



ELEKTRONIKA

dla wszystkich

nr 10/2024 (345) • październik • www.elportal.pl

DIY PLUS
tylko dla prenumeratorów

Klasyczne metronomy LED

PROJEKTY dla elektroników

- ▶ Wzmacniacz mocy 500 W, część 2
- ▶ SMD Tweezers – ulepszona pęseta pomiarowa SMD
- ▶ Prosty limiter do subwoofera

DIY dla wszystkich

- ▶ Przystawka do quiz-ów dla ośmiu graczy
- ▶ Uniwersalny sposób zabezpieczenia bagażu przed kradzieżą

TUTORIALE

- ▶ Audio OUT: Wokoder analogowy, część 1
- ▶ Historia i technologia wyświetlaczy wideo, część 2
- ▶ FnrSi SWM-10 – inteligentna zgrzewarka punktowa
- ▶ Ekscytacje Maxa: Migające diody LED i śliniacy się inżynierowie (13)
- ▶ Know-how: Obwody tłumiące zakłócenia



Pomocna dłoń



automatykaB2B.pl

EP.com.pl

Największy
portal dla
elektroników
konstruktorów

eprasa.pl 4e05f27edc

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



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



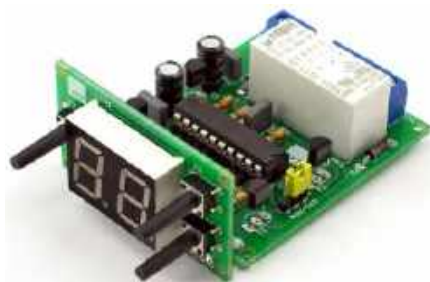
AVT5540 Radio FM z RDS
<https://sklep.avt.pl/avt5540.html>



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



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



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



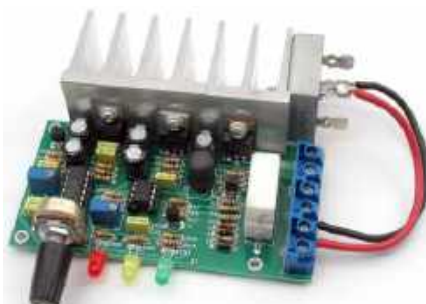
AVT990 Automatyczny włącznik światła
<https://sklep.avt.pl/avt990.html>



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



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



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



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



Pełna oferta na: sklep.avt.pl

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

–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 równolegle 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:

Klasyczne metronomy LED	8
Wzmacniacz mocy 500 W, część 2	21
SMD Tweezers – ulepszona pęseta pomiarowa SMD	30
Prosty limiter do subwoofera	38

Tutoriale:

Audio OUT: Wokoder analogowy, część 1	40
Historia i technologia wyświetlaczy wideo, część 2	44
Fnirsi SWM-10 – inteligentna zgrzewarka punktowa	54
Ekscytacje Maxa:	
• Migające diody LED i śliniacy się inżynierowie (13)	56
• Sprytnie porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania ..	60
Obwody tłumiące zakłócenia	65
Edukacja w EdW dla szkół i uczelni:	
Wykład 23 – filtry dolnoprzepustowe	72



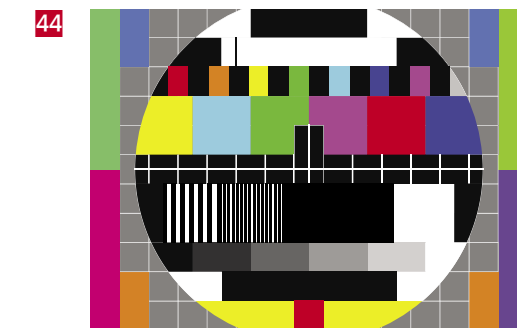
21



30

DIY dla wszystkich:

Przystawka do quiz-ów dla ośmiu graczy	77
Uniwersalny sposób zabezpieczenia bagażu przed kradzieżą	80



44

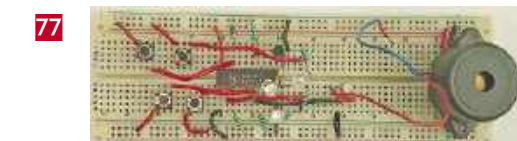
Elektronika dla Wszystkich – Junior:

Gutek. Eko-Czarodziej w krainie (zużytej) elektroniki	79
Czwarte spotkanie z najmłodszymi pasjonatami elektroniki	83

Na zdjęciu na okładce Tymek, koło zainteresowań „Młodych Entuzjastów Elektroniki”, Wrocław

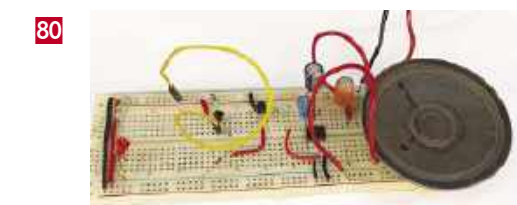
DIY PLUS

tylko dla prenumeratorów zamawiających prenumeratę na www.UlubionyKiosk.pl



77

2-kanałowy moduł przekaźnikowy Wi-Fi wykorzystujący ESP8266 NodeMCU	91
8-kanałowy pilot zdalnego sterowania ESP32	91



80

Rubryki stałe:

Prenumerata	3
Od redakcji	5
Pocztą	6

A za miesiąc w listopadowym EdW



✱ Inteligentny podwójny zasilacz hybrydowy, część 1

Zaawansowany zasilacz warsztatowy: dwa niezależne wyjścia 0–25 V DC, wydajność 5 A na kanał. Możliwość łączenia wyjść w szereg. Prezentacja wyników na wyświetlaczu graficznym. Kontrola za pomocą dwóch obrotowych selektorów i dwóch przycisków.

✱ 500 W wzmacniacz audio, część 3

Ostatnia część opisu budowy wzmacniacza audio o mocy 500 W RMS. W tej części omówimy zasilacz, procedury uruchomienia i testowania wzmacniacza a także podamy szczegółowe informacje na temat zabudowania wzmacniacza w wentylowanej, aluminiowej obudowie.

✱ Sterownik wentylatora chłodzącego i zabezpieczenie głośnika

Sterownik trzech wentylatorów, który załącza je przy zadanej temperaturze a następnie kontroluje ich prędkość obrotową. Pozwala efektywnie i w miarę bezgłośnie wentylować nagrzewające się urządzenie. Może również chronić głośniki przed uszkodzeniem, a także zapobiegać trzaskom przy włączaniu i wyłączaniu zasilania.

✱ Czujniki jakości powietrza

Na rynku pojawiło się wiele przystępnych cenowo mierników jakości powietrza. Zaprezentujemy pobieżnie ich przegląd oraz omówimy ich parametry i zasadę działania.

✱ Zwyczajowa garść ciekawych projektów DIY

✱ Kolejne artykuły w ramach cyklicznych Tutoriali.

✱ Elektroników Juniorów zapraszamy na następną ciekawą przygodę, podczas której znów postaramy się łączyć przystępnie podaną teorię z jeszcze bardziej atrakcyjną praktyką.

**W kioskach
od 30 października**

Kasztany? A gdzie tam, październik w pełnym rozkwicie!

Wrzesień przyniósł nam nieoczekiwane lato, a jaki będzie październik? Trudno cokolwiek przewidzieć, gdyż pogoda zdaje się mieć ostatnimi czasy własne zdanie na temat każdej z nadchodzących pór roku! Nawet spadających kasztanów ciężko dzisiaj być pewnym. Szczęśliwie, gdyby tych miało zabraknąć, kolega Gutek, którego niebawem poznacie, przyniósł nam i trzyma na tę okoliczność prawdziwego asa w rękawie!

Tymczasem w Redakcji, nie zważając na pogodę i urodzaj, ani przez moment nie zwolniliśmy naszego tempa. Przygotowaliśmy dla Was projekty i materiały, którymi mamy ambicję zachwycić wszystkich Czytelników (i każdego z osobna!). Mając na względzie szerokie spektrum możliwych zainteresowań, przygotowaliśmy treści skrótowo opisane poniżej.

Jesień to pora nieco melancholijna. A na melancholię dobry być może kubek dobrej herbaty z imbirem albo gorąca kąpiel. Można też dla Was aż dwa projekty metronomów, konstrukcyjnie stanowiących piękny przykład harmonii nowoczesności z tradycją, wprost idealnych na taką okazję. W numerze kontynuujemy też montaż potężnego wzmacniacza mocy audio, na wypadek, gdybyśmy chcieli nasze domowe kółko wokalne przekształcić w nieco głośniejszą kapelę. By zespół mógł czuć się bezpieczny o własne sprzęty dodatkowo polecamy prosty limiter poziomu sygnału audio dla subwoofera.

Dla naszych kochanych dusz bratnich – introvertów, którzy jesienne wieczory postanowili spędzić przy montażu bądź serwisowaniu nowoczesnej elektroniki mamy projekt niezwykle poręcznej pęsety pomiarowej idealnej do sprawdzania komponentów SMD. Ekstrawertykom polecamy kolejną część cyklu o diodach LED, kibicując tym samym autorowi, który po raz kolejny postanowił porwać Czytelnika w świat własnych fantazji w pogoni za wiedzą i umiejętnościami przemieszanymi z entuzjazmem i dobrym humorem.

Dla Czytelników, którzy cenią sobie spokojny relaks i nie wyobrażają sobie jesiennych wieczorów bez wygodnego fotela oraz ulubionych programów TV, przygotowaliśmy drugą część opowieści o historycznych dziejach ich ukochanego wynalazku.

Październik jest miesiącem uważanym za czas wzmożonych przeziębienia. Zmieniające się warunki pogodowe – spadek temperatur, większa wilgotność oraz częste deszcze – sprzyjają osłabieniu odporności naszego organizmu. Jednak brak odporności nie jest wyłącznie domeną ludzką. O odporność na przeciążenia należy zadbać również projektując układy elektroniczne. Szczególnej uwagi będą wymagały urządzenia zasilane wprost z sieci elektrycznej, jak również obwody pomiarowe z różnej maści czujnikami, gdzie wejścia mogą być narażone na działanie różnorodnych czynników zewnętrznych. Zwiększona ochrona jest szczególnie ważna w środowiskach narażonych na wahaniami napięcia zasilającego oraz różnego typu zakłócenia w tym zjawiska indukcyjne oraz interferencje elektromagnetyczne. W bieżącym numerze zamieściliśmy tekst na temat obwodów niwelujących zakłócenia EMC, a w przyszłym miesiącu opublikujemy dalszy ciąg tej opowieści, szerzej opisując również elementy ratujące elektronikę w przypadku wystąpienia poważniejszych przeciążeń elektrycznych.

Tymczasem najmłodszy entuzjasta elektroniki w ramach cyklu EdW Junior, tym razem będą mogli nie tylko uczestniczyć w prowadzonych specjalnie dla nich zajęciach teoretyczno-praktycznych (zbudować kolejną zabawkę, która przybyła na naszą planetę z dalekich przestworzy), ale też poznać kolegę Gutka, który zabierze ich do swojej czarodziejskiej krainy i pokaże, jak w kreatywny sposób dać drugie życie zużytemu lub uszkodzonym komponentom elektronicznym.

Przyjemnej lektury!

Mariusz Ciszewski

Nowe życie elektronicznego złomu

Nowa rubryka EdW Junior zaowocowała potokiem listów do naszej Redakcji od młodzieży, nauczycieli i instruktorów. Na wezbranej fali tej poczty redakcyjnej naszą uwagę zwrócił wątek nie całkiem elektroniczny, lecz raczej artystyczny. Wybraliśmy tę historię do zaprezentowania jako unikatową, zaskakującą a nawet nieco relaksującą ciekawostkę. W dodatku świetnie nadająca się do numeru z kasztanem w tle.

W sekcji EdW Junior niniejszego numeru zamieściliśmy tekst pod tytułem: „Gutek. Eko-Czarodziej w krainie (zużytej) elektroniki”. Do jego powstania przyczyniła się wcześniejsza korespondencja, pomiędzy kilkorgiem osób. Ponieważ była to całkiem konstruktywna komunikacja, postanowiliśmy podzielić się i zaciekać fragmentami tej korespondencji naszych Czytelników.

Ta proekologiczna historia miała swój początek w sklepie stacjonarnym AVT SPV, do którego zgłosił się nauczyciel techniki ze Szkoły Podstawowej 257 im. prof. Mariana Falskiego w Warszawie z pytaniem o złom elektroniczny na potrzeby kreatywnych zajęć z młodzieżą. Postanowiliśmy chwycić tę nić i dotrzeć do kłębka.

Jakub Sobański, pracownik AVT SPV: (...) Jakiś czas temu przez sklep stacjonarny zgłosił się do mnie Nauczyciel Techniki z SP257 z pytaniem o złom elektroniczny. Przygotowałem dla niego trochę „materiału dydaktycznego”, który ma już stan złomowy. Wiedząc, że uczeń stworzył rzeźbę „perkustysty”, ale i wiele innych podobnych dzieł uznałem, że dobrze byłoby wesprzeć zapał młodego człowieka i pomóc w popularyzacji elektroniki wśród uczniów szkół podstawowych. Warto, aby wraz z nauczycielem opisał swoje dzieła i osiągnięcia w branżowym czasopiśmie. (...)

Tadeusz Lichuta, nauczyciel SP257: (...) Dwa tygodnie temu z firmy AVT-SPV brałem złom elektroniczny na zajęcia do Szkoły Podstawowej 257, aby uczniowie polutowali z niego różne postacie. Dostałem wizytówkę z adresem do kontaktu i z sugestią, aby powstał artykuł do gazety branżowej o wykorzystaniu zużytych części do tworzenia form plastycznych i gadżetów o charakterze ozdobnym. Autorem wspomnianej pracy jest uczeń klasy szóstej, którego umiejętności zostały wykorzystane podczas Pikniku Szkolnego gdzie prowadził warsztaty wprowadzające do umiejętności lutowania.

W szkole działa Koło Techniczne „Inżynierium” w oparciu o stwarzane dwa lata temu pracownię w ramach programu Laboratoria Przyszłości. W to miejsce Gutek również zaglądał, ale profil prac obejmuje raczej konstrukcje z kartonu i drewna więc robił to rzadko.

Można się zastanowić czy nie dało by rady zorganizować aktywności o profilu elektronicznym w formie systematycznych zajęć, realizacji projektu (cykl kilku zajęć) lub jednorazowych (powtarzalnych) warsztatów – dzieci z rodzicami o charakterze np. szkolnym i międzyszkolnym. Najchętniej w oparciu o bazę wiedzy, kontakty i zasoby EdW i AVT.

Mogę tu służyć pomocą związaną z organizacją. Pod koniec roku szkolnego proponowałem Gutkowi wspólny projekt modelu samochodu w ramach którego można by połączyć świat drewna i kartonu z elektronicznym sterowaniem. Gutek proponował, że mógłby nauczyć mnie jak wykorzystać i zaprogramować taki sterownik. No pewnie!

Od września zaczynam pracę w nowo organizowanej szkole gdzie podobny temat można rozwinąć, a od czternastu lat pracuję w Pałacu Młodzieży w Warszawie gdzie starsze dzieciaki stworzyły pod opieką p. Andrzeja „Samoróba” frezarki CNC, a jeszcze inny instruktor prowadzi zajęcia projektowe i wytwórcze na wypalanie laserowej – odpowiadam tematycznie do artykułów. (...)

Wierzmy, że to dopiero początek przygód chłopca ale też załączek współpracy z Panem Tadeuszem, na którą chyba i my mamy parę pomysłów.

Tymczasem zapraszamy Czytelników do zapoznania się z tekstem, który zaistniał na skutek korespondencji, która się wydarzyła. Gorąco zapraszamy również do dzielenia się z Redakcją i za jej pośrednictwem z pozostałymi Czytelnikami informacjami na temat wspólnych inicjatyw, które dzieją się w Waszych szkołach.

Problemy z działaniem wideodomofonu IP z PoE po kablu domofonowym

Mam pytanie dotyczące nietypowego sposobu podłączenia urządzeń wykorzystujących PoE. Chodzi o wideodomofon IP, podłączony w miejsce analogowego. Niestety, w ścianie nie było skrętki, tylko sześcioprzewodowy przewód domofonowy (6×0,5 mm²). Podłączyłem panel

zewnątrzny i panel wewnętrzny, zarabiając wtyki RJ45 z uwzględnieniem poprawnej kolejności żył. Z powodu braku dwóch żył (siódmej i ósmej) w przewodzie w ścianie, podałem zasilanie PoE tylko jedną parą przewodów. Odległość między panelami a switchem wynosi 4 metry. Urządzenia działają prawidłowo, jednak po pewnym czasie (dzień lub dwa) panel wewnętrzny traci kontakt z panelem zewnętrznym. Wówczas wystarczy zrestartować switch PoE znajdujący się pośrodku tej instalacji. Proszę o odpowiedź, czy niestabilność całego systemu może być związana z brakiem żył w obwodzie PoE, brakiem skrętki, czy może ze spadkiem napięcia wynikającym z za małego przekroju przewodów? Co Państwa zdaniem jest największym błędem w tej instalacji? Oczywiście mam świadomość, że jest to rozwiązanie niestandardowe i niedozwolone, jednak chodzi mi raczej o teorię, która to potwierdzi i wskaże najlepszy punkt instalacji.

Konrad

Red. Na pytanie o możliwe przyczyny niestabilności jednego z urządzeń, w zasadzie już Pan sobie odpowiedział. Trzeba uczciwie powiedzieć, że nie ma tu jednoznacznej odpowiedzi, a dodatkowo przyczyny mogą się kumulować.

Brak przepływu (skręcenia każdej z czterech par) niweluje właściwości ekranowania, co przy intensywnej transmisji danych (przesył obrazu) może destabilizować płynność, generować artefakty obrazu bądź całkowicie uniemożliwić transmisję. W dłuższej perspektywie (wspomniane kilka dni/godzin) może to całkowicie unieszkodliwić urządzenie, w zależności od tego, jak bardzo jest ono podatne na zakłócenia oraz w jaki sposób radzi sobie z nimi podczas transmisji danych. Zastosowane (lub nie) przez producenta rozwiązania typu watchdog mogą również wpływać na stabilność urządzenia (ciężko cokolwiek przewidzieć, każdy model, a czasem każda sztuka w ramach modelu, może zachowywać się nieco inaczej).

Spadki napięć na kablach Ethernet, nawet na pojedynczych metrach, potrafią być znaczące (i to na dwóch parach żył!), zwłaszcza że wideodomofon może mieć stosunkowo duże zapotrzebowanie na moc w chwilach aktywnej transmisji wideo (o ewentualnym sterowaniu ryglem nie wspominając). Na kilku czy kilkunastu metrach zjawisko to będzie mocno odczuwalne. Warto zmierzyć, jakie napięcie PoE jest na początku toru (przy switchu PoE), a jakie dociera do domofonu w trakcie aktywnej transmisji wideo. W rozwiązaniach PoE stosuje się nieco wyższe napięcia, aby spadki o kilka czy kilkanaście woltów mimo wszystko pozwalały przetwornicy po stronie odbiornika energii poprawnie pracować. Najczęściej stosowanym napięciem jest 48 V, ale można spotkać też urządzenia dostarczające i/lub wymagające niższych napięć, na przykład 24 V. Warto upewnić się, jakie jest największe dopuszczalne napięcie PoE dla tego domofonu i zastosować switch PoE, który wystawi największe napięcie akceptowane przez domofon. Można też zrezygnować z zasilania przez switch PoE, a zamiast tego zastosować injector PoE (w mieszkaniu), podłączyć go do zasilacza o nieco wyższym napięciu, a po stronie domofonu zastosować zewnętrzną przetwornicę, która obniży napięcie do poziomu wymaganego przez domofon. Z podstawowego wzoru $P=U \times I$ wynika prosta zależność, że stosując wyższe napięcie, można przesłać do urządzenia na drugim końcu większą moc przy takim samym bądź mniejszym prądzie płynącym przez kabel, co w przypadku pojedynczych żył 0,5 mm² może mieć olbrzymie znaczenie. Przy stosowaniu tego typu okablowania (kable domofonowe, skrętka Ethernet) należy stosować napięcia nie wyższe niż tzw. „bezpieczne”, czyli zgodnie z normami takimi jak PN-EN 61140:2016-07 (Ochrona przed porażeniem elektrycznym – Zasady ogólne) w odniesieniu do prądu stałego (DC) napięcia nie wyższe niż 60 V (napięcie to jest uznawane za bezpieczne pod warunkiem odpowiedniego ograniczenia prądu).

Oczywiście to, co powyżej, to dywagacje teoretyczne i podpowiadanych rozwiązań nie polecam wdrażać, zarówno w kontekście bezpieczeństwa (również sprzętu), jak i stabilności/niezawodności rozwiązania. Eksperymentować można dla siebie, kierując się ostrożnością i rozsądkiem, w celach stricte edukacyjnych. Najprostszym i najpewniejszym rozwiązaniem (a już na pewno godnym polecenia) będzie zastosowanie/pociągnięcie odpowiedniego okablowania, zgodnego ze specyfikacją urządzeń, co z perspektywy czasu i uzyskania najlepszych wyników może się też okazać rozwiązaniem najbardziej efektywnym z ekonomicznego punktu widzenia.

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

Dwa opisane w artykule projekty metronomów symulują klasyczny metronom mechaniczny z odwróconym wahadłem, którego wskaźnik w kształcie paterki kołysze się w lewo i w prawo, wytwarzając charakterystyczne kliknięcia na każdym krańcu. W obu projektach zostały zastosowane tylko elementy dyskretnie i proste układy logiczne, dzięki czemu są łatwe do zrozumienia i zbudowania. Ponadto oba metronomy są świetnymi projektami dla początkujących.

Klasyczne metronomy LED

Typowe „nowoczesne” metronomy elektroniczne klikają i/lub migają tylko raz na takt. Dla użytkowników metronomów tradycyjnych nie są więc zbyt atrakcyjnym rozwiązaniem konstrukcyjnym. Opisane w artykule projekty zostały opracowane, aby lepiej symulować metronomy mechaniczne, które są znane i lubiane. Oba projekty zapalają serię diod LED, którym towarzyszą dźwięki uderzeń wydobywające się z głośnika.

W pierwszym projekcie zastosowano osiem diod LED, a całe urządzenie mieści się w standardowej plastikowej obudowie. Drugi projekt natomiast jest nieco bardziej skomplikowany, ma 10 diod LED, a ponadto urządzenie zostało zamknięte w niestandardowej obudowie wykonanej z drewna. Można go więc polecić czytelnikom, którzy mają pewne doświadczenie w obróbce drewna.

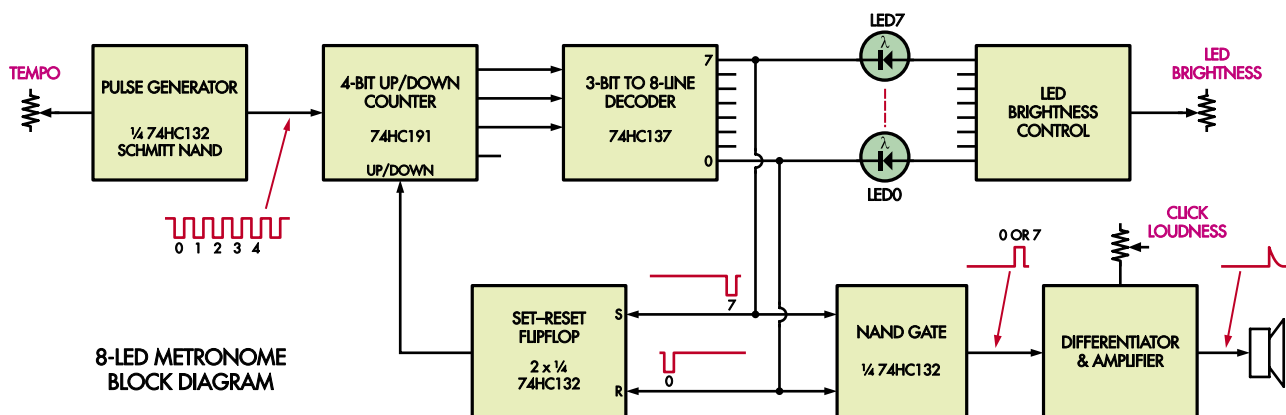
W obu przypadkach diody LED są ułożone w łuk i zapalają się sekwencyjnie, do przodu i do tyłu, naśladując ruch odwróconego wahadła. Kliknięcie na każdym końcu łuku dodatkowo symuluje mechaniczny metronom. Typowy zakres tempa metronomu wynosi od 40 do 208 uderzeń na minutę, co daje stosunek 5,2 do jednego. W opisywanych metronomach zakres ten jest rozszerzony do 36...216 uderzeń na minutę, co daje stosunek sześć do jednego.



Oba projekty są doskonałe dla początkujących. Nie występują w nich sygnały o wysokiej częstotliwości, wysokie napięcia i skomplikowane okablowanie. Nie ma też potrzeby programowania jakiegokolwiek układu. Biorąc pod uwagę oczekiwane tolerancje elementów, konieczne będą jednak pewne pomiary i regulacje w celu skalibrowania przyrządów po zakończeniu budowy.

Dwa projekty

W nieco prostszym metronomie z 8 diodami LED zastosowano układy logiczne serii 74HC. Może być on zasilany bateryjnie. W metronomie z 10 diodami LED zastosowano układy CD4000. Jest on przeznaczony do zasilania z zasilacza wtyczkowego. Oba układy działają podobnie. Generator



Rysunek 1. 8-diodowy metronom jest oparty na trzech cyfrowych układach logicznych serii 74HC. Układ 74HC132 generuje impulsy z możliwą do ustawienia częstotliwością. Impulsy te taktują licznik 74HC191, a jego trzybitowe wyjście steruje ośmioma diodami LED za pośrednictwem układu dekodera 74HC137. Pozostałe trzy bramki w układzie poczwórnej bramki NAND 74HC132 są używane do utworzenia przerzutnika RS używanego do odwracania kierunku wirtualnego LED-owego wahadła metronomu. Dzieje się to za każdym razem, gdy zostanie osiągnięty dowolny kraniec. Układ wytwarza ponadto impulsy sterujące pracą głośnika



impulsów taktuje scalony licznik rewersyjny inkrementując go lub dekrementując z szybkością wymaganą dla żadanego tempa. Inny układ scalony dekoduje wartość licznika, aby sekwencyjnie zapalać diody LED.

Gdy którakolwiek z końcowych diod LED jest zapalona, przerzutnik RS (set/reset, SR-FF: Set-Reset Flip-Flop) przełącza kierunek zliczania, co przy użyciu diod LED generuje animację imitującą ruch wahadła w jednym bądź w drugim kierunku. Kliknięcie jest generowane przez logiczne sumowanie sygnałów z końcowych diod LED. Tak uzyskany impuls jest podawany do układu różniczkującego odpowiednio go skracającego,

a następnie kierowany jest do wzmacniacza jednotranzystorowego obciążonego małym głośnikiem.

Schematy blokowe obu metronomów pokazano na **rysunkach 1 i 2**. W projekcie z 8 diodami LED impuls tempa jest generowany przez bramkę Schmitta typu NAND (część układu 74HC132). Impuls ten taktuje czterobitowy licznik rewersyjny IC2 (74HC191). Do sterowania dekoderek 3...8 linii IC3 (74HC137) zostały użyte tylko trzy z czterech wyjść binarnych. Dekoder ten zapala diody LED w sekwencji (od osiem do limitu dekodera 137). Przerzutnik RS jest wykonany z dwóch kolejnych bramek NAND układu 74HC132.

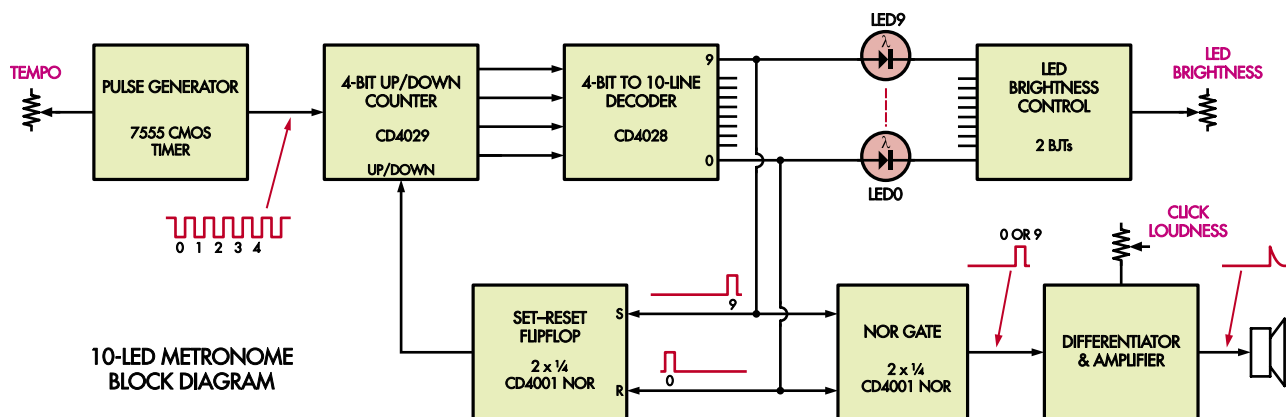
W projekcie z 10 diodami LED impuls jest generowany przez CMOS-ową wersję popularnego timera 555. Zegar taktuje czterobitowy licznik rewersyjny IC3 (CD4029), który steruje dekoderek IC4 (CD4028) z 10 wyjściami. Przerzutnik RS jest zbudowany również na dwóch bramkach układu IC1, chociaż jest to inny układ logiczny.

Obie wersje umożliwiają jaśniejsze miganie diod LED na każdym końcu łuku. Układ ten może być stosowany również w innych aplikacjach, w których występują kolejno rozświecane diody LED.

Opcje LED

W prezentowanych metronomach istnieje kilka opcji dla diod LED. W wykazie elementów podane są sugerowane typy diod LED, ale można je zastąpić innymi, zarówno pod względem rozmiarów, kształtów, jak i kolorów. Dwie końcowe diody LED mogą nawet różnić się od środkowych. Wszystkie diody powinny mieć jednak wysoką intensywność świecenia, najlepiej co najmniej 4000 mCd (diody takie są często określane jako „superjasne”). Ma to na celu zmniejszenie zużycia energii. W metronomie z 8 diodami LED uzyskujemy w ten sposób zwiększenie żywotności baterii, a w przypadku konstrukcji z 10 diodami LED ograniczamy obciążenie układu scalonego CD4028 do bezpiecznego poziomu.

Oba metronomy zostały wykonane przy użyciu owalnych diod LED o średnicy 5 mm: zielonych w wersji 8-diodowej i czerwonych w wersji 10-diodowej. Diody owalne zostały użyte, ponieważ świecą one w linii, a nie w punktach, zapewniając ciekawszy wygląd wyświetlacza. Oczywiście można zastosować również okrągłe diody LED o typowych średnicach 3 mm lub 5 mm. Najlepiej wyglądają soczewki matowe, rozpraszające światło. Można też wybrać inne diody LED



Rysunek 2. 10-diodowy metronom jest zbudowany w oparciu o układ scalony timera 555 zamiast oscylatora opartego na bramce logicznej. Pozostałe układy logiczne pochodzą z serii 4000: 4029 działa jako licznik rewersyjny (w górę/w dół), a 4028 jest dekoderek 4...10, który steruje diodami LED. Dwie z bramek układu scalonego 4001 (poczwórny NOR) tworzą przerzutnik RS, a pozostałe dwie bramki generują impuls kliknięcia

Impuls z wyjścia IC1c jest zbyt długi i spowodowałby kliknięcie na końcu impulsu, jak również na jego początku, a prąd byłby duży przez cały czas włączenia. Aby tego uniknąć,

impuls ten przechodzi przez kondensatory C4 i/lub C5. Dioda D1 bocznikuje ujemny impuls powodując, że tylko dodatni impuls jest podawany do bazy tranzystora Q1.

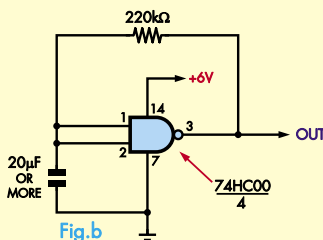
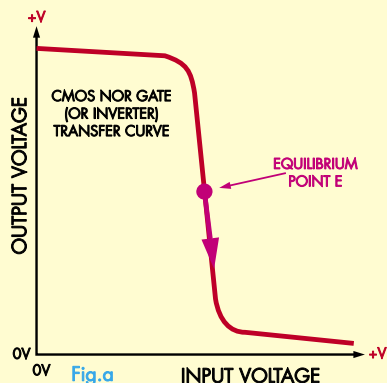
Zasilanie

Układ jest zasilany z czterech ogniw AAA. Zamiast układów serii 74HCT lub 74LS zastosowane są układy serii 74HC. Pozwala

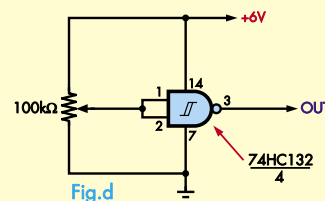
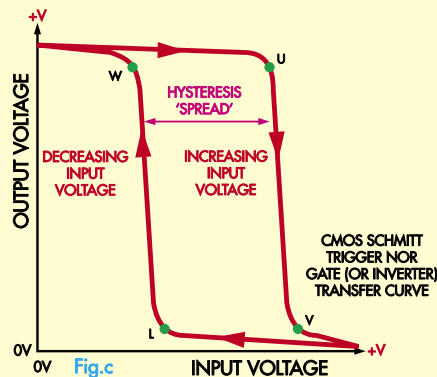
Zastosowanie bramki Schmitta jako generatora impulsów

Bramka Schmitta IC1d układu 74HC132 generuje impulsy taktujące licznik (IC2). Czym więc jest bramka Schmitta i dlaczego jej używamy?

