, gdy oba uzwojenia transformatora są prawidłowo nawinięte. Stąd dwie czarne kropki w pobliżu uzwojeń, które wskazują, że oba uzwojenia powinny być nawinięte w przeciwnych kierunkach.

Tranzystor zaczyna się zamykać, a utrata prądu w uzwojeniu pierwotnym generuje już kolejne wysokie napięcie w uzwojeniu wtórnym, teraz o przeciwnej polaryzacji. Ujemne napięcie, które dotarło do bazy przez pojemność pasywną, teraz szybko się rozprzyszy. Tranzystor powraca do przewodzenia i cykl się powtarza.

Wniosek: obwód oscyluje z dość wysoką częstotliwością. W uzwojeniu wtórnym generowane jest wysokie napięcie, które może wydostawać się do powietrza przez igłę lub kuliste górne połączenie cewki wtórnej pod wpływem generowania wyładowań gazowych.

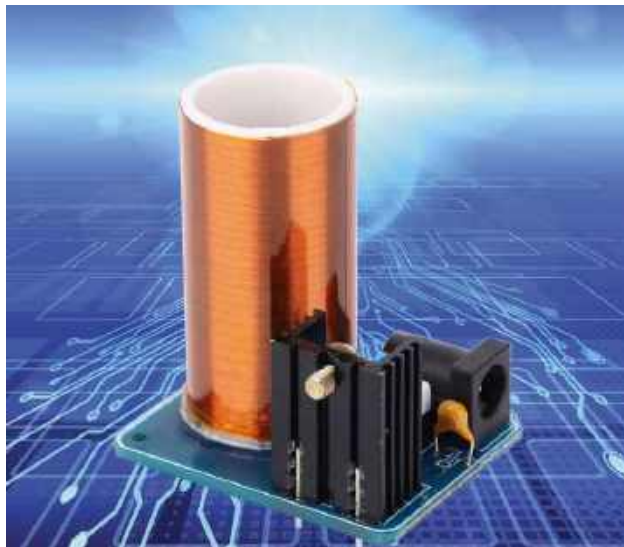
Obwód ten budową i sposobem działania przypomina przetwornice samolawne RCC. W tych przetwornicach występują trzy uzwojenia nawinięte na rdzeniu dławika: pierwotne, wtórne i pomocnicze. Uzwojenie pomocnicze włączone jest między bazę tranzystora, a masę, a uzwojenie wtórne nie ma elektrycznego połączenia ze stroną pierwotną. Układ rezonuje z naturalną częstotliwością rezonansową dławika. Ponieważ uzwojenie pomocnicze ma połączenie z masą, z układu znika dioda. Przyp. tłum.

Miliony woltów. Profesjonalne cewki Tesli pozwalają na generowanie napięć rzędu milionów woltów. Na poniższej fotografii wyładowania mają długość ponad metra. Wytwarza je „Electrum”, największa cewka Tesli zbudowana jako inżynierskie dzieło sztuki. Cewka ta została zaprojektowana przez nowozelandzkiego artystę Leonarda Charlesa Huia Lye i można ją podziwiać w skansenie rzeźbiarskim Gibbs Farm w Kaipara Harbour, 47 km na północ od Auckland. Cewka ta ma uzwojenie wtórne o wysokości 11,5 m, zużywa 130 kW i dostarcza napięcie 3 000 000 V.

Generator Tesli i samoorganizacja. Możemy wybierać spośród dziesiątek zestawów cewek Tesli we wszystkich przedziałach cenowych. Poniższe zdjęcie przedstawia bardzo taną konstrukcję, kosztującą niecałe 50 złotych, od której nie należy oczekiwać zbyt wiele pod względem generowania iskry.

Zupełnie innego kalibru jest model, który kosztuje 1250 euro, ale może generować wyładowania o długości ponad pięćdziesięciu centymetrów.

Wielu konstruktorów, również polskich, budowało cewki Tesli wytwarzające wyładowania mające ponad metr, wydając przy tym znacznie mniej. Przyp. tłum. ■



Jos Verstraten

20. Bardzo tani zestaw cewki Tesli (© Amazon)



Co potrafi współczesny oscyloskop cyfrowy?

Trwający od kilkunastu lat wyścig „zbrojeń” między różnymi producentami oscyloskopów cyfrowych doprowadził do sytuacji, w której urządzenie w zasięgu budżetu hobbysty ma funkcje, które jeszcze dekadę temu były odpłatnymi dodatkami w profesjonalnych urządzeniach kosztujących wiele tysięcy złotych. Ten spadek cen i postęp technologii sprawia, iż hobbysci często nie zdają sobie sprawy, co ich urządzenia potrafią, a nawet jeśli wiedzą o opisanych w tym artykule funkcjach, to nie wiedzą, jak je efektywnie wykorzystać.

W artykule wykorzystany jest oscyloskop cyfrowy Siglent SDS1104X-U. Urządzenie to jest dość drogie, jak na budżet amatora, gdyż kosztuje około dwóch tysięcy złotych, ale do niedawna był to najlepszy dostępny na rynku oscyloskop cyfrowy dla hobbystów. Czytelnik może stwierdzić „Hola, hola, autorze! Mam inny model innej firmy. Co robić, jak żyć?”. Odpowiedź jest prosta: Czytelniku, zajrzyj do instrukcji i sprawdź. Dla ułatwienia zostaną podane angielskie nazwy tych funkcji. Ponadto jako teoretyczny przykład posłuży też oscyloskop Hantek DSO2D10 kosztujący niecałe 1200 złotych, a dokładniej jego instrukcja obsługi. Model ten nie jest może tak dobry, jak oscyloskop Siglenta, i ma też tylko dwa kanały, ale w zamian ma wbudowany generator funkcyjny, a w większości sytuacji to wystarczy. Zresztą, większość oscyloskopów cyfrowych ma jedną użyteczną funkcję, która pozwala uzyskać więcej kanałów, niż oscyloskop posiada, i od tego właśnie zaczniemy.

Wykresy referencyjne, tryby wyświetlania i inne triki

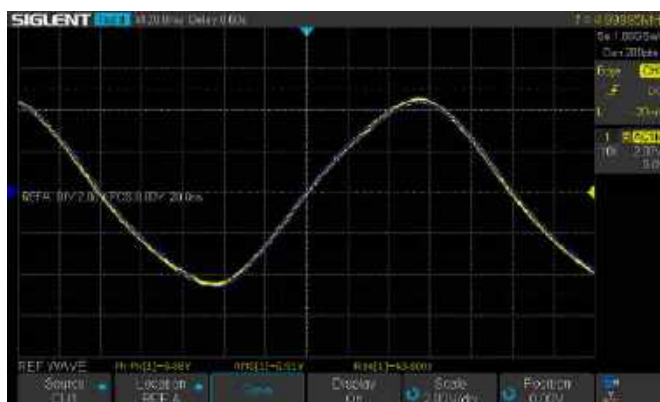
Niewielu ludzi o tym pamięta, ale oscyloskop cyfrowy ma dwa systemy pamięci. Pierwszym jest, oczywiście pamięć próbek. Ale poza nią oscyloskop ma też pamięć systemową, w której „siedzi” bufor ekranu, bieżące parametry, wyniki pomiarów i inne rzeczy, o których będzie wspomniane przy okazji. Pamięć próbek zawiera wszystkie zmierzone w trakcie procesu akwizycji wartości chwilowe sygnału. W pamięci systemowej jednak znajduje się przetworzony obraz sygnału wyświetlany na ekranie. W efekcie pamięć systemowa może przechowywać obraz bieżącego wykresu do późniejszego wykorzystania. Przy okazji gdy wykonujemy zrzut ekranu, do zewnętrznej pamięci USB zapisywana jest kopia bufora ekranu. Można też zapisać zawartość pamięci akwizycji w formie pliku CSV. Pliki te można otwierać w kilku różnych programach i wyświetlać dane w formie wykresów sygnałów.

Załóżmy, iż musimy zaobserwować więcej sygnałów, niż mamy wejść oscyloskopu. Dla przykładu powiedzmy, iż chcemy zobaczyć, czy po włączeniu zasilania w różnych punktach układu pojawiają się potrzebne sygnały, ale tych punktów pomiarowych jest więcej, niż mamy dostępnych sond. Metoda jest dość prosta:

1. Podłączamy dostępne sondy do różnych punktów układu, ale jedna z nich powinna być tam, gdzie pojawia się pierwszy sygnał. Zwykle może to być linia zasilania układu. Ustawiamy to wejście jako wyzwalające.
2. Wyłączamy zasilanie układu i uzbrajamy wyzwalacz oscyloskopu.
3. Włączamy zasilanie. Oscyloskop „złapie” wzrost napięcia zasilania układu oraz wybrane punkty pomiarowe.
4. By zapisać i wyświetlić wybrane kanały należy nacisnąć przycisk Ref (Siglent) lub Save/Recall (Hantek). W tym menu wybiera się kanały do zapisania i w których pamięciach będą zapisane. Przebiegi należy zapisywać do pamięci referencyjnej, by móc je potem wyświetlić, używając tego samego menu.
5. Przelączamy sondy, z wyjątkiem pierwszej, do innych punktów układu i powtarzamy kroki 2...4. Hantek ma tutaj przewagę nad Siglentem, gdyż posiada aż 9 pamięci wykresów.
6. Po ostatniej akwizycji wracamy do menu Ref lub Save/Recall i wyświetlamy zachowane wykresy.

Ważna uwaga, w trakcie pozyskiwania wykresów nie można zmieniać podstawy czasu. Wykresy referencyjne są niemodyfikowalne. Jednym z praktycznych zastosowań poza zwiększaniem liczbeości kanałów jest porównywanie sygnału uzyskanego do sygnału pożądanego, gdy na przykład testujemy lub kalibrujemy jeden układ względem innego, w pełni sprawnego. **Rysunek 1** prezentuje wykres sygnału sinusoidalnego (żółty) oraz wykres referencyjny (niebieski), jak widać, ślady się w tym przypadku pokrywają.

Oscyloskopy posiadają wiele różnych trybów wyzwalania, niektóre, jak „Runt”, przeznaczone są do identyfikowania występujących sporadycznie zakłóceń w trakcie normalnej pracy układu. Problem z tymi zaawansowanymi trybami jest taki, iż trudno jest ustawić parametry wyzwalania dla losowego, bliżej nieznanego zjawiska. Ale i na to jest relatywnie prosty sposób – funkcja nazywana „Persistence”. Funkcja ta w pewnym sensie symuluje zachowanie tradycyjnej lampy oscyloskopowej lub kineskopowej. W obu tych lampach obraz jest rysowany przez skupioną wiązkę elektronów, która uderzając w pokrywający front luminofor powoduje, iż ten zaczyna emitować światło widzialne. W lampach oscyloskopowych luminofor był dość „powolny”, tj. emitował światło przez jakiś czas po trafieniu przez wiązkę elektronów. To zachowanie przez lata było trudne do emulacji w oscyloskopach cyfrowych, i dlatego nie potrafiły one dobrze wyświetlać na przykład sygnału modulacji AM. Współczesne





są uśredniane, przez co ich amplituda spada. Tryb ten pozwala łatwiej zaobserwować sygnały o małej amplitudzie tonące w silnym szumie, ale częstotliwość odświeżania znacząco spada. Jeśli ustawimy uśrednianie z 64 próbek, to tyle kolejnych akwizycji jest sumowanych i uśrednianych by wyświetlić jeden obraz przebiegu. Ostatni tryb występuje pod różnymi nazwami: Siglent nazywa go „Eres”, a Hantek „HR” od słów „High Resolution”. Tryb ten to inna forma uśredniania – tym razem sumowane są kolejne próbki, a nie kolejne przebiegi. Zwiększa to rozdzielczość pionową, a dokładniej pozorną rozdzielczość przetwornika ADC. Uśredniając cztery kolejne próbki zyskujemy jeden bit rozdzielczości, uśredniając szesnaście próbek – dwa bity. 64 próbki dadzą 3 bity rozdzielczości. Ceną jest pasmo przenoszenia – każdy dodatkowy bit rozdzielczości dzieli maksymalną częstotliwość przez cztery. Inaczej pisząc dodanie jednego bita do ADC tą metodą ucina pasmo z 100 MHz do 25 MHz. Z trzema bonusowymi bitami pasmo spada do ~1,5 MHz! Jako ciekawostkę dodam, iż ta technika, zwana oversamplingiem, dostępna jest dla każdego elektronika. Zainteresowanych hobbystów odsyłam do znakomitej noty aplikacyjnej AVR121.

Skoro jesteśmy już w ustawieniach akwizycji, to warto wspomnieć też o trybie XY. W oscyloskopach analogowych w trybie tym sygnał jednego kanału przemieszczał plamkę w pionie, a drugiego w poziomie. Pozwalało to w łatwy sposób obrazować zależności częstotliwości, amplitudy i fazy między dwoma sygnałami. Dla przykładu możemy zaobserwować przesunięcie fazowe wprowadzane przez wzmacniacz operacyjny podając na wejście X sygnał z wejścia wzmacniacza, a na wejście Y sygnał z wyjścia wzmacniacza. Przy okazji należy pamiętać o odpowiednim ustawieniu napięć na działkę. Załóżmy, że na wejście podajemy sygnał 100 mVp-p, a badany układ ma wzmocnienie 50×, to kanał X powinien być ustawiony na 20 mV/dz. A Y na 1 V/dz – w ten sposób wystarczy nam po pięć działek w pionie. Gdy oba sygnały są w fazie, na ekranie będzie ukośna linia o nachyleniu 45 stopni, zaczynająca się w lewym dolnym rogu. Jeśli nie ma żadnego napięcia niezrównoważenia (offsetu), linia będzie przechodzić przez środek układu współrzędnych. Jeśli wzmacniacz ma inne wzmocnienie, lub jeśli źle ustawimy wejścia oscyloskopu lub/i sondy, kąt linii będzie inny. Wraz ze zwiększaniem częstotliwości generowanego sygnału linia na oscyloskopie powoli stanie się elipsą. Dla sygnałów o przeciwnych fazach wykres przybierze kształt koła. Jeśli do wejść oscyloskopu pracującego w trybie XY podamy sygnały o tej samej fazie, ale jeden z nich będzie miał częstotliwość będącą wielokrotnością częstotliwości drugiego sygnału, na ekranie zobaczymy tzw. figurę Lissajous. Jeśli figura powoli się obraca, częstotliwości nie są wielokrotnościami lub występuje różnica w fazie.

Drugim zastosowaniem dla trybu XY oscyloskopu jest rysowanie zależności prądu od napięcia dla półprzewodników. Dla przykładu testowany komponent włączony jest między wyjście generatora, a masę, jedna sonda mierzy napięcie na tym komponencie, druga (zazwyczaj prądowa lub różnicowa) mierzy prąd płynący przez komponent. Kiedyś budowano przystawki oscyloskopowe, które integrowały w sobie prosty generator i wzmacniacz różnicowy do pomiaru prądu płynącego przez badany komponent.

Jest jedna rzecz, którą potrafi droższy Siglent, ale tańszy Hantek już nie, a która jest bardzo użyteczna, zwłaszcza gdy pracujemy z transmisją szeregową. Nie, nie chodzi o dekodowanie protokołów szeregowych – oba oscyloskopy to potrafią. W ustawieniach akwizycji ukryta jest funkcja o nazwie „Sequence Mode”. Jej szczególnym zastosowaniem jest badanie właśnie protokołów komunikacyjnych lub serii impulsów oddzielonych długimi interwałami nicości. Dla przykładu założymy, iż raz na 100 ms przesyłany jest pakiet danych

trwający ledwo 10 µs. W normalnym trybie pracy jeśli chcemy uchwycić więcej tych pakietów, musimy wydłużyć podstawę czasu do jednej sekundy lub więcej. Uchwycimy ledwo dziesięć pakietów, a między nimi będzie 9,99 ms bezużytecznych informacji. Używając pamięci sekwencyjnej możemy kazać oscyloskopowi złapać na przykład tysiąc kolejnych pakietów danych, ignorując całkowicie martwe czasy między nimi. Nie dość, że zbierzemy o wiele więcej danych, to jeszcze rozdzielczość pozioma takiej akwizycji będzie większa, niż w normalnym trybie. W normalnym trybie pracy rozdzielczość czasowa wynosi maksymalnie 200 ns na sekundę akwizycji. W trybie pamięci sekwencyjnej bez problemu uzyskamy rozdzielczość 2 ns dla dwóch kanałów. Dzięki temu nawet protokół SPI z zegarem 25 MHz nie będzie nam straszny. Co więcej, częstotliwość akwizycji dla Siglenta rośnie z 100 tysięcy przebiegów na sekundę do 400 tysięcy przebiegów. Zależność częstotliwości próbkowania i wielkości pamięci działa też w drugą stronę – w normalnym trybie dwa kanały próbkowane z maksymalną częstotliwością 500 MSa/s dadzą maksymalną długość rekordu wynoszącą 14 ms. Weźmy nasz przykład z pakietami trwającymi 10 µs i powtarzającymi się co 100 ms:

1. Ustawiamy kanały tak, by dobrze było widać wszystkie pożądane sygnały, w tym wyzwalenie.
2. Ustawiamy podstawę czasu tak, by było widać cały pakiet. Od tego parametru zależy, jak wiele segmentów będziemy mogli zapisać. W naszym przypadku może wystarczyć 2 µs/dz. Wciskamy przycisk „Run/Stop” by zatrzymać akwizycję.
3. Otwieramy menu „Acquire” przyciskiem na panelu przednim, a następnie wybieramy opcję „Sequence”.
4. Opcja „Segments” pozwala nam wybrać liczbę próbek. Możemy ustawić wartość kręcąc pokrętką uniwersalną lub kliknąć nim by wyświetlić klawiaturę numeryczną.
5. Gdy jesteśmy gotowi, rozpoczynamy akwizycję.
6. Po zakończeniu akwizycji możemy przejść do menu „History” by tam przeglądać zebrane próbki.