„Zwykła” bramka lub inwerter jest w rzeczywistości wzmacniaczem o dużym wzmocnieniu, w dużym zakresie liniowym. W rezultacie przejście sygnału wyjściowego z wysokiego do niskiego lub z niskiego do wysokiego odbywa się w wąskim zakresie napięcia



Rysunek a i b. Charakterystyka przejściowa standardowej bramki NOR. Wyjście jest w stanie niskim, gdy wejście jest w stanie wysokim i odwrotnie, ale jeśli napięcie wejściowe jest pośrednie, napięcie wyjściowe może znajdować się w dowolnym punkcie pomiędzy



Rysunek c i d. Inwerter z wejściem Schmitta ma histerezę, więc gdy jego napięcie wejściowe jest wystarczająco wysokie, wyjście przełącza się w stan niski i pozostaje w tym stanie, dopóki napięcie wejściowe znacznie nie spadnie. Podobnie, gdy napięcie wejściowe spada, a wyjście przechodzi w stan wysoki, pozostaje ono w tym stanie, dopóki napięcie wejściowe znacznie nie wzrośnie

wejściowego, jak pokazano na **rysunku a**, który przedstawia wykres napięcia wyjściowego w funkcji napięcia wejściowego.

W rezultacie obwód RC z ujemnym sprzężeniem zwrotnym pokazany na **rysunku b** zazwyczaj osiąga równowagę w pewnym punkcie (E), a wyjście „utknie” w napięciu pośrednim. Jeśli napięcie wejściowe wzrośnie, jak wskazuje strzałka, wyjście zareaguje spadkiem i przywróci obwód do punktu E ze stałą czasową określoną przez R i C. Odwrotna sytuacja ma miejsce, jeśli napięcie wejściowe spadnie.

Można to przetestować samodzielnie na płytce prototypowej, jeśli mamy w swoich zapasach układ 74HC00. Trzeba pamiętać, aby podłączyć wszystkie nieużywane wejścia do plusa zasilania lub do masy. Multimetr pokaże, że napięcie na nóżce 3 jest stabilne.

Natomiast bramka Schmitta zapewnia oscylację dzięki wbudowanej histerezie i powiązanemu dodatniemu sprzężeniu zwrotnemu. Jest to zilustrowane na **rysunku c**, który przedstawia wykres równoważny do rysunku a, ale dla bramki Schmitta.

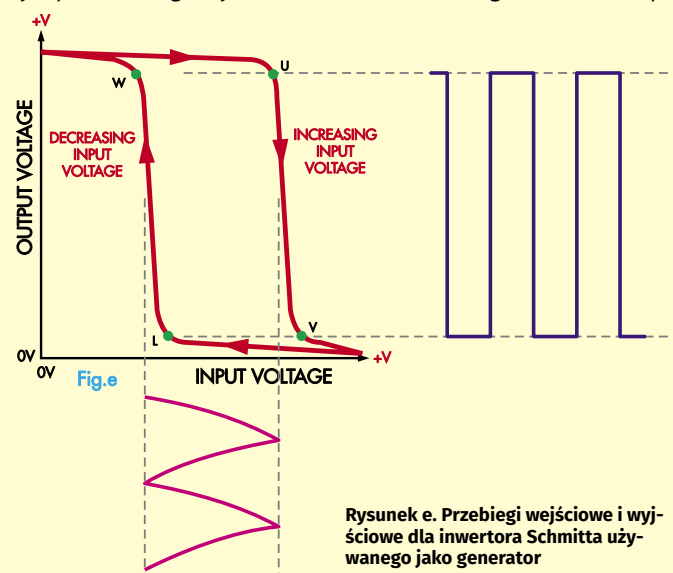
Gdy napięcie wejściowe wzrośnie powyżej górnego progu napięcia wejściowego (VT+, punkt U), wyjście natychmiast „przełącza się” na niski poziom (punkt V). Pozostaje tam do momentu, gdy napięcie wejściowe spadnie poniżej dolnego progu napięcia wejściowego (VT-, punkt L), a wyjście „zaskoczy” na wysoki poziom (punkt W).

Można to zademonstrować na przykładzie obwodu pokazanego na **rysunku d**. Zaczynamy od ustawienia potencjometru w skrajnym położeniu zgodnym z ruchem wskazówek zegara (nóżki 1 i 2 przy +6 V) i podłączamy zasilanie. Dioda LED powinna pozostać wyłączona. Powoli zmniejszamy napięcie wejściowe za pomocą potencjometru, aż dioda LED zacznie świecić. Notujemy napięcie wejściowe. Teraz stopniowo zwiększamy napięcie wejściowe, aż dioda LED zgaśnie. Powinna wystąpić różnica kilku woltów; jest to rozrzut histerezy (VT+ – VT-).

Wykonaliśmy jeden ruch zgodnie z ruchem wskazówek zegara wokół krzywej histerezy, jak wskazują strzałki na **rysunku c**.

Jeśli układ 74HC132 z wejściem Schmitta zostanie zastąpiony układem 74HC00 na **rysunku b**, zauważymy, że oscyluje on, generując na wyjściu przebieg prostokątny. Na wejściu pojawia się wykładniczy przebieg pseudo-trójkątny o amplitudzie równej rozrzutowi histerezy, jak pokazano na **rysunku e**.

Jedną z kluczowych kwestii do rozważenia jest to, jak szybkość oscylacji będzie się zmieniać w zależności od napięcia zasilania (zwłaszcza w układzie zasilanym bateryjnie). Jak się okazuje, zmniejszony prąd ładowania kondensatora jest w pewnym stopniu kompensowany przez spadek rozrzutu histerezy (ponieważ jest on w pewnym stopniu proporcjonalny do napięcia zasilania układu scalonego). Wynika z tego, że częstotliwość impulsów zmienia się tylko o kilka procent przy zmianach zasilania z 6 V do 5,5 V (zmiana napięcia o 8%).



Rysunek e. Przebiegi wejściowe i wyjściowe dla inwertera Schmitta używanego jako generator

ona na nieco wyższe napięcie zasilania, do 6 V. Świeże standardowe ogniwa alkaliczne AAA dostarczają nieznacznie więcej niż 6 V, więc linia zasilająca dla układów scalonych jest ograniczona za pomocą rezystora 47 Ω i diodę Zenera 6 V (ZD1). Można stosować ogniwa alkaliczne, suche ogniwa, akumulatory NiMH lub litowo-jonowe AAA.

Kondensatory filtrujące 4700 μF i 470 μF , w połączeniu z rezystorem szeregowym 47 Ω , niwelują wahania napięcia zasilania które bez nich zawierałoby zakłócenia od głośnika załączonego w momencie osiągnięcia przez odwrócone wahadło maksymalnego wychyłu.

Ponieważ potencjometry do montażu na płytce drukowanej z wbudowanymi przełącznikami są trudno dostępne, zastosowano oddzielny przełącznik zasilania. Zamiast baterii można użyć regulowanego zasilacza wtyczkowego 6 V DC. Trzeba upewnić się, że jest on regulowany, ponieważ w przeciwnym razie jego napięcie wyjściowe mogłoby być zbyt wysokie dla metronomu.

Wersja z 10 diodami LED

Schemat dla tej wersji pokazano na **rysunku 4**. Jest on podobny do wersji z 8 diodami LED, ale w całym układzie zastosowano logikę dodatnią. Impulsy są generowane przez timer CMOS-owy 555 IC2. Zegar taktuje czterobitowy licznik rewersyjny IC3. Wyjścia IC3 są dekodowane do 10 indywidualnych wyjść przez IC4, zapalając kolejno 10 diod LED.

Gdy świeci się dioda LED na którymś z krańców wychyłu (LED0 lub LED9), przerzutnik RS utworzony przez bramki IC2a i IC2c jest ustawiany lub zerowany, przełączając w ten sposób IC3 w tryb zliczania w górę lub w dół, odwracając sekwencję generowaną z użyciem diod LED.

Potencjometr VR5 reguluje jasność świecenia diod LED. Zamiast techniki zastosowanej w projekcie z 8 diodami LED, aby rozjaśnić końcowe diody LED, w tej wersji zastosowano lustro prądowe składające się z tranzystorów Q1 i Q2 z potencjometrem VR4, potencjometrem sterującym VR5 i kilkoma rezystorami stałymi. Potencjometr VR4 reguluje jasność

środkowych diod LED w stosunku do jasności dwóch diod końcowych. W sąsiedniej ramce został zamieszczony opis działania układu.

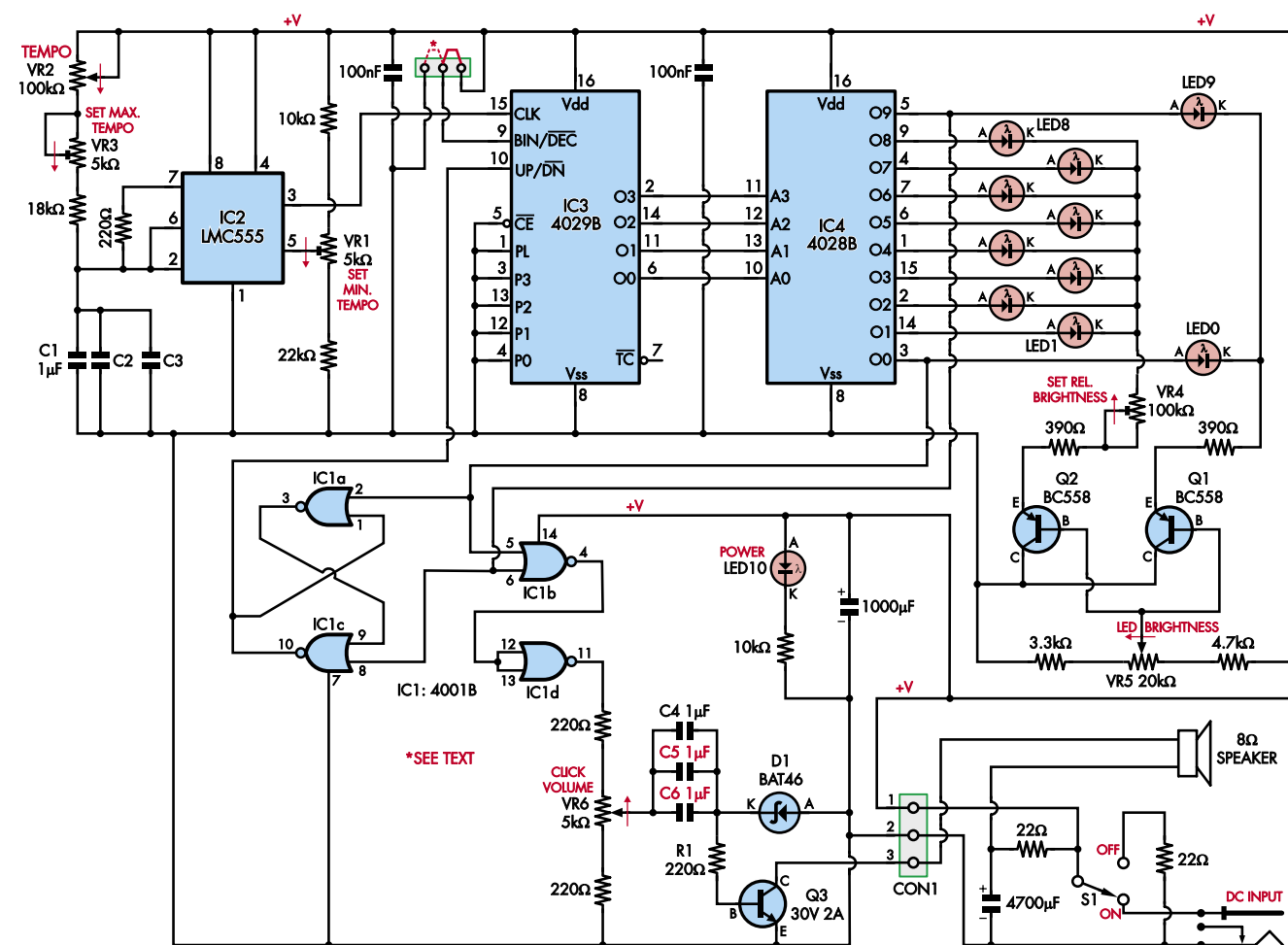
Generowanie kliknięć i warianty układu są takie same jak w przypadku konstrukcji z 8 diodami LED. Wyższe napięcie zasilania w tej wersji zapewni głośniejsze kliknięcie.

Budowa

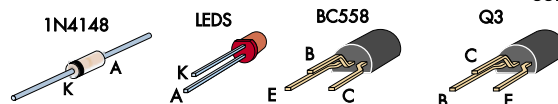
Na **rysunku 5** przedstawiono widok płytki drukowanej dla wersji z 8 diodami LED, a na **rysunku 6** dla wersji z 10 diodami LED. Większość elementów jest montowana na płytkach. Kilka z nich może wymagać zmiany wartości, dlatego niektóre części nie mają przypisanej wartości.

Niezależnie od wersji, proces budowy jest początkowo podobny. Zaczynamy od zamontowania wszystkich rezystorów o podanych stałych wartościach, używając multimetru do sprawdzenia wartości przed ich wlutowaniem.

Lutując diody należy upewnić się, że paski oznaczające katodę są skierowane tak,



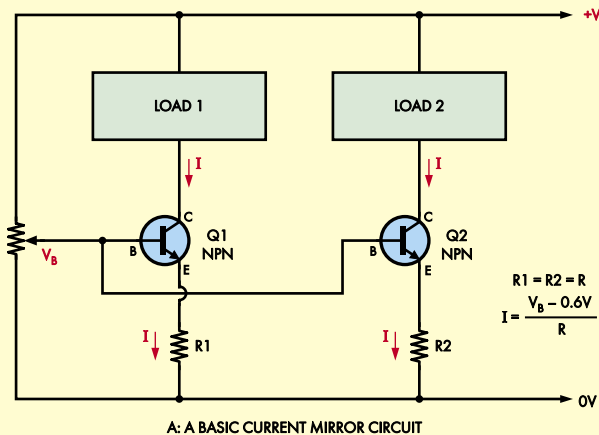
Rysunek 4. Metronom z 10 diodami LED ma bardziej skomplikowany schemat sterowania jasnością diod LED z tranzystorami PNP Q1 i Q2 tworzącymi lustro prądowe. Dzięki temu jasność ośmiu środkowych i dwóch zewnętrznych diod LED jest wybierana w szerokim zakresie regulacji. Dioda LED10 podświetla pokrętkę regulacji uderzeń na minutę. Poza tymi różnicami i zastosowaniem timera CMOS 555 jako generatora impulsów, obwód jest dość podobny do wersji z 8 diodami LED



Jak działa lustro prądowe

Układ lustro prądowego służy do dopasowania dwóch lub więcej prądów w zmiennych warunkach. Na **rysunku f** pokazano podstawowy przykład. Transzystory bipolarnie NPN Q1 i Q2 o jednakowych parametrach (w praktyce są one tylko do siebie zbliżone) mają potężne bazy, na których występuje napięcie sterujące V_B . W ten sposób ich emitory będą miały równe napięcia, około 0,6 V niższe niż V_B .

Jeśli rezystory emiterowe, R1 i R2, mają jednakową rezystancję, tranzystory będą przewodzić jednakowe prądy o natężeniach ok. $I_C = (V_B - 0,6 \text{ V}) / R$.



A: A BASIC CURRENT MIRROR CIRCUIT

Rysunek f. Podstawowy obwód lustro prądowego. Ponieważ Q1 i Q2 są tranzystorami o zbliżonych parametrach i dzięki ujemnemu sprzężeniu zwrotnemu zapewnianemu przez rezystory emiterowe, zmiana ich napięć na bazie za pomocą potencjometru skutkuje ściśle dopasowanymi prądami przez dwa niezależne obciążenia

Zakładając wystarczająco wysokie wzmocnienie prądowe tranzystorów (≈ 100), a tym samym pomijalne prądy bazy, prąd kolektora każdego z nich będzie taki sam jak jego prąd emitera. Innymi słowy, prądy płynące przez LOAD 1 i LOAD 2 będą takie same. Jeśli napięcie na bazie (V_B) jest zmieniane, prądy emitera i kolektora też będą się zmieniać, ale pozostaną takie same w obu tranzystorach.

Podobnie, jeśli jedno lub oba obciążenia różnią się rezystancją – w pewnych granicach – ich prądy będą nadal równe i określone powyższym równaniem.

W przypadku regulacji jasności 10-diodowego metronomu chcemy, aby prąd środkowych diod LED był ułamkiem prądu końcowych diod LED i był taki sam w szerokim zakresie prądów. Gdybyśmy zastosowali powyższy schemat, obwód wyglądałby jak na **rysunku g**, z R będącym ułamkiem $VR + R$, tj. nierównymi rezystorami emiterowymi.

Istnieje jednak pewien problem: ponieważ każda grupa diod LED jest naprzemiennie wyłączana, rezystancja obciążenia staje się niezwykle wysoka. W rezultacie tranzystor w wyłączonej części obwodu nie ma prądu kolektora, a prąd bazy staje się duży, ponieważ złącze baza-emiter jest diodą spolaryzowaną w kierunku przewodzenia. Zmniejsza to napięcie na bazie, a tym samym prąd kolektora drugiego tranzystora.

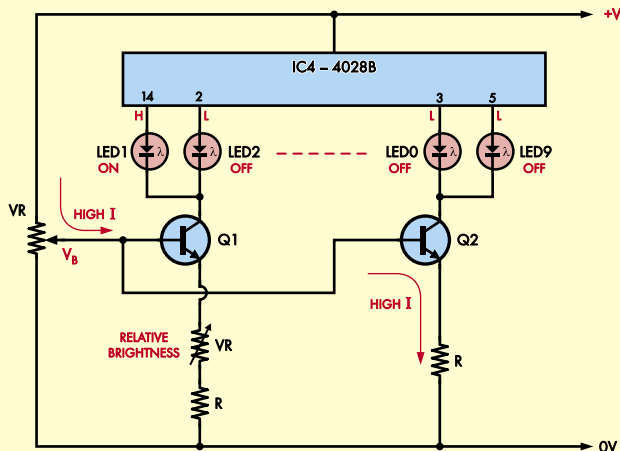
Na przykład, na **rysunku g** pokazano jedną ze środkowych diod LED włączoną, podczas gdy obie końcowe diody LED są wyłączone. Skutkuje to wysokim prądem bazy płynącym przez ich tranzystor (Q2).

Aby tego uniknąć, autor opracował inny schemat dla 10-diodowego metronomu, który jest pokazany na **rysunku h**. Układ ten działa, ponieważ dwa obciążenia są w praktyce prawie stałe i równe, a każde z nich składa się z jednej diody LED na raz.

Obwód lustro prądowego został odwrócony, zastosowane są tranzystory PNP zamiast NPN. Każda grupa diod LED jest częścią obwodu emitera, połączonego szeregowo z rezystorem, który określa prąd i względną jasność.

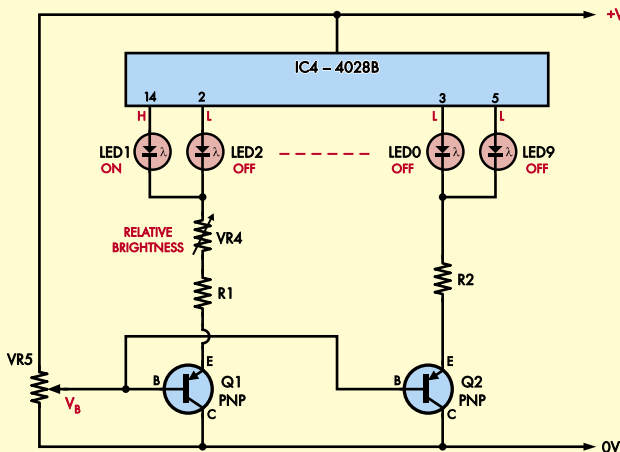
Po zapaleniu w sekwencji, każda środkowa dioda LED jest połączona szeregowo z $R1 + VR4$, który jest większy niż $R2$ połączony szeregowo z każdą końcową diodą LED. Prąd kolektora Q1 będzie ułamkiem prądu Q2 ($R2 / [R1 + VR4]$). Ułamek ten – stosunek prądów – będzie utrzymywany w zakresie V_B sterowanym przez VR5, a zatem jasność ośmiu środkowych diod LED będzie ułamkiem jasności dwóch końcowych diod LED w szerokim zakresie.

Sytuacja ta ulegnie załamaniu, jeśli V_B przekroczy $12 \text{ V} - V_{LED} - 0,6 \text{ V}$ lub około 10 V. Można tego uniknąć, dołączając do końcówek potencjometru VR5 rezystory stałe. Potencjometr montażowy VR4 umożliwia ustawienie stosunku rezystancji $R2 / [R1 + VR4]$ zgodnie z potrzebami, ustawiając w ten sposób różnicę jasności.



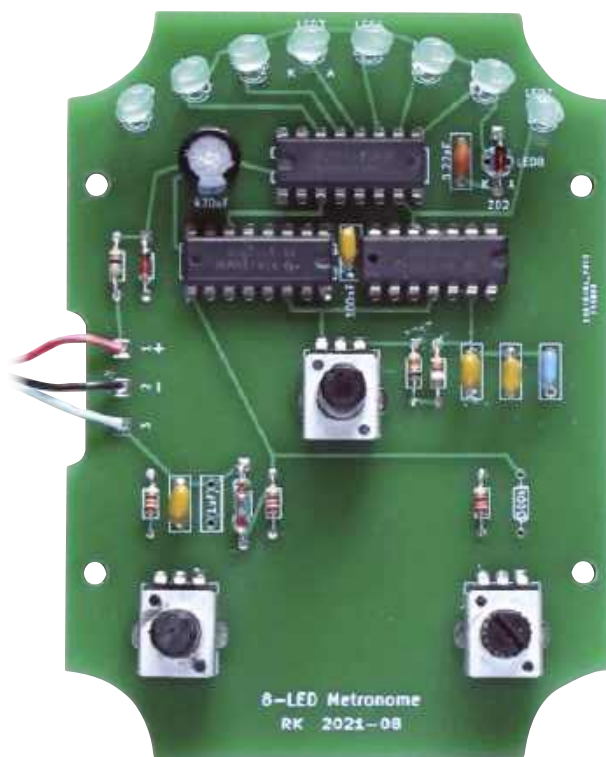
B: WHY THE SCHEME IN (A) DOESN'T WORK FOR TWO GROUPS OF LEDS

Rysunek g. Aby prądy obciążenia były różne, mogą być użyte różne wartości rezystora emiterowego, ale zachowują podobne proporcje prądu, gdy zmienia się napięcie na bazie tranzystorów

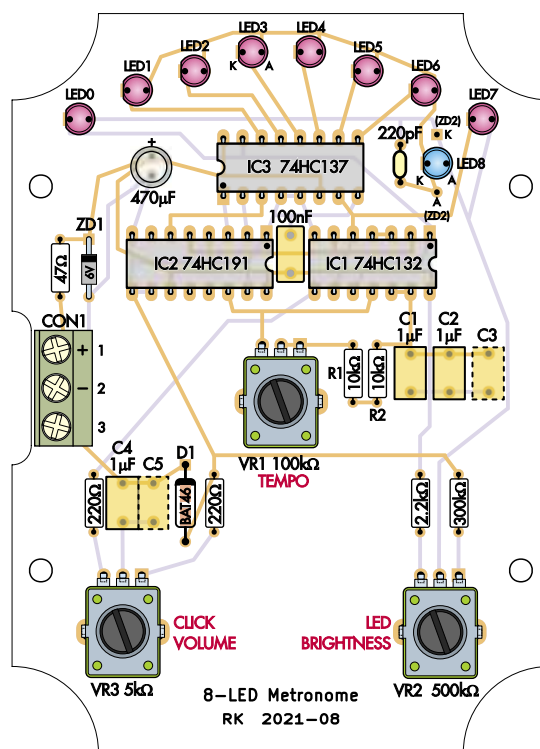


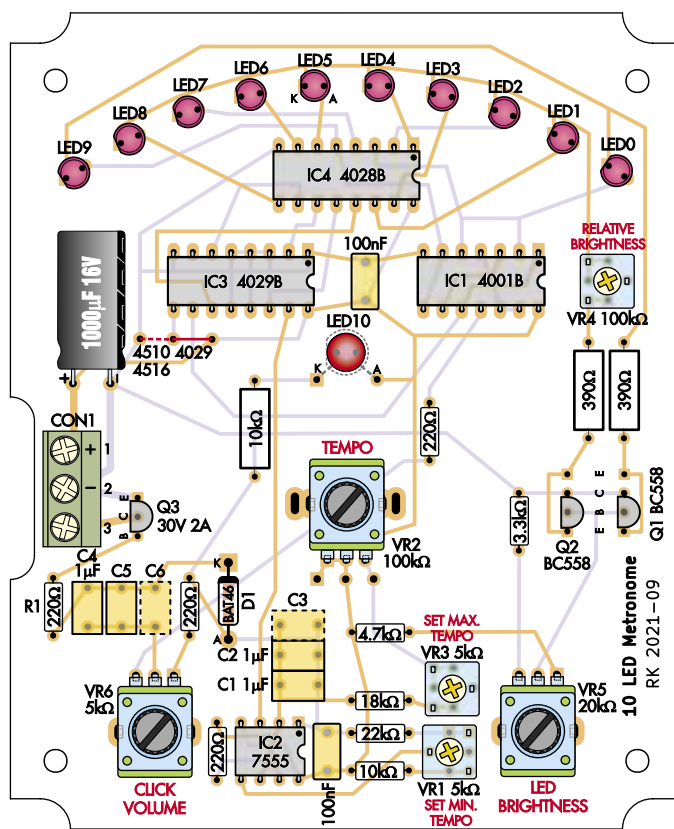
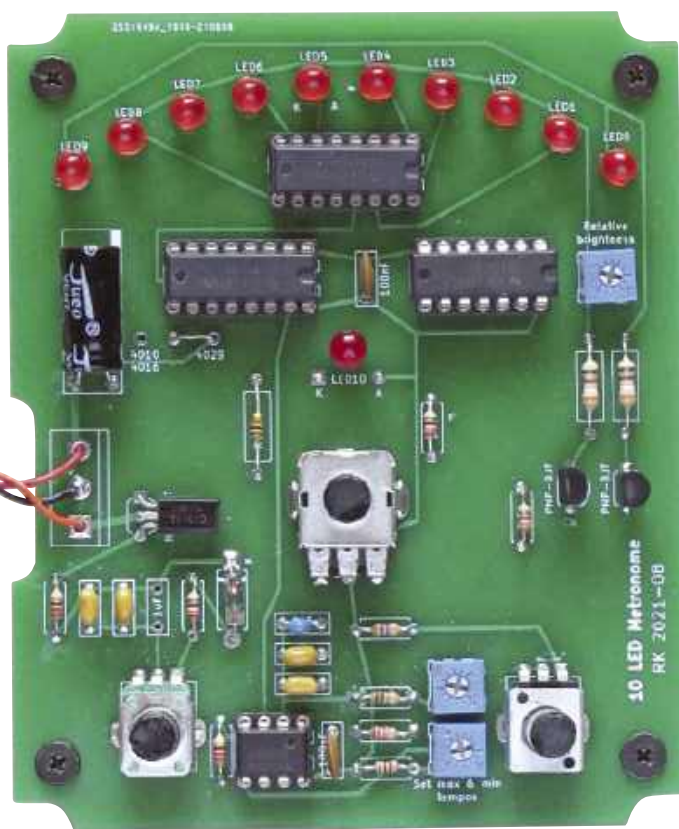
C: MODIFIED CURRENT MIRROR FOR THE 4000 VERSION METRONOME

Rysunek h. Obwód pokazany na rysunku g może działać wadliwie z powodu problemów z nadmiernym prądem płynącym do baz tranzystorów, gdy obciążenia mogą być włączane i wyłączane niezależnie. Przedstawiony obwód rozwiązuje ten problem, zamieniając tranzystory NPN na PNP i utrzymując połączenia rezystorów ustawiających prąd na emiterach tranzystorów



Rysunek 5. Większość elementów wersji 8-ledowej jest montowana na płytce drukowanej, jak pokazano tym rysunku. Montaż jest prosty, ale należy uważać, aby prawidłowo umieścić układy scalone, diody LED i diody zgodnie z rysunkiem. Nie należy też pomylić potencjometrów, ponieważ wszystkie mają różne wartości (sprawdź kody wydrukowane na ich obudowach)





Rysunek 6. Metronom z 10 diodami LED jest nieco bardziej skomplikowany niż wersja z 8 diodami. Należy zwrócić uwagę na elementy zamontowane na leżąc i upewnić się, że tranzystory znajdują się w odpowiednich otworach

należy wywiercić otwory na diody LED i potencjometr w przedniej części obudowy. **Rysunek 7** może posłużyć jako szablon do precyzyjnego wykonania otworów.

Można wydrukować szablon na kartonie, wyciąć w szablonie otwory montażowe 5 mm i tymczasowo przykleić szablon do wewnętrznej strony przedniej połowy obudowy.

W przypadku korzystania z zalecanych owalnych diod LED konieczne będzie ostrożne wydłużenie otworów po ich wywierceniu. Należy pamiętać, że linia podświetlenia owalnej diody LED jest prostopadła do większego wymiaru korpusu diody LED. Trzeba zdecydować, którą orientację wybrać i odpowiednio ustawić diody LED i otwory. Podczas wiercenia lub dopasowywania otworów do użytych diod LED należy sprawdzić, czy diody pasują do otworów, ale jednocześnie czy ich zamontowanie nie wymaga użycia nadmiernej siły.

Kolejną czynnością jest montaż trzech potencjometrów na płytce drukowanej bez ich lutowania, mocujemy też płytkę drukowaną do przedniej części obudowy. Aby zapewnić miejsce na komponenty na płytce drukowanej, mogą być potrzebne podkładki dystansowe na śrubach. Trzeba sprawdzić długości wałków potencjometrów i skrócić je w razie potrzeby.

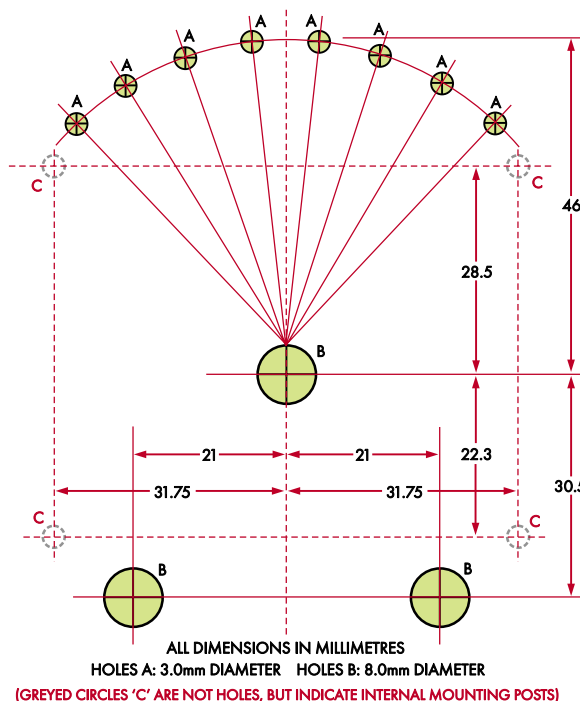
Po dopasowaniu wysokości pokręteł z założonymi gałkami do wymiarów obudowy

tak, aby nie było zbyt dużych luzów i traci można przylutować trzy potencjometry, po wcześniejszym sprawdzeniu, czy mają prawidłowe wartości.

Teraz, zwracając uwagę na ich orientację (patrz oznaczenia A i K na płytce drukowanej),

wkładamy nóżki diod LED do płytki bez ich lutowania. Jeśli używane są diody owalne, należy je lekko przekręcić, aby dopasować do łuku.

Ponownie mocujemy płytkę drukowaną do przedniej części obudowy i umieszczamy każdą diodę LED w odpowiednim



Rysunek 7. Szablon do wiercenia otworów w panelu przednim wersji z 8 diodami LED. Należy wywiercić jedenaście otworów: osiem na diody LED (rozmiar i kształt dostosowany do używanych diod LED, oznaczone jako „A”) i trzy otwory 8 mm na wałki potencjometrów, oznaczone jako „B”. Przerywane okręgi wyznaczają pozycje słupków montażowych w określonej obudowie. Nie należy ich wiercić

Otwory o średnicy 3 mm są określone jako „A”, ale ich rozmiar zależy od typu używanych diod LED

otworze z przodu obudowy. Najlepiej, aby lekko wystawały. Sprawdzamy wygląd diod LED i dostosowujemy je w razie potrzeby, po czym można już je przylutować do płytki drukowanej. Płytką musi znajdować się w docelowym położeniu. Diody LED prawdopodobnie nie będą przylegały do płytki drukowanej, ale będą oddalone od niej o kilka milimetrów.