Jak widać, procedura jest dość prosta, a przy tym pozwala pozbyć się długich odstępów czasu między pożądanymi wykresami.

Tryby wyzwiania: co potrafią i do czego można je wykorzystać

Tradycyjne oscyloskopy analogowe generalnie miały jeden tryb wyzwiania: zboczem narastającym lub opadającym. Obecnie jednak nawet budżetowy oscyloskop oferuje szereg różnych opcji w menu „Trigger”, czyli właśnie wyzwiania. Pamięć akwizycji zachowuje przebieg nie tylko od momentu wyzwolenia, ale też buforuje część informacji sprzed wyzwolenia. Każdy oscyloskop cyfrowy pozwala wybrać kanał, który będzie wyzwalał akwizycję. Część oscyloskopów oferuje dodatkowe wejście zewnętrznego wyzwiania, pozwalające na synchronizację z zewnętrznym generatorem funkcyjnym albo przystawką pomiarową. Niemal wszystkie stacjonarne oscyloskopy cyfrowe mają wyjście sygnału wyzwolenia, co też pozwala na synchronizację, ale też na pomiar częstotliwości akwizycji oscyloskopu. Obie te funkcje (oraz sygnał Pass/Fail) używają zazwyczaj tego samego gniazda. Wracając do samego wyzwiania, domyślnym ustawieniem jest wyzwianie zboczem („Edge”) narastającym („Rising”), przy czym można regulować poziom napięcia, przy którym wyzwianie nastąpi. Dostępne jest też wyzwianie zboczem opadającym („Falling”) lub oboma. Ten typ wyzwiania jest odpowiedni dla znakomitej większości przypadków. Pokrewnym typem wyzwiania jest „Window”. Dla tego wyzwiania wybieramy górny i dolny próg działania. Jeśli sygnał przekroczy tylko dolny próg, nastąpi wyzwolenie zboczem narastającym. Jeśli sygnał przekroczy

tylko górny próg, nastąpi wyzwolenie zboczem opadającym. Jeśli sygnał przekracza oba progi, wyzwolenie nastąpi w obu przypadkach.

Trzecim typem wyzwiania, który może się przydać szczególnie przy testowaniu lub naprawie przetwornic impulsowych i innych sterowników jest wyzwianie „Pulse”. Oscyloskop rozpocznie aktywność wtedy, gdy dodatni lub ujemny impuls trwa dłużej lub krócej niż zdefiniowana wartość. Dzięki temu możemy nie tylko sprawdzić, czy sygnał PWM osiąga pożądaną minimalną lub maksymalną szerokość impulsu, ale też wykryć problemy w łączności szeregowej związane z niespełnianiem granicznych wartości czasów trwania sygnałów. Innym użytecznym typem wyzwiania jest wyzwianie nachyleniem zbocza („Slope”). Użytkownik wyznacza dwa punkty wyzwiania: górny i dolny, oraz czas między nimi. Wyzwalanie nastąpi wtedy, gdy badany sygnał przejdzie od dolnego do górnego limitu (lub od górnego do dolnego) w czasie krótszym lub dłuższym od wyznaczonego. Inaczej pisząc wyzwianie nastąpi wtedy, gdy zbocze sygnału narasta lub opada szybciej niż wolniej od wyznaczonej wartości. Przykładem zastosowania może być badanie, czy sygnał protokołu I²C, lub 1-Wire ma wystarczająco strome zbocza dla poprawnego działania. Długie przewody i pojemności pasożytnicze oraz zbyt duże wartości rezystorów podciągających mogą te czasy spowolnić, co objawi się problemami w łączności.

Wyzwalanie „Dropout” (Siglent) lub „Overtime” (Hantek) przydaje się wtedy, gdy szukamy zaniku sygnału. Wyzwolenie nastąpi gdy między pierwszym zboczem narastającym (lub opadającym), a następnym zboczem opadającym (lub narastającym) upłynie dłuższy czas, niż ustawiono. Inaczej pisząc wyzwolenie nastąpi, gdy sygnał przestanie się zmieniać, co pozwala wykryć sytuację, gdy na przykład oscylator lub timer przestaje generować sygnał, albo układ przestaje wysyłać impulsy synchronizacyjne. Pokrewnym typem wyzwiania jest „Interval”, w którym wyzwianie następuje wtedy, gdy między zboczami narastającym i opadającym minie więcej lub mniej czasu, niż ustawiono. W przeciwieństwie do „Pulse” limit czasu odnosi się do pojawienia się następnego zbocza tego samego typu. Dodatkowo „Interval” daje więcej opcji, niż „Dropout/Overtime”. Wyzwalanie to pozwala zdefiniować dokładnie kryteria czasowe, co przydaje się, gdy badany układ generuje impulsy o ściśle określonych tolerancjach. Dobrym przykładem może być mikrokontroler, który wysyła sygnały resetu do zewnętrznego układu nadzorczego typu Watchdog Timer, i jeśli te impulsy pojawiają się zbyt szybko lub zbyt wolno po sobie, Watchdog resetuje mikrokontroler. „Interval” pozwoli ustawić limity czasowe układu Watchdog i uchwyci moment, gdy mikrokontroler resetuje go za szybko lub zbyt wolno. Ciekawym typem wyzwiania, który może się przydać przy pracy z układami logicznymi jest wyzwianie typu „Pattern”. Dla każdego dostępnego kanału można wybrać, czy aktywnym stanem jest zbocze narastające czy opadające, czy też kanał jest ignorowany. Następnie dla każdego kanału wybierany jest próg działania. Na końcu użytkownik wybiera operację logiczną, która jest wykonywana na wybranych sygnałach by spowolnić wyzwolenie. Siglent ma dodatkowy parametr czasowy, i cztery operacje logiczne do wyboru zamiast dwóch. Praktycznym zastosowaniem jest ustawienie kombinacji sygnałów, która jest nieprawidłowa w danym układzie i sprawdzenie, czy w trakcie pracy badanego układu pojawia się powodując jego złe działanie.

Wspomniana była metoda poszukiwania zakłóceń, które mogą pojawić się w sygnale, na przykład szpilek. Wyzwalanie „Runt” (Siglent) lub „Under Amp Trigger” (Hantek) pozwala wychwycić takie szpilki. Podobnie do wyzwiania „Window” mamy dwa poziomy do ustawienia: górny i dolny. Zależnie od ustawionej polaryzacji oscyloskop uchwyci sygnał, który narastając przekroczył dolny

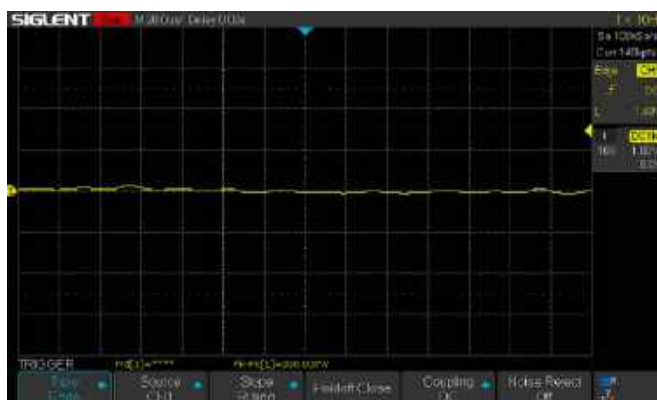
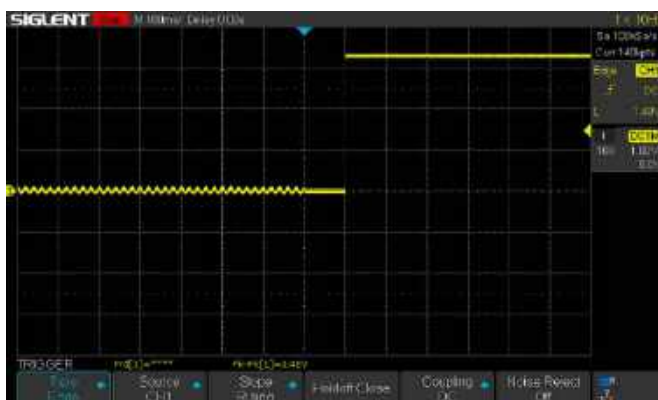
poziom, ale nie osiągnął górnego, albo opadając przekroczył górny poziom, ale nie osiągnął dolnego. Dzięki temu można zobaczyć niepożądane szpilki, które są na tyle duże by układ cyfrowy uznał je za poprawne zmiany stanu.

Chyba najrzadziej używanym typem wyzwiania dla hobbystów jest wyzwianie sygnałem wideo. Zarówno Hantek, jak i Siglent, oraz wiele innych oscyloskopów pozwala na wybór systemu między PAL i NTSC oraz wybór linii obrazu, którą chcemy zaobserwować. Siglent dodaje do tego jeszcze formaty HD oraz własny użytkownika. Wszyscy producenci ignorują system SECAM.

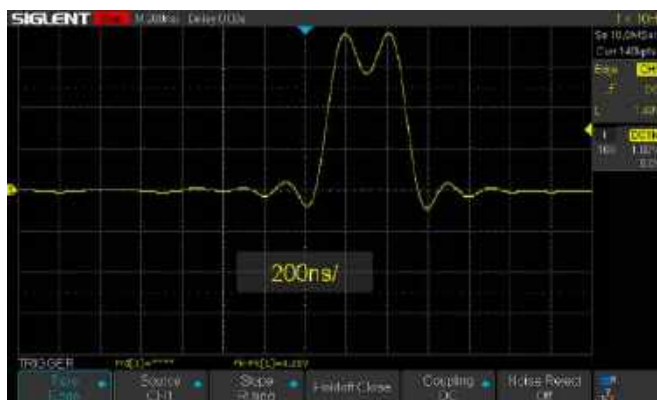
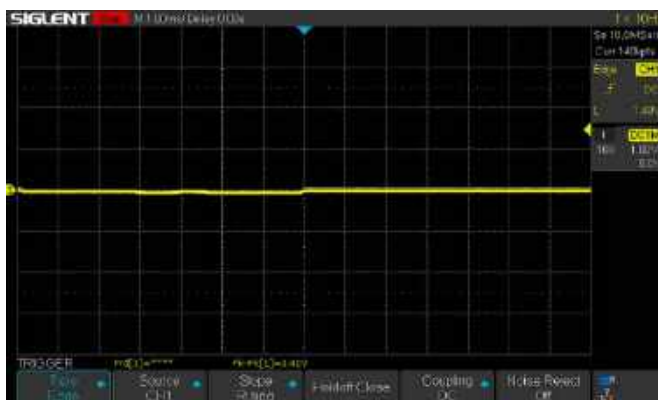
Wyjątkowo cennym dla hobbysty pracującego z mikrokontrolerami jest zestaw typów wyzwiania dla różnych protokołów transmisji szeregowej. Zarówno Siglent, jak i Hantek wspierają protokoły UART, LIN, CAN, SPI oraz I²C. Dla UART wyzwianie może być ustawione dla bitów startu, stopu, specyficznej zawartości pakietu lub dla błędu parzystości. Dla protokołów LIN, CAN i I²C oba oscyloskopy wspierają wszystkie funkcje, jakie one oferują, w tym wyzwianie gdy pojawi się określony adres lub identyfikator albo dane. Siglent dzięki posiadaniu czterech kanałów ma przewagę przy protokole SPI, gdyż można wykorzystać je do połączenia linii CLK, MISO, MOSI i CS. W przypadku oscyloskopu dwukanałowego podłącza się tylko linię CLK oraz MOSI lub MISO. Wraz z wyzwianiem transmisją szeregową oba oscyloskopy udostępniają też ich dekodowanie. Warto zaznaczyć, iż przez długi czas funkcje wyzwiania i dekodowania protokołów sieciowych w tanich oscyloskopach były albo w ogóle niedostępne, albo odpłatne. Nawet swego czasu ekstremalnie popularny Rigol DS1052E nie udostępniał wielu tych funkcji, chyba że użytkownik zhakuje swój oscyloskop. Rigol nie miał z tym problemu tak długo, jak hakowaniem parali się hobbysci i profesjonalisci używający ich oscyloskopu w domu. Firmy jednak nigdy by nie zaryzykowały łamania zabezpieczeń oscyloskopu ze względu na ryzyko bardzo kosztownego pozwu.

Wspomniane było wyzwianie sygnałem zewnętrznym. Jest to zawsze wyzwianie zboczem narastającym lub opadającym, i zarówno Hantek, jak i Siglent to potrafią. Siglent dodatkowo oferuje też wyzwianie napięciem sieciowym. Przydaje się to wtedy, gdy szukamy zakłóceń w układzie pochodzących od zasilacza lub instalacji elektrycznej. Podobnie do wyzwiania sygnałem wideo nie jest to funkcja, z której często się korzysta. Warto jednak wspomnieć o kilku dodatkowych opcjach w systemie wyzwiania. Można wybrać, czy układ wyzwiania ma mieć sprzężenie AC czy DC, filtr szumów, filtr w.c.z. i m.c.z. oraz możliwość opóźnienia wyzwolenia. Opóźnienie, czyli „Holdout” pozwala zablokować wyzwianie na określony przez użytkownika czas. Na przykład jeśli mamy do czynienia z długą serią impulsów, możemy ustawić wyzwolenie na zboczu pierwszego z nich, a potem wstrzymać wyzwianie aż do końca serii. Może się to szczególnie przydać, gdy oglądamy niestandardową transmisję szeregową.

Na zakończenie tej sekcji musimy omówić problem, który może dotknąć każdego z nas: amatora, profesjonalistę, młode osoby, jak i te posunięte w latach. To nie jest temat do śmiechu, gdyż, choć nie zdarza się to zbyt często, to każdy z nas może mieć problem z przedwczesnym wyzwianiem. **Rysunek 3a** przedstawia typowy przykład pomiaru linii zasilania układu przy jego uruchamianiu. Ustawiona jest długa podstawa czasu, 100 ms/dz. by sprawdzić, czy napięcie osiąga pożądaną wartość. Próg wyzwiania ustawiony jest na 1 V. Oscyloskop ustawiony jest w trybie „Single”, po czym włączany jest nasz układ. Ale na rysunku wyraźnie widać, iż wyzwolenie było przedwczesne. Przybliżmy więc uchwyciony obraz skracając podstawę na uchwycionym przebiegu. **Rysunek 3b** pokazuje to zbliżenie, a na nim nie widać przyczyny przedwczesnego wyzwolenia. Co tu jest grane?



Rysunek 3. Przedwczesne wyzwolenie (a), ale etiologia nieznaną mimo powiększenia (b)



Rysunek 4. Przy podstawie czasu 1 ms na działkę nadal nie widać przyczyny (a), ale po powiększeniu wszystko staje się jasne (b)

Odkryjemy problem, gdy ustawimy krótszą podstawę czasu, powiedzmy 1 ms/dz. Pokazuje to **rysunek 4a**, a zbliżenie jest na **rysunku 4b**. Przedwczesne wyzwolenie wywołują szybka szpilka na linii, symbolizującą w tym przypadku zakłócenia elektromagnetyczne od dużej przetwornicy. Dlaczego wcześniej jej nie było widać? Dlaczego wywołała one przedwczesne wyzwolenie pozostając ukrytymi przed naszym wzrokiem?

Układ wyzwolenia oscyloskopu cyfrowego jest oddzielnym obwodem, który można rozpatrywać jako wręcz kolejny kanał oscyloskopu. Przetwornik ADC gromadzi dane w pamięci akwizycji przez cały czas, ale to dopiero sygnał z tego obwodu zapisuje je i nakazuje głównemu układowi obrobienie ich i przeniesienie do bufora obrazu. Mechanizm wyzwolenia opiera się o prosty komparator, a bardziej zaawansowane typy wyzwolenia są realizowane programowo na bazie tego prostego sygnału. Jeśli Czytelnik zwróci uwagę na górny prawy róg ostatnich czterech rysunków, zauważy, iż częstotliwość próbkowania jest niższa, a pamięć próbek wynosi tylko 140 kpts. Dopuszczałem się tu pewnego oszustwa i celowo zredukowałem rozmiar pamięci próbek by ukazać problem przedwczesnego wyzwolenia. Przy wartości domyślnej 14 Mpts. te szpilki wciąż byłyby widoczne. Dlaczego? Gdyż mamy 14 działek w poziomie, więc na jedną działkę wypada milion punktów. Dla podstawy 100 milisekund na działkę daje to 100 ns na pojedynczą próbkę. Dla pamięci ograniczonej do 140 kpts. będzie to już tylko 10 μ s na próbkę – za mało by zarejestrować nasze szpilki. Oczywiście jeśli ustawić tryb akwizycji „Peak Detect”, to szpilki też byłyby widoczne.

W jakich sytuacjach może się pojawić problem przedwczesnego wyzwolenia?

1. Gdy podstawa czasu jest zbyt duża.
2. Gdy używamy pamięci sekwencyjnej.

3. W trybie „Eres/HR”, gdy kolejne próbki są uśredniane, przez co szpilka może zniknąć.
4. Gdy mamy ustawioną za małą pamięć próbek, lub gdy używamy więcej niż jednego kanału, przez co ilość dostępnej pamięci jest automatycznie podzielona między kanały.

Wszystkie te przypadki można podsumować jednym zdaniem: przedwczesne wyzwolenie nastąpi wtedy, gdy na jedną próbkę mamy sporo więcej czasu, niż trwa sama próbka.

Zakończenie

Ten artykuł nie wyczerpuje listy możliwości współczesnych oscyloskopów, a i nie omawia ich wszystkich zbyt szczegółowo. Nie zajęliśmy się w ogóle funkcjami matematycznymi, w tym FFT, ani nie pobawiliśmy się automatycznymi pomiarami czy pomiarami za pomocą kursorów. Moim celem jest zachęcenie Czytelnika do samodzielnego poznania tego niezwykle użytecznego instrumentu. W przyszłości ukaże się jednak artykuł omawiający kilka nieco bardziej zaawansowanych technik pomiarowych, wymagających zbudowania prostych przystawek lub posiadania generatora funkcyjnego. ■

Paweł Kowalczyk



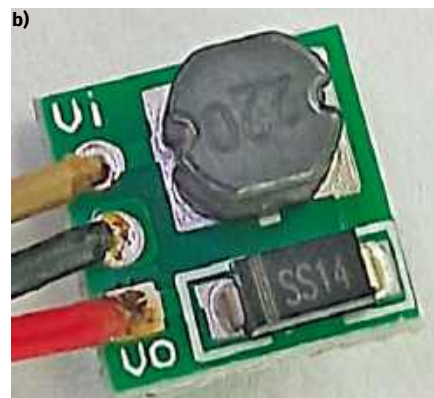
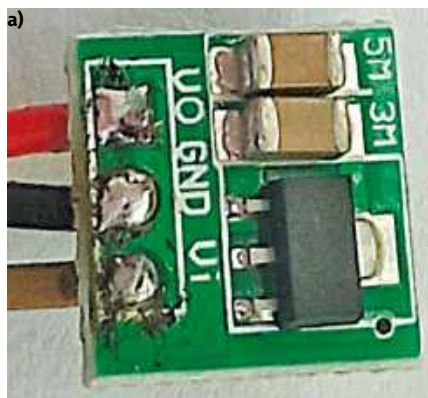
Zasilanie 5 V z pojedynczych ogniw baterii

Wiele urządzeń małej mocy zadowala się zasilaniem z jednej baterijki typu AA lub AAA. To zaledwie 1,5 V, aczkolwiek zastosowane obwody i układy wewnętrznej elektroniki wymagają napięcia wyższego, choć często prąd i wymagana moc pozostaje na bardzo niskim poziomie. Rozsądnym rozwiązaniem zamiast kilku ogniw baterii łączonych w szereg jest przetwornica step-up podwyższająca napięcie.

W przypadku zasilania baterijnego bardzo istotnym parametrem jest sprawność takiego przetwarzania energii. Dobrym przykładem urządzeń o takich właśnie wymagach są bezprzewodowe myszy komputerowe, piloty zdalnego sterowania itp. Wykonanie przetwornicy podwyższającej napięcie z elementów dyskretnych nie jest zadaniem prostym. Oczekujemy bowiem nie tylko dużej sprawności, ale aby cały układ był równocześnie zwarty, niewielkich rozmiarów i mieścił się w niewielkiej obudowie. Na szczęście opracowano układy scalone dedykowane dla takiego celu. Prezentowany w bieżącym projekcie układ wytwarza i stabilizuje napięcie 5 V z jednej baterijki 1,5 V lub innego ogniwa o napięciu niższym od wymaganych 5 V. Ilość wymaganych elementów jest zaskakująco mała, bo to zaledwie 3 elementy. Na **rysunku 1** widzimy płytkę dwustronną na której został zmontowany prototyp wykonany przez autora. Spis elementów potrzebnych do wykonania tego układu zamieszczono w poniższej tabelce.

Opis układu i jego działanie

Zaproponowany układ jest tak oszczędny w ilość komponentów dzięki wykorzystaniu układu scalonego CE8301 (IC1). To specjalizowany kontroler przetwornicy pracujący w klasycznej konfiguracji podwyższającej step-up. Zawiera wszystkie podzespoły sterownika wraz z tranzystorem kluczującym i elementami stabilizującego ujemnego sprzężenia zwrotnego. Zastosowana wersja jest 5 V, aczkolwiek istnieją wersje CE8301 pozwalające uzyskać napięcia od 1,8 V do 6,5 V z krokiem 0,1 V. Wersja IC1 wybrana przez autora projektu zamknięta jest w obudowie SMD o jedynie trzech wyprowadzeniach. Dzięki temu uzyskano tak zwartą budowę całej przetwornicy. Kluczowymi elementami dodatkowymi jest cewka L1 i dioda Schottky'ego D1. Z uwagi na wysoką częstotliwość kluczkowania wartość wymaganej indukcyjności



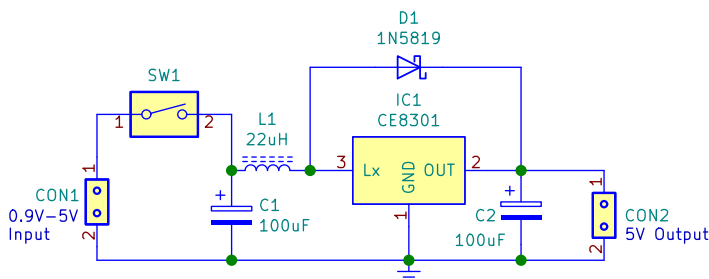
Rysunek 1. Prototyp wykonany przez autora; (a) i (b) – obie strony uzbrojonej płytki PCB

to zaledwie 22 μH , a dioda powinna być szybka, tu – 1N5819. Działanie przetwornicy bazuje na klasycznej modulacji szerokości impulsów, aczkolwiek modulowana jest też częstotliwość PFM. Dzięki temu uzyskano dużą sprawność przetwarzania energii w szerokim zakresie obciążenia mocą odbiornika. Poza wymienionymi wyżej trzema kluczowymi elementami, schemat zawiera jedynie dwa dodatkowe kondensatory elektrolityczne, przycisk SW1 i ew. dwa dwupinowe złącza dla napięcia wejściowego i wyjściowego. Do wejścia CON1 można podłączyć napięcie z zakresu od 0,9 V do 5 V, zaś dedykowane jest 1,5 V z pojedynczego ogniwa baterijki. Zamknięcie przełącznika SW1 natychmiast uruchamia pracę przetwornicy i na złączu CON2 dostępne jest napięcie 5 V. Kondensator C2 na wyjściu filtruje ew.

tętnienia U_{wy} o częstotliwości kluczkowania. Natomiast C1 ulokowany na wejściu jest buforem obniżającym dynamiczną impedancję widzianą od strony wejścia przetwornicy. Istotną cechą w przypadku zasilania baterijnego jest niska wartość prądu zasilania sterownika. CE8301 wykonany jest w technologii CMOS i spełnia ten wymóg. Maksymalna częstotliwość kluczkowania sięga 100 kHz, a przetwornica startuje już przy zasilaniu na wejściu na poziomie 0,9 V. To pozwala na pełne wykorzystanie cennej energii

Wykaz elementów:

CE8301 (IC1) – 1 szt.
1N5819 (D1) – 1 szt.
Przełącznik on/off (SW1) – 1 szt.
cewka 22 μH (L1) – 1 szt.
Kondensator 100 $\mu\text{F}/16\text{V}$ (C1, C2) – 2 szt.
Złącze 2-pinowe – 2 szt.



Rysunek 2. Schemat ideowy układu

w bateryjce. Przy tak niskim zasilaniu, gwarantowany prąd na wyjściu jest na poziomie jedynie 1 mA, co jednak w wielu aplikacjach wystarcza. Spośród wielu dostępnych wersji CE8301X, autor wykorzystał podstawową wersję ze stałym U_{wy} o wartości 5 V.

Konstrukcja i testowanie układu

Zaprojektowany tu obwód zasilania 5 V oferuje niewielką moc i cały układ zawiera raptem kilka elementów. Dlatego nie pokazujemy projektu płytki PCB. Należy oczekiwać, iż w większości aplikacji istotne będą małe rozmiary całej płytki. Dlatego w razie potrzeby należałoby użyć płytkę dwustronną. Warto też wiedzieć, że tego typu gotowe moduły dostępne są w handlu, co zdecydowanie zaoszczędzi nakład pracy dla projektu, a jedynie może być potrzebna niewielka modyfikacja zakupionego modułu.

W przypadku użycia jednostronnej płytki PCB, układ scalony CE8301 montujemy od strony druku, ponieważ jest to element w obudowie SMD. Pozostałe elementy, kondensatory, cewkę i diodę zaleca się montować

możliwie blisko (jak to tylko możliwe) układu scalonego przetwornicy. To ogólne zalecenie w przypadku zasilaczy impulsowych kluczujących z dużą częstotliwością.

Uwaga od Redakcji EFY

W miejsce układu scalonego IC1 można zastosować zamiennik LTC3525-5, ME2108 lub BL8530. Jako kondensator C2 zaleca się zastosować elektrolit tantalowy. Jego pojemność powinna być co najmniej na poziomie 10 μ F, a dopuszczalne napięcie trzykrotnie wyższe od U_{wy} wypracowanego przetwornicą step-up (w tym przypadku 16 V). Oba kondensatory elektrolityczne zastosowane w układzie powinny być typu LOW ESR, a cewkę L1 powinna cechować niska rezystancja drutu z rdzeniem nie nasycającym się przy maksymalnym osiąganym prądzie wyjściowym zasilacza-przetwornicy. ■

A. Samiuddhin

Ten artykuł był wcześniej opublikowany na łamach „EFY”, listopad 2023 (efymag.com)

Od Redakcji EdW: Wydaje się, iż krytycznym elementem w tej prostej przetwornicze jest cewka L1. W każdej przetwornicy konfiguracji step-up występuje duża nieciągłość prądu po stronie wyjścia. To znaczy, przez dużą część czasu energia do obciążenia dostarczana jest przez kondensator na wyjściu (tu C2), a krótkie są odcinki czasu gdy energia jest pompowana przez diodę (tu D1). Ta „relacja czasów” jest tym gorsza im większy jest stosunek napięć wyjściowego do wejściowego. Dlatego nie należy się sugerować, iż przy niskim napięciu baterii gwarantowany prąd dostarczany z 5 V U_{wy} jest stosunkowo niewielki. Wartość maksymalna prądu w cewce L1 jest wielokrotnie wyższa. Fizyczna wielkość cewki wyznaczona jest głównie nie przez jej indukcyjność, a właśnie przez szczytową wartość prądu. Dlatego L1 nie może być tak „smarkata” jak można by sądzić po mocy tej mini-przetworniczki. Te same uwarunkowania występujące w przetwornicy step-up narzucają wymogi pojemności i ESR kondensatora wyjściowego, na którym „spoczywa obowiązek” dostarczania energii do obciążenia przez znaczną część czasu (okresu kluczowania).

REKLAMA

Publikujemy dla projektantów i programistów elektroniki

ELPORTAL.pl



Znajdziesz nas również na Facebooku: facebook.com/ElportalPL

Sięgnij po archiwalne wydania

Przesyłka **GRATIS**

ELEKTRONIKA dla WSZYSTKICH



Prenumeratorzy mają bezpłatny dostęp do e-wydań archiwalnych EdW starszych niż 24 miesiące



Zamów wygodnie na **www.UlubionyKiosk.pl**

eprasa.pl a64b3da2db

Bezpieczna przystawka żelazka do ochrony odzieży

Czy komuś zdarzyło się pozostawić na dłuższą chwilę nieruchome żelazko bez nadzoru podczas prasowania odzieży? Konsekwencje mogą być całkiem poważne, łącznie z wypaleniem dziury w prasowanym materiale. Sytuacje awaryjne zdarzają się, gdy np. zadzwoni telefon lub nasza uwaga zostanie przekierowana na coś ważniejszego. Bieżyący projekt wychodzi na przeciw takiemu zagrożeniu.

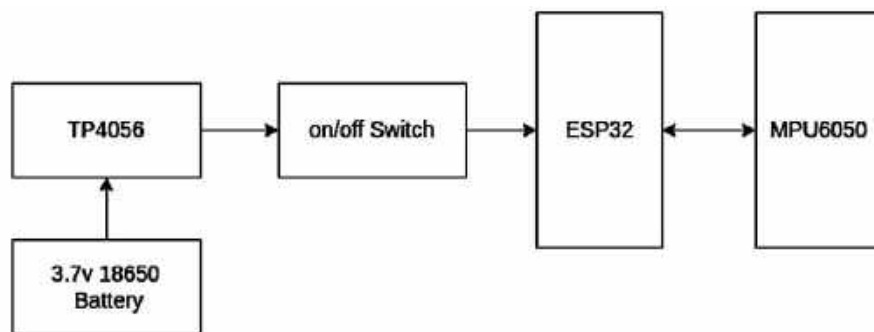
Działanie układu bazuje na rozpoznawaniu ruchu, i jeśli sensor stwierdzi przez dłuższą chwilę bezruch żelazka, odłączy je od zasilania sieci energetycznej. Podczas prasowania czujnik żyroskopowy lub akcelerometr rozpoznaje ciągły ruch, jak również pozycję żelazka (poziomą podczas prasowania lub pionową gdy żelazko poprawnie odstawimy). Układ wyposażono w łącze bezprzewodowe po którym czujnik ruchu przesyła informację do jednostki centralnej. Zasilanie pozostaje włączone tylko w sytuacji, gdy pozycja żelazka jest pozioma i rozpoznawany jest ciągły ruch. Ze względu na wykorzystanie bezprzewodowej łączności między czujnikiem a wykonawczą jednostką centralną, nasza przystawka składa się z dwu odrębnych części. Spis materiałów dla nadajnika i odbiornika zebrano oddzielnie w tabeli 1 i 2. Na rysunku 1 pokazano prototyp wykonany przez autora. W lewej części rysunku znajduje się nadajnik, a po prawej część odbiorczą proponowanej przystawki.



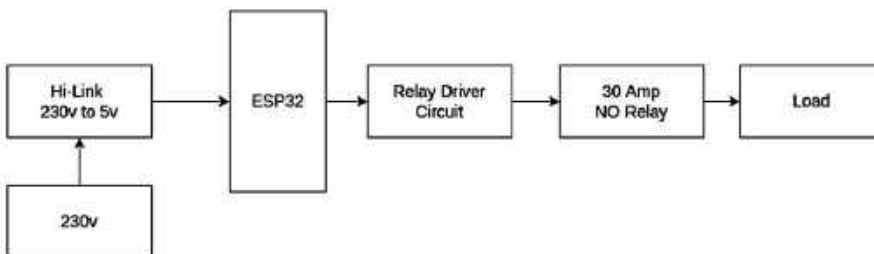
Rysunek 1. Prototyp wykonany przez autora (po lewej – część nadawcza; po prawej stronie – odbiornik i wykonawcza jednostka centralna)

Budowa układu i jego działanie

Na rysunkach 2 i 3 pokazano schemat blokowy obu części projektu. W części nadawczej znajduje się czujnik ruchu i położenia żelazka. Odbiornik steruje wysokoprądową częścią, gdzie elementem wykonawczym jest przełącznik. Kluczowym podzespołem nadajnika jest MPU6050. To akcelerometr potrafiący rozpoznać ruch (przyspieszenie) w trzech kierunkach przestrzeni x, y i z. Dane przesyłane są na bieżąco do mikrokontrolera ESP32 za pośrednictwem magistrali I²C. Głównym zadaniem mikrokontrolera jest dalszy przesył danych do odbiornika, również opartego na ESP32. Transmisja odbywa się zgodnie z protokołem ESP-NOW. W części nadawczej trzeba wykorzystać zasilanie bateryjne. Jest nim akumulator litowo-jonowy 3,7 V. Akumulator ten ładowany jest z typowej ładowarki 5 V USB. Układ kontrolera ładowania Aku zintegrowano z płytką nadajnika i jest nim powszechnie stosowany układ scalony



Rysunek 2. Schemat blokowy nadajnika



Rysunek 3. Schemat blokowy części odbiorczej

Tabela 1. Spis materiałów nadajnika

TP4056 (MOD1): ładowarka akumulatora z układem scalonym TP4056 – 1 szt.
 Bateria: akumulator litowo-jonowy 3,7 V typu 18650 – 1 szt.
 Switch (SW1): wyłącznik zasilania – 1 szt.
 ESP32 DEVKITC-32D (MOD2): płytka mikrokontrolera ESP32 – 1 szt.
 MPU6050 (MOD3): żyroskop i akcelerometr 3-osiowy xyz – 1 szt.

Tabela 2. Spis elementów odbiornika

Złącze śrubowe 01x02: złącze zasilania od strony wejścia oraz wyjściowe styki wykonawcze – 2 szt.
 Przetwornica 230 V AC na 5 V DC: zasilacz płytki odbiornika – 1 szt.
 1N4007 (D1): dioda prostownicza – 2 szt.
 1 kΩ (R1): rezystor 1 kΩ/0,25 W – 1 szt.
 BC547 (T1): tranzystor NPN – 1 szt.
 ESP32 DEVKITC-32D (MOD1): płytka mikrokontrolera ESP32 – 1 szt.
 RAYEX-L90 (RL1): przekaźnik typ-T 30 A – 1 szt.

TP4056. Część odbiorcza połączona jest z siecią napięcia przemiennego 230 V, zatem tutaj zasilanie stanowi 5-woltowa przetwornica niewielkiej mocy. Mimo wykorzystania mikrokontrolerów o sporych możliwościach, tutaj protokół transmisji jest bardzo ubogi. To w istocie zero-jedynkowa informacja binarna. Zasilanie żelazka ma być włączone lub wyłączone, w zależności od tego czy „jest w ruchu” czy żelazko spoczywa (w pozycji

poziomej lub pionowej). ESP32 w odbiorniku angażuje tylko jedną linię GPIO na której wystawia stan wysoki lub niski. Sygnał ten poprzez prosty driver na tranzystorze NPN steruje przekaźnikiem wykonawczym. Styki tego przekaźnika włączają lub odłączają zasilanie grzałki żelazka.