Na **rysunku 8** przedstawiono tarczę pokrętki tempa. Można je pobrać ze strony internetowej Silicon Chip. Dobrym pomysłem jest wydrukowanie go na papierze fotograficznym. Osiągniemy w ten sposób dobry wygląd. Zakłada się, że VR1 jest odpowiednikiem typu podanego w wykazie elementów. Potencjometr musi mieć kąt obrotu równy 280° .

Dopasowujemy pokrętkę do wału potencjometru tempa i przyklejamy je do przedniej części obudowy. Dopasowujemy pokrętkę do potencjometru tempa w taki sposób, aby jego obrót rozciągał się równo poza linie tempa 36 i 216. Jest to konieczne, ponieważ potencjometry zwykle mają martwą strefę na każdym z krańców, gdzie rezystancja zmienia się bardzo nieznacznie.

Tranzystor NPN Q1, kondensator $4700 \mu\text{F}$, przełącznik S1, głośnik i pojemnik baterii nie są zamontowane na płytce drukowanej, ale przymocowane do tylnej części obudowy (patrz zdjęcie poniżej). Otwory na głośniki mogą mieć dowolny wzór. Autor użył perforowanej blachy, wybrał wiertło o średnicy otworów, zacisnął blachę do wewnętrznej strony tylnej części obudowy i użył jej jako przewodnicy do wiercenia.

Przełącznik suwakowy mocujemy do panelu za pomocą małych śrub i nakrętek. Pojemnik baterii i głośnik są utrzymywane na miejscu za pomocą zacisków wykonanych z dużego, wytrzymałego spinacza do papieru.

Tranzystor Q1 i kondensator $4700 \mu\text{F}$ są montowane blisko głośnika i podłączone

bezpośrednio do zacisków głośnikowych, co minimalizuje rezystancję pasożytniczą. Nie są one przełączane przez S1, również w celu zmniejszenia rezystancji pasożytniczej. Może to być ważne, ponieważ napięcie zasilania jest stosunkowo niskie, a impedancja głośnika wynosi 8Ω .

Gdy metronom jest wyłączony, przez te elementy przepływa jedynie niewielki prąd upływowy. Jeśli jednak metronom jest nieużywany przez dłuższy czas, najlepiej jest wyjąć ogniwa.

Przed zamontowaniem w obudowie tranzystora Q1 i kondensatora należy przylutować nóżkę emitera i ujemne wyprowadzenie kondensatora. Konieczne jest dokładne sprawdzenie wyprowadzeń tych elementów. Błąd może spowodować przepływ nadmiernego prądu i uszkodzić tranzystor Q1 lub spalić cewkę i membranę głośnika.

Kolejna czynność, to wycięcie drewnianej podstawy tak, aby pasowała do obudowy i przymocowanie do niej tylnej części za pomocą śrub, nadając obudowie lekkie nachylenie do tyłu.

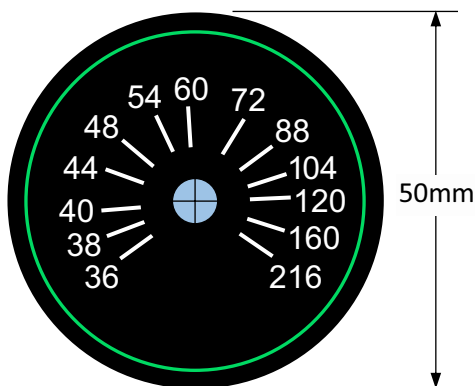
Na koniec podłączamy części poza PCB do 3-pinowego złącza zaciskowego, jak pokazano na zdjęciach. Jeśli czytelnik zrezygnuje z listwy zaciskowej, można przylutować przewody bezpośrednio do padów PCB.

Regulacja

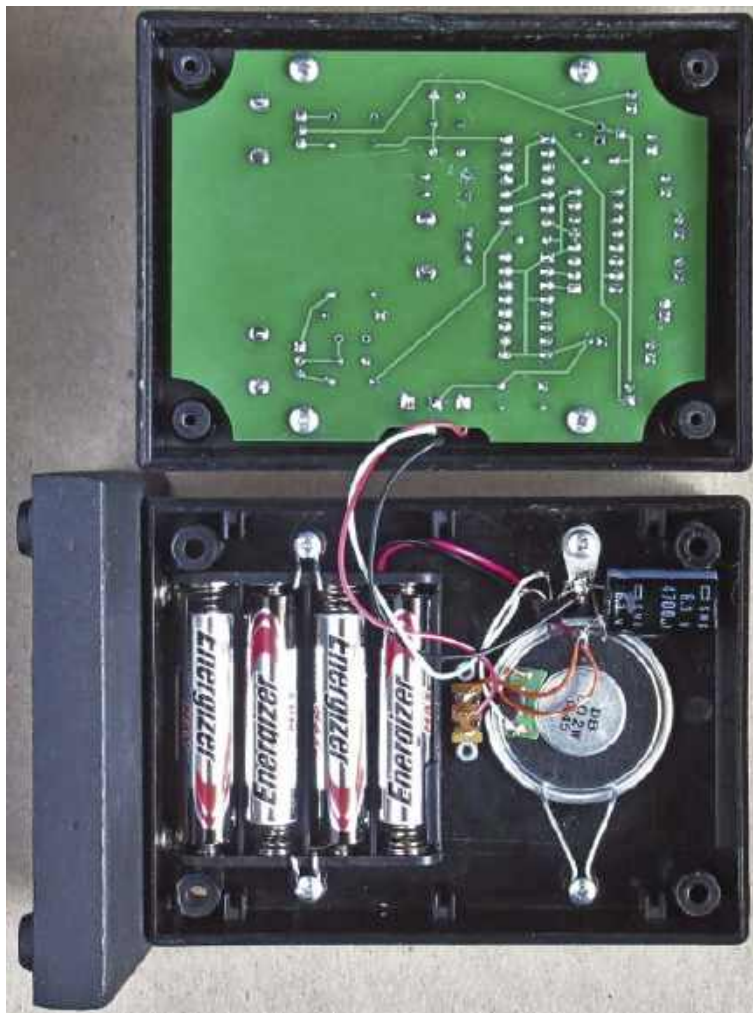
Tempo taktowania i jego zakres prawdopodobnie będą wymagały regulacji. Ośmiu różnych układów scalonych 74HC132 wykazało kilkuprocentowy rozrzut, z czego jeden o około 7% powyżej średniej. Tempo może z tego powodu nie odpowiadać oznaczeniom na pokrętkę, mogą ponadto występować różnice w wartościach kondensatorów i rezystorów VR1 i R1/R2. Potencjometry mogą różnić się nawet o 20%.

Jeśli uzyskane wyniki nie będą zadowalające, do dokładnej kalibracji konieczne będzie użycie miernika częstotliwości. Taki tryb pomiarowy ma wiele tanich multimetrów.

Mierzymy częstotliwość impulsów na nóżce 11 układu IC1 lub na nóżce 14 układu IC2. Przy VR1 ustawionym



Rysunek 8. Przedstawioną grafikę tarczy dla wersji 8-LED należy wydrukować na papierze fotograficznym, wyciąć ją i przykleić z przodu obudowy. Dokładna średnica nie jest krytyczna, ale powinna być zbliżona do 50 mm. Jest ona dostępna do pobrania ze strony internetowej w formacie PDF



W przypadku metronomu LED niektóre elementy, takie jak głośnik i uchwyt baterii, nie są montowane na płytce drukowanej, ale z tyłu obudowy. Zdjęcie przedstawia układ metronomu z 8 diodami LED

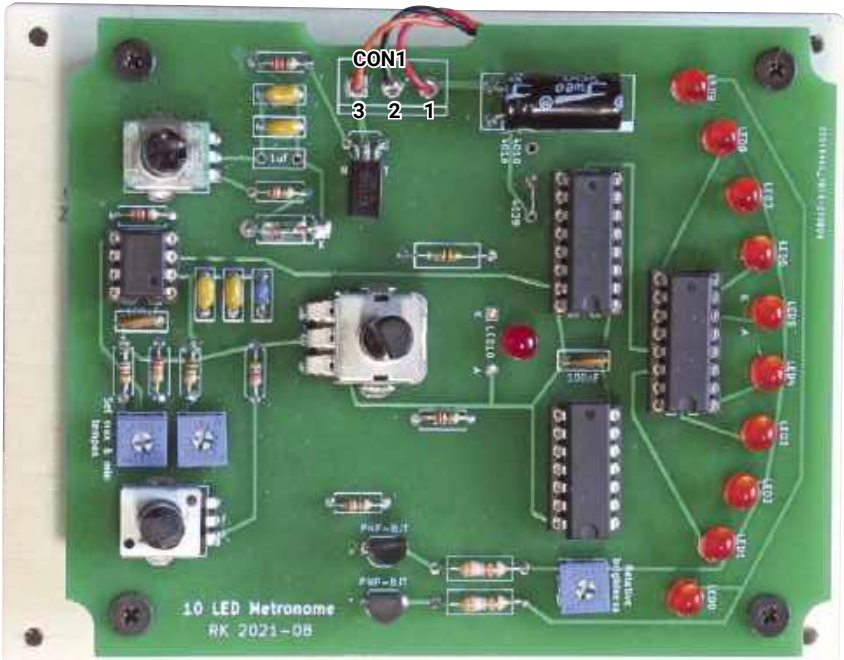
na najniższe tempo (wyrównane z oznaczeniem pokazującym 36 uderzeń/minutę), należy dopasować kondensator lub kondensatory równolegle dla C1...C3, które dadzą częstotliwość impulsów 4,2 Hz (tabela 1).

Jeśli nie dysponujemy wystarczającym zasobem kondensatorów do wypróbowania (lub chcemy zaoszczędzić czas), można obliczyć procentowy błąd częstotliwości (rzeczywisty vs oczekiwany) i zmierzyć pojemność C1...C3. Następnie mnożymy odczytaną pojemność przez wartość procentową i dzielimy przez 100. Jest to pojemność, którą należy dodać (jeśli drgania są zbyt szybkie) lub odjąć (jeśli jest są zbyt wolne).

Aby odjąć pojemność, należy wymienić C1 i/ lub C2 na kondensatory o niższej pojemności, a następnie ponownie sprawdzić i ewentualnie dodać nieco większą pojemność, dopasowując C3, aby uzyskać odpowiednią częstotliwość.

Zakładając, że wybrano R1 i R2 jako 1/5 wartości VR1, przy VR1 ustawionym na 216 uderzeń/min, odczyt częstotliwości powinien wynosić 25,2 Hz (patrz tabela 1). Jeśli wartość ta znacznie odbiega od normy, warto dostosować wartości rezystorów, zmniejszając je, aby przyspieszyć, lub zwiększając, aby spowolnić tempo. Obliczenie procentowego błędu częstotliwości w stosunku do 25,2 Hz wskazuje procent, o jaki należy zmienić całkowitą rezystancję. Wynik końcowy powinien zawierać wszystkie tempa od 36 do 216 (a tym samym częstotliwości impulsów w tabeli 1) zgodnie z oznaczeniami na tarczy.

Na koniec sprawdzamy głośność i barwę kliknięcia. Sprawdzamy, czy VR3 płynnie



Podobnie jak w przypadku metronomu 8-LED, metronom 10-LED również ma elementy zamontowane na tylnym panelu, a nie na płytce drukowanej, jak widać na poniższym zdjęciu, ma przewody wychodzące z górnej części

zmienia głośność kliknięcia od zera do maksimum. Jeśli jest zbyt miękkie przy maksimum, potrzebny jest tranzystor o wyższym współczynniku h_{FE} . Jeśli kliknięcie nie zmienia się płynnie, należy zwiększyć rezystor 220 Ω , aż zmiana głośności będzie zadowalająca.

Barwę kliknięcia można zmieniać, dodając kondensator w miejscu oznaczonym C5. Dodanie kondensatora powinno zapewnić bardziej „łagodne” kliknięcie.

Można także wypróbować głośnik w obu polaryzacjach, ponieważ może to wpłynąć na barwę dźwięku.

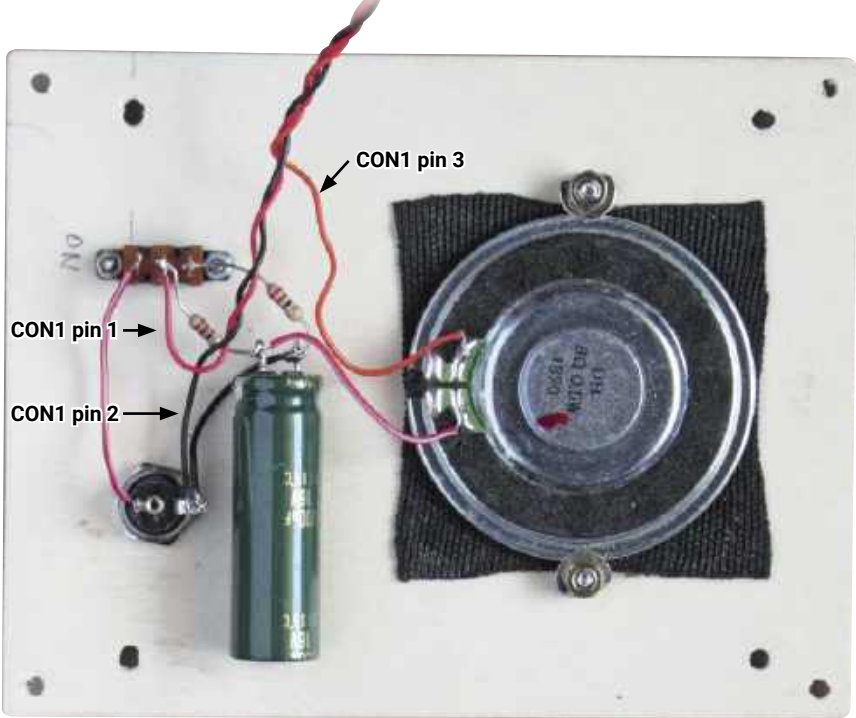
Rozwiązywanie problemów

Jeśli diody LED nie świecą lub zachowują się dziwnie, należy sprawdzić orientację wszystkich elementów.

Sprawdzamy, czy nie ma mostków lutowniczych lub zimnych lutów, i czy nóżki układu scalonego nie są wygięte. Sprawdzamy również,

Tabela 1. Zależność między częstotliwością impulsów (Hz) a tempem (uderzenia/minutę)

Tempo [bpm]	wersja 8 LED [Hz]	wersja 10 LED [Hz]
36	4,20	5,40
38	4,43	5,70
40	4,67	6,00
44	5,13	6,60
48	5,60	7,20
54	6,30	8,10
60	7,00	9,00
66	7,70	9,90
72	8,40	10,8
80	9,33	12,0
88	10,3	13,2
104	12,1	15,6
120	14,0	18,0
160	18,7	24,0
216	25,2	32,4



czy na nóżce 11 układu scalonego IC1 i nóżce 14 układu scalonego IC2 występują impulsy.

Finalizacja metronomu 10-diodowego

Otwory montażowe dla VR2 są przystosowane do potencjometru o kącie obrotu 280° (pasującego do pozostałych) lub większego potencjometru 300°. Choć różnica jest subtelna, potencjometr 300° pozwala na nieco większe rozłożenie wartości tempa. Przed zamontowaniem VR2 należy zmierzyć jego rezystancję, podzielić przez 5,5 i sprawdzić, czy jest ona bliska 18 kΩ. Jeśli nie, może być konieczne zastąpienie rezystora 18 kΩ inną wartością, która jest do niej zbliżona.

Jeśli posiadany układ IC3 to 4029, należy dopasować czerwone połączenie pokazane

na rysunku 6. W przeciwnym razie dopasujemy połączenie przewodowe tam, gdzie znajduje się przerywana czerwona linia. W obu przypadkach można użyć odciętego wyprowadzenia elementu.

Wkładamy diodę LED10 przez płytkę drukowaną od tyłu (patrz zdjęcie na stronie). Lutujemy jej wyprowadzenia z tyłu PCB do punktów „K” i „A”.

Zamiast wydruku umieszczonego na tarczy, jak w wersji z 8 diodami LED, w metronomie z 10 diodami LED zastosowano przezroczystą plastikową tarczę z czarnym nadrukiem i wyraźnymi cyframi tempa (**rysunek 9**). Wybieramy wzór odpowiedni dla posiadanego typu potencjometru VR2. Jeśli nadrukowana tarcza nie jest wystarczająco sztywna, konieczne może być zastosowanie przezroczystej tarczy nośnej.

Przyklejamy dysk do tylnej części plastikowej tulei zamontowanej do wału VR2 na wcisk lub z użyciem kleju. Tuleja ta może być wykonana z przyciętego pokrętkła. Dioda LED10 podświetla liczbowe oznakowanie tempa na tarczy, aby były widoczne przez plastikowy panel. Jasność diody LED10 jest określana rezystorem R2 (przyjęto 10 kΩ). Jeśli nie są zadowalająca, należy zmniejszyć rezystancję R2, aby była jaśniejsza lub zwiększyć R2, aby była ciemniejsza.

Dla wersji z 10 diodami LED, diody te nie wystają przez panel przedni, ale są widoczne, więc otwory na diody LED nie są potrzebne. Otwory dla dwóch najniższych potencjometrów (VR5/6) mogą mieć 8 mm, ale VR2 będzie wymagał większego otworu, aby pomieścić tuleję przytrzymującą pokrętko tempa.

Wykaz elementów

metronom 8-diodowy

- 1 płytka dwustronna kod PCB 23111211, 71 mm × 98 mm
- 1 3-pinowa listwa zaciskowa (CON1)
- 1 obudowa plastikowa Serpac 131-BK, 111 mm × 82,5 mm × 38 mm [Mouser 635-131-B, Digi-Key SR131B-ND].
- 1 drewniana podstawa, 75 mm × 90 mm × 12,5 mm (DIY)
- 4 ognia AAA, najlepiej akumulatory NIMH (BAT1)
- 1 pojemnik na baterie 4×AAA (BAT1) [Keystone Electronics 2482; Mouser 534-2482, Digi-Key 36-2482-ND]
- 1 głośnik 8 Ω, średnica 36 mm (SPK1) [DB Unlimited SM360608-1; Mouser 497-SM360608-1, Digi-Key 2104-SM260608-1-ND].
- 1 100 kΩ liniowy potencjometr pionowy 9 mm/10 mm (VR1) [Mouser 652-PTV09A4025UB104].
- 1 500 kΩ liniowy potencjometr pionowy 9 mm/10 mm (VR2) [Mouser 652-PTV09A4020UB504].
- 1 5 kΩ potencjometr liniowy pionowy 9 mm/10 mm (VR3) [Mouser 652-PTV09A4030UB502, Digi-Key PTV09A4030UB502-ND].
- 1 przełącznik suwakowy SPST lub SPDT (S1) [Alpha SS60012F-0102-4V-NB; Mouser SS60012F-0102-4V-NB].
- 1 podstawa 14-pinowa DIL IC (opcjonalne dla IC1)
- 2 podstawki 16-pinowa DIL IC (opcjonalnie dla IC2 i IC3)
- 3 pokrętki pasujące do VR1...VR3
- 4 samoprzylepne nóżki gumowe
- 2 małe śruby i nakrętki (do montażu przełącznika suwakowego)
- 1 duży, wytrzymały spinacz do papieru
- 8 wkrętów samowinujących 4 × 6 mm
- 2 małe, krótkie (~10 mm) wkręty do drewna z tłem walcowym (do montażu obudowy do podstawy)
- 1 końcówka lutownicza z otworem o średnicy ~3,25 mm
- przewody połączeniowe o różnych długościach i kolorach

Półprzewodniki

- 1 74HC132 poczwórna 2-wejściowa bramka Schmitta NAND, DIP-14 (IC1)
- 1 74HC191 wstępnie ustawiany 4-bitowy binarny licznik rewersyjny, DIP-16 (IC2)
- 1 74HC137 lub 74HC138 dekodery 3 do 8 linii, DIP-16 (IC3)
- 1 tranzystor NPN 30 V 1 A, TO-92 (Q1) [KSD471ACYTA, KSC2328AYTA lub ZTX690B].
- 8 „superjasnych” diod LED, okrągłych lub owalnych (LED0...LED7); [Broadcom HLMP-HM74-34CDD (zielony, owalny), Kingbright WP7083ZGD/G (zielony, 5 mm), Jameco 2169846 (zielony, 3 mm)].
- 1 dioda Zenera 6,0 V 500 mW (ZD1) [1N5233 lub odpowiednik]
- 1 niebieska/biała dioda LED lub dioda Zenera 4,7 V (LED8/ZD2) [1N5231]
- 1 dioda Schottky’ego 150 mA; np. BAT46/BAT48/BAT85 (D1); [Jaycar ZR1141, Altronics Z0044, Mouser 511-BAT46].

Kondensatory

- 1 4700 µF 6,3V elektrolityczny [Mouser 667-EEU-FS0J472]
- 1 470 µF 6,3V elektrolityczny niskoprofilowy [Mouser 232-63AX470MEFC8X75]
- 4 1 µF 50 V wielowarstwowy ceramiczny
- 1 100 nF 50V ceramiczny
- 1 220 pF 50V ceramiczny

Rezystory

- (1% 1/4 W, 1/8 W lub 1/16 W, metalizowane, chyba że podano inaczej)
- 1×300 kΩ (opcjonalnie)
- 2×10 kΩ
- 1×2,2 kΩ
- 2×220 Ω
- 1×47 Ω

Zalecane jest stosowanie rezystorów 1%, ale w razie potrzeby można użyć rezystorów 5%. Uzyskanie prawidłowego zakresu tempa może wymagać dokładniejszej regulacji.

metronom 10-diodowy

- 1 płytka dwustronna – kod PCB 23111212, 108 mm × 89 mm
- 1 zasilacz wtyczkowy 12 V DC 100 mA+
- 1 3-pinowa listwa zaciskowa (CON1)
- 1 gniazdo zasilaczowe do montażu w obudowie, pasujące do wtyczek (CON2)
- 1 głośnik 8 Ω o średnicy 50 mm (SPK1) [DB Unlimited SM500208-1; Mouser 497-SM500208-1].
- 2 miniaturowe potencjometry montażowe ze śrubką od góry 5 kΩ (VR1, VR3)
- 1 100 kΩ 280° liniowy potencjometr pionowy 9 mm/10 mm (VR2) [Mouser 652-PTV09A4025UB104] LUB
- 1 100 kΩ 300° liniowy potencjometr pionowy 9 mm/10 mm (VR2) [Mouser 652-PDB12-M4251-104BF] (patrz tekst)
- 1 miniaturowy potencjometr montażowy ze śrubką od góry 100 kΩ (VR4)
- 1 liniowy potencjometr pionowy 9 mm/10 mm 20 kΩ (VR5) [Mouser 652-PTV09A-4030UB203, Digi-Key PTV09A-4030U-B203-ND].
- 1 liniowy potencjometr pionowy 9 mm/10 mm 5 kΩ (VR6) [Mouser 652-PTV09A4030UB502, Digi-Key PTV09A4030UB502-ND].
- 1 przełącznik suwakowy SPDT (S1) [Alpha SS60012F-0102-4V-NB; Mouser SS60012F-0102-4V-NB]
- 1 podstawa 14-pinowa DIL IC (opcjonalne; dla IC1)
- 1 podstawa 8-pinowa DIL IC (opcjonalne; dla IC2)
- 2 podstawki 16-pinowa DIL IC (opcjonalne; dla IC3 i IC4)
- 1 drewniana podstawa, 75 mm × 150 mm × 12 mm (DIY)
- 1 rama drewniana, 110 mm × 130 mm × 40 mm (DIY)
- 1 przyciemniany na czerwono przezroczysty panel akrylowy o wymiarach 100 mm × 125 mm × 2,5...3 mm (lub pasujący do ramki)
- 3 pokrętki dopasowane do VR2, VR5 i VR6
- 1 przezroczysty, nadający się do zadrukowania plastik i plastikowy podkład pod pokrętko tempa, tuleję lub przycięte pokrętko do montażu na VR2
- 4 podkładki dystansowe 15 mm z gwintem M3 (do montażu płytki drukowanej do panelu)
- 8 śrub z tłem walcowym M3 × 6 mm (do montażu płytki drukowanej do panelu)
- 2 małe śruby i nakrętki (do montażu przełącznika suwakowego)
- 2 śruby z tłem walcowym M3 × 10 mm, płaskie podkładki i nakrętki (do montażu głośnika)
- 1 duży, wytrzymały spinacz do papieru
- 6 małych, krótkich (~10 mm) wkrętów do drewna z tłem walcowym
- przewody połączeniowe o różnych długościach i kolorach
- 1 kwadratowa tkanina głośnikowa 55 mm × 55 mm

Półprzewodniki

- 1 CD4001BE poczwórna 2-wejściowa bramka NOR, DIP-14 (IC1)
- 1 układ czasowy CMOS TLC555IP lub LMC555CN, DIP-8 (IC2)
- 1 CD4029BE, CD4510BE lub CD4516BE 4-bitowy binarny licznik rewersyjny, DIP-16 (IC3)
- 1 CD4028BE dekodery binarny 4-10, DIP-16 (IC4)
- 2 tranzystory PNP BC558 30 V 100 mA, TO-92 (Q1, Q2)
- 1 tranzystor 30 V 2 A NPN, TO-92 (Q3) [KSC2328AYTA lub ZTX690B].
- 1 dioda Schottky’ego 150 mA, np. BAT46/BAT48/BAT85 (D1); [Jaycar ZR1141, Altronics Z0044, Mouser 511-BAT46].
- 10 „superjasnych” diod LED, okrągłych lub owalnych (LED0-LED9); [Zalecane Cree C566D-RFE-CV0X0BB1 (czerwone, owalne)].
- 1 superjasna czerwona dioda LED 5 mm (LED10) [zalecane Kingbright WP7113SRD/J4].

Kondensatory

- 1 4700 µF 16 V elektrolityczny, średnica 13 mm [Mouser 232-16PK4700MEFC125X]
- 1 1000 µF 16 V elektrolityczny, średnica 8 mm [Mouser 232-16ZLH1000MEFC8X2]
- 5 1 µF 50 V wielowarstwowy ceramiczny
- 2 100 nF 50 V ceramiczne

Rezystory (1% 1/4 W, 1/8 W lub 1/16 W, metalizowane, chyba że podano inaczej)

- 1×22 kΩ 1×18kΩ 2×10 kΩ 1×4,7 kΩ
- 1×3,3 kΩ 2×390 Ω 4×220 Ω 2×22 Ω

Konieczne będzie wycięcie cienkiego panelu, na którym zostanie zamontowana płytką drukowaną za pomocą gwintowanych kołków 15 mm. Panel ten będzie pasował do tylnej części drewnianej ramy. Zawiera on ponadto głośnik, gniazdo zasilania, przełącznik zasilania, elektrolit 4700 μF i dwa rezystory 22 Ω (patrz zdjęcie). Gniazdo zasilania powinno być dopasowane do zastosowanej wtyczki.

Głośnik jest montowany za pomocą śrub i nakrętek, z podkładkami, które zostały wygięte w dół po jednej stronie. Do zabezpieczenia dużego kondensatora elektrolitycznego można użyć kleju na gorąco lub silikonu.

Panel, wraz z płytką PCB, jest przymocowany do ramy za pomocą małych wkrętów do drewna, które są przykręcane do czterech kawałków drewna o wymiarach około 8 mm $\times$ 8 mm $\times$ 12 mm przyklejonych do wewnętrznych narożników ramy.

Obudowa metronomu jest widoczna na zdjęciu. Tworzą ją cztery sklejone z 45-stopniową fazą kawałki drewna. Ci, którzy nie lubią stolarki, prawdopodobnie będą mogli znaleźć odpowiednie plastikowe pudełko z przezroczystą pokrywą, które będzie wystarczająco duże, aby pomieścić płytkę drukowaną i inne elementy.

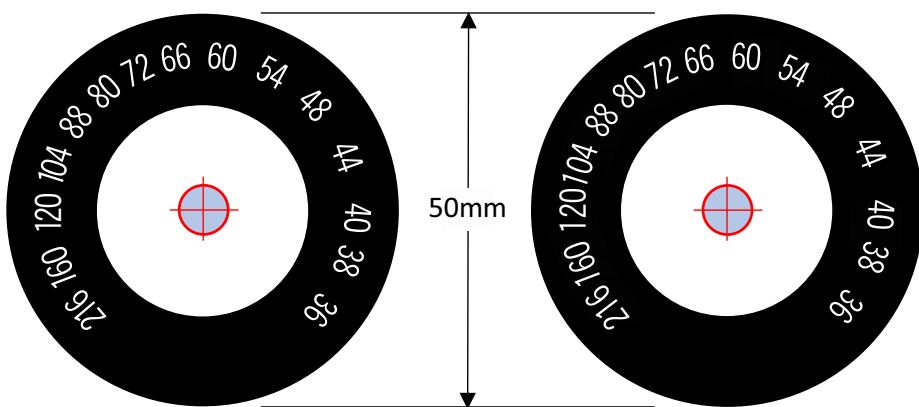
Sterowanie i regulacja

Przed podłączeniem zasilania należy dokładnie sprawdzić okablowanie do elementów zewnętrznych. Błąd w tym miejscu może spowodować nadmierny prąd i uszkodzenie tranzystora Q3, a także spalanie cewki lub membrany głośnika. Warto porównać swoje okablowanie z pokazanym na zdjęciach.

Ze względu na różnice w elementach, tempo będzie prawdopodobnie wymagało kalibracji. Będzie do tego pomocny miernik częstotliwości (nawet bardzo podstawowy, jaki można znaleźć w wielu multimetrach).

Ustawiamy VR2 na najwolniejsze tempo (36 uderzeń/min) i mierzymy częstotliwość impulsów na nóżce 3 IC2 lub na nóżce 15 IC3. Aby uzyskać częstotliwość 5,4 Hz (patrz tabela 1), za pomocą potencjometru montażowego VR1 regulujemy napięcie sterujące (nóżka 5) timera IC2. Jeśli regulacja potencjometrem montażowym VR1 nie pozwala uzyskać częstotliwości do 5,4 Hz, należy dodać kolejny kondensator równoległe do C1 i C2 (w pozycji C3), aby ją spowolnić, lub zmniejszyć wartość C1 i/lub C2, aby ją przyspieszyć.

Gdy uzyskana częstotliwość jest prawidłowa, ustawiamy tempo na 216 uderzeń/min i regulujemy VR3, aby uzyskać 32,4 Hz. Jeśli VR3 nie pozwala uzyskać



Rysunek 9. Metronom 10-LED posiada dwa pokręta, pasujące do większego potencjometru 300° (po lewej) lub standardowego potencjometru 280° (po prawej). W przeciwieństwie do wersji z 8 diodami LED, są one wydrukowane na przezroczystej folii i połączone z obracającym się wałem potencjometru. Dzięki temu dioda LED10 z tyłu może oświetlać wybrane tempo

Przerzutniki RS

Obie wersje metronomu zawierają przerzutnik RS, układ logiczny z dwoma stanami: set i reset (ustawiony/wyzerowany). Podanie wysokiego poziomu na wejście S przy jednoczesnym utrzymaniu niskiego poziomu na wejściu R powoduje przełączenie przerzutnika w stan ustawiony, w którym pozostaje do momentu wyzerowania. Analogicznie, podanie wysokiego poziomu na wejście R przy jednoczesnym utrzymaniu niskiego poziomu na wejściu S powoduje przejście przerzutnika RS w stan wyzerowania i pozostanie w nim do momentu ponownego ustawienia.

Zgodnie z dzisiejszymi standardami nazewnictwa, RS jest przezroczystym zatraskiem, a nie przerzutnikiem, ponieważ nie ma wejścia zegarowego, ale tradycyjny termin „przerzutnik” jest nadal używany. Można również powiedzieć, że jest to 1-bitowa komórka pamięci lub układ bistabilny.

Przerzutnik RS jest prostym typem sekwencyjnego obwodu logicznego, co oznacza, że stan wyjściowy zależy od jego „historii” – czyli układ ma pamięć. Można to porównać z logiką kombinatoryczną, w której wyjścia zależą tylko od wartości wejść, nie ma historii, ani pamięci.

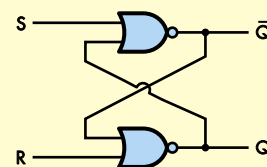
Przerzutnik RS może być wykonany z dwóch bramek NOR, jak pokazano na sąsiednim schemacie, można też uzyskać gotowe układy scalone flip-flop. W 10-diodowym metronomie używamy już bramek NOR do innych celów, więc realizacja przerzutnika na bramkach pozwala uniknąć konieczności stosowania dodatkowego układu scalonego.

Zasada działania jest następująca. Wyobraźmy sobie, że oba wejścia R i S są w stanie niskim. Obwód może początkowo znajdować się w dowolnym stanie: ustawionym, z Q w stanie wysokim i $\bar{Q}$ w stanie niskim, lub wyzerowanym, odwrotnie. Podawanie impulsów do wejścia S spowoduje, że $\bar{Q}$ przejdzie w stan niski lub pozostanie w stanie niskim, a to spowoduje, że Q będzie w stanie wysokim. Jest to stan ustawiony. Dalsze impulsy podawane do wejścia S nie będą miały żadnego wpływu, ponieważ wyjście Q utrzymuje górne wejście bramki NOR w stanie wysokim, zakładając, że R pozostaje w stanie niskim.