Schemat ideowy nadajnika i odbiornika pokazano odpowiednio na **rysunkach 4 i 5**. W nadajniku widzimy 3 moduły: MOD1

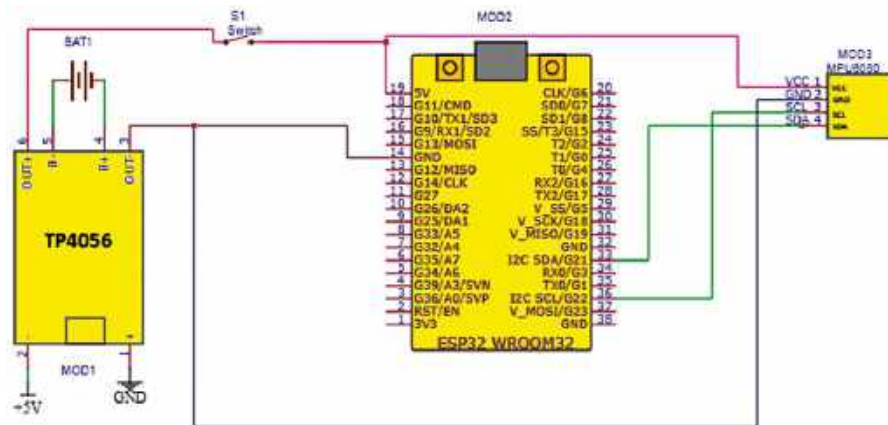


Rysunek 6. Jednokierunkowa komunikacja między modułami ESP32

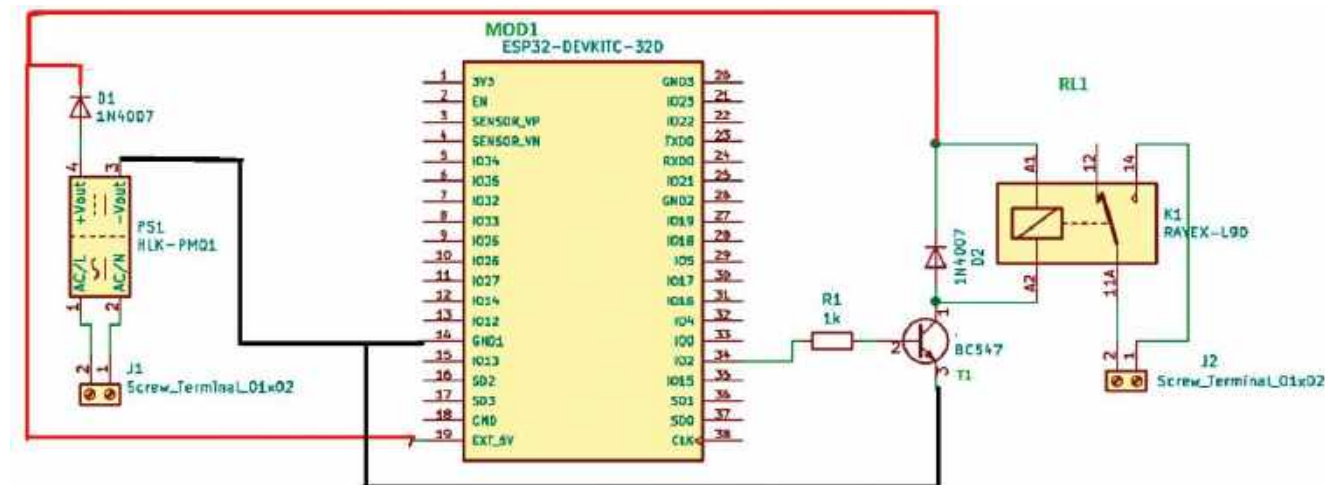
to moduł ładowarki akumulatora z elementem TP4056, MOD2 to płytka z mikrokontrolerem ESP32, a MOD3 to MPU6050 właściwy czujnik rozpoznający ruch lub spoczynek i kierunek montażu położenia żelazka. Część odbiorcza wykonana jest bardzo podobnie. MOD1 to ESP w wersji DEVKITC-32D. Kolejnym modulem jest przetwornica Power Source PS1 230 V AC na 5 V DC. Ponadto na schemacie z rysunku 5 widzimy jedynie wykonawczy przekaźnik RL1 oraz sterujący nim tranzystor T1 typu BC547.

Mimo prostoty układu, konieczny jest jakiś protokół transmisji bezprzewodowej. Wykorzystano protokół ESP-NOW, co pozwoliło nie angażować żadnego routera, jak również nie ma potrzeby łączności z internetem. Równocześnie komunikacja jest szybka i bezpieczna. ESP-NOW pozwala bezpośrednio łączyć inteligentne urządzenia i nie jest wymagana duża moc dla tej transmisji. W naszym przypadku komunikacja jest jednokierunkowa. Jeden ESP32 przesyła dane do drugiego ESP32 jak pokazano na **rysunku 6**.

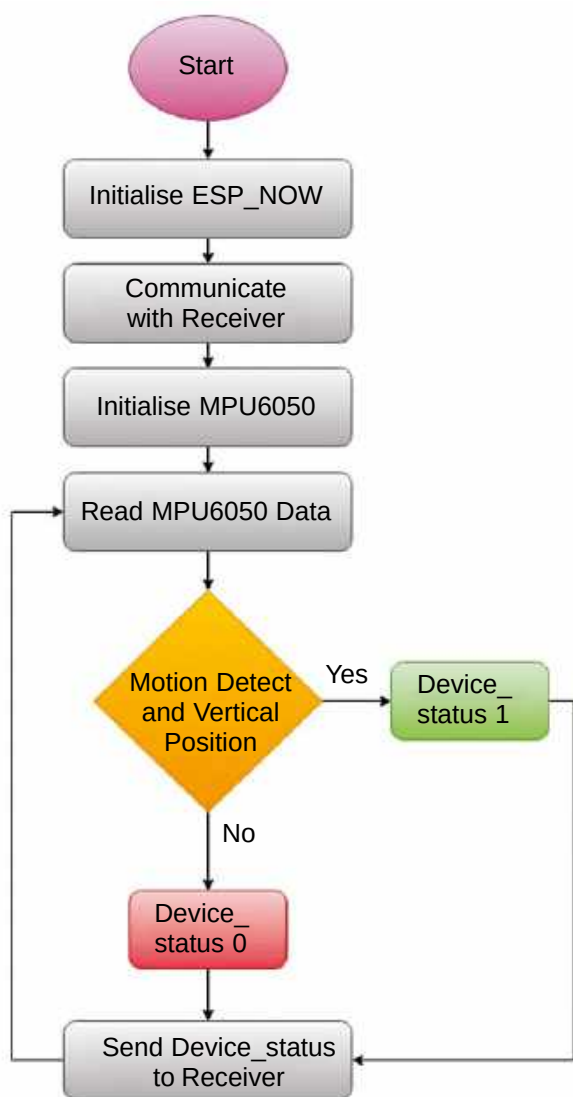
Implementacja protokołu ESP-NOW jest względnie prosta. Nasz projekt nie wymaga przesyłu dużej ilości danych. Całość sprowadza się do kontroli jednego wyjścia GPIO skonfigurowanego jako wyjście cyfrowe.



Rysunek 4. Schemat ideowy nadajnika



Rysunek 5. Schemat ideowy odbiornika



Rysunek 7. Flow chart kodu źródłowego nadajnika

Aczkolwiek przyjęty protokół pozwala na realizację zdalnego sterowania z prawdziwego zdarzenia. Jak również jeden nadajnik ESP może przysyłać dane do wielu bliźniaczych odbiorników.

Protokół ESP-NOW wymaga znajomości adresu MAC odbiornika. Ponieważ każdy moduł ESP32 wyposażony jest w swój własny niepowtarzalny adres MAC, w ten sposób łatwo jest identyfikować moduły mikrokontrolera dla których wysyłane dane są przeznaczone.

Oprogramowanie

Bieżący projekt angażuje dwa mikrokontrolery i wymagane jest oddzielne oprogramowanie dla każdego z nich. Oprogramowanie zostało wrzucane do modułów ESP32 za pomocą oprogramowania Arduino IDE. Dlatego należy rozpocząć od pobrania i zainstalowania tego IDE (<https://www.arduino.cc/en/software>). Następnie załaduj biblioteki dla protokołu ESP-NOW oraz dla MPU6050 (https://github.com/adafruit/Adafruit_MPU6050). Wybierz typ modułu mikrokontrolera jako ESP32 i przygotuj kod źródłowy dla nadajnika. W nadajniku kluczowym elementem jest czujnik żyroskopowy i akcelerometr MPU6050. Na nim spoczywa zadanie rozpoznania, czy aktualnie „machasz” żelazkiem prasując, czy spoczywa ono nieruchome. Flow chart struktury oprogramowania pokazano na rysunku 7.

Rysunek 8. Fragment kodu źródłowego nadajnika ustalający konfigurację ESP-NOW

Rysunek 9. Fragment kodu źródłowego odbiornika

Pierwszą czynnością po włączeniu zasilania jest nawiązanie komunikacji między nadajnikiem i odbiornikiem. Następnie po inicjalizacji czujnika MPU6050, program pracuje w pętli, sprawdzając, czy żelazko jest w ruchu, czy też spoczywa w bezruchu.

Na bieżąco też przesyłany jest status 0 lub 1 do odbiornika. Softwareowa część projektu sprowadza się do wczytania kodu źródłowego, uzupełnienia go o potrzebne biblioteki programów oraz do konfiguracji adresu MAC odbiornika do którego dane będą przesyłane.



Rysunek 10. Przystawka umocowana na żelazku

Zrzut ekranowy fragmentu kodu źródłowego, który zawiaduje nadajnikiem, pokazano na **rysunku 8**. Po skompletowaniu całości kodu, należy go załadować do ESP. W tym celu, w IDE należy ustawić właściwy typ płytki mikrokontrolera oraz odpowiedni numer portu po którym software ma być przesłany.

W celu przygotowania kodu źródłowego dla odbiornika postępujemy podobnie. Należy rozpocząć od inicjalizacji biblioteki ESP-NOW. Potwierdzeniem poprawności będzie odbiór wiadomości z nadajnika. Fragment kodu źródłowego sterującego odbiornikiem pokazano na **rysunku 9**. Tu zasadnicza część programu także polega na pracy w prostej pętli programowej. Zadanie sprowadza się do ustawienia stanu wysokiego lub niskiego na wybranej linii portu GPIO, w zależności od wiadomości odebranej z nadajnika.

Wgranie szkicu programu do odbiornika odbywa się w ten sam sposób jak w przypadku nadajnika. Kończącą częścią projektu jest montaż modułów nadajnika i odbiornika w odpowiednio przygotowanych obudowach. Rysunek 1 pokazywał prototyp wykonany przez autora. Na **rysunku 10** pokazano ostateczny wygląd nadajnika umocowanego na rękojeści żelazka. Odbiornik umieszczono w pobliżu gniazda zasilającego 230 V AC.

Teraz możesz być spokojny. Nawet jeśli zapomnisz wyłączyć żelazko lub pozostawisz je na prasowanej tkaninie, przystawka wyłączy je automatycznie. ■

S. Maheshwaran
K. Vairamani

Ten artykuł był wcześniej opublikowany na łamach „EFY”, czerwiec 2023 (efymag.com)

REKLAMA



m.technik
Ciekawi świata są zawsze młodzi

w prezencie na każdą okazję przejrzyś i kupisz na www.ulubionykiosk.pl

WYPRAWA NA ZIEMIĘ
Odkrywanie nieznanego planety



Krystian, Młodzi Entuzjaści Elektroniki,
Szkoła Podstawowa nr 86, Wrocław

Codziennie stajemy przed wyborami. Większość z nich to drobiazgi – kawa czy herbata, schody czy winda, papier czy plastik. Decyzje w takich sprawach wydają się błahе, ale to one kształtują naszą codzienność. Chcemy tego, czy nie, z reguły znacznie wcześniej, niż tego się spodziewamy, każdy z nas staje przed jakimś wyborem życiowej rangi, daleko wybiegającym ponad rodzaj napoju, środek transportu czy typ opakowania. Jednym z takich wyborów będzie z pewnością wybór typu szkoły ponadpodstawowej, później uczelni wyższej, lub decyzja o podjęciu jakiegoś rodzaju pracy zarobkowej.



Fotografia 1. Montaż zestawu „Wspomagacz wyboru” (AVTEDU639). Od lewej: Kornel, Kacper, Bartek, Błażej i Krystian. Młodzi Entuzjaści Elektroniki, Szkoła Podstawowa nr 86, Wrocław

To nie koniec życiowych wyborów, bo czekają Cię również i takie, które będą miały jeszcze poważniejsze konsekwencje. Na przykład wybór życiowej partnerki albo partnera. Może się również okazać, że nie mając parcia na biologiczne potomstwo postanowisz żyć poza konwenansami, na przykład, przepędzając czas na wolontariacie (choćby montując z dziećmi elektronikę), co przyniesie Ci największą w życiu frajdę i radość. Gdyby tak było pędz za tą myślą. Cudze definicje szczęścia nigdy nie będą kluczem do Twojego własnego, albowiem szczęście nigdy nie miało uniwersalnej recepty – każdy z nas musi odnaleźć je na swój sposób. Każdy musi odnaleźć własne.

Bez wątpienia układ elektroniczny, który zbudujesz za chwilę, nie sprawdzi się jako pomoc przy podejmowaniu życiowych decyzji. Te będziesz musiał podjąć sam i im

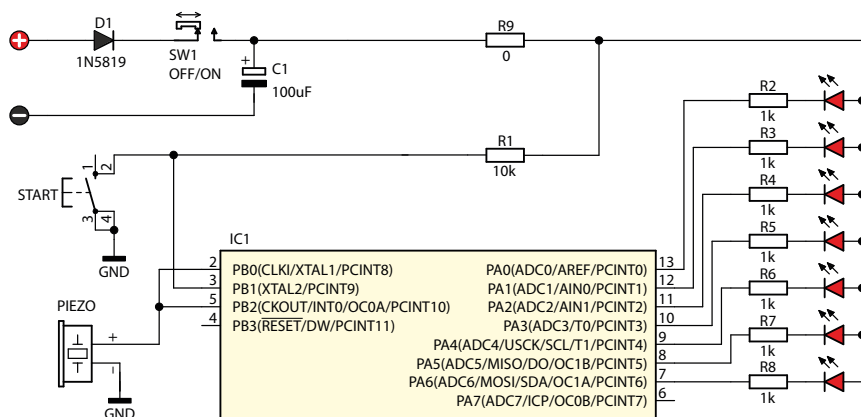
więcej w tym wyborze będzie Twojej samodzielności, tym lepiej dla Ciebie. Chybnym wyborów też pewnie nie unikniesz, trzeba się będzie wtedy podnieść, otrząsnąć i przeć dalej naprzód. Jak w tej piosence – Show Must Go On.

W życiowych decyzjach **wspomagacz wyboru** (AVTEDU639) z pewnością Ci nie pomoże. Sprawdzi się on natomiast wszędzie tam, gdzie konsekwencje wyboru będą w zasadzie żadne. Na przykład mając do wykonania kilka porządków domowych czasem ciężko się zdecydować, od czego zacząć. Odkurzanie, czy zmywanie podłogi? A może jednak mycie naczyń? Chociaż w dzisiejszych czasach zmywa najczęściej zmywarka, a podłogę czyści mopujący robot sprząający, i tak trzeba to wszystko odpowiednio uruchomić, nadzorować i serwisować. Czym więc zająć się w pierwszej kolejności? Ciężki orzech

do zgryzienia. Na taką ewentualność **wspomagacz wyboru** świetnie się nada!

Schemat elektryczny **wspomagacza wyboru** zaprezentowano na **rysunku 1**.

To już drugi podczas spotkań Juniorów układ elektroniczny wykorzystujący mikrokontroler, czyli układ scalony z wgranym programem (firmware), odpowiedzialnym za realizację zaplanowanych przez autora projektu zadań. Tym wcześniejszym układem, zbudowanym również z wykorzystaniem mikrokontrolera, był zestaw AVTEDU632 UFOLEDek, który opisany został w EdW 10/2024. Omawiając tamten zestaw postanowiłem przekonać Cię, że wkraczanie coraz głębiej w świat elektroniki wcale nie musi komplikować tematu. Jeśli pamiętasz, zastosowanie mikrokontrolera pozwoliło uniknąć realizowania wielu zadań za pomocą mniej i bardziej skomplikowanych obwodów analogowych,



Rysunek 1. Schemat ideowy układu

złożonych z dziesiątków pojedynczych elementów dyskretnych, takich jak rezystory, kondensatory, diody i tranzystory, i niemal wszystko to zastąpić pojedynczą kosteczką – mikrokontrolerem z wgranym oprogramowaniem (firmware) realizującym wszystkie zadania w sposób cyfrowy.

Podobieństwa układów budowanych z użyciem mikrokontrolerów

We **wspomagaczu wyboru** sytuacja wygląda dokładnie tak samo. Jeśli masz pod ręką egzemplarz EdW 10/2024 zerknij na stronę 83, gdzie, na rysunku 1 znajdziesz schemat UFOLEDka, bardzo podobny do dzisiejszego układu **wspomagacza wyboru**. Zwróć uwagę na wszystkie podobieństwa. Podpowiem, że w obu zestawach wykorzystano mikrokontroler tego samego producenta, firmy Microchip Technology Inc. W zestawie UFOLEDek zastosowano mikrokontroler ATtiny2313 a we **wspomagaczu wyboru** ATtiny24A. Początkowo mikrokontrolery te były produkowane przez firmę Atmel, ale Microchip przejął firmę Atmel w 2016 roku i obecnie to on kontynuuje produkcję oraz rozwój tych mikrokontrolerów.

Ponieważ przebiegi czasowe w obu wspomnianych zestawach (UFOLEDek i **wspomagacz wyboru**) nie są krytyczne, w żadnym z tych układów **nie zastosowano zewnętrznego rezonatora kwarcowego**. Zamiast tego, w roli generatora taktu zegarowego, dzięki któremu krok po kroku mogą być wykonywane kolejne zadania programu, wykorzystano (włączono programowo) wewnętrzny oscylator RC, co jest dość charakterystyczną cechą tych mikrokontrolerów, po którą konstruktorzy chętnie sięgają. Zawsze to jeden element mniej do zamontowania. W zasadzie trzy komponenty, bo odpada rezonator kwarcowy i dwa wspomagające go kondensatory.

Jest zatem prościej i taniej a proste układy będą działały tak samo, czy to na zewnętrznym rezonatorze kwarcowym, czy też wewnętrznym oscylatorze RC. Możliwość załączenia wewnętrznego oscylatora RC to jedno z podobieństw, którego nie widać wprost na schematach, a o którym warto było wspomnieć.

Kolejną sprawą, którą warto zauważyć, jest fakt, że żaden z układów **nie zawiera stabilizatora napięcia**, który obniżałby je do wartości bezpiecznej dla zastosowanych w obu przypadkach mikrokontrolerów, czyli do napięcia o wartości 5 V. Z użycia stabilizatorów zrezygnowano najpewniej dlatego, że oba te układy przeznaczone są do zasilania z dołączonego do zestawu koszyeczka na trzy baterie AAA (R03) o napięciu 1,5 V każda. Wszystkie trzy baterie są połączone wewnątrz koszyeczka w sposób szeregowy, to jest jedna za drugą. Kabelek w kolorze czarnym podłączony jest do minusa pierwszej baterii, plus pierwszej baterii podłączony jest z minusem drugiej baterii, plus drugiej baterii z minusem trzeciej baterii a plus trzeciej baterii z kablem czerwonym. Napięcia szeregowo połączonych baterii sumują się, dlatego pomiar dokonany za pomocą multimetru ustawionego w tryb woltomierza napięć stałych powinien wskazać napięcie około: $3 \cdot 1,5 \text{ V} = 4,5 \text{ V}$. Należy jednak pamiętać, że jest to napięcie nominalne, a świeżo zakupione pojedyncze ogniwo AAA będzie miało najprawdopodobniej około 1,6 V. Tym samym odczyt pomiaru napięcia rzędu 4,8 V na koszyeczku, do którego włożyłeś trzy świeżo zakupione baterie AAA nie powinien być sporą niespodzianką. Przesunie co najwyżej nieco w górę wartości w obliczeniach, których dokonamy poniżej.

Stosując trzy baterie o sumarycznym napięciu 4,5 V nigdy nie przekroczymy bezpiecznego dla wspomnianych mikrokontrolerów nominalnego napięcia

zasilania o wartości 5 V. Stąd też stosowanie stabilizatora napięcia nie było potrzebne. Musisz o tym bezwzględnie pamiętać, gdybyś mierzył się z pomysłem zasilania tych układów z zewnętrznego zasilacza wtyczkowego. Do tematu wrócę nieco później, w akapicie „Baterie, a może... zasilacz?”.

Na obu schematach widzimy identyczną diodę D1, która **zabezpiecza mikrokontrolery obu układów przed omyłkowym podaniem na nie napięcia o błędnej polaryzacji**, co dla nich samych z dużym prawdopodobieństwem skończyłoby się tragicznie. Ponieważ dioda przewodzi prąd w jednym kierunku (od anody w stronę katody katody) prąd z baterii do układu elektronicznego popłynie tylko wtedy gdy zostanie zachowana odpowiednia polaryzacja podanego zasilania. Gdy podłączymy bieguny baterii na odwrót, prąd nie popłynie wcale. Dla dociekliwych: w rzeczywistości przez diodę spolaryzowaną zaporowo (co będzie miało w miejsce w przypadku odwrotnej biegunowości podanego zasilania) popłynie minimalny prąd wsteczny (rzędu mikroamperów lub nawet mniej, w zależności od napięcia wstecznego i temperatury), jednak ten prąd będzie na tyle mały, że nie będzie miał on wpływu na działanie układu ani nie spowoduje jego uszkodzenia). Warto zauważyć, że w obu układach zastosowano diodę 1N5819, która jest nieco mniej znana, niż na przykład klasyczna dioda 1N4001 (tudzież 1N4004 czy 1N4007). Analizując notę katalogową diody 1N5819 zauważysz, że jest to dioda Schottky'ego, która różni się od klasycznej diody krzemowej (np. 1N4001) między innymi tym, że ma ona mniejszy spadek napięcia w kierunku przewodzenia. W przypadku użycia diody krzemowej spadek napięcia wyniósłby około 0,7 V, co oznacza, że w przypadku zasilania układu z koszyeczka z bateriami o sumarycznym napięciu 4,5 V za diodą D1 byłoby już tylko 3,8 V i napięcie to będzie spadało wraz z ze stopniem rozładowania kompletu baterii. Dlatego lepszy rezultat da w roli zabezpieczenia przed podaniem napięcia o odwrotnej polaryzacji dioda Schottky'ego o podobnych parametrach, na przykład 1N5819, na której spadek napięcia wyniesie około 0,3...0,4 V. Wówczas za diodą D1 napięcie zasilające układ (w tym mikrokontroler) powinno wynieść około 4,1...4,2 V. Różnica pomiędzy 3,8 V i 4,2 V może nie jest spektakularna, ale można założyć, że przy zastosowaniu diody Schottky'ego, na jednym komplecie baterii układ podziała dłużej.

Włącznik zasilania przed czy za elementem?

Czyżby kolejna decyzja do podjęcia? Przydałby się już ten wspomagacz wyboru (ciagle jakieś wybory i decyzje). Podpowiem, że tu sprawdziłby się wyśmienicie, bo konsekwencja wyboru będzie tu w zasadzie żadna. Dylemat dotyczy oczywiście widocznej różnicy w sposobie umiejscowienia przełącznika SW1 w obu układach. W przypadku UFOLEDka przełącznik ten znajduje się przed diodą D1, a dla odmiany we **wspomagaczu wyboru** przełącznik ten znajduje się dopiero za diodą D1. Nie ma znaczenia w którym miejscu nastąpi przerwa w obwodzie, bo w każdym z tych przypadków prąd przestanie płynąć. Za diodą D1, przed diodą D1 nie ma tu żadnego znaczenia. Równie dobrze przełącznik można byłoby umieścić pomiędzy czarnym („minusowym”) przewodem koszyczka z bateriami a dowolnym padem masy (GND) na płycie drukowanej. Gdy przełącznik znajdzie się w pozycji otwartej, prąd również, ani z baterii nie wypłynie, ani do niej nie wróci.

Kondensatory filtrujące

W obu układach znajdują się kondensatory filtrujące napięcie zasilające. W przypadku UFOLEDka rolę taką pełniły aż trzy kondensatory C1...C3 o wartości 100 μF połączone równolegle. Ponieważ pojemności łączone równolegle sumują się, napięcie filtrowane było pojemnością o wartości 300 μF . W układzie **wspomagacza wyboru** zastosowano tylko jeden kondensator C1 o pojemności 100 μF . Skąd taka różnica? Po pierwsze pojemność filtrująca nie jest wartością krytyczną. Potrzebna pojemność zależy od jakości stabilizacji napięcia zasilającego oraz od poziomu zakłóceń generowanych przez sam układ elektroniczny, w tym przez mikrokontroler wykonujący swoją pracę. Oba układy są domyślnie zasilane z koszyczka z bateriami, stąd napięcie wejściowe jest idealnie stabilne. Niemniej oba układy załączają i wyłączają sporą liczbę diod LED, co też powodować będzie minimalne skoki napięcia w obwodzie zasilania. Dla dociekliwych: skąd te skoki napięcia? Otóż wraz ze wzrostem liczby zasilanych elementów (na przykład diod LED załączanych przez mikrokontroler) układ elektroniczny (jako całość) będzie pobierał raz mniej raz więcej prądu z baterii. Bateria nie jest idealnym źródłem napięciowym, i wcale nie dostarcza zawsze tego samego napięcia. Napięcie na jej biegunach będzie malało nie tylko na skutek rozładowywania się baterii w czasie, ale również w sposób

chwilowy, wraz ze wzrostem wartości pobieranego z niej prądu. „Brakujące” napięcie, czyli różnica pomiędzy napięciem nominalnym baterii a napięciem, które można zmierzyć na jej biegunach po podłączeniu do niej obciążenia (np. serii kilkunastu diod LED) odłoży się na tzw. rezystancji wewnętrznej baterii. Rezystancję wewnętrzną można sobie wyobrazić jako rezystor szeregowo połączony z idealną baterią. Napięcie w obwodach zasilania (na przykład pomiędzy nóżkami VCC i GND mikrokontrolera) będzie zawsze różnicą pomiędzy napięciem nominalnym z baterii i napięciem, które odłoży się na tym rezystorze (rezystancji wewnętrznej) i, jak wspomniałem, będzie się ono zmieniało wraz ze zmianami obciążenia generowanymi przez pracujący układ. Takim wahaniom ma właśnie przeciwstawić się zastosowana pojemność filtrująca, którą w przypadku **wspomagacza wyboru** stanowi kondensator elektrolityczny C1 o pojemności 100 μF . W prostych układach wielkość tej pojemności często odbiera się doświadczalnie. Im większa pojemność tym lepsze filtrowanie tętnień o stosunkowo małej częstotliwości.

Rezystory

W obu zestawach znajdziemy po kilka rezystorów. Te o wartości 0 Ω , czyli wszystkie oznaczone na obu schematach cyfrą „0” można by było zastąpić kawałkiem dowolnego przewodnika. Więcej o roli tego typu rezystorów opowiedziałem na spotkaniu poświęconym UFOledkowi w EdW 10/2024 na stronie 84 w akapicie **Zwory na płytkach drukowanych (jumper wire)** więc nie tutaj tych treści powtarzać. Zignorujemy je zupełnie pamiętając o tym, że stanowią one połączenie (na schemacie można by je było zastąpić linią, a na płycie PCB ścieżką, oczywiście pod warunkiem, że w danym obszarze PCB nie wystąpiła by kolizja z inną miedzianą ścieżką).

Zewnętrzny Pull-Up na linii RESET, stosować czy nie?

Znów jakiś wybór. Znów trzeba podjąć decyzję, i ponieść jego konsekwencje. Spójrzmy na oba schematy. Oba mikrokontrolery mają sygnał RESET aktywny w stanie niskim (bliskim 0 V). Oznacza to, że do poprawnej pracy należy pin $\sim$ RESET (pojedyncza fala przez frazę RESET wskazuje, że mam na myśli sygnał aktywowany stanem niskim) ustawić w stan wysoki, na przykład połączyć za pomocą rezystora z „plusem” zasilania. Pewną niespodzianką jest brak rezystora podciągającego pin $\sim$ RESET

mikrokontrolera ATTiny24A w układzie **wspomagacza wyboru** do napięcia zasilania (**rysunek 1**). O ile w układzie UFOledka funkcję tę pełnił rezystor R1, w układzie **wspomagacza wyboru** pin 4 mikrokontrolera ATTiny24A, którego jedną z funkcji jest właśnie funkcja $\sim$ RESET „wisi w powietrzu”. Wiemy już, jakiego wyboru dokonała osoba projektująca układ. Czy to oznacza, że popełniła błąd konstrukcyjny i zapomniiała dodać ten ważny dla poprawnego startu układu element? I tak, i nie. Z jednej strony oba mikrokontrolery posiadają na linii $\sim$ RESET funkcję wewnętrznego pull-up’a, czyli podciągania do zasilania za pomocą wewnętrznego rezystora. Problem w tym, że funkcja ta jest albo domyślnie włączona albo domyślnie wyłączona. Gdyby się okazało, że funkcja ta jest domyślnie wyłączona, wówczas należałoby ją programowo aktywować. Aby to uczynić, mikrokontroler musiałby zostać najpierw uruchomiony i wykonać odpowiedzialne za to instrukcje. Ten paradoks sprawiłby, że zewnętrzny rezystor i tak należałoby zastosować, a ewentualne programowe załączenie (później) wewnętrznego rezystora pull-up na pinie $\sim$ RESET niejako traciłoby większy sens dla samego wystartowania programu. Sięgając po notę katalogową, szczerze przyznam, że trudno mi jest, tak na szybko, odnaleźć jednoznacznie brzmiącą odpowiedź na to, czy pull-up na pinie $\sim$ RESET jest domyślnie załączony, czy wyłączony, jednak w innym źródle trafiłem na informację, która mówi, że nawet gdy wewnętrzny rezystor pull-up jest załączony warto zastosować zewnętrzny rezystor o mniejszej rezystancji rzędu 10 k Ω w celu zwiększenia stabilności sygnału $\sim$ RESET, a tym samym zwiększenia odporności układu na zakłócenia elektromagnetyczne w eterze. Suma-summarum wychodzi na to, że również w przypadku ATTiny24A rezystor podciągający pin $\sim$ RESET należy podciągnąć do zasilania. Koniec i kropka!

Wcześniej wspominałem, że w układzie **wspomagacza wyboru** tego rezystora zabrakło. Czy to oznacza, że układ działał nie będzie? Będzie! Przekonasz się o tym jak tylko zbudujesz swój układ. Co najwyżej będzie on mniej odporny na zakłócenia i może się okazać, że zresetuje się samoczynnie, lub nawet zawiesi, na przykład po 10 minutach, godzinie, albo dopiero po roku, w zależności do zakłóceń na linii zasilania lub w eterze. Tyle, że nasz układ będzie zazwyczaj pracował kilka chwil, od załączenia przycisku SW1, wygenerowania

(wylosowania) wyboru w jakimś temacie po użyciu przycisku **START**, aż do ponownego wyłączenia przycisku SW1. Ryzyko płynące z braku tego rezystora jest więc w zasadzie żadne.

Na schemacie **wspomagacza wyboru** znajduje się rezystor R1 o wartości 10 kΩ, i owszem podciąga on do zasilania ale zupełnie inny pin: PB1 mikrokontrolera. Zamiast tego rezystora można by było załączyć wewnętrzny rezystor pull-up, tak jak miało to miejsce w układzie UFOledka, gdzie również zastosowano przyciski (S1...S3) ściągające po ich naciśnięciu napięcia obecne na wybranych pinach mikrokontrolera do masy (GND). Nie stosowano tam zewnętrznych rezystorów podciągających te piny do napięcia zasilania. Zamiast tego programowo (w kodzie programu) załączono po prostu rezystory pull-up wbudowane w strukturę mikrokontrolera, co dało dokładnie taki sam efekt. W układzie **wspomagacza wyboru** autor zdecydował o zastosowaniu zewnętrznego rezystora (nie musiał ale miał do tego prawo). Pin PB1 podłączony jest również do przycisku **START**, który po naciśnięciu zwiera ten pin mikrokontrolera do masy oznaczonej na schemacie jako GND, czyli do minusa zasilania. W wyniku powyższego na pinie PB1 będzie panował albo, za sprawą rezystora R1, stan wysoki (napięcie bliskie napięciu zasilania) gdy przycisk start jest rozarty, lub stan niski (napięcie bliskie GND, czyli „minusa” zasilania), gdy przycisk **START** jest naciśnięty. Ponieważ mikrokontroler może odczytywać stan panujący na pinie PB1 i rozróżniać stany wysoki i niski, będzie mógł również, w zależności od odczytanego stanu, wykonać w bloku decyzyjnym szereg instrukcji zaszytych we wgranym programie. Spoglądając na schemat, zapewne już się domyślasz, że będzie to się wiązało z generowaniem dźwięku za pomocą **buzzera PIEZO** oraz sterowaniem **siedmioma świecącymi diodami LED**.

Do omówienia zostały właśnie te diody LED. W przypadku UFOledka, wszystkie anody diod LED były połączone razem a następnie za pomocą pojedynczego rezystora podłączone do plusa zasilania. Takie rozwiązanie owszem, działa, ale nie do końca jest ono eleganckie. W EdW 10/2024 na stronie 86 w akapicie o tytule „Gdzie podziały się szeregowo rezystory diod LED?” dokładnie wyjaśniłem dlaczego tak jest, a teraz tylko przypomnę, że w przypadku gdy do pojedynczego rezystora szeregowego podłączone zostanie wiele diod LED, będzie odkładało się na nim napięcie zależne

od płynącego prądu, czyli zależne będzie od tego, ile diod LED będzie w danym momencie włączonych. Tym samym jasność świecenia każdej pojedynczej diody LED będzie zależała od liczby diod LED załączonych przez mikrokontroler jednocześnie.

Takiego efektu nie będzie w układzie **wspomagacza wyboru** ponieważ rozwiązanie tutaj przyjęte jest dużo bardziej eleganckie. Jak możesz zauważyć na schemacie (rysunek 1) każda dioda LED ma teraz swój własny rezystor szeregowy R2...R8 o wartości 1 kΩ każdy. Ponieważ wszystkie anody diod LED są trwale podłączone do „plusa” zasilania, gdy tylko mikrokontroler ustawi stan niski na dowolnej liczbie linii PA0...PA6 odprowadzi tym samym prąd z odpowiadających tym liniom gałęzi rezystor 1 kΩ – dioda LED i diody te zaczną świecić. I w drugą stronę. Każda z tych diod zgaśnie, gdy mikrokontroler na odpowiadającej jej linii ustawi stan wysoki. Dlaczego? Ano dlatego, że po obu stronach takiej gałęzi z rezystorem szeregową i diodą LED będzie panowało wówczas napięcie bliskie „plusa” zasilania, stąd nie wystąpi różnica potencjałów i nie popłynie żaden prąd.