Podobnie, podawanie impulsów do wejścia R spowoduje, że wyjście Q przejdzie w stan niski, a tym samym $\bar{Q}$ przejdzie w stan wysoki. Jest to stan wyzerowania. Dalsze impulsy R nie będą miały żadnego wpływu, zakładając, że S pozostanie w stanie niskim.

W przypadku metronomu 8-diodowego, przerzutnik RS jest zbudowany z dwóch bramek NAND zamiast bramek NOR. Oznacza to, że RS używa logiki ujemnej. W logice ujemnej bramki NAND stają się bramkami NOR, a przerzutnik RS jest ustawiany i zerowany stanami niskimi impulsów, w szczególności z LED7 (ustawianie) i LED0 (zerowanie). Wyjście Q przerzutnika RS jest dołączone do licznika 74HC191 w celu zmiany kierunku zliczania.

W metronomie 10-diodowym przerzutnik RS jest zbudowany z dwóch bramek NOR układu CD4001. Używana jest logika dodatnia, a RS działa w sposób opisany powyżej.



Przerzutnik RS zbudowany z dwóch bramek NOR. Wyjścia Q i $\bar{Q}$ mają zawsze przeciwną polaryzację. Q przechodzi w stan wysoki, gdy wejście S przechodzi w stan wysoki, a przejście wyjścia Q w stan niski, jest następstwem przejścia wejścia R w stan wysoki. Oba wejścia nie mogą być jednocześnie w stanie wysokim

częstotliwości do 32,4 Hz, należy zmienić wartość rezystora szeregowego 18 kΩ, a następnie powtarzamy regulację dla najwolniejszego i najszybszego tempa.

Na koniec regulujemy jasności między dwiema końcowymi i środkowymi diodami LED. Służy do tego potencjometr VR4.

Barwę i głośność kliknięcia można również zmodyfikować dla wersji z 10 diodami LED. Do uzyskania płynnej zmiany głośności dobieramy rezystor R1, jak opisano dla wersji z 8 diodami LED.

Aby zmienić barwę kliknięcia, należy poeksperymentować z łączną wartością C4...C6. Większa pojemność powinna zapewnić łagodniejsze kliknięcie. Głośnik może również

wpływać na ton, więc można wypróbować podłączenie głośnika w obu polaryzacjach.

Rozwiązywanie problemów

Jeśli metronom nie działa, należy sprawdzić orientację układu IC2 i spolaryzowanych elementów dyskretnych. Sprawdzamy też, czy wszystkie nóżki układu scalonego są prawidłowo włożone. Czasami są one wygięte i nie wchodzi do gniazda lub płytki drukowanej. Sprawdzamy, czy występują impulsy na nóżce 3 układu IC2 i nóżce 15 układu IC3.

Jeśli sekwencja LED działa tylko w jednym kierunku, prawdopodobnie nie działa przetrzutnik RS, lub nie odbiera impulsów S i R z układu IC4.

Obsługa

Obsługa obu wersji jest prosta. Włącz metronom i dostosuj głośność kliknięć, jasność diod LED i tempo zgodnie z potrzebami.

Prąd zasilania metronomu z 8 LED-ami wynosi około 2...4 mA, w zależności od jasności diod LED, głośności kliknięć i tempa. Ogniwa AAA mają zwykle pojemność około 900 mAh, więc, zakładając, że urządzenie jest używane przez około pół godziny dziennie, ogniwa alkaliczne lub akumulatory powinny wystarczyć w przypadku wersji 8-LED nawet na miesiąc. ■

Randy Keenan

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

REKLAMA

Prenumeratorzy mają bezpłatny dostęp do e-wydań archiwalnych
EdW starszych niż 24 miesiące

Sięgnij po archiwalne wydania ELEKTRONIKI dla WSZYSTKICH

Przesyłka **GRATIS**

Zamów wygodnie na www.UlubionyKiosk.pl

eprasa.pl 4e05f27e00



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/dp3bh>
Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania

Wzmacniacz mocy 500 W, część 2

W poprzednim odcinku został opisany moduł wzmacniacza o mocy 500 W. Znalazły się w nim szczegóły dotyczące jego parametrów wraz z opisem części układowej. W tym odcinku zajmiemy się budową wzmacniacza, zaczynając od montażu płytki drukowanej (modułu wzmacniacza). W kolejnym odcinku zbudujemy kompletny wzmacniacz wraz z aktywnym chłodzeniem, zabezpieczeniem głośników i wskaźnikiem przesterowania.

Wzmacniacz mocy 500 W składa się z czterech głównych elementów: modułu wzmacniacza, zasilacza, sekcji aktywnego układu chłodzenia i ochrony głośników oraz układu wskaźnika przesterowania.

W tym artykule skoncentrujemy się na montażu modułu wzmacniacza, którego układ został opisany w pierwszym odcinku. W następnym odcinku, zostanie szczegółowo opisany zasilacz, obudowa oraz końcowy montaż i okablowanie, które połączy wszystkie te części razem. A teraz zajmiemy się budową najważniejszego modułu wzmacniacza.

Budowa

Moduł wzmacniacza o mocy 500 W jest zbudowany na dwustronnej, płytce drukowanej (kod 01107021) o wymiarach 402 mm × 124 mm. Na **rysunku 6** przedstawiono rozmieszczenie elementów, z którego warto korzystać podczas montażu.

Przed rozpoczęciem prac należy dokładnie sprawdzić płytkę. Pozwoli to bliżej zapoznać się z jej koncepcją montażu ale ujawni też ewentualne wady (jakkolwiek mało prawdopodobne).

Budowę należy rozpocząć od zamontowania tranzystorów Q1 i Q2. Są to małe tranzystory SOT-23/TO-236 do montażu powierzchniowego. Są one stosunkowo łatwe

do przylutowania ze względu na szeroko rozstawione piny, ale może być potrzebna lupa i dobre oświetlenie.

Najpierw należy wypośrodkować tranzystor Q1 względem padów, przytrzymując go pęsetą i przylutować jeden z pinów do PCB. Sprawdzamy, czy jest prawidłowo ustawiony względem pozostałych padów i w razie potrzeby ponownie podgrzewamy punkt lutowniczy celem wykonania poprawek. Następnie lutujemy pozostałe piny. Tranzystor Q2 jest montowany w podobny sposób.

Nie należy przejmować się, jeśli dodamy tak dużo cyny, że złącza na tych elementach w obudowie SOT-23 będą wyglądać jak małe srebrne kulki. Jest mało prawdopodobne, aby stwarzało to jakiegokolwiek problemy. Chcemy, aby połączenia były błyszczące, ale sytuacja, w której cyny jest zbyt dużo jest lepsza niż wtedy, gdy jest jej za mało!

Jeśli jednak zostanie podjęta decyzja o usunięciu nadmiaru cyny, należy dodać odrobinę topnika i dotknąć połączenia czystym grotiem lutownicy.

Teraz montujemy małe (1/4 W lub 1/2 W) rezystory. Przed wlutowaniem sprawdzamy każdą wartość za pomocą multimetru ustawionego na funkcję omomierza. W celu określenia rezystancji raczej nie należy polegać wyłącznie

na kodach paskowych, ponieważ mogą być mylące i trudne do dokładnego odczytu.

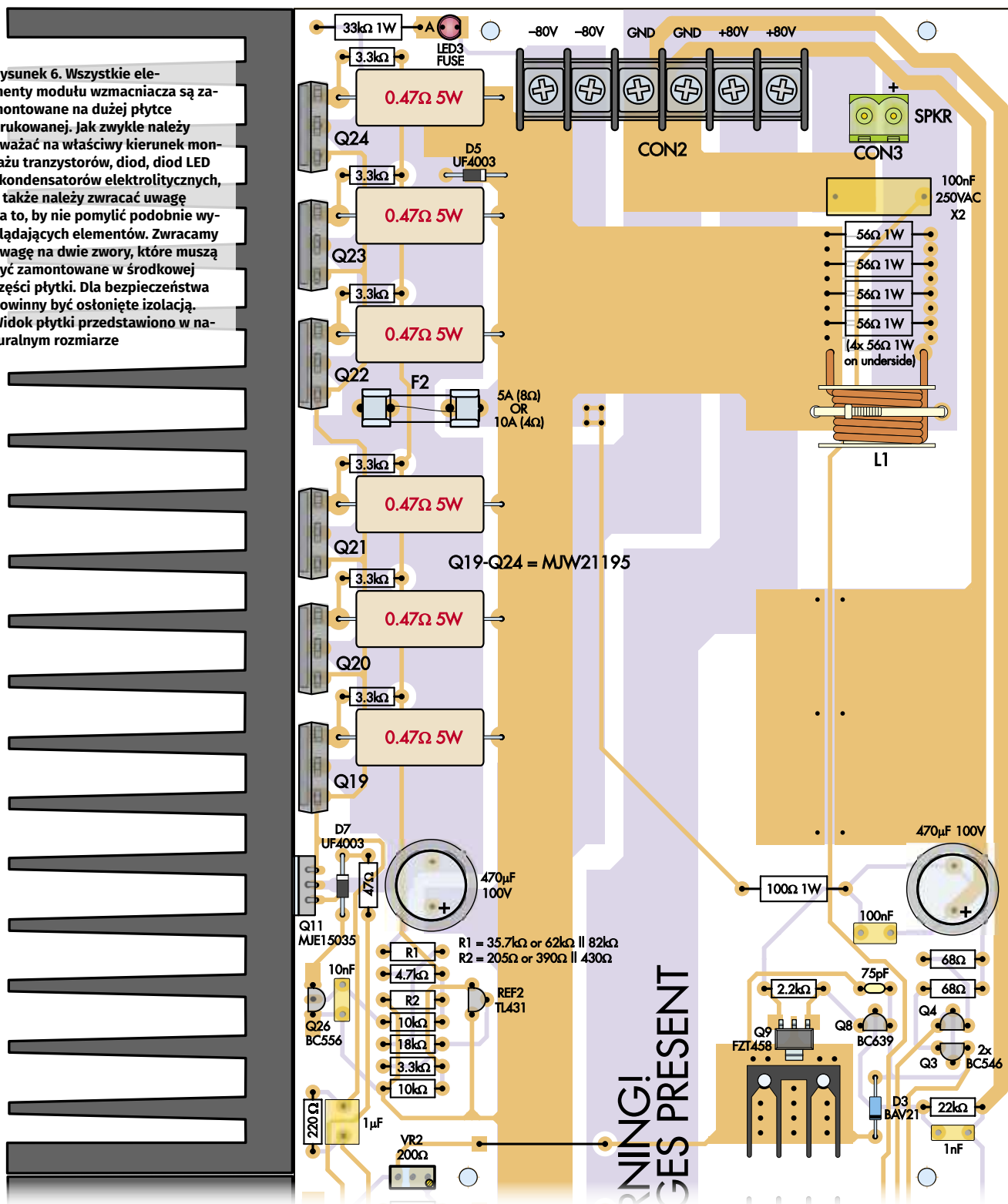
Jak widać, na płytce drukowanej znajdują się dwie pary rezystorów oznaczone jako R1 i R2; nie mają one przypisanych wartości. Nominalne wartości wymagane dla tych rezystorów (które definiują krzywe obszarów bezpiecznej pracy SOA) to $R1=35,328\text{ k}\Omega$ i $R2=204,8\ \Omega$. W praktyce takich wartości nie da się zakupić, ale istnieją dwa sposoby, aby się do nich zbliżyć.

Możemy użyć rezystorów o wartości z szeregu E96, z $R1=35,7\text{ k}\Omega$ (+1%) i $R2=205\ \Omega$ (+0,1%). W praktyce może się okazać, że łatwiej będzie zakupić cały zestaw rezystorów w rzeczonym szeregu, w którym te wartości się znajdują, niż pojedyncze elementy o takich wartościach.

Nieco bardziej precyzyjną metodą na znalezienie wartości R1 i R2 jest użycie równoległych par rezystorów, z których jedna jest zamontowana na górnej stronie płytki drukowanej w normalny sposób, a druga jest przylutowana do padów pod spodem. Są to $62\text{ k}\Omega||82\text{ k}\Omega$ dla R1 dające $35,3\text{ k}\Omega$ (-0,08%) i $390\ \Omega||430\ \Omega$ dające $204,5\ \Omega$ dla R2 (+0,15%).

Nie sądzimy, aby błąd +1% przy użyciu $35,7\text{ k}\Omega$ dla R1 miał znaczenie. Rezystory pomiarowe $0,47\ \Omega$ mają 5% tolerancję, a krzywe SOA mają wbudowany margines

Rysunek 6. Wszystkie elementy modułu wzmacniacza są zamontowane na dużej płytce drukowanej. Jak zwykle należy uważać na właściwy kierunek montażu tranzystorów, diod, diod LED i kondensatorów elektrolitycznych, a także należy zwracać uwagę na to, by nie pomylić podobnie wyglądających elementów. Zwracamy uwagę na dwie zwory, które muszą być zamontowane w środkowej części płytki. Dla bezpieczeństwa powinny być osłonięte izolacją. Widok płytki przedstawiono w naturalnym rozmiarze



bezpieczeństwa. Jeśli stanowi to jakiś problem, można zamiast tak dobranych rezystorów zastosować pary równoległe.

Montujemy teraz te rezystory w ośmiu lokalizacjach, używając dowolnej metody.

Następnie zgodnie z rysunkiem oraz warstwą opisową płytki PCB montujemy dwie małe diody 1N4148 (D1 i D2), których katody

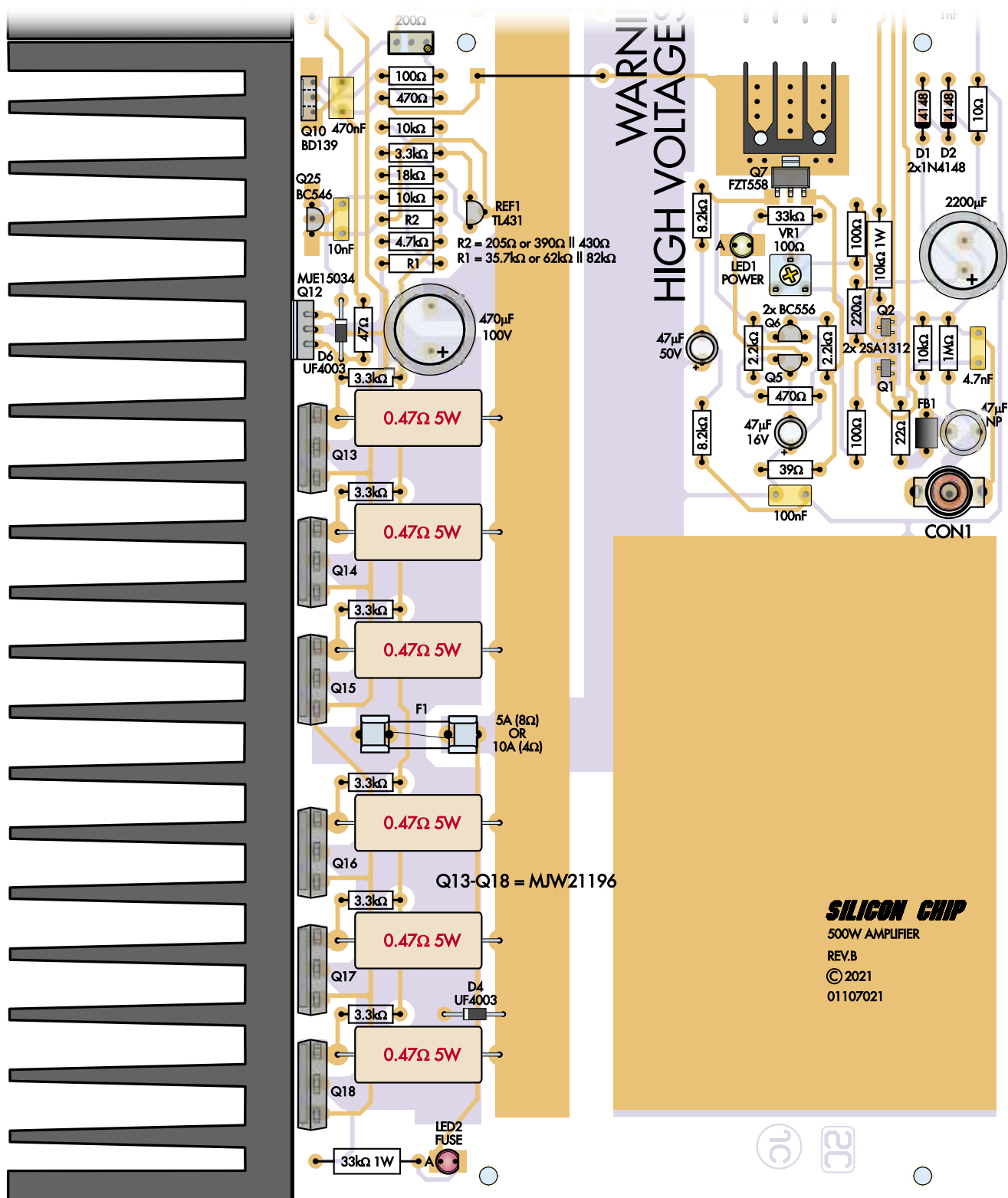
zostały oznaczone paskiem. Następnie montujemy diodę BAV21 (D3) z katodą skierowaną w tę samą stronę.

Następnie ponownie zgodnie z rysunkiem 6 oraz opisem na PCB należy zamontować diody UF4003 (D4...D7).

Pośrodku płytki drukowanej nad Q7 i Q9 należy zamontować dwie zwory. Wykonujemy

je przy użyciu ocynowanego drutu miedzianego o średnicy 0,7 mm pokrytego rurką termokurczliwą o średnicy 1 mm i długości tak dobranej, aby odsłonięte były tylko same końce.

Kontynuujemy montaż rezystorów 1 Ω , ponownie przed montażem sprawdzając ich wartości. W przypadku rezystorów 56 Ω w pobliżu złącza głośnikowego CON3, cztery z nich należy



zamontować na górnej stronie płytki drukowanej, a pozostałe cztery na spodzie. Opis na PCB pokazuje pozycje rezystorów po obu stronach.

Następnie montujemy tranzystory małosygnałowe w obudowach TO-92. Są to Q3 i Q4 (BC546) oraz Q5 i Q6 (BC556). W tej chwili pomijamy Q25 i Q26, ponieważ należy je zamontować na radiatorze. Można jednak

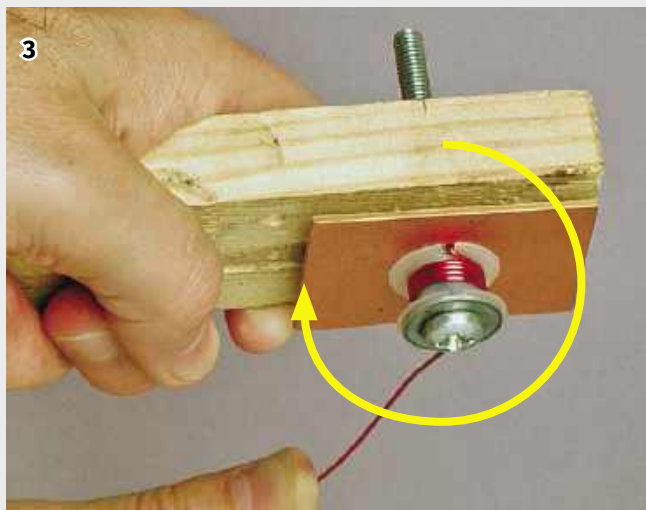
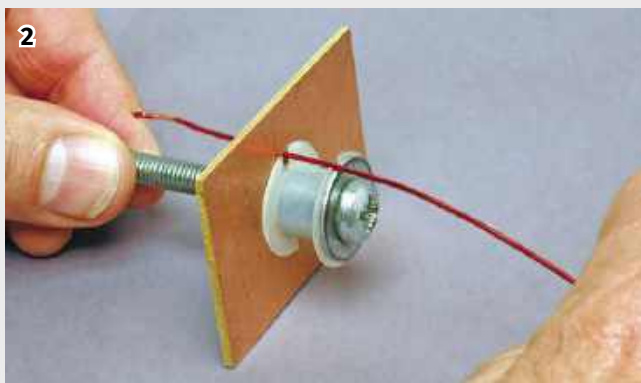
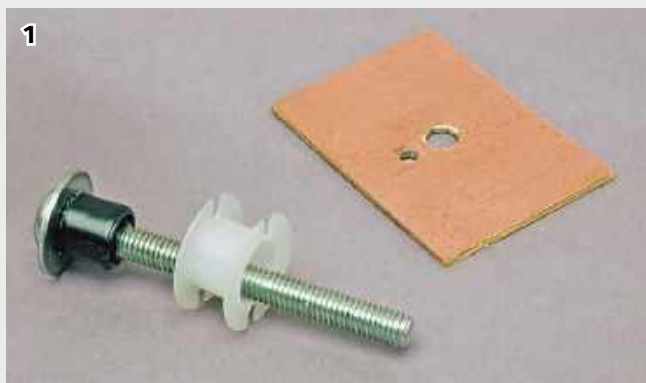
teraz zamontować dwa układy napięcia referencyjnego TL431, również w obudowach TO-92 (REF1 i REF2). Trzeba uważnie przyrzeć się oznaczeniom tranzystorów i upewnić się, że w każdym miejscu zainstalowany jest właściwy typ.

Trzy diody LED są montowane około 5 mm nad płytką drukowaną. Należy

zadbać o ich prawidłowy kierunek montażu. Dioda LED1 powinna być koloru zielonego. Dłuższa nóżka jest anodą, a jej pozycja jest oznaczona na płytce literą „A”.

Montujemy teraz kondensator 75 pF 200 V wraz z kondensatorami 1 nF, 10 nF, 100 nF, 470 nF i 1 μF MKT. Kolejne elementy to potencjometry montażowe

Przyrząd do nawijania L1



taśmy izolacyjnej na śrubę lub na rurkę. Kołnierz musi mieć wystarczającą średnicę, aby karkas dobrze przylegał, ale zestawienie nie było zbyt ciasne.

Wiercimy otwór o średnicy 5 mm przez środek PCB oraz otwór wyjściowy o średnicy 1,5 mm w odległości około 8 mm, który będzie wyrównany z jedną ze szczelin w karkasie.

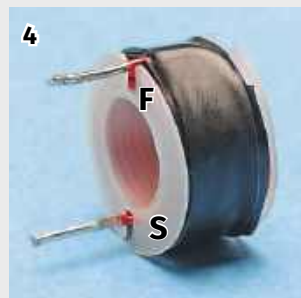
Karkas można wsunąć na kołnierz, po czym koniec używanej płytki drukowanej jest nasuwany na śrubę, tj. karkas jest umieszczany w pozycji między podkładką a płytką.

Wyrównujemy karkas tak, aby jedna z jego szczelin pokrywała się z otworem wyjściowym, a następnie przykręcamy pierwszą nakrętkę i mocno ją dokręcamy. Następnie montujemy uchwyt, wierząc otwór o średnicy 5 mm na jednym końcu, wsuwając go na śrubę i przykręcając drugą nakrętkę.

Zdjęcia pokazują, w jaki sposób należy użyć przyrządu do wykonania cewki 2,2 μH .

Najpierw należy nasunąć karkas na kołnierz śruby (1), a następnie założyć płytkę końcową i przewlec drut przez szczelinę wyjściową (2). Następnie mocujemy uchwyt i nawijamy cewkę ciasno na szpulę, nawijając 13,5 zwoja emaliowanego drutu miedzianego o średnicy 1,25 mm (3). Na koniec zabezpieczamy gotowy zwoj (4) od zewnątrz za pomocą rurki termokurczliwej o średnicy 20 mm.

Drut należy nawinąć na karkas zgodnie z ruchem wskazówek zegara.



Przyrząd do nawijania składa się ze śruby M5 70 mm, dwóch nakrętek M5, płaskiej podkładki M5, kawałka zbędnej płytki drukowanej lub podobnego materiału o wymiarach około 40 mm $\times$ 50 mm oraz kawałka drewna (około 140 mm $\times$ 45 mm $\times$ 20 mm) na uchwyt.

Podczas użytkowania, płaska podkładka przylega do tła śruby, po czym na śrubę nakładany jest kołnierz, który ma pomieścić karkas. Kołnierz ten powinien być nieco mniejszy niż wewnętrzna średnica karkasu i może być wykonany poprzez nawinięcie

VR1 i VR2. VR2 ma śrubę regulacyjną skierowaną ku dolnej części płytki (prawa krawędź na rysunku 6).

Następne są cztery zaciski bezpieczników M205. Należy umieścić je całkowicie na płycie przed lutowaniem i upewnić się, że klipsy mocujące znajdują się na zewnątrz. Najlepszym sposobem gwarantującym właściwy montaż podstawki bezpiecznika (klipsów) jest wstępne zamontowanie bezpiecznika w gnieździe, a następnie włożenie do płytki i przylutowanie całości takiego złożenia. Pozwoli to utrzymać klipsy bezpieczników w prawidłowej pozycji.

Teraz można wlutować 12 rezystorów 0,47 Ω 5W. Powinny być zamontowane około 2 mm od płytki drukowanej, aby powietrze

mogło je efektywnie chłodzić. Do zapewnienia prawidłowych odstępów od podłoża może być zastosowana tekturowa przekładka wsunięta pod korpusy rezystorów przed lutowaniem.

Teraz należy zamontować złącza, tj. gniazdo RCA (CON1), dwupinowe gniazdo do podłączenia głośnika (CON3) i 6-pinowe złącze zasilania (CON2). W przypadku CON3 najpierw należy włożyć wtyczkę bloku zacisków do gniazda, a następnie włożyć gniazdo w otwory płytki drukowanej z wejściami przewodów w kierunku zewnętrznej krawędzi płytki drukowanej.

Teraz montujemy kondensator 100 nF X2 znajdujący się w pobliżu CON3. Następnie można zamontować kondensatory elektrolityczne 47 μF , 470 μF i 2200 μF . Kondensator

elektrolityczny 47 μF NP (niespolaryzowany) może być umieszczony w dowolny sposób, ale dla pozostałych należy zastosować właściwy kierunek montażu.

Należy pamiętać, że kondensator 47 μF powyżej Q5 i Q6 musi być przystosowany do pracy z napięciem co najmniej 50 V (np. typ 63 V byłby akceptowalny).

Mini radiatory

Przed zamontowaniem Q7 i Q9 należy najpierw przymocować radiatory. W tym celu wkładamy słupki montażowe do otworów w płytce drukowanej i lutujemy je do spodu płytki. Będą one wymagały silnego podgrzewania lutownicą zanim cyna stopi się. Należy uważać, aby nie poparzyć się gorącymi radiatorami.

Trzeba poczekać, aż ostygną przed zamontowaniem Q7 i Q9.

Teraz zajmijmy się tranzystorem Q7 (FZT558). Przed umieszczeniem elementu pomocne będzie rozproszanie niewielkiej ilości topnika w postaci pasty na wszystkich czterech padach. Wyśrodkowujemy tranzystor względem padów PCB i lutujemy jeden z pinów do płytki. Sprawdzamy wyrównanie i jeśli to konieczne ponownie podgrzewamy cynę, aby poprawić ustawienie elementu. Następnie lutujemy pozostałe piny.

Radiator komponentu należy przylutować do płytki drukowanej tuż obok dodatkowego radiatora zewnętrznego. Ponownie konieczne będzie podgrzanie jej lutownicą przez dłuższy czas, ponieważ radiator odprowadza ciepło. Gdy cyna rozpuści się, należy przylutować wyprowadzenie elementu tak szybko, jak to możliwe, aby uniknąć przegrzania tranzystora. Teraz, w ten sam sposób instalujemy tranzystor Q9 (FZT458).

Cewka L1

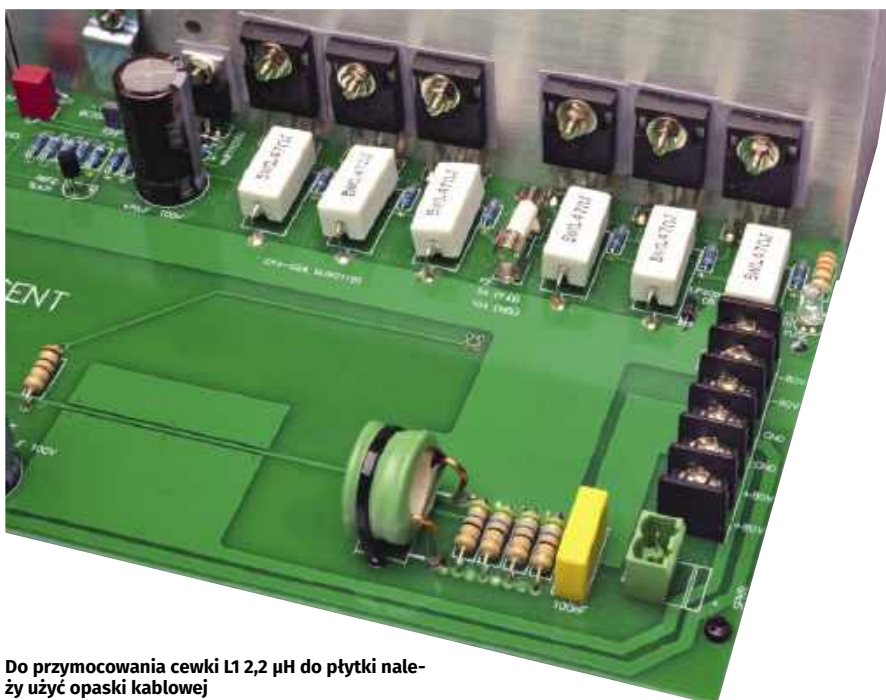
Cewka L1 jest nawinięta przy użyciu około 900 mm emaliowanego drutu miedzianego o średnicy 1,25 mm na plastikowym karkasie. Należy użyć przyrządu do nawijania, jak pokazano na zdjęciach powyżej. Bez tego procedura jest znacznie trudniejsza i istnieje ryzyko uszkodzenia stosunkowo delikatnego karkasu. Karkas należy przymocować do przyrządu, a następnie nawinąć 13,5 zwoja drutu o średnicy 1,25 mm zgodnie z ruchem wskazówek zegara, pozostawiając około 20 mm wolnego miejsca na każdym końcu.

Po zakończeniu, trzeba zabezpieczyć uzwojenie wąskim paskiem taśmy izolacyjnej, a następnie nasunąć 15 mm rurkę termokurczliwą o średnicy 20 mm na karkas i delikatnie ją podgrzać (trzeba uważać, aby nie stopić karkasu). Następnie używamy małego, ostrego nożyka do zeszkrobania emalii z wystających odcinków drutu na całym obwodzie i cynujemy odsłoniętą miedź na końcach, upewniając się, że cyna pokryje drut.

Cewka może być następnie zamontowana na płycie drukowanej. Powinna być ustawiona w pozycji leżącej (pod kątem 90°) zgodnie z rysunkiem. Przymocujemy ją do płytki drukowanej za pomocą opaski kablowej (trytytki). Opaskę układamy wokół obwodu cewki L1, następnie przewlekamy przez otwory w płycie PCB i zaciskamy.

Przygotowanie głównego radiatora

Następnym krokiem jest wywiercenie otworów w radiatorach przy użyciu dostarczonych szablonów (rysunek 7). Ważne jest,



Do przymocowania cewki L1 2,2 μH do płytki należy użyć opaski kablowej

aby dokładnie rozmieścić otwory tak, aby były wyśrodkowane między żeberkami radiatora. W ten sposób lby śrub będą idealnie mieściły się między żebrami.

Przed wywierceniem otworów w radiatorze należy ostrożnie zaznaczyć miejsca otworów za pomocą bardzo ostrego ołówka, a następnie użyć punktaka (lub młotka i gwoźdźcia), aby zaznaczyć środki otworów. Następnie wiercimy otwory o średnicy 3 mm we wszystkich zaznaczonych miejscach.

Początkowo należy użyć małego wiertła pilotującego (np. 1,5 mm), a następnie zwiększamy rozmiar wiertła do 2,5 mm lub 3 mm. Podczas wiercenia otworów należy używać odpowiedniego smaru. Dla aluminium zalecany środkiem smarującym jest nafta, ale okazało się, że olej maszynowy (np. Singer lub 3 w 1) również będzie odpowiedni.

Otwory muszą znajdować się między żebrami, dlatego przed ich wywierceniem należy sprawdzić, czy pozycje otworów są prawidłowe.

Należy pamiętać, że otworów nie należy wiercić za pierwszym podejściem. Podczas wiercenia w aluminium ważne jest, aby regularnie wyjmować wiertło z otworu i usuwać metalowe opiłki. Jeśli się tego nie robi, opiłki aluminium mają nieprzyjemny zwyczaj zakleszczania wiertła i łamania go. Mogą również zarysować powierzchnię radiatora. Za każdym razem przed wznowieniem wiercenia należy ponownie nasmarować otwór i wiertło.

Na tym etapie można wywiercić otwory o średnicy 2,5 mm w dolnej krawędzi radiatora. Będą one gotowe do gwintowania gwintownikiem M3. Należy to zrobić w dwóch miejscach wzdłuż dolnej krawędzi każdego radiatora. Otwory te posłużą



Zbliżenie pokazuje sposób montażu tranzystorów do radiatora

do późniejszego montażu radiatorów do obudowy.

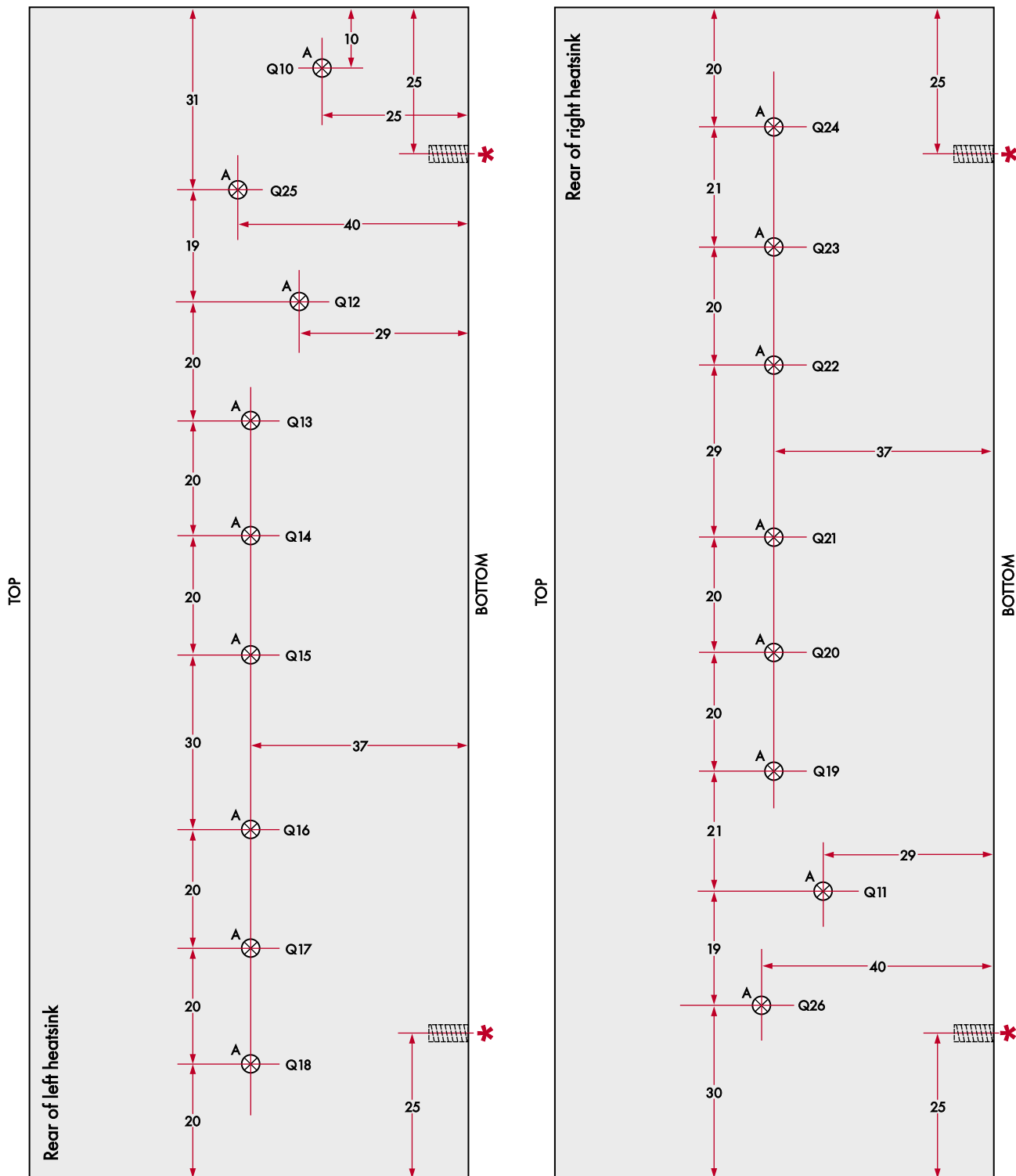
Gwintowanie

Do gwintowania otworów montażowych od spodu potrzebny będzie gwintownik

pośredni M3 (lub początkowy, nie końcowy). Cała sztuka polega na tym, aby robić to delikatnie i powoli. Podczas gwintowania trzeba utrzymywać smar w górze i regularnie wyciągać gwintownik. Umożliwimy w ten sposób usuwanie opiłków z otworu. Przed każdym

wznowieniem pracy należy ponownie nasmarować gwintownik.

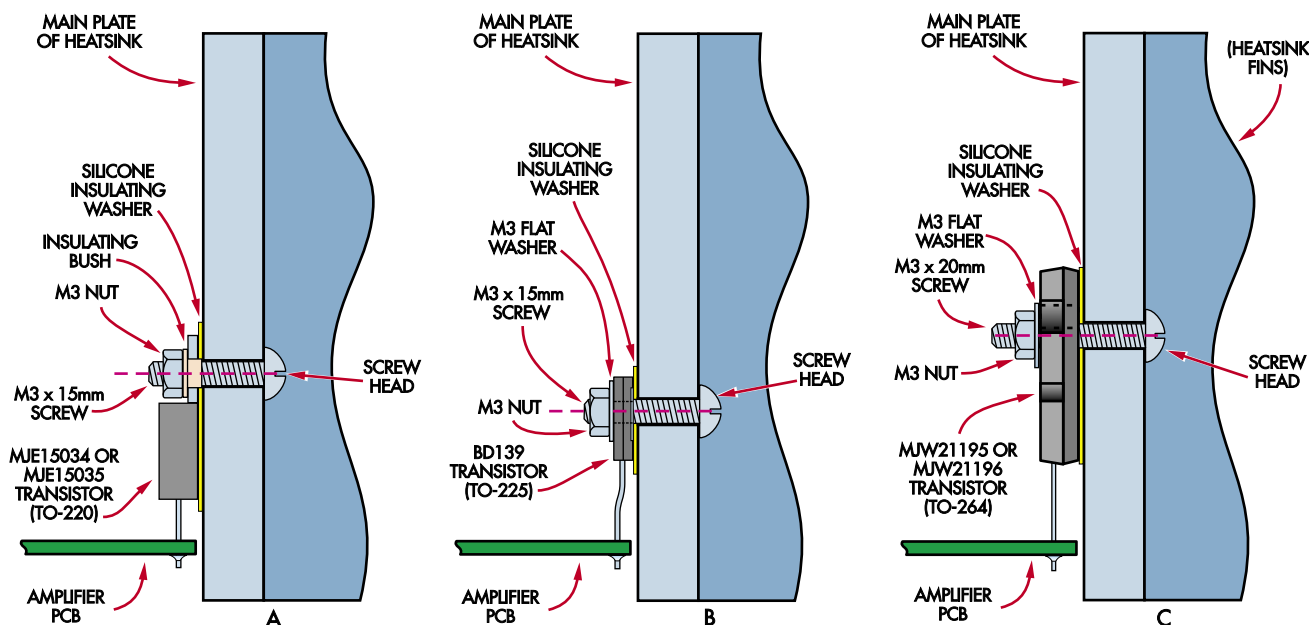
W czasie gwintowania nie należy przykładać nadmiernej siły do gwintownika, gdyż bardzo łatwo można go złamać. Podobnie, jeśli napotkamy jakikolwiek opór należy gwintownik



ALL HOLES 'A' ARE 3mm IN DIAMETER

* DRILL & TAP FOR M3 THREAD IN UNDERSIDE EDGE OF HEATSINK BASEPLATE (2 PER HEATSINK)

Rysunek 7. Trzeba wiercić dwa radiatorzy obok siebie, jak to pokazano na tym rysunku. Można wywiercić otwory montażowe tranzystorów przez radiator wiertłem 3 mm, a następnie zamontować tranzystory za pomocą śrub, nakrętek i podkładek. Spodnia krawędź jest nawiercona do 2,5 mm i nagwintowana do M3 w dwóch miejscach na każdym radiatorze, dzięki czemu można go zamontować do obudowy (wszystkie wymiary w milimetrach)



Rysunek 8. Rysunek pokazuje sposób montażu elementów na radiatorze. Należy zwrócić uwagę na użycie silikonowych podkładek izolacyjnych dla wszystkich dużych tranzystorów (nie ma potrzeby stosowania podkładek mikowych, biorąc pod uwagę, jak rozłożone jest obciążenie cieplne). Konieczne będą również plastikowe tuleje dla tranzystorów TO-220 z catkowicie odstłoniętymi metalowymi radiatorami

delikatnie obróć do przodu i do tyłu, co powinno umożliwić jego wyjęcie. Krótko mówiąc, użycie siły jest zabronione.

Na koniec należy lekko usunąć zadziory na krawędziach otworów za pomocą wiertła o dużo większym rozmiarze i usunąć wszelkie cząstki aluminium lub opiłki. Sprawdzamy, czy obszar wokół otworów jest idealnie gładki, w przeciwnym razie podkładki izolacyjne mogą ulec uszkodzeniu. Pozostało jeszcze dokładne wyszorowanie radiatora wodą z detergentem i pozostawienie go do wyschnięcia.

Montaż końcowy

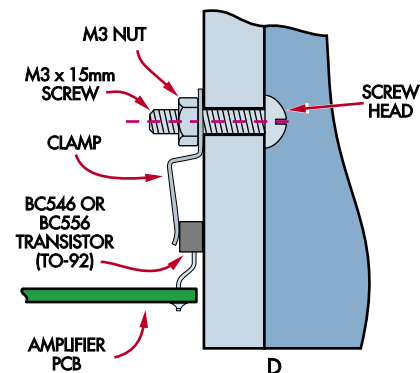
Na **rysunku 8** przedstawiono szczegółowy montaż tranzystorów. Zaczynamy od zamontowania tranzystorów Q13 do Q24, zwracając uwagę, że Q13...Q18 to tranzystory MJW21196, a Q19...Q24 to tranzystory MJW21195. Q13...Q18 są zamontowane na lewym radiatorze, a Q19...24 na prawym radiatorze. Lokalizacje tych elementów pokazano na **rysunku 7** (można również odnieść się do rysunku 6).

Wszystkie te elementy są montowane z silikonową podkładką izolacyjną między każdym tranzystorem a powierzchnią radiatora. Są one mocowane za pomocą śrub M3 x 20 mm włożonych między żebra radiatora oraz płaskiej metalowej podkładki i nakrętki M3 na powierzchni tranzystora. Nie dokręcamy jeszcze śrub, aby było możliwe przesunięcie podkładek izolacyjnych i tranzystorów w celu umożliwienia montażu na płytce drukowanej.

Q12 (MJE15034) na lewym radiatorze i Q11 (MJE15035) na prawym radiatorze wymagają silikonowych podkładek izolacyjnych TO-220 i tulei izolacyjnej włożonej w otwór zakładki tranzystora przed przymocowaniem za pomocą śruby M3 x 15 mm i nakrętki M3. Należy je na razie również pozostawić luźno.

Q10, BD139, montuje się metalową powierzchnią w kierunku radiatora i z podkładką silikonową TO-220 między radiatorem a tranzystorem. Mocujemy go za pomocą śruby M3 x 15 mm i nakrętki M3, ponownie pozostawiając luźne połączenie śrubowe.

Teraz montujemy płytkę drukowaną na sześciu podkładkach dystansowych 9 mm z gwintem M3 i umieszczamy ją na płaskiej powierzchni. Montujemy radiatory



do płytki PCB wkładając wyprowadzenia tranzystorów przez właściwe otwory. Po ich włożeniu, dopychamy je do płytki tak, aby wszystkie cztery elementy dystansowe (i radiator) stykały się z blatem stołu.

Dostosowuje to długości wyprowadzeń tranzystorów i zapewnia, że spód PCB znajduje

REKLAMA

Certyfikat
Underwritten
Laboratories

MA 985-0
E400140
TYPE 1

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

OBWODY DRUKOWANE

Produkcja, Projektowanie, Montaż

Płytki jednowarstwowe Płytki dwuwarstwowe Płytki na podłożu aluminium	Serie maszyn Prototypy Maksymalny wymiar płytki 30 x 300 mm	Dokumentacja technologiczna Dokumentacja konstrukcyjna	Montaż elektroniczny Hotci modelowe produkcyjne
Aktywny kalkulator prototypów na stronie internetowej Polaryzacja i testy SMD Maski, opisy montażowe w różnych kolorach	Płyty szlifierskie FR4 Trasowanie szablony SMD	Krótkie terminy Wykonania super ekspresowe	

Wykaz elementów:

kompletny wzmacniacz 500 W

- 1 zmontowany moduł wzmacniacza 500 W (Silicon Chip, kwiecień i maj 2022 r.)
- 1 zmontowany wskaźnik przesterowania wzmacniacza ustawiony na zasilanie ± 80 V DC (Silicon Chip, marzec 2022)
- 1 zmontowany sterownik wentylatorów i osłona głośników z trzema wentylatorami PWM 120 mm (patrz Silicon Chip, luty 2022)
- 1 zasilacz sieciowy 12 V 15 W z przetwornikiem trybu pracy [Jaycar MP3296, Altronics M8728]

Obudowa

- 1 aluminiowa obudowa rack 3U, 558,80 mm $\times$ 431,80 mm $\times$ 133,35 mm, wykonana z:
- 1 element obudowy RM-14222 Rackmount Chassis Kit (przód, tył i boki) [Digi-Key 377-1392-ND]
- 1 element obudowy TBC-14253 Solid Rackmount Cover (do podstawy) [Digi-Key 377-1396-ND]
- 1 element obudowy TBC-14263 Perforowana pokrywa do montażu w szafie (do pokrywy) [Digi-Key 377-1397-ND]
- 4 nożyki do montażu w podstawie [Jaycar HP0830/HP0832, Altronics H0890]
- 1 kątownik aluminiowy 20 mm $\times$ 20 mm $\times$ 3 mm o długości 400 mm [ze sklepu z narzędziami]
- 1 etykieta na panel przedni 220 mm $\times$ 60 mm

Zasilanie

- 1 toroidalny transformator sieciowy 800 VA z uzwojeniami 2 $\times$ 115 V AC i 2 $\times$ 55 V AC [RS Components 1234050]
- 1 krążek montażowy transformatora toroidalnego (wywiercić otwór o średnicy 8 mm) [RS Components 6719202]
- 2 podkładki neoprenowe do transformatora toroidalnego [RS Components 6719218]
- 1 mostek prostowniczy 35 A 400 V (BR1) [MB354, KPC3504 lub podobne]
- 1 arkusz izolacyjny 208 mm $\times$ 225 mm $\times$ 0,8 mm (Prespahn, Elephantide lub podobny) [Jaycar HG9985]
- 1 plastikowy arkusz o wymiarach 295 mm $\times$ 125 mm $\times$ 3 mm (Perspex, poliwęgiel, PVC, akryl lub podobny)
- 1 złącze wejściowe IEC z bezpiecznikiem [Jaycar PP4004, Altronics P8324]
- 1 osłona izolacyjna złącza sieciowego IEC [Jaycar PM4015]
- 1 przewód zasilający IEC
- 1 bezpiecznik M205 3.15A (F3)
- 1 wyłącznik sieciowy DPDT z czerwoną lampką neonową (S1) [Jaycar SK0982, Altronics S3242B]
- 13-pinowa listwa zaciskowa 6 A z zasilaniem sieciowym [Jaycar HM3194, Altronics P2130A]
- 8 kondensatorów elektrolitycznych 10000 μ F 100 V [Jaycar RU6712 z uchwyty montażowymi]
- 6 rezystorów 15 k Ω 1 W
- 2 diody LED 5 mm (LED4, LED5)
- Żółte izolowane zaciskane oczka 6 mm [Jaycar PT4714, Altronics H2061B]
- 6 żeńskich złączy zaciskanych 6,3 mm z niebieską izolacją [Jaycar PT4625, H1996B]
- 10 opasek kablowych 150 mm
- 7 samoprzylepnych kotew kablowych do montażu panelowego
- zestaw rurek termokurczliwych

Przewody i kable

- 300 mm przewodu uziemiającego 7,5 A lub 10 A (zielony/żółty w paski) [może być odizolowany od trójżyłowego przewodu sieciowego]
- 1 dwużyłowy przewód zasilający w osłonie 7,5 A o długości 1,5 m
- 5 m drutu miedzianego o średnicy 0,5 mm (np. miedziany drut do ram obrazów)
- 400 mm dwużyłowego ekranowanego kabla mikrofonowego (lub jednożyłowego, jeśli używane jest gniazdo wejściowe RCA)
- 2 m czerwonego przewodu połączeniowego 25 A, 2,9 mm² [Jaycar WH3080]
- 2 m czarnego przewodu połączeniowego 25 A, 2,9 mm² [Jaycar WH3082]
- 1 m przewodu ósemkowego, 2,93 mm² na żyłę [Jaycar WB1732]
- 1 m przewodu ósemkowego o przekroju 2,5 mm² na żyłę [Jaycar WB1712]
- 2 m przewodu ósemkowego, 0,76 mm² na żyłę [Jaycar WB1708]
- 1 m przewodu ósemkowego, 0,44 mm² na żyłę [Jaycar WB1704]

Sprzęt, w tym śruby

- 2 wkręty samogwintujące nr 4 $\times$ 6 mm (lub dwa wkręty maszynowe M2 $\times$ 6 mm i dwie nakrętki M2)
- 1 śruba M8 $\times$ 75 mm, nakrętka sześciokątna M8 i podkładka do transformatora [sklep z narzędziami]
- 8 śrub M4 $\times$ 50 mm
- 1 śruba M4 $\times$ 20 mm
- 3 śruby M4 $\times$ 15 mm
- 22 śruby M4 $\times$ 10 mm
- 4 łączniki z gwintem M4
- 39 nakrętek sześciokątnych M4
- 3 podkładki gwiazdkowe M4
- 2 śruby M3 $\times$ 15 mm
- 4 śruby z tłem stożkowym M3 $\times$ 12 mm
- 10 wkrętów maszynowych M3 $\times$ 10 mm
- 11 nylonowych wsporników M3 $\times$ 9 mm
- 2 śruby maszynowe M3 $\times$ 6 mm
- 22 śruby maszynowe M3 $\times$ 5 mm
- 12 nakrętek sześciokątnych M3

Inne części

- 1 przełącznik SPDT 30 A, cewka 12 V (RLY1) [Altronics S4211]
- 3-pinowe żeńskie złącze panelowe XLR [Jaycar PS1054, Altronics P0903] (lub izolowane gniazdo panelowe RCA)
- 1 para wytrzymałych zacisków głośnikowych do montażu panelowego [Jaycar PT0457, Altronics P9257A]
- 1 wtyczka liniowa RCA
- 1 maskownica do montażu panelowego dla diody LED 5 mm [Jaycar SL2610, Altronics Z0220]
- 3 żeńskie złącza zaciskane 6,3 mm z żółtą izolacją [Jaycar PT4725, Altronics H1842A]
- 1 kondensator MKT 560 nF 100 V
- 2 termistory NTC 10 k Ω [Altronics R4112]

się dokładnie 9 mm nad dolną krawędzią radiatora.

Sprawdzamy, czy w każdej pozycji znajduje się prawidłowy tranzystor i regulujemy zespół PCB w poziomie tak, aby każdy z nich wystawał o 1 mm poza bok radiatora. Teraz dokręcamy wszystkie śruby tranzystorów na tyle, aby utrzymać je na miejscu, jednocześnie utrzymując prawidłowo wyrównane podkładki izolacyjne. Tył każdego radiatora powinien przylegać płasko do krawędzi montażowej tranzystora na płytce drukowanej.

Następnym krokiem jest wstępne przyłutowanie wyprowadzeń tranzystorów od góry płytki PCB, a przynajmniej tyłu wyprowadzeń, do których można łatwo uzyskać dostęp od góry. Następnie ostrożnie odwracamy cały zespół do góry nogami i podpieramy przednią krawędź płytki, umieszczając książki lub jakieś inne wsporniki pod płytką. Płytkę drukowaną powinna być utrzymywana pod kątem prostym do radiatora.

Jako tymczasowych podpór można też użyć kilku kartonowych rolek (np. po papierowych ręcznikach papierowych) przyciętych do 63 milimetrów.

Jeśli taka podpora nie zostanie zastosowana, wyprowadzenia zegną się pod ciężarem napierającej konstrukcji i pozostaną w tym stanie po ich przyłutowaniu.

Teraz można przyłutować pozostałe wyprowadzenia tranzystorów i dodać więcej cyny do każdego pola lutowniczego, które tego wymaga. Upewniamy się, że połączenia są odpowiednio dobrej jakości, ponieważ przy pełnej mocy może przez nie przepływać prąd o natężeniu wielu amperów. Po zakończeniu trzeba jeszcze przyciąć nadmiarową długość wyprowadzeń i ponownie obrócić płytkę prawą stroną do góry.

Następnie dokręcamy śruby mocujące tranzystory, zapewniając dobry transfer ciepła między elementami i radiatorem. Muszą być mocno dokręcone, ale nie ma potrzeby sięgania po sprzęt pneumatyczny czy udarowy.

Sprawdzanie izolacji tranzystorów

Przechodzimy teraz do sprawdzenia, czy wszystkie tranzystory są elektrycznie odizolowane od radiatora. W tym celu należy przełączyć multimetr na najwyższy zakres pomiaru rezystancji i zmierzyć rezystancję między powierzchnią montażową radiatora a kolektorami tranzystorów zamontowanych na radiatorze.

W przypadku tranzystorów Q13...Q24 jest to prosta czynność – sprawdzamy czy nie ma zwarcia między każdym z zacisków bezpieczników znajdujących się najbliżej radiatora a samym radiatorem po każdej stronie

wzmacniacza. Dzieje się tak, ponieważ kolektory tranzystorów stopnia wyjściowego są połączone ze sobą i podłączone do odpowiednich bezpieczników. Odczyt powinien być powyżej 10 MΩ a najprawdopodobniej „OL”, ponieważ powinien on być zbyt wysoki dla użycia do pomiaru miernika.

Testowanie zwarć dla tranzystorów Q10 (mnożnik napięcia V_{BE}), Q11 i Q12 będzie wyglądało inaczej. W tym przypadku należy sprawdzić, czy nie ma zwarcia między wyprowadzeniem środkowym (kolektorowym) każdego tranzystora a radiatorem.

Jeśli zwarcie zostanie wykryte, należy odkręcić po kolei każdą śrubę mocującą tranzystor, aż zwarcie zniknie. Następnie wystarczy zlokalizować przyczynę problemu i wymienić uszkodzony lub poprawić wcześniej wadliwie zamontowany tranzystor. Pamiętajmy, aby wymienić podkładkę izolacyjną, jeśli została w jakikolwiek sposób uszkodzona (np. przebita).

Teraz można zamontować tranzystory Q25 (BC546) i Q26 (BC556). Są one utrzymywane za pomocą zacisków tranzystorowych przymocowanych do radiatora za pomocą śrub M3 15 mm i nakrętek.

Teraz nakładamy niewielką ilość pasty termoprzewodzącej na płaską powierzchnię każdego z nich, montujemy zaciski tranzystorów i umieszczamy każdy tranzystor tak, aby zaciski utrzymywały je w miejscu mniej więcej pośrodku korpusu tranzystora. Następnie dokręcamy śruby. Odwracamy PCB do góry nogami, lutujemy i przycinamy nadmiar wyprowadzeń uprzednio przylutowanych tranzystorów.



Po zakończeniu prac montażowych, wzmacniacz 500 W będzie miał wentylatory przymocowane za pomocą metalowego wspornika do podstawy obudowy z tyłu radiatora

Teraz należy usunąć trzy wsporniki dystansowe z krawędzi płyty przylegającej do radiatora. Ta krawędź płytki musi być podtrzymywana tylko przez wyprowadzenia tranzystorów radiatora. Pozwala to uniknąć ryzyka pęknięcia ścieżek PCB i padów wokół tranzystorów zamontowanych na radiatorze z powodu rozszerzalności cieplnej i kurczenia się, gdy zespół nagrzewa się i stygnie.

Za miesiąc

Na tym kończy się montaż modułu wzmacniacza. W następnym odcinku zostanie opisany

zasilacz, sposób uruchomienia i testowania wzmacniacza oraz zostaną podane szczegółowe informacje na temat umieszczenia wzmacniacza w wentylowanej aluminiowej obudowie (pokazanej obok ze zdjętą pokrywą). To rozwiązanie powinno zapewnić utrzymywanie poprawnej temperatury pracy wzmacniacza, nawet w warunkach pełnego obciążenia. ■

John Clarke

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

REKLAMA

ELEKTRONIKA PRAKTYCZNA

tylko Prenumeratory

NOWOCZESNE PRAKTYCZNE IMPULSOWE

ZESTAWY EDUKACYJNE I EWALUACYJNE

przejrzysz i kupisz na www.ulubionykiosk.pl



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

SMD Tweezers – ulepszona pęseta pomiarowa SMD

Projekt pęsety pomiarowej SMD, który był opublikowany w październiku 2021 roku w Silicon Chip cieszył się ogromną popularnością. Trudno się temu dziwić, biorąc pod uwagę, że pęseta była poręczna, kompaktowa, łatwa w użyciu, łatwa do wykonania, a koszt zestawu był rozsądny. Mimo tak wielu zalet, zawsze można zastanowić się, czy można ten projekt ulepszyć? Odpowiedź znajdziesz w artykule.

Pęseta SMD to prosta, ale pomysłowa konstrukcja. Ośmionóżkowy mikrokontroler PIC12F1572 zasilany z baterii pastylkowej służy do sprawdzania rezystorów, diod i kondensatorów, a następnie wyświetlania wyników na małym wyświetlaczu OLED.

Układ PIC12F1572 radzi sobie przyzwoicie, ale oprogramowanie Tweezers zajmuje niemal całą pamięć Flash, pozostawiając jedynie 42 wolne bity. Nie daje to nadziei na jakąkolwiek rozbudowę.

W oryginalnej pęsecie SMD zastosowany był mikrokontroler PIC12F1572. W tym czasie był to najtańszy dostępny układ. Konkurencyjny mógł być jeszcze układ PIC12F1571, ale ma on jeszcze mniejszą pamięć. Do niedawna PIC12F675 i później PIC12F617 były idealnymi mikrokontrolerami do tego zastosowania, ale Microchip wciąż wprowadza nowe wersje swoich produktów o większej wydajności i większej liczbie funkcji w niższych cenach, więc warto śledzić nowości i za nimi podążać.

Układ PIC12F1572 jest bardziej wydajny niż starszy układ PIC12F675, co było wspomniane w artykule zamieszczonym w Silicon Chip w listopadzie 2020 roku (siliconchip.com.au/Article/14648).

Po pewnym czasie pojawił się więc pomysł, aby zmodernizować Tweezers o pewne ulepszenia oprogramowania. Okazało się wówczas,

że do dodania nowych funkcji lub ulepszenia istniejących układ PIC12F1572 nie ma wystarczającej wolnej pamięci. W tym celu konieczne by było przejście na najnowszą generację mikrokontrolera PIC. Kiedy więc pojawiła się nowa rodzina układów PIC, została przeprowadzona analiza, jakie funkcjonalności można by było dodać przy ich zastosowaniu.

Nowy PIC

Oryginalna pęseta SMD nie wykorzystuje żadnych egzotycznych układów peryferyjnych mikrokontrolera. Przetwornik analogowo-cyfrowy (ADC) i zegar watchdog, których wymaga oprogramowanie, znajdują się w większości mikrokontrolerów PIC. Tryb uśpienia o niskim poborze mocy jest również dość standardowy i jest niezbędny do pracy w trybie gotowości, gdy jest zasilany z małej baterijki. Pozwala to na pozostawienie pęsety w stanie bezczynności, ale gotowej do pracy w każdej chwili.

Interfejs I²C do wyświetlacza OLED jest emulowany programowo poprzez przełączanie pinów GPIO (wejście/wyjście ogólnego przeznaczenia), techniką często znaną jako „bitbanging”. Wszystko to oznacza, że prawie każdy 8-pinowy mikrokontroler z większą pamięcią programu może być użyty w pęsetach SMD.

Pod koniec 2021 roku, już po opracowaniu oryginalnej pęsety autor dowiedział się o serii mikrokontrolerów PIC16F152xx. Seria

ta obejmuje układy mające od ośmiu do czterdziestu nóżek. Chociaż seria ta ma funkcje, które są skromne jak na obecne standardy, to są one nadal bardziej wydajne niż starsze mikrokontrolery, takie jak PIC12F675.

PIC16F15213 i PIC16F15214 to 8-nóżkowe układy z tej serii. Są one tańsze niż PIC12F1572, chociaż obserwowane od pewnego czasu niedobory części na rynku oznaczają, że ich dostępność może być niska.

Ważne jest, że mikrokontroler PIC16F15214 jest dostępny w obudowie SOIC i ma dwa razy więcej pamięci Flash niż PIC12F1572.

Rozeznanie wykazało, że PIC16F15214 jest zarówno tańszy, jak i w pełni zdolny do zastąpienia PIC12F1572 jako kontrolera dla pęsety SMD, a jednocześnie ma większą pamięć programu potrzebną do dodania nowych funkcji.

Pęseta 2.0

W nowej wersji pęsety zostały wprowadzone trzy aktualizacje.

Po pierwsze, został rozszerzony zakres pomiaru pojemności. Obecnie można mierzyć zarówno większe, jak i mniejsze pojemności.

Po drugie, została dodana procedura kalibracji i konfiguracji.

Wreszcie, została poprawiona użyteczność dla osób leworęcznych (lub tych, którzy w przymym ręku trzymają na przykład lutownicę), umożliwiając obrót ekranu o 180°.

Cechy i specyfikacja:

- Wykorzystuje identyczny sprzęt jak oryginalny Tweezers (październik 2021, siliconchip.com.au/Article/15057) z wyjątkiem mikrokontrolera PIC
- Identyfikuje typ elementu (rezystor, kondensator, dioda lub LED) i mierzy wartości krytyczne
- Rezystory: wartość od 10 Ω do 1 M Ω
- Diody: napięcie przewodzenia do około 3 V
- Kondensatory: wartość od (około) 10 pF do 150 μ F
- Napięcie baterii bez podłączenia
- Niski pobór mocy w trybie uśpienia w stanie beczynności pozwala uniknąć konieczności stosowania przycisku on/off
- Natychmiastowe wybudzanie przez zetknięcie końcówek sondy
- Opcja wyboru wyświetlacza dla osoby lewo- lub praworęcznej
- Kalibracja rezystancji wewnętrznej i kontaktowej

Wszystkie te ulepszenia zostały wprowadzone w oprogramowaniu, więc poza zmianą Mikrokontrolera IC12F1572 na PIC16F15214, sprzęt jest identyczny, a ogólne działanie jest takie samo.

Szczegóły układu

Na **rysunku 1** został pokazany schemat, który różni się od poprzedniego jedynie układem IC1. Wszystkie komunikaty i wyniki pomiarów są wyświetlane na małym wyświetlaczu OLED dołączonym do złącza CON2.

Układ IC1 steruje pinami RA5 (pin 2, IOTOP) i RA4 (pin 3, IOBOT) w stanie wysokim i niskim i mierzy napięcie występujące na pinie 5 za pomocą wewnętrznego przetwornika analogowo-cyfrowego. Może na przykład określić rezystancję rezystora podłączonego między punktami CON+ i CON-, korzystając z równania dzielnika napięcia.

W wyniku pomiarów diod będzie wyświetlane napięcia przewodzenia i napięcie wsteczne występujące pomiędzy CON+ i CON–, gdy mikrokontroler przyłoży napięcie. Następnie określa on orientację diody, pokazując jej polaryzację i napięcie przewodzenia.

Kondensatory są najpierw ładowane przez ustawienie IOTOP w stan wysoki i IOBOT w stan niski. Następnie wyprowadzenie IOTOP jest ustawiane w stan niski i badana jest szybkość rozładowania określająca pojemność badanego elementu. Pęseta może nawet mierzyć własne napięcie zasilania, odczytując napięcie wewnętrznego napięcia odniesienia 1,024 V względem napięcia zasilania.

Funkcje te są już zaimplementowana w oryginalnym programie pęsety, a szczegóły zostały opisane w z oryginalnym artykule Tweezers z październik 2021 r., zamieszczonym w Silicon Chip. Zawiera on opis korzystania z nowych funkcji.

Ulepszenia

Górna granica zakresu pojemności była ograniczona przez użycie 8-bitowego licznika do pomiaru czasu rozładowania. Nowy

mikrokontroler pozwala na korzystanie z licznika 16-bitowego.

Teoretycznie rozszerza to zakres 256-krotnie, ale w praktyce wykorzystanie całego zakresu nie jest możliwe. Górna granica wynosi obecnie około 150 μF , co odpowiada około 12 bitom lub 16-krotnemu zwiększeniu zakresu.

Pierwszym powodem jest to, że wyższe wartości przepełniłyby wymagane 32-bitowe obliczenia matematyczne. Drugim jest to, że czas potrzebny do naładowania i rozładowania większego kondensatora staje się nieracjonalnie długi, rzędu kilku sekund między odczytami. Jedynym sposobem na pokonaniu tego problemu byłaby zmiana szeregowego rezystora testowego, co wpłynęłoby również na inne odczyty.