Dla wprawy w prostych obliczeniach znów sprawdzmy autora i policzmy sobie czy prąd płynący w gałęziach z diodami LED jest dla tych diod (i dla wyjść mikrokontrolera) bezpieczny. Pamiętajmy, że potrzebujemy znać prąd i napięcie przewodzenia użytej diody LED, przy jakich dioda może w sposób bezpieczny pracować. Do zestawu dołączono czerwone diody LED. Nie znamy nazwy ani kodu katalogowego dostarczonych diod, przyjmijmy więc, tak jak to uczyniliśmy w poprzednim miesiącu (EdW 1/2025, strona 85), że przez diodę może popłynąć prąd 20 mA i powinno się wtedy na niej odłożyć napięcie o wartości 2,2 V. Wartość rezystorów R2...R8 dołączonych do zestawu znamy (1 kΩ każdy). Napięcie zasilania to trzy baterie o wartości 1,5 V każda (zakładając, że baterie nie są rozładowane), czyli sumarycznie 4,5 V. Spadek napięcia na diodzie Schottky’ego D1 to powiedzmy 0,3 V. Dla uproszczenia pominiemy ewentualną różnicę pomiędzy minusem zasilania a stanem niskim, które wystawi na linii PA0...PA6 mikrokontroler. W powyższych warunkach każda z gałęzi dioda LED – rezystor zasilona będzie napięciem 4,5 V – 0,3 V czyli 4,2 V. Sprawdźmy zatem, jaki przez diodę (i porty mikrokontrolera) popłynie prąd, pamiętając jednocześnie o tym, że nie powinien on przekroczyć 20 mA, by nie uszkodzić diody LED. Nota katalogowa mikrokontrolera ATtiny24A podaje,

że dopuszczalny prąd na pojedynczym pinie I/O (wejścia/wyjścia) to aż 40 mA, więc tu bardziej trzeba się martwić o diodę LED niż sam mikrokontroler.

Sięgamy zatem po ulubiony wzór każdego elektronika, czyli po prawo Ohma, i wyliczamy co trzeba:

$$I = \frac{U}{R} = \frac{4,2V - 2,2V}{1000\Omega} = \frac{2V}{1000\Omega} = 0,002A = 2mA$$

Zauważyłeś już pewnie, że mamy tu spory zapas. Biorąc pod uwagę wcześniejsze założenia, przez diodę LED mógłby popłynąć prąd nawet 20 mA, a więc 10-krotnie większy niż ten wyliczony. Pewnie intuicja Ci podpowiada, że można by użyć rezystorów nawet 10-krotnie mniejszych, czyli 100 Ω w miejsce 1 kΩ:

$$I = \frac{U}{R} = \frac{4,2V - 2,2V}{100\Omega} = \frac{2V}{100\Omega} = 0,02A = 20mA$$

Zostań jednak przy wartości 1 kΩ (z resztą w zestawie dołączono tylko takie). Układ z jednego kompletu baterii podziała wtedy zdecydowanie dłużej, a intensywność świecenia diod LED nie ma w tym przypadku większego znaczenia. Chodzi przecież o wskaźnik z podpowiedzią wyboru, a nie efekt dyskotekowy. Warto też wspomnieć, że intensywność świecenia dołączonych do zestawu diod LED, nawet przy rezystorach 1 kΩ, i tak jest naprawdę imponująca.

Uruchomić, czy kupić zmontowany?

Kolejny wybór, choć raczej prosty, bo wydaje mi się, że zestawy serii AVTEDU dostępne są w sklepie AVT wyłącznie w wersji do samodzielnego montażu. Gdyby jednak było inaczej odpowiedź jest chyba oczywista. Pewnie że montujemy samodzielnie! Ot, cała frajda przecież! Do podejmowania tego typu decyzji nawet **wspomagacz wyboru** nam nie jest potrzebny. Ach, żeby wszystkie wybory w naszym życiu były tak proste! A zatem, do dzieła!

Montaż rezystorów

W zestawie znajdują się rezystory o różnej wartości. Siedem spośród tych rezystorów ma wartość 1 kΩ, jeden z nich ma wartość 10 kΩ a jeden 0 Ω. Rezystor o wartości 0 Ω ma dosyć charakterystyczny wygląd i rozpoznasz go po pojedynczym czarnym pasku na jego obudowie. Rezystor o wartości 10 kΩ ma kod paskowy złożony z kolorów ułożonych w następującej kolejności:

brązowy-czarny-pomarańczowy-żółty. Pozostałe rezystory powinny mieć wartość 1 kΩ i mają prawie taki sam układ kolorów jak rezystor 10 kΩ, tyle, że na trzeciej pozycji zamiast paska w kolorze pomarańczowym znajduje się pasek w kolorze czerwonym. Dla pewności każdy rezystor przed wlutowaniem do płytki warto zmierzyć. Gdybyś miał jakikolwiek problem z pomiarem rezystorów za pomocą multimetru, przypomnę, że na stronie <https://elportal.pl/do-pobrania-jako-materiał-dodatkowy-do-numeru-EdW-11/2024> dostępna jest do pobrania instrukcja robocza: **Pomiar wartości rezystorów za pomocą multimetru**. Gorąco zachęcam do jej pobrania, wydrukowania i postępowania każdorazowo wedle zamieszczonych tam instrukcji.

Po zamocowaniu i przylutowaniu wszystkich rezystorów płytka powinna wyglądać podobnie jak ta z **fotografii 3**.

Montaż diody Schottky'ego

Jako kolejną zamontuj proszę diodę D1, pamiętając o zachowaniu właściwej polaryzacji, czyli w taki sposób, by biała (a może raczej szary?) pasek na obudowie diody zwrócony był w stronę paska zaznaczonego białą farbą na obrysie tego komponentu na płytce drukowanej (**fotografia 4**). Gdybyś zamontował ją na odwrót, a koszyk z bateriami zamontował poprawnie, mimo załączenia przełącznika SW1 prąd w obwodzie nie popłynie. Jeśli natomiast niepoprawnie zamontowałeś diodę D1 oraz, dodatkowo niepoprawnie zamontowałeś również koszyk z bateriami, to prąd popłynie i uszkodzi mikrokontroler oraz po kilku chwilach spowoduje również dezintegrację obudowy (wybuch) kondensatora C1.

Montaż podstawki pod mikrokontroler

Jako następną w kolejności zamontuj proszę podstawkę pod mikrokontroler. Dzięki

niej, w przypadku pomyłki kierunku montażu albo konieczności zaprogramowania w przyszłości mikrokontrolera nowym kodem (wsadem, firmwarem) albo wymiany uszkodzonego mikrokontrolera na nowy zrobisz to bez najmniejszego wysiłku, podważając osadzony w niej mikrokontroler, na przykład płaskim śrubokrętem i wyciągając go z podstawki, chwilę później osadzając w jego miejsce nowy. Demontaż trwale przylutowanego bezpośrednio do płytki, mającego 14 nóżek mikrokontrolera, stanowiłby nie lada wyzwanie i wymagał dodatkowych narzędzi, takich jak odsysacz do cyny albo najlepiej stację rozlutowniczą z kompresorem. W przypadku płytki z miedzią po jednej stronie (ang. *single sided PCB*) jak ma to miejsce w przypadku płytki **wspomagacza wyboru** istniało by również spore prawdopodobieństwo zerwania (oderwania od ścieżek i płytki) pierścieni miedzi wokół otworów nóżek mikrokontrolera i koniecznością naprawy/odtworzenia zerwanego druku (ścieżek i padów).

Podstawkę należy przylutować do płytki w kierunku wyznaczonym przez wycięcie na jednym z krótszych boków mikrokontrolera oraz obrysu na płytce PCB dla tego komponentu. Jak zwykle, lokalizacja wycięć musi być taka sama na płytce i komponentencie (**fotografia 5**).

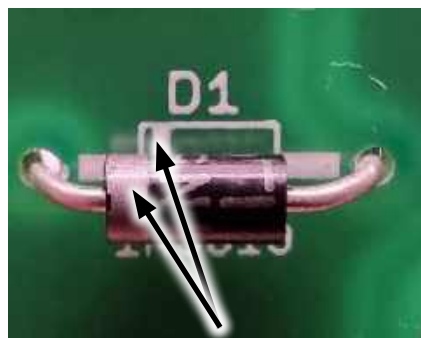
Montaż diod LED

Jako następne zamontuj proszę wszystkie diody LED opisane na płytce PCB jako LED1...LED7. Podczas ich montażu pamiętaj o zachowaniu właściwej polaryzacji. W przeciwnym wypadku diody LED nie będą świeciły. Do zestawu dołączono nowe diody LED z jeszcze nie przyciętymi wyprowadzeniami. W takich diodach LED, jak już pewnie bardzo dobrze pamiętasz, krótsze wyprowadzenia to katody i należy je zamontować w otworach oznaczonych literą „K” lub znakiem „-”, z kolei wyprowadzenia

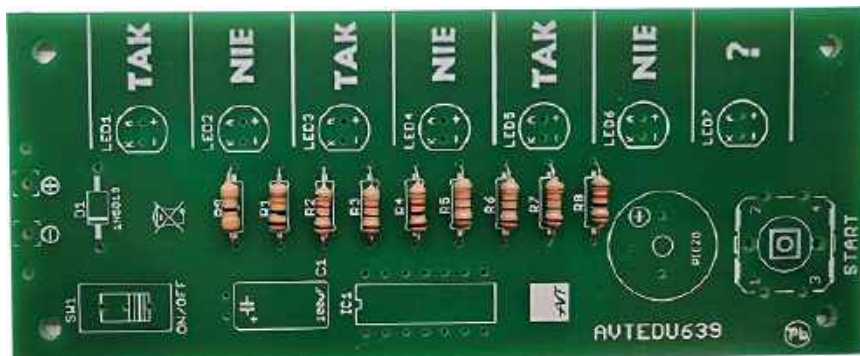


Fotografia 2. Kacper uzbraja płytkę wspomagacza wyboru w rezystory, Kornel już rozpoczął proces ich lutowania. Młodzi Entuzjaści Elektroniki, Szkoła Podstawowa nr 86, Wrocław

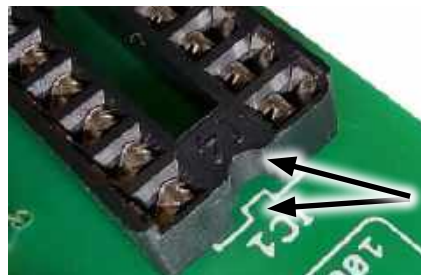
dłuższe to oczywiście anody, które należy umieścić w otworach oznaczonych literą „A” lub znakiem „+” (**rysunek 2a i b**). Jeśli czujesz „informacyjny niedosyt” nieco więcej szczegółów dotyczących diod LED oraz ich budowy znajdziesz w EdW 12/2024



Fotografia 4. Poprawny, zgodny z polaryzacją montaż diody D1. Strzałki na fotografii wskazują, że paski polaryzacji, jeden namalowany na obudowie diody, drugi namalowany na obrysie komponentu na płytce PCB znajdują się po tej samej stronie



Fotografia 3. Płytką z zamontowanymi rezystorami. Dostrzegasz różnicę w kolorze paska na trzeciej pozycji pomiędzy rezystorami R1 i R2? Jeśli „nie bardzo” to wiesz teraz, dlaczego preferuję pomiary za pomocą multimetru w miejsce postępowania się kodem paskowym



Fotografia 5. Poprawne ustawienie podstawki pod mikrokontroler względem płytki PCB. Wycięcie na jednym z krótszych boków podstawki musi być skierowane tak, by pokrywało się z wcięciem narysowanym na jednym z krótszych boków obrysu mikrokontrolera na płytce PCB



Fotografia 6. Kacper lutuje podstawkę pod mikrokontroler do płytki PCB. Zagięcie po włożeniu do płytki ale jeszcze przed rozpoczęciem procesu lutowania skrajnych dwóch wyprowadzeń podstawki (po przekątnej), np. nóżki 1 i 8, pozwala manewrować płytką, bez obawy o to, że podstawka wypadnie. Młodzi Entuzjaści Elektroniki, Szkoła Podstawowa nr 86, Wrocław

na stronach 84 i 85. Wspominałem tam o tym, że w przypadku większości diod LED, gdy spojrzeć na nie pod światło, większa elektroda (nie mylić z dłuższym wyprowadzeniem) na której zlokalizowany jest również chip półprzewodnikowy i odbłyśnik to z reguły katoda (**rysunek 2c**). „Z reguły”, ponieważ zdarzają się pewne odstępstwa, w związku z czym, kierowanie się tą ostatnią cechą diody LED bywa zwodnicze. W przypadku wątpliwości warto diodę LED przetestować za pomocą multimetru skonfigurowanego w tryb

pomiaru ciągłości obwodu. Jeśli przewody multimetru są poprawnie zamontowane (**fotografia 13**), po przyłożeniu czerwonej sondy multimetru do anody diody LED, a czarnej sondy multimetru do katody diody LED, dioda LED powinna się zaświecić.

W EdW 12/2024 wspominałem też, że cechą rozpoznawczą ujemnego wyprowadzenia (katody) jest również płaskie ścięcie na średnicy diody LED, widoczne zarówno na średnicy diody jak i na obrysie diody LED na płytce PCB (rysunek 2b i c).

Montaż kondensatora elektrolitycznego

W zestawie **wspomagacza wyboru** do zamontowania jest tylko jeden kondensator elektrolityczny, na schemacie i płytce opisany jako C1. Powinien mieć on wartość 100 μF , co, dla pewności, możemy odczytać na korpusie (obudowie) tego kondensatora (**fotografia 7**). Na jego korpusie możemy odczytać też dopuszczalne napięcie pracy tego kondensatora. Ja w zestawie znalazłem kondensator o pojemności 100 μF na napięcie 16 V. Co mówi nam wartość 16 V i czy należy się nią przejmować? O ile pojemność (w mikrofaradach) jest zazwyczaj wielkością krytyczną, od której wartości zależy zazwyczaj zachowanie układu opartego o różnego rodzaju zależności czasowe, o tyle dopuszczalne napięcie (w voltach) mówi tylko tyle, jakiego napięcia nie można przekroczyć w docelowym zastosowaniu. Nasz układ zasilany jest napięciem 4,5 V a po spadku napięcia na diodzie D1 będzie to niecałe 4,2 V. Oznacza to, że nigdy nie przekroczymy bezpiecznego dla tego kondensatora napięcia, czyli 16 V. Nic nie stoi

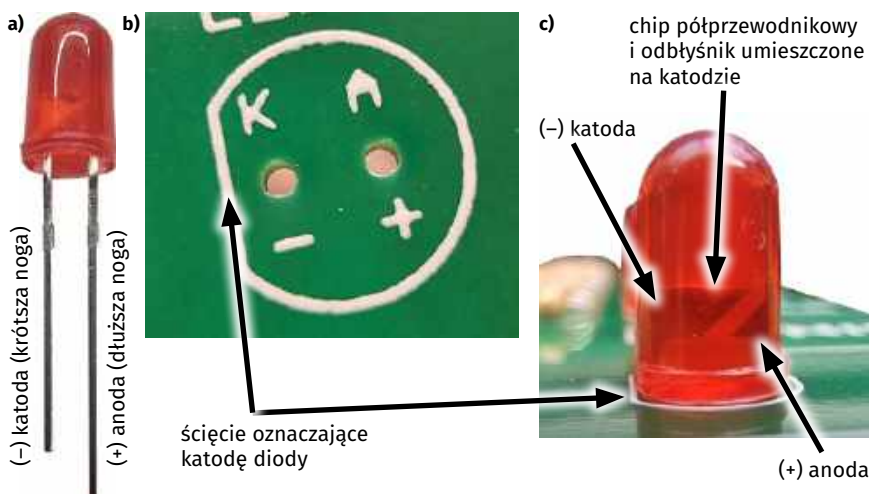
na przeszkodzie, by wykorzystać kondensator o pojemności 100 μF na wyższe napięcie, np. 100 μF na 25 V lub 100 μF na 50 V i często tak właśnie się robi, gdy nie mamy pod ręką kondensatora na sugerowane napięcie. Zawsze można wybrać kondensator na wyższe napięcie. Na niższe już nie. Dopóki napięcie w układzie nie przekracza tego jakie dopuszcza oznakowanie kondensatora elektrolitycznego, zachowanie tego kondensatora w układzie będzie dokładnie takie samo. Ważne jest jedynie aby pojemność była zgodna z wymaganą. Jedyny problem, na jaki możemy natrafić stosując w układzie kondensator elektrolityczny na nieco wyższe napięcie to brak wystarczającej ilości miejsca na płytce PCB, należy bowiem wiedzieć, że, zazwyczaj, im większe dopuszczalne napięcie pracy kondensatora elektrolitycznego, tym większy jego fizyczny gabaryt. Kondensator elektrolityczny, jak zapewne pamiętasz, jest elementem spolaryzowanym, i podobnie jak ma to miejsce w przypadku diod LED, tu również dłuższa nóżka nowego (nieprzyciętego jeszcze) elementu jest wyprowadzeniem dodatnim (+) a krótsza ujemnym (-). Poprawny montaż kondensatora elektrolitycznego C1 na płytce **wspomagacza wyboru** pokazano na **rysunku 3** oraz fotografii 7.

Obrys komponentu (prostokąt zamiast okręgu) sugeruje, że zamysłem projektanta był montaż poziomy kondensatora, dlatego, zgodnie z takim zamysłem, kondensator ten warto zamontować w pozycji leżącej (**fotografia 7**).