Stosunkowo wysoka wartość szeregowego rezystora testowego oznacza również, że odczyty pojemności mogą być zniekształcone przez prąd upływu. Upływ jest zwykle wyższy w kondensatorach o wyższej pojemności, zwłaszcza w kondensatorach elektrolitycznych, więc dokładność i użyteczność tych wyższych zakresów jest mniejsza niż teoretycznie możliwa.

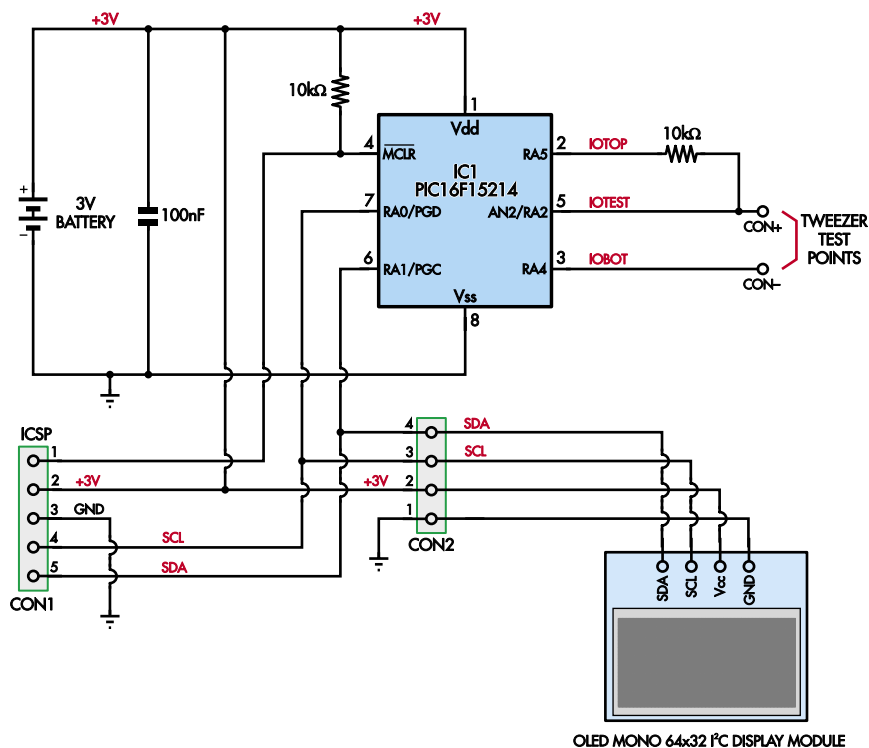
Kondensatory o większych pojemnościach mogą być mierzone, ale odczyt pojawi się z krótkim opóźnieniem, a dokładność nie będzie tak dobra, jak w przypadku niższych pojemności.

Pomiary niskiej pojemności

Wartości niższe niż 1 nF są mierzone zupełnie inną metodą. Jest ona tak czuła, że może mierzyć pojemność dotyku dłoni, rzędu pikofaradów.

Jest to metoda współdzielonego wykrywania pojemności i została zastosowana do wykrywania dotknięć palców w ATtiny816 Breakout Board ze stycznia 2019 r. (siliconchip.com.au/Article/11372).

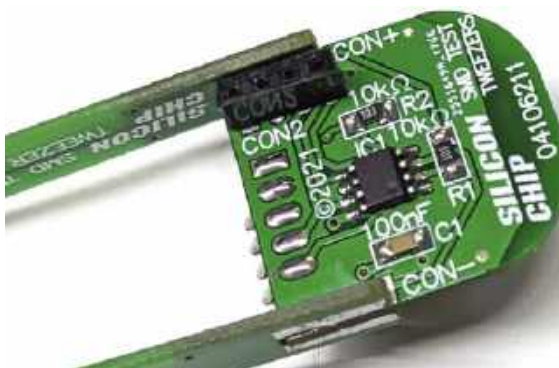
Pomiar polega na porównaniu względnej wielkości dwóch kondensatorów poprzez początkowe naładowanie jednego i rozładowanie drugiego, jak pokazano na **rysunku 2**. Gdy kondensatory są połączone, cały ładunek jest dzielony między nimi proporcjonalnie do ich pojemności. Stosunek napięcia



Rysunek 1. Układ zaktualizowanej pęsety jest praktycznie taki sam jak w starej wersji, z wyjątkiem tego, że IC1 to teraz PIC16F15214. Może wykonywać wszystkie swoje testy, przykładając różne napięcia do pinów IOTOP i IOBOT oraz testując napięcie na pinie IOTEST



Kilku konstruktorów miało trudności ze znalezieniem mosiężnych pasków, które były zalecane dla oryginalnych pęset. Substytutem są typowe piny łączówek kołkowych i gdy znajdują się w plastikowych osłonkach można je łatwo wyrównać podczas lutowania. Dostępne są nawet wersje pozłacane



Aby moduł wyświetlacza OLED mógł być wyjmowany podczas uruchamiania prototypu, zostało zastosowane niskoprofilowe gniazdo kołkowe (Altronics P5398). Ponadto dostępne stały się piny używane do programowania mikrokontrolera. Wymagało to również wycięcia pinów gniazda na spodzie wyświetlacza i usunięcia plastikowego bloku dystansowego. Alternatywą jest po prostu przylutowanie wyświetlacza bezpośrednio do głównej płytki drukowanej

początkowego i końcowego odnosi się bezpośrednio do stosunku pojemności. Teoria i matematyka są wyjaśnione bardziej szczegółowo w artykule ATtiny816 Breakout Board.

W przypadku nowej pęsety, kondensator podłączony do CON+ i CON- jest ładowany przez rezystor 10 kΩ. Drugi kondensator jest w rzeczywistości małym kondensatorem wewnętrznym, który służy do próbkowania i utrzymywania napięć odczytywanych przez przetwornik ADC mikrokontrolera. Kondensator ten ma nominalną pojemność 5 pF i jest rozładowywany przez próbkowanie kanału ADC połączonego do masy.

Na szczęście przetwornik ADC mikrokontrolera ma możliwość wyboru wewnętrznego połączenia z masą, więc nie wymaga to dodatkowego pinu. Zewnętrzny kondensator jest odłączany od rezystora, a dwa kondensatory są łączone poprzez odczyt ADC z zewnętrznego kondensatora.

Do obliczenia wartości zewnętrznego kondensatora używane jest równanie podobne do równania dzielnika napięcia w odniesieniu do wyniku przetwornika ADC. Kondensatory dzielą ładunek analogicznie jak rezystory dzielą napięcie w rezystorowym dzielniku napięciowym.

Oprogramowanie dokonuje również drobnych korekt w celu uwzględnienia niektórych pojemności rozproszonych występujących w układzie, które są znaczące przy tych wielkościach. W celu dostrojenia tych odczytów zostały wykonane testy na rzeczywistych kondensatorach w zakresie pikofaradów.

Dolna granica jest przyjęta dość arbitralnie. Została wybrana w celu uniknięcia wykrycia przez pęsetę pojemności rozproszonej jako mierzonego elementu, co mogłoby spowodować nieprawidłowe wyłączenie zasilania. Przy tej skali, nawet sposób trzymania pęsety może znacząco zmienić odczyt.

Gdy odczyt ADC zbliża się do górnego limitu dla większych pojemności, rozdzielczość jest słaba około 1 nF, a kroki rosną nawet do 100 pF. Dlatego odczyty przy użyciu tej metody są zawsze wyświetlane jako pF. Inne metody są używane do pomiarów w nF lub μF.

Szczegółowy proces kalibracji i konfiguracji zostanie przedstawiony w dalszej części.

Budowa

Pęseta pomiarowa SMD jest zbudowana z trzech płytek PCB, z których główna ma kod 04106211 i wymiary 28 mm × 26 mm. Podczas budowy należy zapoznać się z opisami na PCB (rysunki 3 i 4).

Główna płytkę drukowaną nie jest trudna do zbudowania, nawet jeśli wszystkie części są montowane powierzchniowo. Do montażu zalecana jest lutownica z cienkim grot, lupa, pasta topnikowa, plecionka lutownicza i pęsetka.

Podczas montażu małej płytki drukowanej warto umieścić ją w małym imadle. Innym rozwiązaniem może być użycie kleju, na przykład takiego jak Blu-Tack.

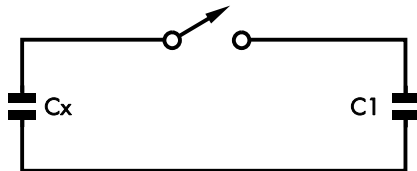
Jeśli to możliwe, należy również ustawić wyciąg oparów, który zabezpieczy nas przed

wydychaniem szkodliwego dymu powstającego podczas pracy z topnikiem. Ostatecznym rozwiązaniem powinna być praca w pobliżu otwartego okna a nawet na zewnątrz. Przydatna jest również gąbka do czyszczenia końcówek.

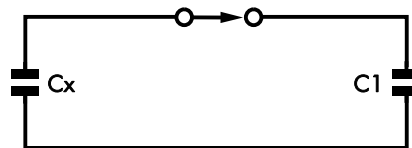
Zaczynamy od nałożenia topnika na górne pady PCB dla IC1 i trzech elementów pasywnych. Następnie umieszczamy IC1 na miejscu za pomocą pęsety, upewniając się, że kropka lub skos pinu 1 znajduje się w kierunku zaokrąglonego końca płytki. Wyrównujemy element na padach, czyścimy grot lutownicy, nakładamy świeżą cynę i lutujemy jedno wprowadzenie w lokalizacji docelowej.

W razie potrzeby układ scalony należy lekko skorygować, aby upewnić się, że przylega płasko do płytki drukowanej i jest wyśrodkowany z padami. Następnie lutujemy pozostałe piny, czyszcząc grot lutownicy i dodając cynę w razie potrzeby.

Do ewentualnego usuwania mostków lutowniczych używamy plecionki lutowniczej poprzez dodanie większej ilości topnika, a następnie docisnięcie plecionki do nadmiaru cyny za pomocą lutownicy. Gdy cyna zostanie wchłonięta ostrożnie odciągamy zarówno grot, jak i opłot.



$$\begin{aligned} V_{\text{across } Cx} &= V_{\text{initial}}, C = Cx \\ \Rightarrow Q_{\text{across } Cx} &= V_{\text{initial}} \times Cx \\ V_{\text{across } C1} &= 0, C = C1 \\ \Rightarrow Q_{\text{across } C1} &= 0 \text{ (because } V_{\text{across } C1} = 0) \end{aligned}$$



$$\begin{aligned} C &= Cx + C1 \text{ (parallel capacitors)} \\ Q &= Q_{\text{across } Cx} + Q_{\text{across } C1} = V_{\text{initial}} \times Cx \text{ (conservation of charge)} \\ \text{measure } V_{\text{final}} \\ \Rightarrow Cx &= (V_{\text{final}} \times C1) \div (V_{\text{initial}} - V_{\text{final}}) \end{aligned}$$

Rysunek 2. Cx jest testowanym urządzeniem (DUT) podłączonym do sond pęsety, natomiast C1 jest kondensatorem próbkującym i podtrzymującym wewnętrzny ADC mikrokontrolera. Kondensatory są połączone poprzez próbkowanie Cx za pomocą ADC. Jeśli wartość Cx jest równa C1, wynikowe napięcie jest równe połowie napięcia początkowego. Jest to analogiczne do rezystancyjnego dzielnika napięcia, a wzory są takie same, z ładunkiem kondensatora zastępującym napięcie na rezystorach

Pozostałe trzy elementy nie są spolaryzowane, więc ich orientacja nie ma znaczenia. Kondensator znajduje się w pobliżu CON-, natomiast dwa identyczne rezystory oznajdują się w pobliżu układu IC1 na jego drugim końcu i boku.

Podobną technikę stosujemy do IC1. Lutujemy jedną nóżkę, korygujemy położenie elementu, a następnie lutujemy drugą nóżkę. Można także poprawić wyprowadzenia, jeśli połączenie nie wygląda prawidłowo. Powinno być gładkie i błyszczące. Aby każdym etapie poprawić jakość lutowania, można dodać więcej topnika.

Następnie lutujemy pojedynczy element do tylnej części płytki drukowanej. Aby wyrównać dwa zewnętrzne piny do ich padów, należy wyśrodkować uchwyt ogniwa. Jeśli lutownica ma regulację temperatury, można zwiększyć ją podczas lutowania większego elementu. Należy również upewnić się, że szerszy otwór w uchwycie baterii jest skierowany w stronę zaokrąglonej krawędzi płytki drukowanej. Ułatwi to dostęp do baterii w celu jej zamontowania i wyjęcia.

Tak jak poprzednio, nakładamy topnik, lutujemy jedną nóżkę i korygujemy położenie. Następnie lutujemy drugą nóżkę. W przypadku znacznie większych padów pomocne może być nałożenie dodatkowej cyny bezpośrednio na pad. Powstanie dzięki temu zaokrąglenie, które można zobaczyć na zdjęciach.

Po zamontowaniu elementów montowanych powierzchniowo można wyczyścić płytkę PCB za pomocą środka do czyszczenia topnika. Dobrymi środkami alternatywnymi do czyszczenia wielu topników jest spirytus etylowany lub alkohol izopropylowy. Są również dostępne uniwersalne środki do czyszczenia topników (działają one zwykle lepiej niż zwykły alkohol).

Przed przejściem do kolejnych kroków należy upewnić się, że łatwopalny rozpuszczalnik całkowicie wyparował.

Wykaz elementów:

- 1 dwustronna płytka drukowana, kod – 04106211, 28 mm × 26 mm (główna płytka drukowana)
 - 2 dwustronne płytki PCB – kod 04106212, 100 mm × 8 mm (ramiona pęsety)
 - 1 mikrokontroler 8-bitowy PIC16F15214-I/SN lub PIC16F15214-E/SN zaprogramowany kodem 0410621B.HEX, SOIC-8 (IC1)*
 - 1 0,49-calowy moduł wyświetlacza OLED 64×32 (Silicon Chip Online Shop nr kat. SC5602)
 - 1 gniazdo baterii pastylkowej do montażu powierzchniowego (BAT1) [Digi-key BAT-HLD-001-ND, Mouser 712-BAT-HLD-001 lub podobny]
 - 1 pastylkowa bateria litowa CR2032 lub CR2025
 - 1 5-stykowe złącze kątowe męskie (CON1 – opcjonalne, do programowania układu IC1 w układzie)
 - 1 kondensator ceramiczny 100 nF SMD 50 V X7R, rozmiar 3216/M1206 [Altronics R9935]
 - 2 rezystory SMD 10 kΩ 1%, rozmiar 3216/M1206 [Altronics R8188]
 - 2 krótkie kawałki cienkiej (np. 1 mm) blachy miedzianej 15 × 2 mm na końcówki (opcjonalnie) LUB
 - 2 pozłacane piny do końcówek (patrz tekst)*
 - 1 przezroczysta rurka termokurczliwa o długości 40 mm i średnicy 30 mm (opcjonalnie)
 - 2 odcinki 100 mm rurki termokurczliwej o średnicy 10 mm (opcjonalnie)
 - 1 4-pinowa niskoprofilowa listwa kołkowa żeńska (opcjonalna, dla CON2, może być wycięta z Altronics P5398)*
- * te części zostały zmienione w porównaniu do oryginalnej pęsety.
- Kompletny zestaw (SC5934) jest dostępny na stronie siliconchip.com.au/Shop/20/5934.

Programowanie IC1

Do projektu można użyć wstępnie zaprogramowanego mikrokontrolera PIC, albo czystego układu IC1, który trzeba będzie zaprogramować wsadem odpowiednim dla tego projektu. Ten krok można pominąć, jeśli mikrokontroler został już zaprogramowany.

Jak już wiemy, układ PIC16F15214 jest znacznie nowszy niż PIC12F1572, więc potrzebny będzie nowy programator i nowa wersja MPLAB X IPE (zintegrowane środowisko programowania) firmy Microchip. Można je pobrać jako część MPLAB X IDE ze strony siliconchip.com.au/link/abd2.

Autor przetestował wersje v5.40 i nowsze. Konieczne może być również pobranie pakietu DFP (device family pack). Można go pobrać z poziomu IDE, a następnie IPE wykryje, że DFP jest zainstalowany. Należy szukać rodziny PIC16F152xx.

Będzie także potrzebny najnowszy programator, taki jak Snap lub PICKit 4, ponieważ starszy PICKit 3 nie obsługuje tych mikrokontrolerów.

Programator dołączamy do płytki drukowanej przez CON1, parując ze sobą strzałki oznaczające pin 1. Złącze można przylutować, ale praktyka pokazuje, że zwykle wystarczające

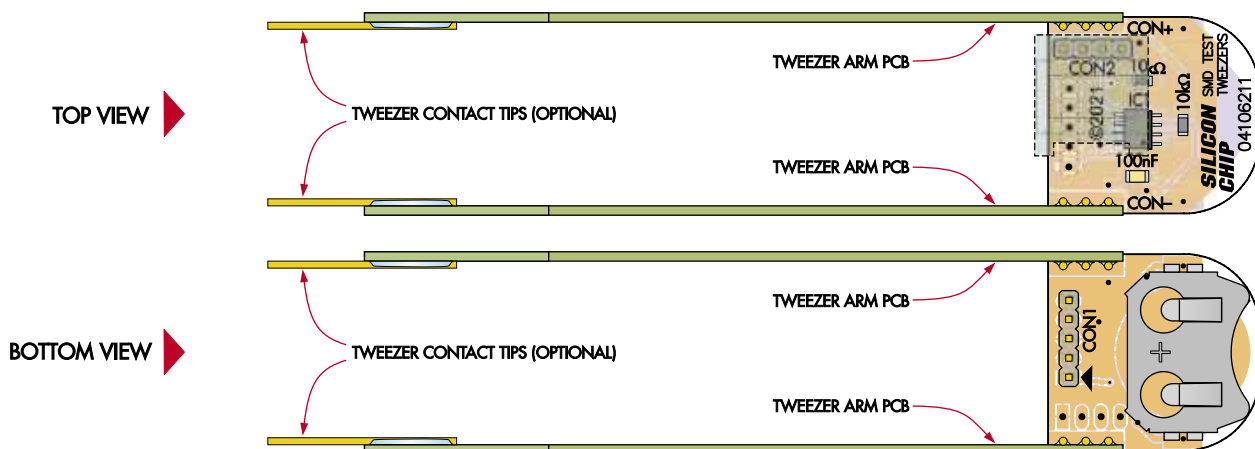
jest przytrzymanie złącza i mocne dociśnięcie go do padów w celu uzyskania kontaktu.

Wybieramy układ PIC16F15214 i otwieramy plik 0410621B.HEX. Aby umożliwić programatorowi podłączenie zasilania, może być konieczna zmiana ustawień. Następnie klikamy „Program” i sprawdzamy, czy układ programuje się i weryfikuje poprawnie.

Ramiona pęsety

Następnie należy przymocować płytki PCB z dwoma ramionami, ponieważ moduł OLED zakrywa większość głównej płytki drukowanej, ograniczając dostęp.

Aby nadać ramionom pęsety delikatniejsze końcówki niż tylko gołe płytki PCB w pierwszej wersji pęsety były zastosowane małe kawałki mosiężnego paska. W przypadku trudności z ich znalezieniem można zastosować rozwiązania alternatywne. Wskazane jest, aby pęseta miała pozłacane końcówki! W tej roli dobrze sprawdzą się szpilki listew kołkowych. Inną ich zaletą jest to, że dobrze pasują do płytek prototypowych i przewodów połączeniowych, dzięki czemu bardzo łatwo jest dołączyć pęsetę do innych elementów w celu odczytu wyników pomiaru bez użycia rąk.



Rysunek 3. Konstrukcja pęsety jest taka sama jak poprzedniej (październik 2021), z wyjątkiem mikrokontrolera (IC1), który jest kompatybilny z poprzednikiem, ale ma większą pamięć. D montażu nie ma wielu elementów, trzeba tylko upewnić się, że IC1 jest prawidłowo zorientowany



Można zalecić dopasowanie ramion mniej więcej do krawędzi płytki drukowanej, ale lekko przechylonych do wewnątrz, z około 15 mm odstępem na końcach końcówek. Podobnie jak w przypadku elementów SMD, należy z grubsza przylutować ramiona na miejscu i dostosować je do własnych upodobań.

Preferowany jest montaż ramion z napisem i główną ścieżką styku biegnącą od wewnątrz. Pomaga to ekranować i izolować ścieżkę od kontaktu zewnętrznego lub pojemności pasożytniczych.

Przed rozpoczęciem korzystania z pęsety należy przetestować działanie i nacisk, gdy ramiona są ustawione. Jeśli wynik takiego testu będzie zadowalający, można ramiona pęsety zamocować nakładając dużą ilość cyny po obu stronach głównej płytki drukowanej.

Aby dopasować końcówki, można użyć np. 7-pinowego gniazda kołkowego (co pozwoli uzyskać 15 milimetrový rozstaw ramion) i umieścić piny pęsety w plastikowym gnieździe, po czym przylutować końcówki ramion do krótkich końcówek pinów. Użycie uchwytu sprawi, że piny będą równoległe i równe.

Gdy zostanie osiągnięta zadowalająca położenie końcówek, ponownie nakładamy dużą ilość cyny, a następnie ostrożnie i równomiernie wyciągamy ramiona i ich końcówki z plastikowego uchwytu. W tej sytuacji przydadzą się spiczaste szczytce.

Wyświetlacz OLED

Ostatnim krokiem jest zamontowanie modułu wyświetlacza OLED. Moduł można przylutować bezpośrednio do głównej płytki drukowanej. Z uwagi na to, że konieczne było przeprowadzenie rozszerzonych testów dla tej nowej wersji pęsety, zostało zastosowane niskoprofilowe gniazdo umożliwiające usuwanie wyświetlacza. Jest to konieczne, ponieważ piny programowania są również używane do połączenia z wyświetlaczem OLED.

Do wielokrotnego programowania mikrokontrolera PIC w trakcie prowadzenia

testów został zastosowany programator Programming Helper z czerwca 2021 roku (siliconchip.com.au/Article/14889). Użyte zostało niskoprofilowe (5 mm wysokości) gniazdo słuchawkowe pozwalające zachować kompaktowość urządzenia. Można to zaobserwować na zdjęciach.

Taka metoda montażu była stosowana w celu umożliwienia prowadzenia testów, w ogólnym przypadku zalecany jest jednak montaż bezpośredni. Wyjątkiem będą ci użytkownicy, którzy rozważają opracowanie własnego oprogramowania układowego.

Jeśli gniazdo wyświetlacza OLED nie jest podłączone, należy je teraz przylutować pod kątem prostym do płytki drukowanej modułu. Następnie montujemy wyświetlacz na płytce drukowanej. Może okazać się, że tył modułu wyświetlacza OLED dotyka IC1. W takim przypadku na czas lutowania należy rozdzielić te dwa elementy za pomocą masy mocującej Blu-Tack lub tekturowej podkładki. Długie szpilki powinny wystawać z tyłu płytki drukowanej. Można je ostrożnie przyciąć ostrym nożykiem do cięcia bocznego.

Testowanie

Do wykonania testów konieczne jest umieszczenie baterii pastylkowej CR2025 lub CR2032 3 V w koszyku, zwracając uwagę na polaryzację pojemnika na baterię. Po około sekundzie na wyświetlaczu powinien być wyświetlony napis „R HAND” (ekran 1). Jeśli tak nie jest, należy sprawdzić lutowanie i upewnić się czy napięcie 3 V występuje na stykach gniazda modułu wyświetlacza lub stykach 2 i 3 CON1. Jeśli napięcie 3 V nie zostanie w tych punktach wykryte, może to świadczyć o problemach z baterią – brak kontaktu lub zwarcie. Trzeba pamiętać o sprawdzeniu odwrotnej strony płytki drukowanej, ponieważ koszyk baterii, ramiona i gniazdo OLED znajdują się bardzo blisko siebie.

Przed przystąpieniem do korzystania z pęsety pomiarowej SMD warto przejść przez procedurę kalibracji opisaną w ramce.

Działanie

Po zakończeniu kalibracji i konfiguracji rozpocznie się normalna praca. Na wyświetlaczu powinien pojawić się wskaźnik napięcia baterii poprzedzony literą B. Po pięciu sekundach pęseta przejdzie w tryb uśpienia, z którego można ją wybudzić poprzez zwarcie końcówek.

W tym momencie nowa pęseta działa tak samo jak starsza wersja, z wyjątkiem rozszerzonego zakresu pomiaru pojemności.

Jeśli końcówki zostaną zwarte w celu zmierzenia rezystancji zwarcia i wszystko działa poprawnie, na wyświetlaczu powinno pojawić się skaczące wskazanie pomiędzy 0 Ω a 1 Ω .

Nowa pęseta pobiera prąd około 4 mA podczas pracy i 5 μ A w stanie uśpienia. Żywotność baterii będzie więc zależeć głównie od tego, jak często pęseta jest używana.

Finalizacja

Można rozważyć dodanie do pęsety koszulki termokurczliwej, która zapewniłaby dodatkową ochronę i zapobiegła wyjęciu baterii. **Koszulki termokurczliwe 10 mm można nałożyć na ramiona, przedstawiając odsłonięte tylko końcówki. Przed obkurczeniem koszulek opalarką należy ją mocno docisnąć do głównej płytki drukowanej.**

Szersza koszulka termokurczliwa pasuje do głównej płytki drukowanej i powinna wystawać na tyle, aby uniemożliwić usunięcie ogniwa. Oczywiście, aby wymienić ogniwo, folię należy usunąć i wymienić.

Należy również uważać, aby nie obkurczyć dużej rurki termokurczliwej zbyt mocno wokół wyświetlacza, ponieważ jego szklany ekran może być kruchy. Gorące powietrze powinno być skierowane wzdłuż krawędzi, aby uniknąć nagrzania wyświetlacza i baterii.

Dostępność i aktualizacja zestawu

Nowa sonda jest w zasadzie niemal identyczna jak pierwsza jej wersja. Konieczne jest jednak zaktualizowanie oprogramowania



Do uzyskania solidnego zaokrąglenia pomocne jest nałożenie dodatkowej cyny bezpośrednio na podkładkę ramion pęsety

Końcówki mogą wyglądać nieco dziwnie, ale gdy ramiona są ściśnięte, aby je potączyć, stają się równoległe w odległości, w jakiej zwykle są używane (wystarczająco szerokie, aby pomieścić typowy element SMD)



Ekran 1. Pierwszy ekran po włączeniu pęsety to komunikat HAND, zorientowane zgodnie z ustawieniem. Pozostawienie otwartych końcówek oznacza obsługę prawą ręką



Ekran 2. W różnych momentach podczas konfiguracji może zostać wyświetlone polecenie zwolnienia pęsety („RELEASE”) poprzez otwarcie jej końcówek. Zapobiega to przypadkowemu wprowadzeniu wielu ustawień



Ekran 3. Pierwsze polecenie dotyczy zakończenia procesu kalibracji i towarzyszy mu nominalny pięciosekundowy licznik czasu. Jeśli końcówki pozostaną otwarte w tym czasie, kalibracja zostanie pominięta

do najnowszej wersji. Zestaw nadal zawiera wspomnianą wcześniej rurkę termokurczliwą. W nowej wersji natomiast do końcówek zostały dodane dwie pozłacane szpilki.

Koszt zestawu jest na tyle niski, że warto jednak zbudować nową pęsetę całkowicie od początku. Jeśli jednak ma być tylko zmodernizowana poprzednia wersja, konieczna będzie aktualizacja oprogramowania. Mikrokontroler z aktualnym oprogramowaniem można zamówić.

Ewentualna decyzja o zamianie mikrokontrolera z najnowszym oprogramowaniem wiąże się z umiejętnością usuwania i montowania układów SMD. Najłatwiej jest to zrobić za pomocą stacji gorącego powietrza, chociaż jest to również możliwe za pomocą zwykłej lutownicy.

Kolejne ulepszenia

Pęseta SMD jest nieco ograniczona przez posiadanie tylko jednego rezystora do przyłożenia napięcia do elementów, co z kolei jest ograniczone przez 8-pinowy układ PIC. Rezystor 10 kΩ ogranicza przyłożony prąd do około 300 μA, a to oznacza, że napięcia przewodzenia diod są znacznie niższe niż oczekiwano, a diody LED nie świecą bardzo jasno. Rozważany jest więc bardziej skomplikowany projekt pęsety SMD, w którym znajdzie się chip z, powiedzmy, 14 pinami. Mogłoby to pozwolić

na dodanie nowych trybów testowych i ulepszenie istniejących.

Dodatkowe piny oznaczają, że można używać wielu rezystorów ograniczających prąd, a tym samym możliwy byłby wybór prądu testowego. Oprócz rozszerzenia zakresu, dodatkowe rezystory testowe poprawią również ogólną dokładność.

Konfiguracja i kalibracja

Dokładność nowoczesnych rezystorów do montażu powierzchniowego jest doskonała. Pęseta SMD będzie w stanie rozróżnić rezystory wystarczająco dobrze dla większości zastosowań. Mimo to, dodatkowa przestrzeń programowa dostępna w mikrokontrolerze PIC16F15214 stwarza miejsce na dołożenie kilku procedur pozwalających na dodanie niektórych ustawień i stałych kalibracyjnych.

Jak się okazało, autorzy projektu będący pracownikami Silicon Chip są praworęczni. W trakcie prac zauważono, że został przeoczony istotny aspekt, który prawdopodobnie sprawia, że oryginalna pęseta jest bardzo trudna w użyciu dla osób leworęcznych. Pierwszym nowym ustawieniem jest więc opcja odwrócenia wyświetlacza tak, aby był czytelny, gdy pęseta jest trzymana w lewej ręce.

Istnieje również opcja ustawienia wartości rezystora szeregowego o nominalnej

rezystancji 10 kΩ między pinami 2 i 5 układu IC1. Zamiast próbować zmierzyć jego wartość, zalecane jest przetestowanie zewnętrznego elementu o znanej wartości i przeprowadzenie kalibracji, aż pęseta zmierzy ją poprawnie.

Oporność szeregową jest po prostu dostosowywana proporcjonalnie do żądanej zmiany obliczonej rezystancji. Na przykład, jeśli wyświetlana rezystancja rezystora testowego jest o 1% za niska, należy zwiększyć wartość rezystora szeregowego o 1%.

Metody tej nie stosuje się dla kalibracji rezystancji ścieżki i styku. W tym celu należy wykonać osobny krok kalibracji. Mimo to, jeśli używana pęseta nie będzie różniła się znacząco od egzemplarza autora, wstępnie ustawiona wartość wpisana do oprogramowania układowego pęsety, będzie dość dokładna.

Nietrudno zauważyć, że pęseta nie ma żadnych przycisków, więc różne ustawienia są konfigurowane za pomocą jednego dostępnego urządzenia wejściowego – jej końcówek!

Możemy przejść przez konfigurację i kalibrację, otwierając i zamykając końcówki pęsety w różnych momentach. Jest to powolny, ale skuteczny proces, ułatwiony dzięki ekranowi reagującemu na poszczególne akcje.