Montaż buzzera PIEZO

Przyszła pora na montaż pasywnego buzzera piezoelektrycznego, opisanego jako PIEZO. Słowo „pasywny” oznacza, że element ten sam z siebie nie generuje dźwięku. Nie ma on zaimplementowanej elektroniki, która generowałaby przebieg akustyczny, dlatego element ten pełni funkcję analogiczną do głośnika magnetoelektrycznego, choć jest on zbudowany zupełnie inaczej. Kilka słów na temat różnicy pomiędzy pasywnym buzzerem piezo a głośnikiem magnetoelektrycznym znajdziesz w EdW 9/2024 na stronie 88, gdzie omawiałem budowę Konsoli Audiochaos (AVTEDU624).

Buzzery piezo są z definicji elementami bez polaryzacji, bardzo często próżno więc szukać na ich obudowie tego typu oznaczeń. Nie ma ich także buzzer dołączony do zestawu (**fotografia 8a i b**) dlatego montażu tego elementu możesz dokonać w dowolnym kierunku. Niemniej zdarzają się też

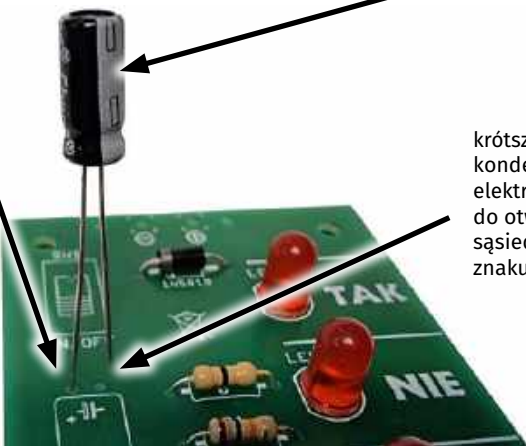


Rysunek 2. a) polaryzacja wyprowadzeń diody LED według długości wyprowadzeń, b) sposób oznakowania polaryzacji diod LED na płytce PCB (w otworach oznaczonych literą „K” lub znakiem „-” należy umieścić krótszą nóżkę diody LED, z kolei wyprowadzenia dłuższe należy umieścić w otworach oznaczonych literą „A” lub znakiem „+”); c) polaryzację w większości diod LED można rozpoznać również po wielkości elektrod wewnątrz obudowy

dłuższe wyprowadzenie kondensatora elektrolitycznego trafia do otworu oznaczonego znakiem „+”

oznakowanie „minusa” na korpusie kondensatora elektrolitycznego

krótsze wyprowadzenie kondensatora elektrolitycznego trafia do otworu sąsiedniego, bez znaku

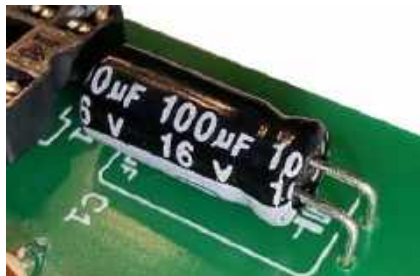


Rysunek 3. Poprawny, zgodny z polaryzacją sposób włożenia kondensatora elektrolitycznego C1 do płytki PCB. Na płytce znakiem „+” oznaczono otwór, w którym należy umieścić dodatnie (dłuższe) wyprowadzenie kondensatora elektrolitycznego a na korpusie (obudowie) kondensatora w sposób bardzo wyraźny zaznaczono wyprowadzenie ujemne kondensatora, które jest przy okazji wyprowadzeniem krótszym. Należy je zamontować do sąsiedniego otworu kondensatora elektrolitycznego (bez znaku)

niedobrze egzotyczne buzzery pasywne z zaimplementowaną diodą ochronną, redukującą napięcie wsteczne, bądź w których przyjęta konstrukcja powoduje, że buzzer działa optymalnie przy określonej polaryzacji. Wówczas z pewnością znajdziemy na takim elemencie znacznik polaryzacji. Dla takich elementów można wówczas skorzystać ze znacznika polaryzacji (znak „+”) widniejącego na warstwie opisowej płytki drukowanej (**fotografia 8c**).

Montaż przycisku SW1

Pora na zamontowanie włącznika zasilania opisanego na schemacie i PCB jako SW1. Przełącznik łączy swój pin środkowy z jednym z dwóch skrajnych, w którego kierunku jest w danym momencie skierowany hebelkę przełącznika. Z uwagi na taką konstrukcję kierunek montażu tego elementu nie ma żadnego znaczenia. Podczas montażu



Fotografia 7. Wygląd kondensatora elektrolitycznego C1 zamontowanego na płytce PCB w pozycji leżącej. Na zdjęciu widać również opisy dotyczące pojemności oraz maksymalnego napięcia pracy tego kondensatora. Widać też, że nóżka przy długim białym pasku na korpusie kondensatora (znak „minusa”) została umieszczona w otworze bez znaku

warto przylutować środkowy pin do płytki PCB, a po upewnieniu się czy komponent dobrze leży na płytce przylutować dwa pozostałe.

Montaż przycisku START

Przycisk start jest przedostatnim elementem do przylutowania. Kierunek montażu tego elementu również jest obojętny, ważne jest tylko to, aby wszystkie cztery nóżki przeszły przez otwory w płytce PCB a sam przycisk dobrze na niej leżał. Cała podstawa przycisku powinna dotykać powierzchni PCB (**fotografia 10**).

Montaż koszyeczka z bateriami

Przyszła pora na przylutowanie koszyeczka z bateriami. W pierwszej kolejności, jeśli końcówki kabelków nie są jeszcze „obrane” z izolacji, należy je odizolować na długości około 5 mm a następnie skrócić i pocynować miedziane końcówki. Jeśli nie bardzo wiesz, jak się za to zabrać, zerknij proszę do EdW 8/2024, gdzie podczas drugiego spotkania Juniorów omawiałem montaż lampki LED (AVTEDU622). Jeśli masz ten numer pod ręką, otwórz go na stronie 87 a następnie postępuj według fotografii 7, 8 i 9. Jeśli kabelki są już przygotowane, lub też były



Fotografia 8 a) i b) – brak oznaczeń polaryzacji na pasywnym buzzersie piezo dołączonym do zestawu, c) znacznik polaryzacji na płytce PCB na wypadek potrzeby użycia bardziej egzotycznego pasywnego buzzera

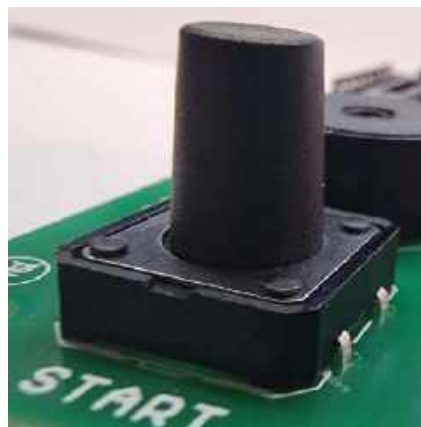


Fotografia 9. Kornel powoli kończy montaż układu. Do przylutowania zostały jeszcze dwa przełączniki. Młodzi Entuzjaści Elektroniki, Szkoła Podstawowa nr 86, Wrocław

one już przygotowane fabrycznie, należy przeprowadzić je w kierunku od strony lutowania płytki PCB do strony z komponentami, a następnie umieścić odizolowane końcówki w docelowych otworach i przylutować. Czerwony kabelk należy przylutować od spodu do pierścienia miedzi wokół otworu przy którym widnieje znak „+” a czarny do pierścienia miedzi wokół otworu ze znakiem „-”. Poprawnie zamontowane kabelki pokazuje **fotografia 11**.

Podsumowanie montażu

Po ukończeniu montażu upewnij się proszę, czy wszystkie połączenia lutowane są błyszczące i nie ma zimnych lutów, oraz, czy żadne sąsiednie pola lutownicze nie są ze sobą połączone, na przykład za pomocą dużej kulki cyny. Poprawnie zmontowany



Fotografia 10. Poprawny montaż przycisku start. Wszystkie nóżki przycisku przeszły przez otwory na PCB oraz cała podstawa przycisku leży bezpośrednio na powierzchni PCB



Fotografia 11. Poprawny sposób montażu kabelków do płytki PCB. Czerwony kabelk zamontowany został w otworze oznaczonym znakiem „+” a czarny w otworze ze znakiem „-”

układ powinien wyglądać jak na **fotografii 12a i b**. Zauważ, że w podstawie nie zamontowałeś jeszcze mikrokontrolera i nie rób tego, do poki nie zrobisz podstawowych pomiarów, które pomogą Ci przeprowadzić już za chwilę.

Uruchomienie

Po zmontowaniu każdego układu elektronicznego zawierającego układy scalone, a zwłaszcza mikrokontroler, po uprzedniej obowiązkowej weryfikacji poprawności lutowania, warto (a nawet trzeba) sprawdzić, czy na pinach zasilających podstawek znajdują się oczekiwane napięcia. Zerkając w notę katalogową (najlepiej), lub też zwyczajnie analizując układ ścieżek po stronie lutowania (przynajmniej), łatwo zauważyć, że „plus” zasilania powinien być dostarczony na pin numer 1 a „minus” zasilania na pin numer 14. Czas więc dokonać pomiaru. W tym celu upewnij się, że masz ubrane gogle ochronne, włóż do koszyczka baterie, ustaw przełącznik SW1, w taki sposób, by jego hebelki znalazł się przy najbliższej krótszej krawędzi płytki PCB (**fotografia 14**). Następnie ustaw multimetr w tryb woltomierza napięć stałych (DCV) na odpowiedni zakres, u mnie jest to zakres do 20 V. Jeśli kable są prawidłowo podłączone do miernika (**fotografia 13**) sondę czerwoną przyłóż do pinu pierwszego a sondę czarną



Fotografia 13. Poprawne podłączenie przewodów na przykładzie multimetru DT-830B ustawionego na funkcję woltomierza napięć stałych. Wolne gniazdo „10 ADC” używane jest wyłącznie do pomiarów prądów w przedziale od 200 mA do maksymalnie 10 A. Oznacza to, że widoczne na zdjęciu podłączenie kabli poprawne jest dla większości wykonywanych na co dzień pomiarów (pomiar ciągłości, omomierz, woltomierz, i amperomierz – ten ostatni wyłącznie w zakresie do 200 mA)

do czternastego (**fotografia 14**). Zmierzone napięcie powinno wynieść nieco poniżej 5 V.

Jeśli tak właśnie jest, można wyłączyć zasilanie, poprzez ustawienia hebelka przełącznika SW1 w kierunku przeciwnym, czyli do wnętrza płytki PCB. Dla pewności można też wyjąć jedną z baterii z koszyczka. Następnie w podstawie umieść mikrokontroler, pamiętając o właściwej polaryzacji (**fotografia 15**) i po upewnieniu się, że każda z nóżek mikrokontrolera trafiła do swojego gniazda w podstawie, ostrożnie dociśnij układ do podstawki, cały czas obserwując, czy podczas tego procesu, żadna z nóżek się nie łamie i nie zagina. Gdyby tak się zadziało (nóżki zaczęły się giąć), należy wyjąć układ z odstawki, podważając go za pomocą np. płaskiego śrubokręta, wyprostować nóżki i podjąć ponowną próbę montażu, cały czas pamiętając o właściwym kierunku montażu. Cały proces należy wykonywać bardzo ostrożnie, ponieważ po kilku wygięciach ograniczona wytrzymałość mechaniczna nóżek

spowoduje, że się one odłamią, i trzeba będzie próbować je jakoś dorabiać (lub zakupić i użyć nowy, odpowiednio zaprogramowany mikrokontroler).

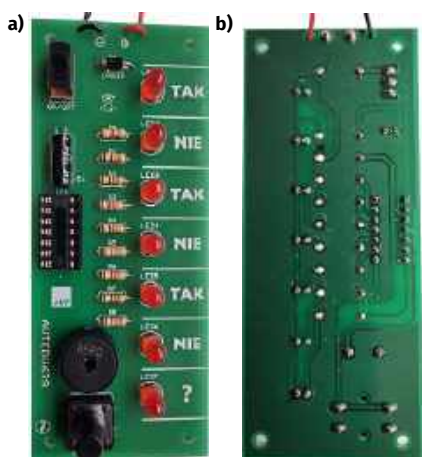
Ponownie upewniając się, że masz gogle ochronne w odpowiednim miejscu (czyli na nosie) ponownie załącz przycisk SW1 oraz naciśnij przycisk START. Powinna uruchomić się efektywna sekwencja wspomagacza wyboru wzbogacona melodyjnymi dźwiękami losowania. Och, jakby to było pięknie, gdyby wszystkie wybory naszego życia dało się powierzyć maszynie losującej. Przy istotnych decyzjach, nigdy nie korzystaj z tego urządzenia! Można by się wtedy ładnie na nim „przejechać”...

Plansze tematyczne do wspomagacza wyboru

Tym razem w zestawie, obok instrukcji montażu, znajdziesz również kolorowe plansze tematyczne na sztywnych kartonikach (papier o fakturze kartek z bloku technicznego), które będziesz mógł wyciąć i za pomocą dodatkowo wyciętych otworków nałożyć na wkręcone od spodu płytki PCB dołączone do zestawu śrubki M3. Otwory umożliwiają wkręcenie gwintu bezpośrednio do płytki, co będzie wymagało dobrego śrubokręta i nieco siły, ale mnie się udało. Po nałożeniu na powstałe w ten sposób bolce kartonika tematycznego można go dodatkowo unieruchomić za pomocą przykręconych od góry, dołączonych do zestawu nakrętek M3. **Wspomagacz wyboru** z przytwierdzonym kartonikiem pokazano na **fotografii 16**.

Baterie, a może... zasilacz?

I zaczynają się wybory, decyzje i... konsekwencje. Owszem, można zastosować zasilacz i oczywiście będzie to miało swoje



Fotografia 12. Zmontowany układ, a) widok od strony komponentów, b) widok od strony lutowania



Fotografia 14. Po ustawieniu przetłacznika SW1 w pozycji ON (hebelki w stronę najbliższej krawędzi PCB) pomiędzy pinem pierwszym (plus zasilania) oraz czternastym (minus zasilania) powinno znajdować się napięcie nieco poniżej 5 V

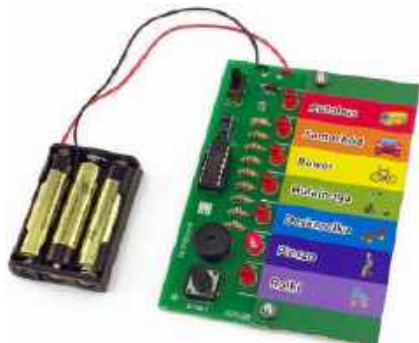


Fotografia 15. Mikrokontroler poprawnie zamontowany w podstawie. Strzałki wskazują, że wycięcia pokazujące kierunek montażu podstawki i mikrokontrolera są zgodne z wycięciem widocznym na warstwie opisowej płytki PCB

konsekwencje. Jeśli wszystko zrobimy z głową, w pełni trzymając się wszystkich zasad, jedyną konsekwencją będzie to, że **wspomagacza wyboru** uruchomimy wyłącznie tam gdzie będzie dostęp do gniazdka napięcia przemiennego sieci 230 V, ale ponieważ nie będziemy korzystali z baterii, nigdy nam się one nie rozładują i nie będzie potrzeby zakup nowych. Trzeba jednak wiedzieć, że zastosowanie zasilacza złego typu, lub też błędna polaryzacja, uszkodzą niektóre komponenty w układzie i trzeba je będzie wymienić na nowe. Jeśli uszkodzimy mikrokontroler, to nie wystarczy potem zakupić identyczny w sklepie, ponieważ te ze sklepu nie są zaprogramowane żadnym wsadem (firmwarem). Trzeba się więc będzie udać do sklepu AVT, i zakupić mikrokontroler zaprogramowany firmware'm do konkretnego zestawu. Jak widać, każdy wybór ma swoje konsekwencje. Zawsze będą jakieś zalety i wady, korzyści i przeciwwskazania.

Zasilacze napięcia stałego

Jeśli do zasilenia któregoś ze wspomnianych wyżej zestawów zdecydujesz się zastosować zasilacz, musisz pamiętać o tym, że do układu potrzebujesz dostarczyć **napięcie stałe, stabilizowane, o wartości możliwie najbliższej 5 V**. Jeśli dysponujesz wyłącznie zasilaczem napięcia stałego, ale



Fotografia 16. Wspomagacz wyboru z przetywierzdaną za pomocą śrubek i nakrętek planszą tematyczną

o wyższej wartości (np. 12 V), wystarczy zastosować dodatkowy stabilizator napięcia, np. najpopularniejszy chyba układ stabilizatora liniowego 7805, najlepiej wraz z kondensatorami filtrującymi, zgodnie z rekomendacją, którą znajdziesz w nocie katalogowej tego elementu. Pamiętaj też o tym, że istnieją zarówno zasilacze dostarczające napięcia stałego jak i zasilacze dostarczające napięcia przemiennego. O zasilaczach napięcia przemiennego dzisiaj niestety nie zdążę Ci opowiedzieć, być może jednak wrócimy do nich innym razem. Musisz jednak wiedzieć, że zasilacz na napięcie przemienną, nawet o wartości 5 V lub mniejszej natychmiast uszkodzą oba nasze układy z mikrokontrolerami! Wróćmy więc do zasilaczy na napięcie stałe. Zasilacze napięcia stałego mają na tabliczkach znamionowych oznaczenie w postaci dwóch poziomych równoległych linii, analogiczne do znaku równości, ale z przerywaną dolną linią (rysunek 4a i b).