Aby zrozumieć przebieg kalibracji należy przeanalizować algorytm pokazany na **rysunku 5**. Procedura konfiguracji działa tylko wtedy, gdy mikrokontroler jest włączony, więc



Ekran 4. Zwarcie ramion pęsety w trakcie wyświetlania informacji pokazanej na tym ekranie spowoduje zwiększenie zapisanej wartości szeregowego rezystora testowego w krokach co 1 Ω. Algorytm wyjaśniający ten proces pokazano na rysunku 5



Ekran 5. Podobnie, ten ekran umożliwia zmniejszenie wartości rezystora testu szeregowego. Jeśli wystąpi jakakolwiek zmiana, te dwa kroki są powtarzane do momentu, aż nie zostanie wykryta żadna zmiana



Ekran 6. Żądanie potwierdzenia chęci zapisania wprowadzonej wartości w pamięci nieulotnej Flash. W tym celu należy złączyć ramiona pęsety



Ekran 7. Ekran wyświetlany podczas zapisywania wartości, potwierdzający ten wybór



Ekran 8. Plecenie zwarcia końcówek w celu skalibrowania ich rezystancji styku. W przeciwnym razie zapisana wartość nie zostanie zmieniona



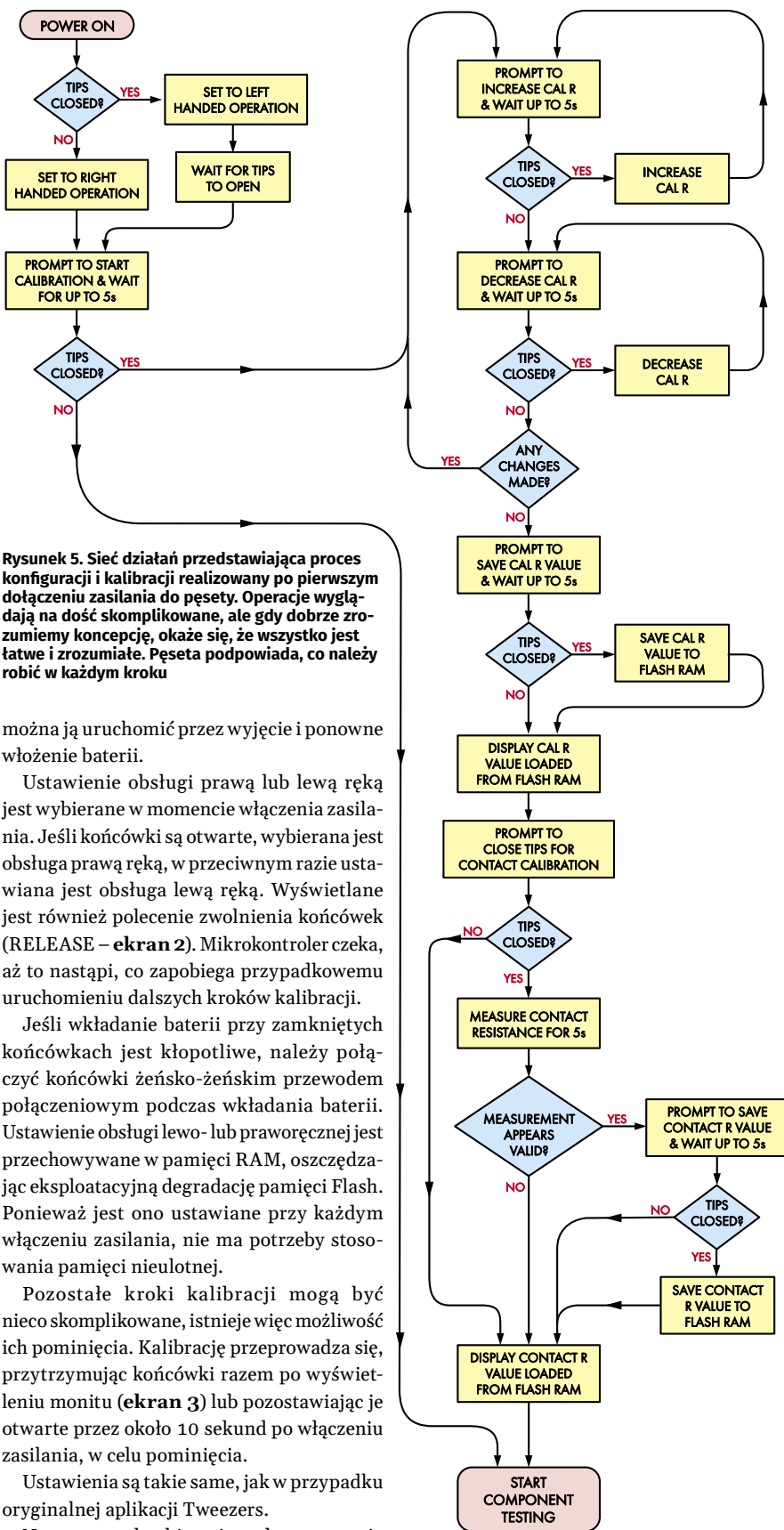
Ekran 9. Rezystancja styków jest uśredniana z około 20 pomiarów. Pokazana tu wartość jest wyższa niż domyślna wartość 16 Ω



Ekran 10. Jeśli nie pojawi się ten komunikat, pęseta wykryła, że końcówki mogły zostać otwarte, więc zmierzona wartość jest niedokładna



Ekran 11. Na koniec rzeczywista wartość zapisana w pamięci Flash zostanie ponownie załadowana, aby można było potwierdzić, że zapisana wartość jest prawidłowa



Rysunek 5. Sieć działań przedstawiająca proces konfiguracji i kalibracji realizowany po pierwszym dołączeniu zasilania do pęsety. Operacje wyglądają na dość skomplikowane, ale gdy dobrze zrozumiemy koncepcję, okaże się, że wszystko jest łatwe i zrozumiałe. Pęseta podpowiada, co należy robić w każdym kroku

można ją uruchomić przez wyjęcie i ponowne włożenie baterii.

Ustawienie obsługi prawą lub lewą ręką jest wybierane w momencie włączenia zasilania. Jeśli końcówki są otwarte, wybierana jest obsługa prawą ręką, w przeciwnym razie ustawiana jest obsługa lewą ręką. Wyświetlane jest również polecenie zwolnienia końcówek (RELEASE – ekran 2). Mikrokontroler czeka, aż to nastąpi, co zapobiega przypadkowemu uruchomieniu dalszych kroków kalibracji.

Jeśli wkładanie baterii przy zamkniętych końcówkach jest kłopotliwe, należy połączyć końcówki żeńsko-żeńskim przewodem połączeniowym podczas wkładania baterii. Ustawienie obsługi lewo- lub praworęcznej jest przechowywane w pamięci RAM, oszczędzając eksploatacyjną degradację pamięci Flash. Ponieważ jest ono ustawiane przy każdym włączeniu zasilania, nie ma potrzeby stosowania pamięci nieulotnej.

Pozostałe kroki kalibracji mogą być nieco skomplikowane, istnieje więc możliwość ich pominięcia. Kalibrację przeprowadza się, przytrzymując końcówki razem po wyświetleniu monitu (ekran 3) lub pozostawiając je otwarte przez około 10 sekund po włączeniu zasilania, w celu pominięcia.

Ustawienia są takie same, jak w przypadku oryginalnej aplikacji Tweezers.

Następnym krokiem jest dostosowanie wartości rezystora testowego 10 kΩ. Na wyświetlaczu pojawi się CAL R+ i licznik czasu (ekran 4). Za każdym razem, gdy końcówki

zostaną zamknięte podczas tej fazy, wyświetlana wartość wzrośnie, a licznik czasu zostanie wyzerowany.

Po tym następuje faza CAL R– (**ekran 5**), która działa podobnie, ale pozwala na zmniejszenie wartości. Jeśli zostaną wprowadzone jakiegokolwiek zmiany, cykl powtarza kroki CAL R+ i CAL R– do momentu, aż nie zostaną wprowadzone żadne zmiany.

Następnie na wyświetlaczu zostanie wyświetlone polecenie zapisania wartości. Potwierdzeniem wykonania zapisu jest ponowne dotknięcie końcówek przed wyświetleniem limitu czasu. Po zgłoszeniu potwierdzenia zapisu zostanie wyświetlony krótki komunikat informujący o tym fakcie (**ekrany 6 i 7**).

Wreszcie, niezależnie od tego, czy jakiegokolwiek zmiany zostały wprowadzone lub zapisane, wartość szeregowego rezystora testowego jest świeżo ładowana z pamięci Flash i wyświetlana w celu potwierdzenia przez użytkownika. Następny krok mający na celu ustawienie rezystancji styku jest prostszy, ponieważ jest ona mierzona, a nie wprowadzana.

Uwaga!

Jak każdy projekt wykorzystujący baterie pastylkowe, pęseta powinna być trzymana z dala od dzieci, które mogą ją połknąć. Pęseta ma również dość spiczaste końcówki, co jest kolejnym powodem, dla którego należy ją trzymać poza zasięgiem ciekawskich palców.

Należy pamiętać, że zegary wyświetlane na tych ekranach nie są bardzo precyzyjne. Wewnętrzne czasy różnią się w zależności od tego, co jest wyświetlane (zwłaszcza zmieniające się liczby, których renderowanie wymaga czasu). Licznik powinien odmierzać interwały o długości ok. 1 sekundy, ale nie są one bardzo dokładne.

Komunikat widoczny na **ekranie 8** oznacza rozpoczęcie kalibracji rezystancji styków. Gdy końcówki są przytrzymywane razem, wyświetlany jest **ekran 9**. Pokazuje on zmierzoną rezystancję styku, uśrednioną z kilku odczytów. Domyślna wartość to 16 Ω , zmierzona na egzemplarzu prototypowym.

Jeśli końcówki zostaną przypadkowo otwarte, proces zostanie przerwany i w celu jego powtórzenia konieczne będzie ponowne uruchomienie kalibracji. W przeciwnym razie wyświetlana jest uśredniona wartość wraz z poleceniem jej zapisanie, jak pokazano na **ekranie 10**. Zwarcie końcówek pęsety potwierdza chęć zapisu, a pozostawienie otwartych pozwoli licznikowi na upływanie czasu. Na **ekranie 11** wyświetlana jest rzeczywista wartość załadowana z pamięci Flash. ■

Tim Blythman

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

REKLAMA

przejrysz i kupisz na www.ulubionykiosk.pl



**świat
radio**
Magazyn wszystkich użytkowników eteru
KRÓTKOFALARSTWO CB RADIOTECHNIKA



Prosty limiter do subwoofera

Niniejsza publikacja zawiera opis techniczny prostego limitera umożliwiającego zabezpieczenie stopnia końcowego subwoofera przed przeciążeniem. Obwód wyzwalania limitera jest odseparowany galwanicznie od części małosygnałowej za pośrednictwem zespołu składającego się z diody LED oraz fotorezystora. Rozwiązanie to może być szczególnie pomocne konstruktorom układów elektronicznych do subwooferów aktywnych.

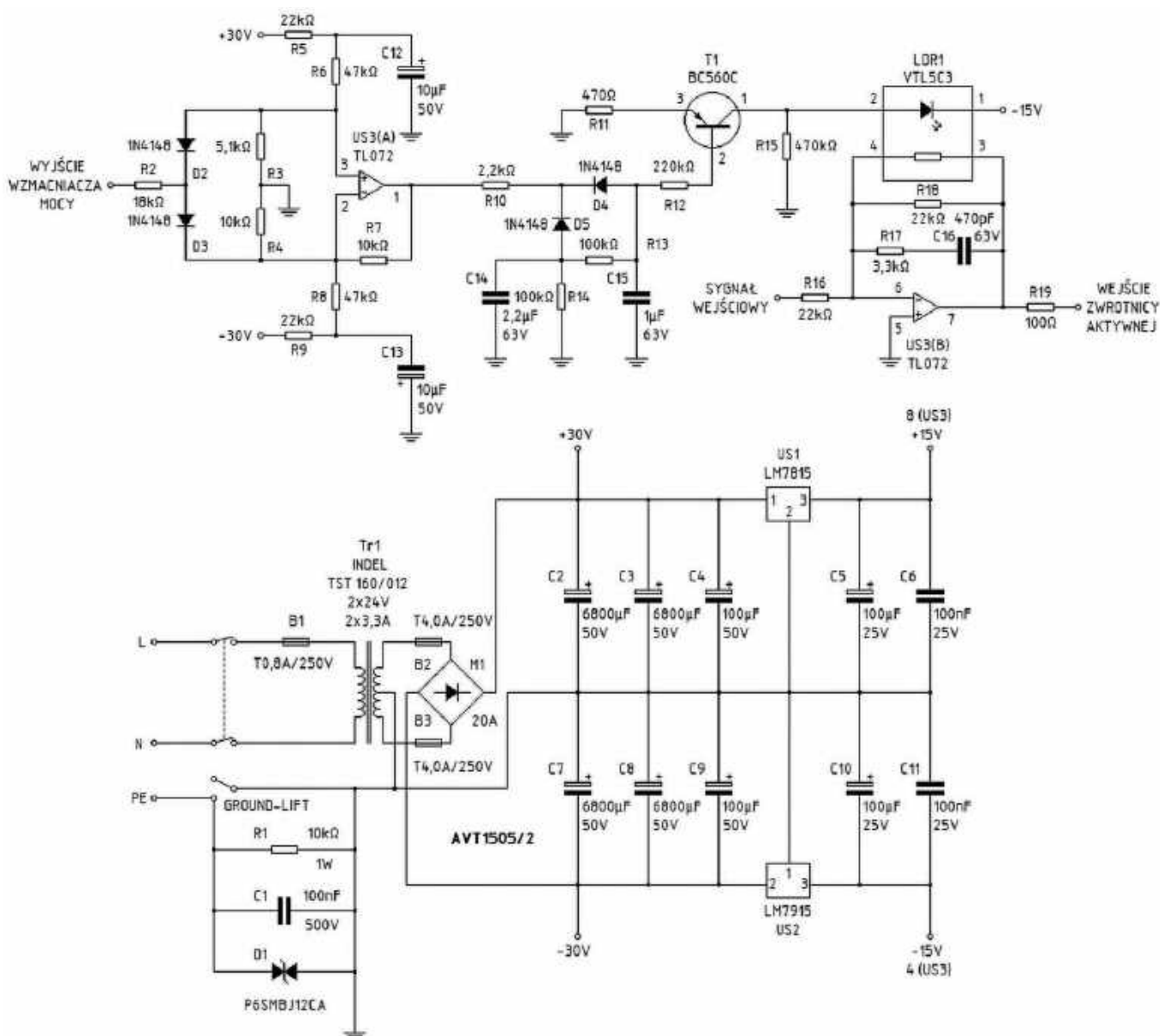
Opis układu

Układy elektroniczne subwooferów aktywnych stosowanych w technice nagłośnieniowej są narażone na przeciążenie spowodowane podaniem na ich wejścia sygnałów o zbyt dużej amplitudzie. Pewne rozwiązanie tego problemu może stanowić zastosowanie w układzie

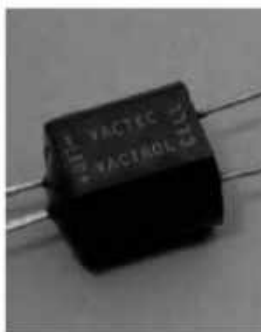
prostego limitera. Zadaniem tego układu będzie nadążna rejestracja poziomu sygnału wyjściowego stopnia końcowego subwoofera i zmniejszanie wzmacnienia pierwszego stopnia przedwzmacniacza w przypadku zarejestrowania sygnału o zbyt dużej amplitudzie. Zaproponowany układ elektroniczny limitera

jest prosty i niezawodny. Bardzo podobne rozwiązanie stosowano swego czasu w projekcie subwoofera typu Aktiv 250 produkowanego przez wrzesińską firmę Tonsil. Schemat ideowy limitera przedstawiono na **rysunku 1**.

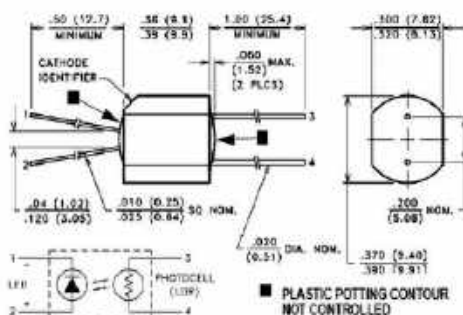
Do zasilania wzmacniacza mocy służy toroidalny transformator sieciowy Tr1



1. Schemat ideowy prostego limitera do subwoofera



PACKAGE DIMENSIONS INCH (MM)



DESCRIPTION

VTL5C3 has a steep slope, good dynamic range, a very low temperature coefficient of resistance, and a small light history memory. VTL5C4 features a very low "on" resistance, fast response time, with a smaller temperature coefficient of resistance than VTL5C3.

ABSOLUTE MAXIMUM RATINGS @ 25°C

Maximum Temperatures Storage and Operating:	-40°C to 75°C	LED Forward Voltage Drop @ 20 mA:	2.0V (1.65V Typ.)
Cell Power:	175 mW	Min. Isolation Voltage @ 70% Rel. Humidity: 2500 VRMS	
Derate above 30°C:	3.9 mW/°C	Output Cell Capacitance:	5.0 pF
LED Current:	40 mA	Cell Voltage:	250V (VTL5C3), 50V (VTL5C4)
Derate above 30°C:	0.9 mW/°C	Input - Output Coupling Capacitance:	0.5 pF
LED Reverse Breakdown Voltage:	3.0 V		

ELECTRO-OPTICAL CHARACTERISTICS @ 25°C

Part Number	Material Type	ON Resistance R_{ON}		OFF Resistance R_{OFF} @ 10 sec. (Min.)	Slope (Typ.) R_{ON} @ 0.5 mA R_{ON} @ 5 mA	Dynamic Range R_{DARK} R_{ON} @ 20 mA	Response Time t_r	
		Input current	Dark Adapted (Typ.)				Turn-on to 63% Final R_{ON} (Typ.)	Turn-off (Decay) to 100 kHz (Max.)
VTL5C3	3	1 mA 10 mA 40 mA	30 k Ω 5 k Ω 1.5 k Ω	10 M Ω	20	75 db	2.5 ms	35 ms
VTL5C4	4	1 mA 10 mA 40 mA	1.2 k Ω 125 k Ω 75 k Ω	400 M Ω	18.7	72 db	5.0 ms	1.5 sec

Refer to Specification Notes, page 41.

ParkinElmer Optoelectronics, 10900 Page Ave., St. Louis, MO 63132 USA.

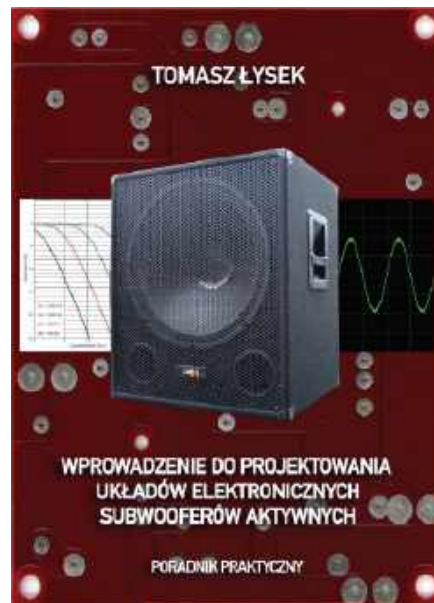
Phone: 314-425-4960 Fax: 314-425-5956 Web: www.parkinvalmer.com/logo

2. Karta katalogowa zespołu VTL5C3 składającego się z diody LED oraz fotorezystora Źródło: [https://www.tme.eu]

dołączany do sieci zasilającej za pośrednictwem wyłącznika dwusekcyjnego. Uzwojenie pierwotne transformatora zabezpieczone jest bezpiecznikiem zwłocznym B1. Uzwojenie wtórne tego transformatora posiada odstęp wyprowadzony ze środka tego uzwojenia. Elementy: B2, B3, M1, C2, C3, C7 oraz C8 wchodzi w skład zestawu do samodzielnego montażu typu AVT1505/2 i są zamontowane na osobnym obwodzie drukowanym tworząc zasilacz symetryczny dostarczający do wzmacniacza napięcia ± 30 V. Układ zasilania uzupełniono o elementy C4, US1, C5, C6, C9, US2, C10 oraz C11, pozwalające na obniżenie symetrycznych względem masy napięć zasilania do poziomu ± 15 V, co z kolei pozwala na wykorzystanie ich do zasilania układu zwrotnicy aktywnej subwoofera oraz układów pomocniczych, do których

zalicza się m.in. układ limitera. Schemat zawiera także układ typu Ground-Lift składający się z elementów R1, C1 oraz D1. Układ ten służy do przerywania pętli masy w przypadku, gdyby źródło sygnału podłączone do subwoofera było jednocześnie podłączone do zacisku ochronnego PE sieci zasilającej. Układ limitera składa się z dwóch odseparowanych od siebie galwanicznie obwodów. Do rejestracji poziomu sygnału wyjściowego stopnia końcowego subwoofera służą elementy: R2, D2, D3, R5, R6, C12, US3(A), R7, R8, R9 oraz C13. Uzyskany w ten sposób sygnał jest w następnej kolejności prostowany i filtrowany przez elementy: R10, D4, D5, C14, R13, R14 oraz C15 aby po przekroczeniu progu zadziałania, ujemny stały prąd mógł za pośrednictwem opornika R12 wprowadzić tranzystor bipolarny T1 w stan przewodzenia. Kiedy tranzystor bipolarny

T1 zostanie wprowadzony w stan przewodzenia, prąd płynie od masy, przez opornik R11, złącze tranzystora T1 oraz diodę LED wchodzącą w skład zespołu LDR1 składającego się z diody LED oraz fotorezystora, do ujemnej szyny zasilania -15 V. Drugi obwód składa się z elementów: R16, R17, R18, C16, US3(B) oraz R19. Jest to w praktyce wzmacniacz odwracający, który w normalnych warunkach pracy dla niskich częstotliwości ma wzmocnienie napięciowe bliskie jedności natomiast dla wysokich częstotliwości tłumi sygnał około ośmiokrotnie. Przepływ prądu przez diodę LED powoduje oświetlenie fotorezystora, zmniejszenie jego rezystancji i w rezultacie zmniejszenie sumarycznego wzmocnienia układu poniżej jedności. Tłumienie sygnału na wejściu zwrotnicy aktywnej służy zabezpieczeniu stopnia końcowego subwoofera i działa na zasadzie sprzężenia zwrotnego. Opornik R19 zabezpiecza wzmacniacz operacyjny US3(B) w przypadku, gdyby na wejściu zwrotnicy aktywnej nastąpiło zwarcie. Na rysunku 2 przedstawiono kartę katalogową zespołu LDR1 składającego się z diody LED oraz fotorezystora.

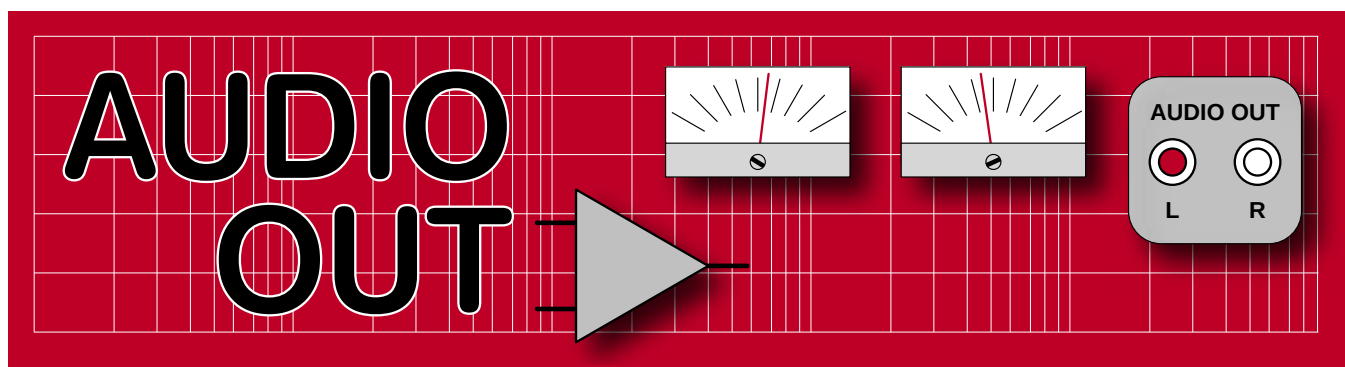


3. Okładka książki pt. „Wprowadzenie do projektowania układów elektronicznych subwooferów aktywnych. Poradnik praktyczny”

Książka o układach elektronicznych do subwooferów aktywnych

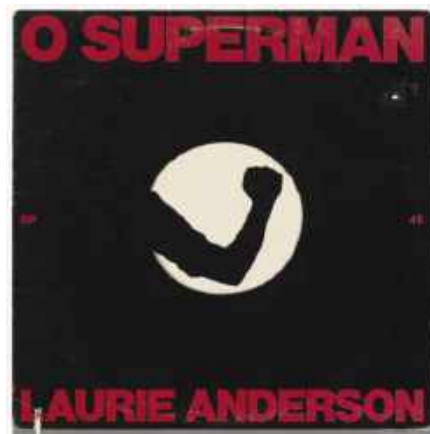
Zapraszam do zapoznania się z moją najnowszą książką pt. „Wprowadzenie do projektowania układów elektronicznych subwooferów aktywnych. Poradnik praktyczny”: <https://youtu.be/KIO1eqxj4AE>, <https://youtu.be/gpQe89R5HEK>. ■

mgr inż. Tomasz Łysek



Wokoder analogowy, część 1

Pora na nowy projekt muzyczny! Wiele z was zapewne kojarzy efekt zwany „wokoderem”, stosowany czasem w muzyce rozrywkowej do przetwarzania głosu wokalisty. Wokoder to fascynujące urządzenie elektroniczne. A my w ciągu najbliższych miesięcy zaprojektujemy i zbudujemy jego przykładową wersję o bardzo wysokiej jakości.



Rysunek 1. „O Superman”, płyta Laurie Anderson; wokaliza w stylu „avant guard minimalism” z użyciem Rolanda VP-330



Rysunek 2. Wokodery programowe to dla większości muzyków najtańszy środek, o ile się ma szybki komputer i zainstalowane oprogramowanie cyfrowej stacji roboczej audio. Jak na ironię, mogą one posłużyć do zoptymalizowania projektu wokodera analogowego przed jego zbudowaniem

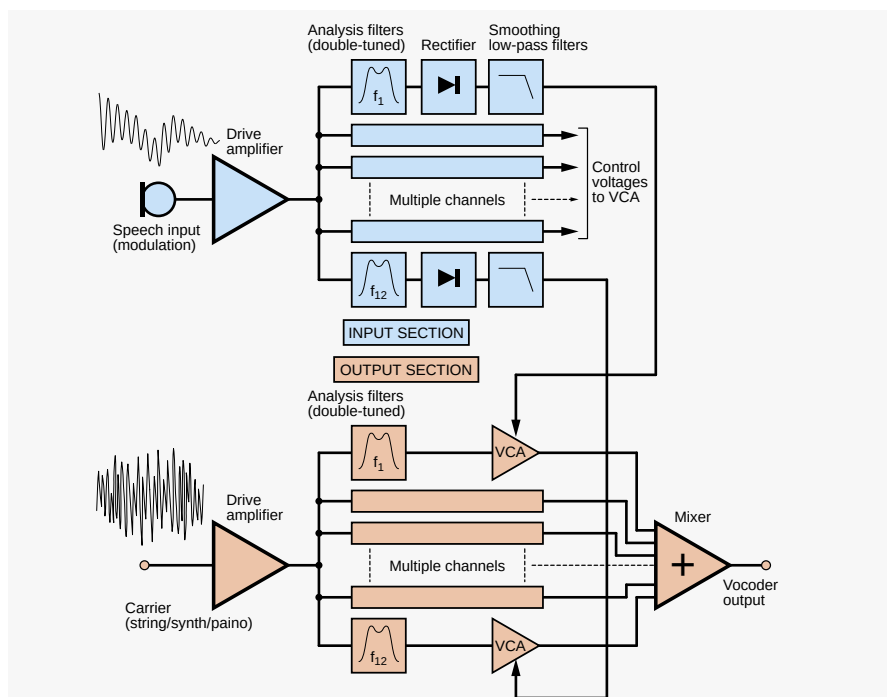
Charakterystyczną cechą wokodera jest to, że głos ludzki jest niejako nakładany na brzmienie jakiegoś instrumentu muzycznego i nadaje temu brzmieniu swój charakter. Przypis redaktora: za historycznie pierwsze użycie wokodera w nagraniu muzycznym uważa się utwór „The Raven” zespołu Alan Parsons Project z 1976 roku. Wersja zremasterowana: <https://youtu.be/N3fHbOsSaE>.

Nowszym przykładem jest „wokoderowa” piosenka „O Superman” (<https://youtu.be/S39NaDPNDtk>), wydana przez Laurie Anderson w 1981 roku (na **rysunku 1** widzimy

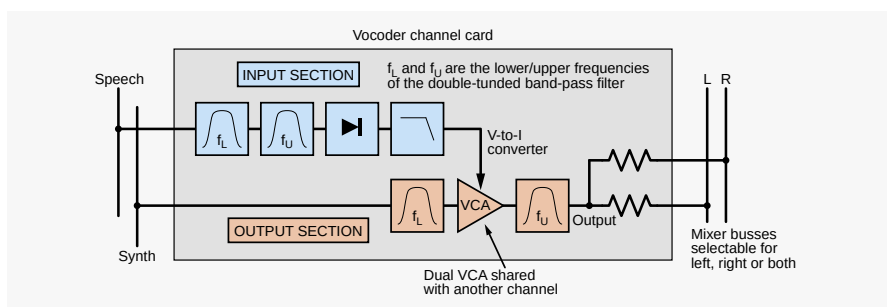
okładkę cennego egzemplarza płyty winylowej w nieskazitelnym stanie ze zbiorów jednego z naszych redaktorów). Utwór został nagrany na analogowym syntezatorze z wokoderem Roland VP-330: https://tiny.pl/c7ytt_56.

Początki

Wokoder, podobnie jak szereg innych urządzeń stosowanych w technice nagraniowej, został pierwotnie stworzony do celów zupełnie niemuzycznych, a mianowicie na potrzeby telekomunikacji. Wynalazł go w 1936 roku Homer Dudley (Bell Labs), a zadaniem urządzenia było



Rysunek 3. Schemat blokowy wokodera analogowego, pokazujący sekcje analizy, sterowania i resyntezy



Rysunek 4. Schemat blokowy każdego z modułów realizujących kanały wokodera

przesyłanie mowy ludzkiej przez długie linie kablowe i inne systemy transmisji o niskiej przepustowości. Wokodery to analogowe systemy kompresji danych, w których informacja o sygnale mowy jest reprezentowana w oszczędnej formie: poprzez poziomy tego sygnału w różnych pasmach częstotliwości. Coś jak słupki w analizatorze widma. Te poziomy zmieniają się stosunkowo powoli i mogą być przesyłane liniami o małej przepustowości. Informacja o poziomach jest odbierana na drugim końcu linii i służy do odtworzenia syntetycznej mowy. Używa się sygnału z generatora elektronicznego. Sygnał ten jest rozdzielany na te same pasma częstotliwości, jak w oryginalnym sygnale mowy, a poziomy w tych pasmach są też takie same. Wokodery używało również wojsko do maskowania głosów i kodowania tajnych wiadomości; stąd właśnie wzięła się nazwa: VOICE enCODER.

Wokoder można traktować jako system do nadawania charakterystyki widmowej sygnału mowy innemu dźwiękowi. **Przypis redaktora: w telekomunikacji dźwiękiem tym był sygnał z generatora piłokształtnego, imitujący ton wytwarzany przez ludzką krtań.** W zastosowaniach muzycznych świetnie sprawdzają się syntezatorowe brzmienia typu „sekcja smyczkowa”, bogate w tony składowe. Jeśli zostanie użyty biały szum, zakodowany głos będzie brzmiał jak szepc.

Wokoder analogowy jest układem skomplikowanym. Prezentowany przeze mnie projekt składa się z 28 precyzyjnych filtrów i 14 wzmacniaczy sterowanych napięciem (VCA). Tzw. „rasowy” inżynier elektronik stwierdziłby zapewne, że działanie tych bloków najlepiej osiągnąć poprzez cyfrowe przetwarzanie sygnału z pomocą szybkiej transformaty Fouriera (FFT) oraz sprzętu i oprogramowania DSP.

Punktem wyjścia dla większości muzyków pragnących uzyskać brzmienie wokodera są komputerowe wtyczki programowe (plug-ins). W moim przypadku był to Vocoder Channel Vocoder w programie FL Studio (**rysunek 2**), w którym eksperymentowałem, szukając najlepszych parametrów filtrów (częstotliwości i dobroci). Cyfrowe wokodery stwarzają jednak problemy. Wymagają dużej mocy obliczeniowej

i często wykazują znaczne opóźnienia czasowe (latencję). Wszystkie wykorzystują podobną matematykę, która nadaje dźwiękowi specyficzną, jakby „podwodną” zabarwienie przy niskich poziomach sygnału, co uważam za nieprzyjemne. Sporo osób może znać ten efekt, bo większość smartfonów wykorzystuje „wokodowanie” jako część algorytmów kompresji. Wokodery programowe są też skomplikowane w użyciu ze względu na mnóstwo funkcji i parametrów do skonfigurowania. Natomiast wokodery analogowe mają o wiele lepszą „grywalność” – są na tyle łatwe i intuicyjne w obsłudze, że nadają się do komponowania i występów na żywo.

Wokodery analogowe konstruowane przez amatorów nie są tanie – ich budowa wiąże się zazwyczaj z wydatkiem rzędu tysiąca złotych. Pamiętajmy jednak, że urządzenia fabryczne, takie jak Roland VP-330 Vocoder Plus lub Electronic Music Studios (EMS) Timbra, to już koszt co najmniej kilkunastu tysięcy złotych.

Po prostu znaleźliśmy się w momencie, w którym analogowe elektroniczne instrumenty muzyczne stały się antykami i jako takie osiągają zaporowe ceny. Wokoder to jeden z niewielu przypadków w elektronice, gdzie zbudowanie urządzenia samemu pozwala zaoszczędzić poważne pieniądze.

Jak na ironię, akurat w momencie, gdy projekt naszego wokodera po pięciu latach pracy trafił do druku (listopad 2021), firma Behringer wprowadziła na rynek kopię wspomnianego już Rolanda VP-330. Mowa o syntezatorze VC340, montowanym w Chinach, a kosztującym ok. 2000 zł. Mój projekt ma jednak przewagę – jest stereofoniczny oraz daje się modyfikować. Są to ważne cechy, ponieważ w muzyce elektronicznej istotny jest własny, oryginalny charakter brzmienia, unikatowa „sygnatura dźwiękowa”.

Architektura układu

Schemat blokowy pokazano na **rysunku 3**. Część wejściowa wokodera („niebieska”) dokonuje analizy zmian sygnału modulującego (zwykle mowy, a w zastosowaniach muzycznych śpiewu) i w telekomunikacji często jest nazywana „modulatorem”. Sekcja ta wytwarza wolnozmienną napięcia sterujące, reprezentujące

poziomy sygnał w różnych pasmach częstotliwości. Część wyjściowa („pomarańczowa”) odpowiada za formowanie dźwięku modulowanego (w telekomunikacji określanego jako „fala nośna” lub „pobudzenie”) czyli syntezy.