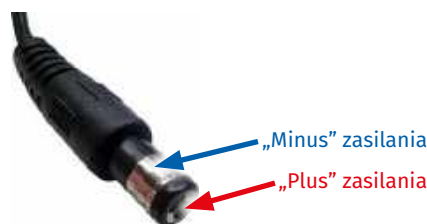
W przypadku zasilacza napięcia stałego, który dostarcza napięcia o stałej wartości, jeden ze styków we wtyczce z napięciem wyjściowym zasilacza zawsze będzie biegunem dodatnim, a drugi, zawsze będzie biegunem ujemnym. W większości zasilaczy napięcia stałego produkcji europejskiej z wtykiem typu **baryłkowego** (barrel jack w języku angielskim) „plus” zasilania dostępny jest najczęściej na wewnętrznej średnicy tego wtyku (rysunek 5), natomiast na zewnętrznej powierzchni dostępny jest „minus” zasilania.

Podłączać, czy nie podłączać?

I znów trzeba dokonać wyboru. Powiedzmy, że trzymasz już w ręku zasilacz opisany na tabliczce znamionowej jako stabilizowany o stałym napięciu wyjściowym 5 V. Na tabliczce znalazłeś też schemat polaryzacji na wtyczce. Teoretycznie wszystko już wiesz. Wiesz, jakiego napięcia dostarcza zasilacz i wiesz gdzie jest dostępny plus i gdzie minus. Jaki masz wybór? Ano możesz podłączyć zasilacz do układu już teraz, albo... mieć ograniczone zaufanie do opisów tekstowych i znakowych na zasilaczu i sprawdzić jak sytuacja wygląda w rzeczywistości. Ponieważ każdy wybór wiąże się z konsekwencjami, jeśli podłączysz zasilacz do układu już teraz (nie mierząc napięcia na zasilaczu i nie upewniając się co do występującej na nim polaryzacji) może się okazać, że albo układ zadziała poprawnie (co jest dość prawdopodobne), albo też, że układ ulegnie uszkodzeniu



Rysunek 4. Na tabliczce znamionowej a) przykładowego zasilacza, b) widoczny jest napis informujący o tym, że zasilacz ten dostarcza (output) napięcie stałe (symbol podobny do znaku równości z przerywaną dolną linią) o wartości 12 V przy dopuszczalnym natężeniu prądu obciążenia o wartości 1 A oraz c) rysunek pokazujący schemat polaryzacji na wtyczce („plus” wewnątrz i „minus” na zewnątrz). Zasilacz na napięcie 12 V bezpośrednio nie nada się do zasilania naszych układów z mikrokontrolerami



Rysunek 5. Rysunek 4c powinien być wystarczająco intuicyjny, niemniej ten rysunek pokazuje przełożenie wspomnianego schematu na rzeczywistość

(też możliwe). Natomiast, jeśli nie porzucasz na zaufaniu względem napisów na zasilaczu i za pomocą multimetru sprawdzisz wartość i polaryzację napięcia, będziesz miał pewność, układ zadziała poprawnie i nic się w nim nie uszkodzi.

Dlaczego warto mieć ograniczone zaufanie do napisów i oznakowań? Otóż wszystko to, co pokazano na fotografiach i rysunkach powyżej to informacja producenta. Nigdy nie mamy pewności, czy zasilacz nie jest przerabiany przez jakiegoś elektronika amatora albo, po prostu, czy zasilacz nie jest wadliwy. Jeśli polaryzacja na wtyczce okaże się inna niż sugeruje to schemat polaryzacji na tabliczce znamionowej, lub też dostarcza wyższego napięcia (bo na przykład uszkodził się stabilizator wewnętrzny tego zasilacza) taki zasilacz, mimo dużo obiecujących, poprawnych napisów i oznakowań uszkodzi nasz układ z mikrokontrolerem (a tego nie chcemy).

Tymczasem do posiadania stuprocentowej gwarancji odnośnie poprawności opisów wystarczy bardzo niewiele. Wystarczy za pomocą multimetru ustawionego w tryb woltomierza napięć stałych (DCV), na przykład w zakresie do 20 V, z prawidłowo



Fotografia 17. Weryfikacja polaryzacji i wartości napięcia stałego na zasilaczu z użyciem multimetru ustawionego na funkcję voltomierza napięć stałych w zakresie do 20 V. Sonda czerwona została umieszczona wewnątrz wtyku (średnica wewnętrzna) a sonda czarna dotyka średnicy zewnętrznej. Brak znaku „-” przed liczbą na wyświetlaczu potwierdza, że „plus” znajduje się na wewnętrznej średnicy wtyku zasilacza

zainstalowanymi w nim przewodami: czarny kabel do gniazda COM a czerwony do gniazda służącego do pomiaru napięć (patrz instrukcja multimetru) przyłożyć sondę czerwoną do spodziewanego „plusa” na wtyku zasilacza oraz sondę czarną do spodziewanego „minusa”. Jeśli na wyświetlaczu multimetru pojawi liczba wskazująca wartość spodziewanego napięcia (bliskie 5 V), bez znaku „-” na przedzie tej liczby, to mamy pewność, że opisy na tabliczce znamionowej zasilacza są poprawne, i możemy śmiało użyć go jako zamiennik baterii w obu wspomnianych wyżej układach z mikrokontrolerami (**fotografia 17**).

Pamiętaj, by nigdy nie podłączać do układu jednocześnie baterii i zasilacza, bo to spowoduje niekontrolowany proces ładowania baterii (zwykłych baterii nie wolno ładować!) a po pewnym czasie ich rozszczelnienie i najpewniej eksplozję, o skutkach trudnych do przewidzenia. Dlatego raz jeszcze powtórzę, nigdy nie podłączaj do układu jednocześnie baterii i zasilacza! Albo zasilacz, albo baterie. Nigdy razem!

Jak widać powyżej udało mi się w moich szpargałach znaleźć stabilizowany zasilacz napięcia stałego o wartości 5 V (**fotografia 17**). Bezpieczny wynik pomiaru napięcia na wyświetlaczu (4,9 V) potwierdza, że nie przekraczam napięcia 5 V więc mogę go użyć do zasilenia **wspomagacza wyboru**. Jednocześnie brak znaku „-” przed cyfrą mówi mi, że sondę czerwoną podłączyłem do „plusa” zasilania a czarną do „minusa”.



Fotografia 18. Gniazdo baryłkowe z przyłutowanymi kabelkami. Kabelki czerwony (plus) przyłutowany do blaszki na tyle gniazda, kabelki czarny (minus) do blaszki na spodzie gniazda. Boczna blaszka pozostaje nielutowana

Ponieważ na płytce PCB nie ma gniazda na zasilacz i do płytki można przyłutować, co najwyżej, kabelki, musiałbym uciąć wtyk zasilacza, zdjąć izolację na odcinku kilku milimetrów oraz pocynować końcówki kabelków, ponownie za pomocą multimetru odnaleźć, który kabelki to „plus”, a który „minus” i odpowiednio przyłutować je do płytki **zamiast** koszyeczka z bateriami. Jednak w ten sposób (odcinając wtyk) zniszczyłbym w pewien sposób w miarę uniwersalny zasilacz. Ponieważ nie lubię niczego psuć, wykorzystałem dodatkowe gniazdko i dodatkowe kabelki (gniazdko i kabelki znalazłem w swoich szpargałach, ale można je kupić w dowolnym sklepie elektronicznym). Blaszka na końcu gniazda odpowiada zazwyczaj wewnętrznej średnicy wtyku baryłkowego, więc tu przyłutowałem kabelki czerwony. Blaszka pod spodem gniazda odpowiada zazwyczaj średnicy zewnętrznej wtyku baryłkowego, więc tutaj przyłutowałem kabelki czarny (**fotografia 18**).

Zmierzymy zatem, czy na kabelku czerwonym jest rzeczywiście „plus” a na czarnym „minus” (**fotografia 19**).

Liczba widoczna na wyświetlaczu miernika pokazuje napięcie stałe o wartości około 4,9 V. To pokazuje, że zasilacz dostarcza napięcia właściwego do zasilania naszych układów z mikrokontrolerami. Z kolei brak znaku „-” na początku wyświetlacza potwierdza, że zmierzone napięcie to 4,9 V (a nie -4,9 V) co oznacza, że czerwony kabelki rzeczywiście dostarcza „plus”

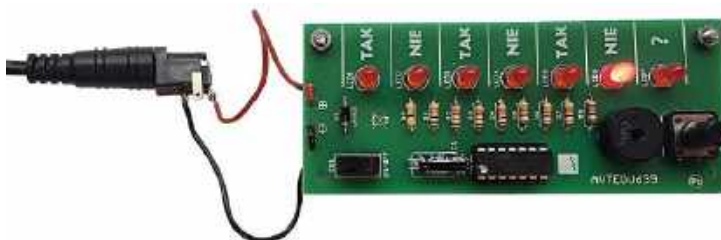


Fotografia 19. Pomiar napięcia na końcach kabelków. Ponownie zmierzono napięcie około 4,9 V a brak znaku przed liczbą świadczy o prawidłowej polaryzacji (plus na czerwonym kabelku). Zwróć też uwagę na zastosowane krokodylki nałożone na sondy przyrządu pomiarowego, które bardzo ułatwiają codzienną pracę

zasilania a czarny stanowi „minus” tego zasilania. Możemy zatem bez obaw zasilić nasz układ za pomocą wspomnianego zasilacza 5 V (**fotografia 20**).

Na koniec życzę Ci dobrych decyzji i przemyślanych wyborów na co dzień. W szkole, w domu, na podwórku, a kiedyś również „w życiu”. Niech powyższy układ na etapie montażu bawi i uczy, a już zbudowany pomaga i cieszy podczas tych najprostszych, codziennych wyborów, które zazwyczaj, na szczęście, nie niosą za sobą żadnych większych dalekoplanowych konsekwencji. Podczas tych trudniejszych wyborów, w miarę możliwości korzystaj z obecności, mądrości i doświadczenia starszego rodzeństwa, rodziców i opiekunów, przyjaciół, póki tylko masz taką możliwość. Czasem warto o coś zapytać, podzielić się jakimś tematem, zapytać, co ktoś myśli bądź sądzi na taki czy inny temat, pamiętając przy tym, że to mimo wszystko cudze myśli, nie Twoje. Te najistotniejsze wybory pozostaw samemu sobie, głęboko wsłuchując się w to, co tak naprawdę myślisz i czujesz. Nigdy się nie spiesz, choćby i cel był szlachetny i piękny. Nierzadko w piękne miejsce da się dojść wieloma ścieżkami. Gdy będziesz gotów, wybierz tę, która współgra z Tobą najbardziej. ■

Mariusz Ciszewski

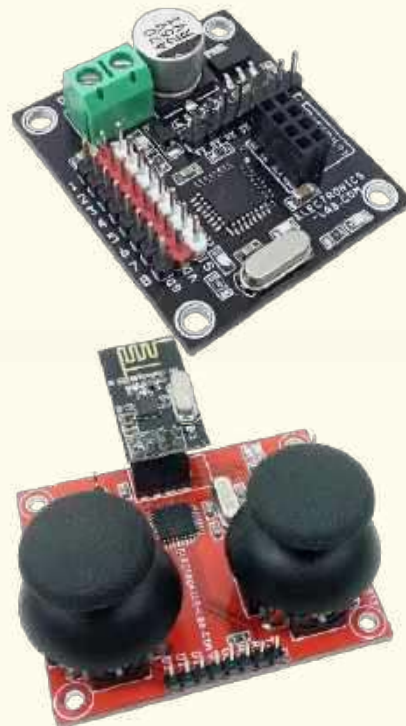


Fotografia 20. Układ pomagacza wyboru zasilony z zasilacza (stabilizowanego) na napięcie stałe o wartości 5 V. Nigdy nie podłączaj do tych układów (wspomagacz wyboru, UFOledek) zasilacza o wyższym (niż 5 V) napięciu. Takie od razu „usmaży” (czytaj: nieodwracalnie zniszczy) mikrokontroler powodując nawet jego wybuch. Pamiętaj, by podczas pracy z elektroniką nosić gogle ochronne!

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**.

8-kanalowy sterownik serwo mechanizmu RC przez łącze RF z modułem RF NRF24L01 – kompatybilny z Arduino

Jest to łatwa do zbudowania płytka kompatybilna z Arduino typu open-source, która umożliwia sterowanie 8 serwo mechanizmami RC za pośrednictwem łącza RF NRF24L01. Projekt może być używany jako samodzielny sterownik serwo mechanizmów RC lub 8-kanalowy zdalnie sterowany odbiornik serwo mechanizmów RC. Opcjonalny wyświetlacz OLED może być wykorzystany do opracowania monitora sygnału RC. Niewielki moduł zawiera mikrokontroler ATmega328, złącza dla 8 interfejsów serwo mechanizmów, złącze zasilania DC, kondensator elektrolityczny C5 na zasilaniu DC, aby zapewnić płynny ruch serwo mechanizmów RC bez zniekształceń. Zasilanie 5 V DC



Nadajnik zdalnego sterowania RF z dwoma joystickami i modułem RF NRF24L01 – sterowanie 2 joystickami

Jest to kompatybilny z Arduino sprzęt o otwartym kodzie źródłowym, który zawiera 2× joysticki, moduł RF NRF24L01, mikrokontroler ATmega328, regulator 3,3 V, diodę LED zasilania, diodę LED funkcji, złącze programowania Arduino i inne wymagane komponenty. Płytkę można używać do tworzenia różnych aplikacji, takich jak gry, zdalny sterownik serwo RC, robotyka i wiele innych. Złącze CN3 służy do programowania mikrokontrolera ATMEGA328 za pomocą Arduino IDE.

Niektóre projekty aktualnie dostępne tylko dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl:

1. Sterowanie prędkością, kierunkiem i zatrzymaniem silnika DC z modułem RF NRF24L01
2. Nadajnik zdalnego sterowania z pojedynczym joystickiem wykorzystujący NRF24L01
3. 8-kanalowy zdalny nadajnik RF z protokołami: Holtek i szeregowym
4. 8-kanalowy zdalny odbiornik RF z protokołami: Holtek i szeregowym
5. Pojemnościowy czujnik wilgotności do konwertera wyjścia analogowego
6. Mostek H dla wysokiej mocy szczotkowego silnika prądu stałego z czujnikiem prądu
7. Przetwornica DC-DC buck 12...75 V na 10 V na wyjściu
8. Czujnik prądu low-side 10 µA...10 mA
9. Kontroler ramienia robota z bezprzewodowym pilotem PS3
10. Termiczny czujnik masowego przepływu powietrza – anemometr stałotemperaturowy
11. Precyzyjny wzmacniacz transimpedancjowy z przełączanym integratorem
12. Kontroler pełnego mostka z przesunięciem fazowym i prostowaniem synchronicznym wykorzystujący UCC28950
13. Wysokowydajny monofoniczny wzmacniacz audio klasy D o mocy 20 W
14. Monitorowanie poziomu cieczy za pomocą czujnika ciśnienia – wyświetlacz stupkowy
15. Sterowanie silnikiem DC za pomocą joysticka
16. 16-kanalowy sterownik serwo mechanizmów RC z interfejsem I²C
17. Programowalny kondycjoner sygnału z czujnika rezystancyjnego mostkowego
18. Choinka z Arduino i pikselowymi diodami
19. 20-segmentowy wyświetlacz stupkowy w rozmiarze jumbo
20. Stacja pogodowa Lilygo ttgo t5-4.7 z wyświetlaczem typu e-papier
21. Półprzewodnikowy przełącznik mocy DC z prądowym sprzężeniem zwrotnym
22. Wyłącznik nadprądowy – przełącznik wyłączający nadprądowy
23. TinyML – Rozpoznawanie ruchu przy pomocy RPI Pico
24. Uniwersalny konwerter napięcia AC – wyjście 18 V DC z wejścia 85...265 V AC
25. Moduł procesora echa głosu – urządzenie opóźniające do efektów dźwiękowych, echo, reverb
26. Sterownik silnika krokowego z joystickiem
27. RPI – stacja pogodowa IoT
28. Niskobudżetowy monitor jakości powietrza IoT oparty o Raspberry Pi 4
29. Automatyczny system ogrodniczy z NodeMCU i Blynk, ArduFarmBot 2
30. Wzmacniacz piezoelektryczny do gitary i skrzypiec
31. Wysokowydajny i niezawodny sterownik bipolarnego silnika krokowego
32. Sonarowy theremin MIDI
33. Sterownik silnika prądu stałego z wykorzystaniem przełącznika i mosfetu – interfejs Arduino
34. Super prosty czuły wykrywacz metali
35. Najlepszy sposób na próbkowanie dźwięku za pomocą ESP32

Miesięcznik „Elektronika dla Wszystkich” (12 numerów w roku) jest wydawany we współpracy z kilkoma redakcjami zagranicznymi



Wydawnictwo:
AVTKorporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Redaktor naczelny:
Mariusz Ciszewski
mariusz.ciszewski@elportal.pl

Adres redakcji:
03-197 Warszawa, ul. Leszczyńska 11
e-mail: edw@elportal.pl, www.elportal.pl

Dział reklamy:
Katarzyna Gugala
katarzyna.gugala@elportal.pl, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobański
jakub.sobanski@elportal.pl

Sekretarz redakcji:
Dariusz Welik
dariusz.welik@elportal.pl

Copyright AVTKorporacja 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)
www.ulubionykiosk.pl

W RUCH S.A., e-mail: prenumerata@ruch.com.pl
tel. 801 800 803, 22 717 59 59, www.prenumerata.ruch.com.pl

Elektor Bestsellers

SAVE UP TO
26% NOW!



www.elektor.com/sale/deals