Dźwięk modulowany – w zastosowaniach muzycznych są to np. akordy z fortepianu lub syntezatora – jest rozdzielany na odpowiednie pasma częstotliwości, takie same jak w części wejściowej. Poziomy sygnałów w tych pasmach są następnie modulowane przez napięcia sterujące pochodzące z analogicznych pasm sekcji wejściowej. Wymaga to użycia wzmacniaczy sterowanych napięciem (VCA), przez które przechodzą sygnały każdego pasma. Efekt końcowy jest taki, że dźwięk fortepianu czy inny sygnał muzyczny zaczyna „mówić” i „śpiewać”.

Możliwe jest również działanie niekonwencjonalne: przemieszanie sygnałów sterujących, czyli sterowanie pasmami częstotliwości w części wyjściowej z zupełnie innych pasm części analizującej. Powstają wtedy dźwięki mocno odbiegające od oryginału, wręcz „szalone”. W taki sposób można skutecznie szyfrować mowę.

Przypis redaktora: w latach 70. firma Moog produkowała wokoder 16-kanałowy, w którym nie było ustalonych połączeń między częścią analizującą a wyjściową. Wszystkie sygnały sterujące, podobnie jak wszystkie wejścia VCA, wyprowadzono na gniazda na panelu czołowym. Użytkownik mógł więc sterować każdym pasmem wyjściowym przez dowolne kombinacje napięć sterujących z różnych pasm wejściowych.

Nasz wokoder składa się z 14 „modułów pasmowych”. Schemat blokowy jednego modułu pokazano na **rysunku 4**. Każdy moduł jest umieszczony na oddzielnej płytce drukowanej. Wszystkie płytki są dołączone do płyty łączącej (magistrali), co widać na zdjęciu prototypu (**rysunek 5**).

12 modułów pasmowych zawiera filtry pasmowoprzepustowe. Ich częstotliwości zostały dobrane pod kątem optymalnego przetwarzania sygnału mowy. Pozostałe dwa moduły przetwarzają skrajne obszary pasma akustycznego: jeden zawiera filtr dolnoprzepustowy 120 Hz, drugi – filtr górnoprzepustowy 8 kHz. Wszystkie filtry są czwartego rzędu i łącznie zawierają 112 kondensatorów o wąskiej

Rysunek 5. Wczesne zdjęcie prototypu wokodera analogowego. Nie ma jeszcze pokręteł





Rysunek 6. Płytkę prototypowego kanału wokodera. W każdym kanale filtry muszą mieć inne wartości kondensatorów. Ostateczna wersja płytki będzie dwustronna, co wyeliminuje mostki drutowe

tolerancji. Te kondensatory stanowiły znaczną część rachunku za elementy! A pomyślnie, że w wokoderze EMS jest takich filtrów aż 22...

Prototypową płytkę kanału pasmowoprzepustowego wokodera widzimy na **rysunku 6**.

Filtry pasmowoprzepustowe

Istnieje wiele sposobów na zbudowanie filtra pasmowoprzepustowego. Jednym z nich, bardzo skutecznym, jest wykorzystanie układu „zmiennej stanu” lub jego alternatywy „bi-quad”; patrz **rysunek 7**. Częstotliwość rezonansową (f_c) „bi-quada” można regulować pojedynczym rezystorem (wyróżnionym przez ramkę z tekstem), a w filtrze „zmiennej stanu” – dwoma rezystorami (R) (**ściślej: potencjometrem podwójnym – przypis redaktora**). Obie te konfiguracje wymagają jednak trzech wzmacniaczy operacyjnych na filtr, a to za dużo w systemie z tak wieloma filtrami, jakim jest wokoder.

Idealne wycinanie pasm częstotliwości można by było osiągnąć, stosując filtry

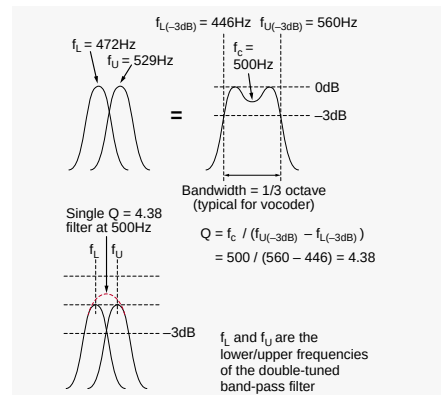
cyfrowe o dużej stromości charakterystyki – ale podejrzewam, że w zastosowaniach muzycznych brzmiałyby one okropnie z powodu „dzwonienia” wywołanego znacznym opóźnieniem grupowym.

Czas na konkrety

Dla przetwarzania mowy, częstotliwości środkowe filtrów pasmowoprzepustowych muszą być rozmieszczone w odstępach około 1/3 oktawy (każda kolejna ok. 1,26 raza większa od poprzedniej – przypis redaktora). Aby skrajne częstotliwości odcięcia (–3 dB) znalazły się we właściwych miejscach, wymagany jest filtr o dobroci Q około 4,38.

Żeby w części przepustowej uzyskać bardziej płaską charakterystykę częstotliwościową, dobrze jest użyć filtrów podwójnych, składających się z dwóch filtrów pasmowych połączonych szeregowo, o zbliżonych (ale nie jednakowych) częstotliwościach środkowych, co daje efekt w postaci nieco „wklęsłej” charakterystyki pokazanej na **rysunku 8**. Oba filtry mają dobroć 8,8, co daje wierzchołek w miarę płaski, nie tak ostry jak w przypadku filtra pojedynczego. Poza tym filtr podwójny ma bardziej strome zbocza (–40 dB na oktawę) i dwukrotnie mocniej tłumi sygnały spoza pasma przepustowego. Technika podwójnych filtrów została pierwotnie opracowana dla transformatorów pośredniej częstotliwości w radiotelefonach superheterodynowych.

Częstotliwości obu filtrów są tak dobrane, aby odpowiadały punktom –3 dB

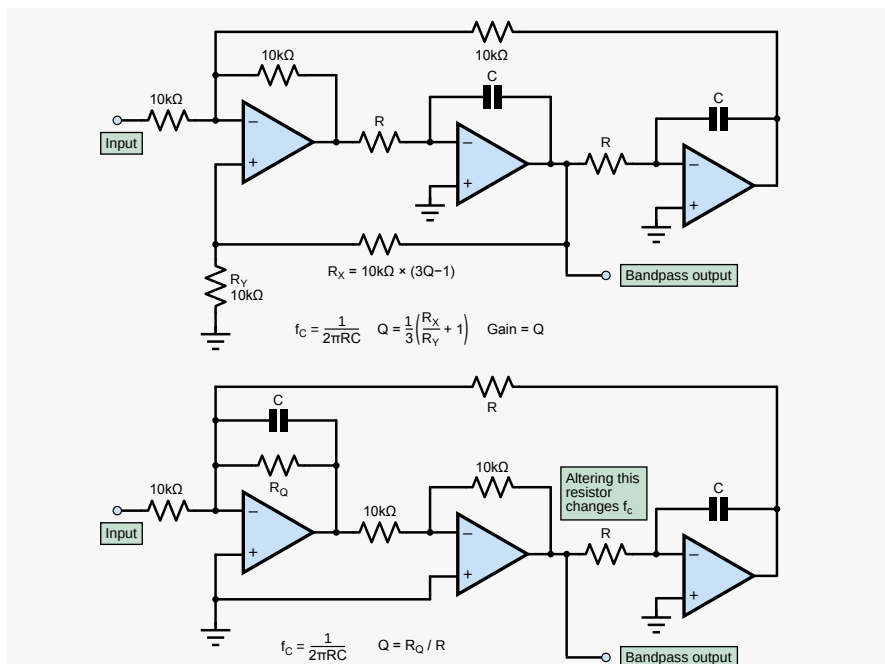


Rysunek 8. Łącząc szeregowo dwa filtry pasmowoprzepustowe o wysokiej dobroci można uzyskać niezłą kompromisową charakterystykę ze stromymi krawędziami i płaską górą. Taki podwójny filtr pasmowoprzepustowy jest znany większości projektantów urządzeń radiowych

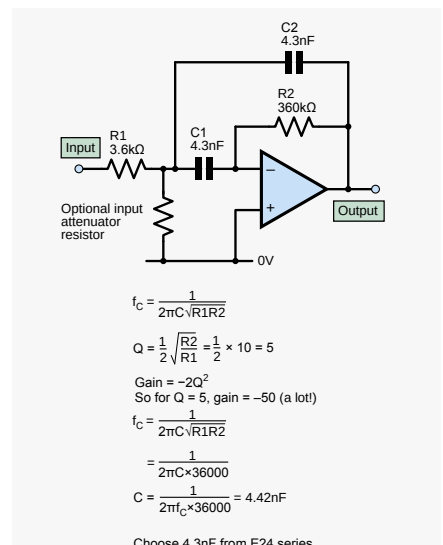
charakterystyki filtra docelowego. Tak więc dla podwójnego filtra o częstotliwość środkową 500 Hz częstotliwości rezonansowe (f_L i f_U) filtrów składowych wyniosą odpowiednio 472 Hz i 529 Hz. Należy zauważyć, że niewielka asymetria wynika z konieczności zapewnienia wykładniczego (w przybliżeniu) skalowania częstotliwości.

Filtr pasmowoprzepustowy z wielokrotnym sprzężeniem zwrotnym

Najprostszym filtrem pasmowoprzepustowym ze wzmacniaczem operacyjnym jest układ pokazany na **rysunku 9**, nazywany



Rysunek 7. Topologie „zmiennej stanu” (na górze) i „bi-quad” (na dole) to wprowadzie idealne rozwiązania konstrukcyjne dla filtrów pasmowoprzepustowych, ale są zbyt złożone i drogie w przypadku tak dużej liczby kanałów występujących w urządzeniu. Układ „bi-quad” można by było ewentualnie zastosować w przypadku wokodera z małą ilością kanałów, których parametry ustawiałoby się na panelu przednim

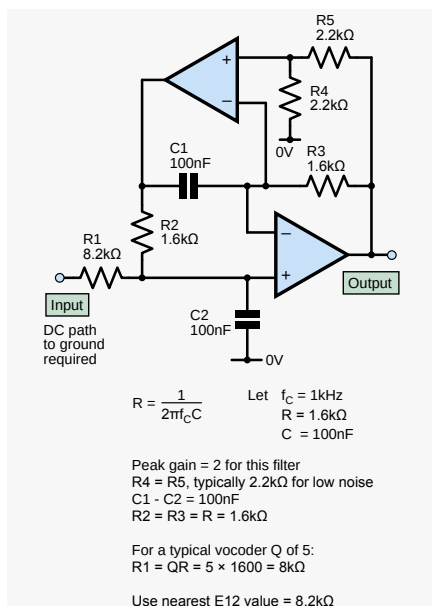


Rysunek 9. Filtr z wielokrotnym sprzężeniem zwrotnym, najprostszy i najtańszy. Wysokie wzmocnienie sprawia, że jest „głośny”. Dwa takie filtry dają filtr podwójny o płaskim wierzchołku. Często dawany jest rezystor do masy, tłumiący sygnał wejściowy; w takim przypadku wartość R1 należy skorygować tak, by równoległa rezystancja obu rezystorów była taka sama jak przedtem wartość R1 – inaczej rozstroimy filtr

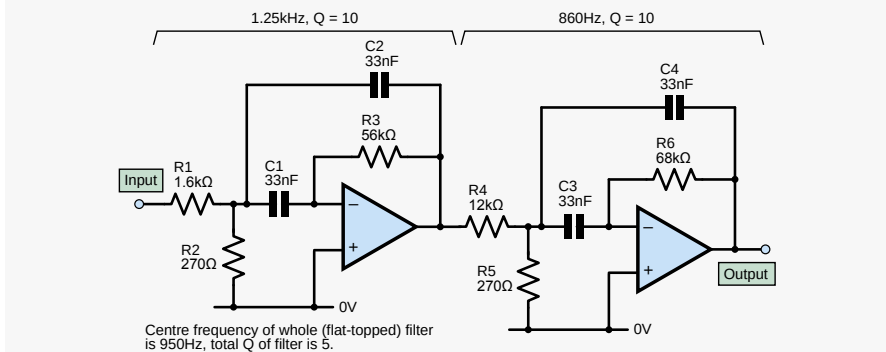
„filtrem z wielokrotnym sprzężeniem zwrotnym”, ponieważ wykorzystuje dwa tory sprzężenia. Wadą tej prostej konstrukcji jest to, że przy Q równym 5 wzmacnienie wynosi aż 50, co stanowi problem, ponieważ przy dużych sygnałach nastąpi przesterowanie wyjścia wzmacniacza operacyjnego. W przypadku 1/3-oktawowych filtrów podwójnych, Q każdej sekcji musi wynosić 8,8, co w praktyce stanowi górną granicę dla tego typu filtrów. Układ filtra podwójnego z wielokrotnym sprzężeniem zwrotnym pokazano na **rysunku 10**.

By uniknąć przesterowania, sygnał na wejściu każdego filtra z rysunku 10 musi być tłumiony, i to mocniej niż w przypadku filtra pojedynczego. W rezultacie stosunek sygnału do szumu nie jest wysoki. Na układ analizy mowy ma to niewielki wpływ, ponieważ po filtrach następują prostowniki pełnookresowe, a one wytwarzają wygładzone, wolnozmiennie napięcia sterujące dla VCA. W sekcji nośnej/wyjściowej ma to większe znaczenie, ale tylko filtry umieszczone za VCA wnoszą stały szum. Szum z filtrów poprzedzających VCA jest wyciszany, gdy VCA się wyłączają.

Alternatywnym typem filtra, który zamierzam wypróbować kiedyś w przyszłości, jest konfiguracja pasmowoprzepustowa z podwójnym wzmacniaczem (DABP), opracowana przez Sedrę i Espinozę, przedstawiona na **rysunku 11**. Ten typ filtra pozwala uniknąć problemu nadmiernego wzmacniania szumów, występującego w filtrach z wielokrotnym sprzężeniem zwrotnym. Wymaga jednak dodania do wokodera 24 wzmacniacze operacyjne,



Rysunek 11. Filtr pasmowoprzepustowy z dwoma wzmacniaczami operacyjnymi. Dwa razy więcej wzmacniaczy, ale lepiej kontrolowane wzmacnienie



Rysunek 10. Układ podwójnego filtra pasmowoprzepustowego. Pasma przenoszenia to 1/3 oktawy, częstotliwość środkowa wynosi 950 Hz

więc raczej zostawimy to na przyszłą aktualizację projektu.

Filtry górno- i dolnoprzepustowe

Topologia modułów z filtrami dolno- i górno-przepustowymi jest zasadniczo taka sama jak topologia modułów pasmowoprzepustowych (rysunek 4), tyle że filtry pasmowoprzepustowe zastąpiono tu filtrami górnoprzepustowym i dolnoprzepustowym czwartego rzędu. Układy tych filtrów pokazano na **rysunku 12**.

Omawiane moduły można potraktować jako dodatkowe. Są potrzebne tylko wtedy, gdy wymagane jest pełne pasmo przenoszenia wokodera. Można z nich zrezygnować, jeśli do sterowania będzie użyty wyłącznie głos ludzki.

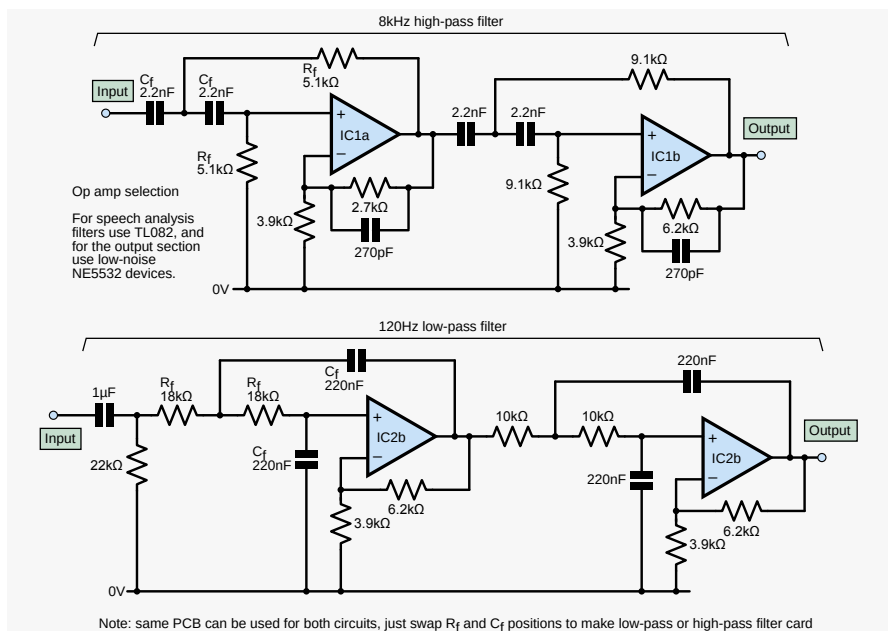
Przykładowo: w zespole, w którym grają basista i perkusista, tonów niskich jest pod dostatkiem i wystarczy, jeśli wokoder pokrywa

tylko pasmo mowy tj. częstotliwości środkowe od około 100 Hz do 8 kHz. W celu uzyskania klarownej „góry”, do sygnału wyjściowego wokodera dobrze jest domiksować trochę oryginalnego sygnału mowy z uwypuklonym wyższym pasmem. A moduły dolno- i górnoprzepustowe można sobie w tej sytuacji darować, uzyskując mniej zagmatwany miks i... zaoszczędzając kilkaset złotych.

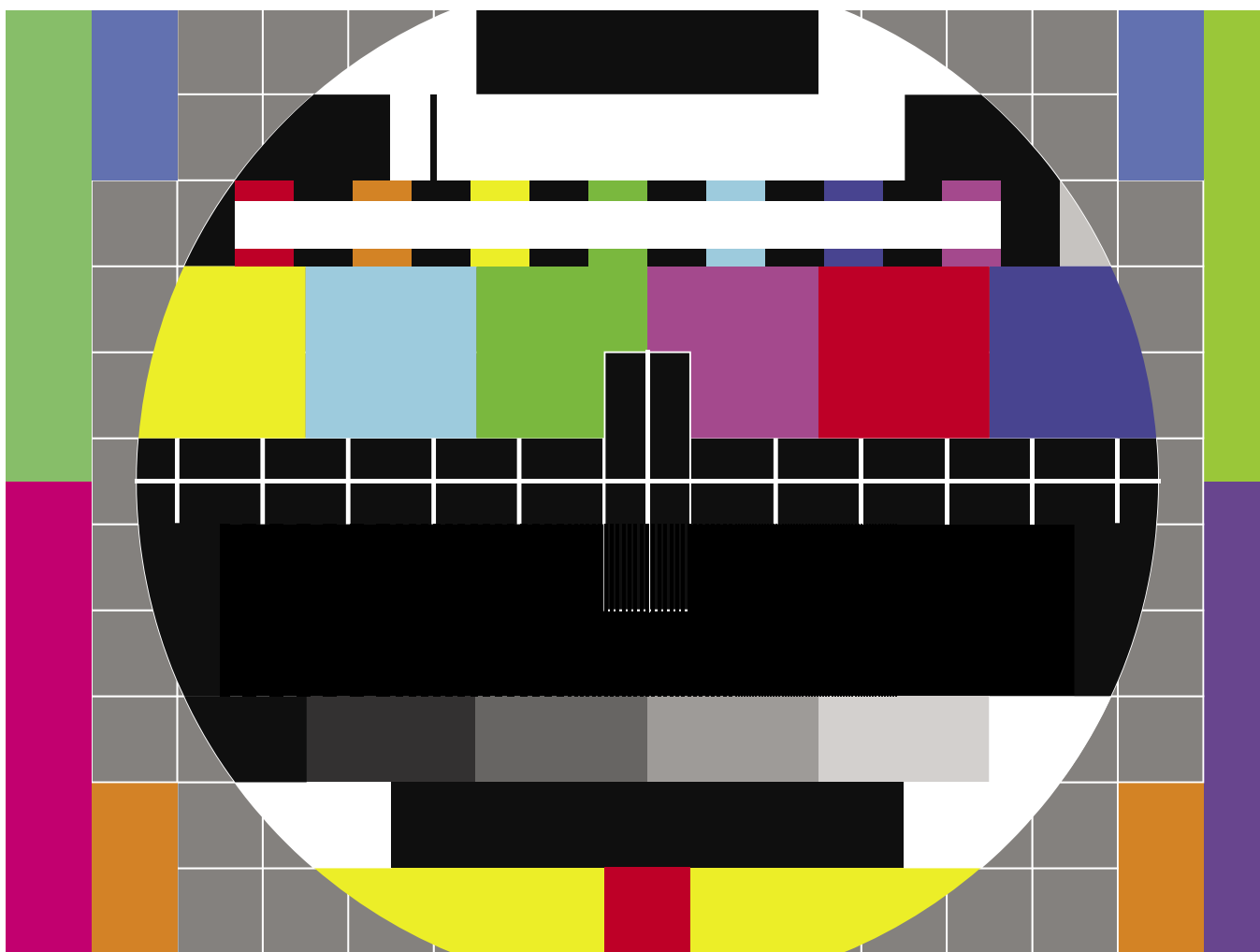
Za miesiąc

Na ten miesiąc to wszystko. W następnym odcinku zaczniemy bardziej szczegółowo przedstawiać przemyślenia towarzyszące powstawaniu projektu. ■

Jake Rothman



Rysunek 12. Dla górnego i dolnego pasma wokodera pełnozakresowego stosowane są filtry dolnoprzepustowe i górnoprzepustowe czwartego rzędu. Wymagają one innej płytki drukowanej i można się bez nich obejść w przypadku używania wokodera tylko z sygnałem mowy



Historia i technologia wyświetlaczy wideo, część 2

Miesiąc temu opisaliśmy rozwój technologii wyświetlania wideo od jej wczesnego powstania do około roku 2000, kiedy to wyświetlacze plazmowe i kineskopowe (CRT) zdominowały przestrzeń konsumencką. W tym artykule opiszemy rozwój ekranów ciekłokrystalicznych (LCD) i najnowsze technologiczne osiągnięcia na tym polu.

Wyświetlacze LCD są obecnie dominującą technologią wyświetlania obrazów statycznych, wykorzystywaną zarówno w monitorach komputerowych jak i ekranach odbiorników telewizyjnych. Na tę sytuację bez wątpienia wpłynęły takie czynniki, jak niski koszt, smukła konstrukcja, niewielka waga, wąskie ramki okalające ekran, doskonałe odwzorowanie kolorów, szerokie kąty widzenia, krótki czas reakcji, wysoki kontrast oraz stosunkowo niskie zużycie energii.

Jednak ekrany LCD nie zawsze takie były. Wczesne wyświetlacze LCD były małe, bardzo prymitywne, powolne w odświeżaniu obrazu

i przydatne tylko w urządzeniach takich jak kalkulatory. Ich rozwój i udoskonalanie zajęło dziesięciolecia, zanim stały się użyteczne przy budowie odbiorników telewizyjnych.

Rzeczony rozwój tej technologii wciąż trwa. Znacznie poprawiono ich podświetlanie, powstały ekrany z nanokryształami półprzewodnikowymi używanymi w ekranach z tzw. kropkami kwantowymi (Quantum dots – QD), ekrany OLED i MicroLED są w powszechnym użytku. Na popularności zyskują również nieco bardziej egzotyczne pod względem technologicznym produkty takie jak telewizory laserowe. Zanim do nich przejdziemy,

zaczniemy od rozwoju technologii wyświetlaczy ciekłokrystalicznych i zasad działania tego typu wyświetlaczy.

Wyświetlacze ciekłokrystaliczne (LCD)

Niektórzy nazywają ciekłe kryształy „czwartym stanem materii”. Czy nie chodzi przypadkiem o plazmę?

To, co obecnie znamy jako ciekłe kryształy, zostało po raz pierwszy zaobserwowane przez Rudolfa Virchowa w 1854 roku. Zauważył on niezwykle zachowanie mieliny (warstwy izolacyjnej wokół wiązań nerwowych).

Następnie w 1857 r. Niemiec Carl von Mettenheimer, również badający mielinę, zauważył, że płynęła ona jak ciecz, ale gdy oglądano ją pod skrzyżowanymi polaryzatorami, światło wykazywało wysoce kolorową dwójłomność, typową dla kryształów. Materiał ten nie został jednak wówczas zidentyfikowany jako ciekły kryształ.

Austriacki botanik Friedrich Reinitzer odkrył ciekłe kryształy w 1888 roku, gdy badał materiał, benzoosan cholesterolu, ekstrahowany z marchwi. Wykazał on specyficzne właściwości pomiędzy dwiema temperaturami („dwie różne temperatury topnienia”, jak je określił), które były charakterystyczne zarówno dla stanu ciekłego (amorficznego), jak i stałego (krystalicznego).

W tym stanie „mezofazy” materiał mógł odbijać spolaryzowane światło i odwracać polaryzację światła. Obserwowanemu zjawisku nadał nazwę ciekłego kryształu. Więcej szczegółów można znaleźć pod poniższymi linkami:

- siliconchip.au/link/abfb,
- siliconchip.au/link/abfc.

W 1922 roku Wsiewołod Fréedericksz i A. Rzepiewa odkryli efekt zwany obecnie przejściem Fréedericksza, który stanowi podstawę technologii ekranów LCD. Gdy ciekły kryształ jest umieszczony między dwiema przezroczystymi szklanymi elektrodami, przepuszczalność światła może być kontrolowana elektrycznie, i stanowić elektronicznie sterowaną barierę optyczną (**rysunek 27**).

Ciekłe kryształy są niezbędne do życia. Błony komórkowe, otoczka mielinowa, która izoluje nerwy i trawienie tłuszczów – wszystko to wymaga ciekłych kryształów.

Zainteresowanie ciekłymi kryształami było niewielkie aż do 1962 roku, kiedy to Richard Williams z RCA Laboratories w USA odkrył elektrooptyczne właściwości tych materiałów. Odkrył on, że ciekłe kryształy tworzą pasiaście wzory po przyłożeniu pola elektrycznego. W 1968 r. George Heilmeyer zademonstrował wyświetlacz ciekłokrystaliczny, choć musiał on pracować w temperaturze 80°C.

Następnie opracowano materiały LCD, które mogły pracować w temperaturze pokojowej. W 1970 roku na międzynarodowej wystawie AICHEMADA zademonstrowano kalkulator wykorzystujący ekran LCD oparty na produktach firmy Merck. Pierwszym dostępnym na rynku kalkulatorem z wyświetlaczem LCD był Sharp EL-805, którego sprzedaż rozpoczęto w 1973 roku.

W latach 1976 i 1978 firma Merck opracowała materiały LCD o krótkim czasie przełączania, skracając czas przejścia z setek milisekund do 20 ms a nawet mniej. Poprawione zostały także właściwości optyczne. W 1980 roku

Rozmiar i proporcje ekranów w odbiornikach TV i monitorach komputerowych

Standardowym sposobem specyfikowania rozmiarów ekranów w odbiornikach TV i monitorach komputerowych jest przekątna ekranu. Jest ona często podawana w calach, choć europejskie i azjatyckie marki zwykle podają również centymetry (wiele japońskich telewizorów CRT było reklamowanych w centymetrach). Ma to tę zaletę, że daje rozsądne wyobrażenie o rozmiarze ekranu dla różnych proporcji obrazu.

Używanie przekątnej do pomiaru rozmiaru ekranu ma swoje historyczne korzenie w czasach, gdy kineskopy były okrągłe, ale musiały wyświetlać prostokątne obrazy, a znaczna część lampy była ukryta za ramką telewizora. Przekątna wskazywała rozmiar wyświetlanego prostokąta, pamiętając, że oryginalny współczynnik proporcji telewizora wynosił 4:3 (1,33:1).

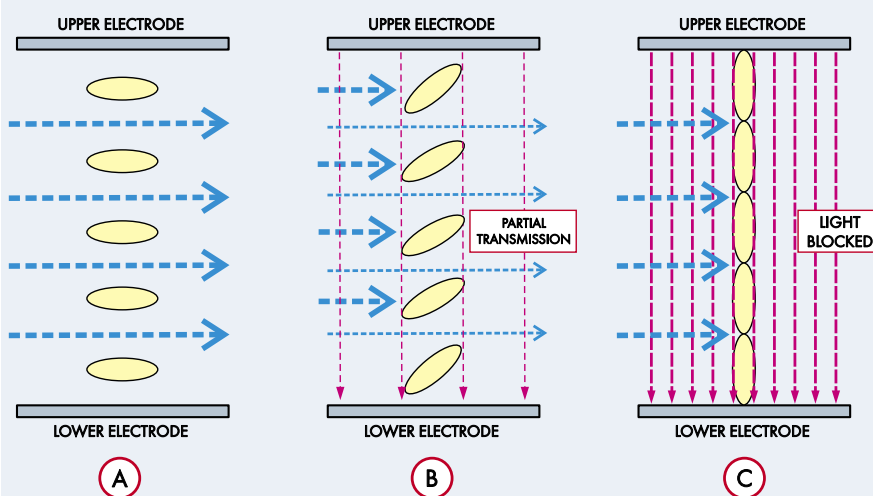
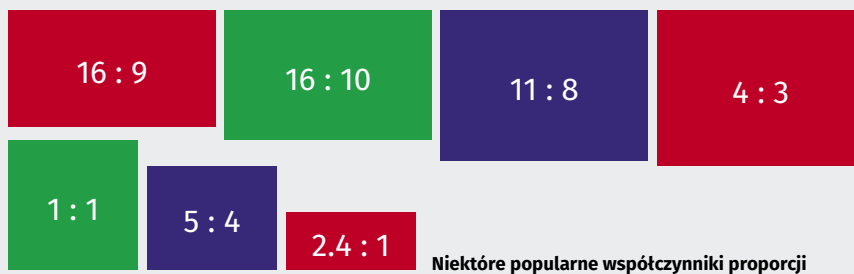
W przypadku płaskich wyświetlaczy przekątna odnosi się do rzeczywistego widocznego obszaru. Filmy są dostępne w wielu proporcjach, ale najpopularniejszym współczynnikiem proporcji telewizora, monitora komputerowego i smartfona jest 16:9 (1,78:1). Jednak niektóre smartfony przekroczyły ten współczynnik, stając się wyższe.

Współczynnik proporcji 16:9 jest standardem Międzynarodowego Związku Telekomunikacyjnego od 1990 roku. Standardowe rozdzielczości HDTV, takie jak 1280×720, 1920×1080 i UltraHD 3840×2160, mają proporcje 16:9, gdy piksele są kwadratowe.

Aby pomieścić inne proporcje materiału źródłowego na ekranie 16:9, obraz jest przycinany lub stosowany jest „letterbox” (czarne paski na górze i na dole), „pillarbox” (czarne paski po bokach) lub, w niektórych przypadkach, „windowbox” z czarną przestrzenią wokół obrazu.

Standardowy współczynnik proporcji filmu ustalony przez Academy of Motion Picture Arts and Sciences to 11:8 (1,375:1), ale filmy były i nadal są produkowane w szerokiej gamie współczynników proporcji, z ultraszerokokątnym 2,35:1, który był dość popularny przez wiele lat w filmach fabularnych. W przypadku monitorów komputerowych również dość powszechnym współczynnikiem proporcji jest 16:10 (jest bardzo zbliżony do złotego podziału, 1,618:1), a 5:4 było również używane w przeszłości (i czasami nadal jest). **Red: „Złoty podział” to matematyczna proporcja, uznawana za estetycznie przyjemną. Jest to stosunek dwóch liczb, równy około 1,618.**

Więcej informacji na temat współczynników proporcji obrazu telewizyjnego i filmowego można znaleźć na stronie https://widescreen.org/aspect_ratios.shtml, a współczynników proporcji obrazu monitora komputerowego na stronie <https://www.wiki/5HTF>.



Rysunek 27. Przejście Fréedericksza jest podstawą technologii wyświetlaczy LCD. Kształty pokazują ułożenie ciekłych kryształów w odpowiedzi na pole elektryczne: a) brak przyłożonego pola elektrycznego, światło przepuszczone; b) przyłożone pośrednie pole elektryczne, światło częściowo przepuszczone; c) przyłożone pełne pole elektryczne, całe światło zablokowane

firma Merck opracowała wyświetlacz typu „viewer independent panel” (panel niezależny od widza), który stał się podstawą wszystkich wyświetlaczy LCD z aktywną matrycą.

W 1982 roku pierwszy telewizor LCD został wypuszczony przez Seiko Epson w formie zegarka na rękę. W 1984 roku firma Citizen wypuściła kolorowy kieszonkowy wyświetlacz LCD o przekątnej 2,7 cala (6,8 cm), który jako pierwszy wykorzystywał aktywną matrycę lub wyświetlacz TFT (tranzystor cienkowarstwowy).

Wyświetlacz LCD były jedną z pierwszych technologii zastępujących telewizory CRT

i ekrany plazmowe. Wczesne wyświetlacze plazmowe mogły generować większy obraz niż wyświetlacze LCD, ale charakteryzowały się niską jasnością i wysokim zużyciem energii.

W 1988 roku Sharp wyprodukował wysokiej klasy 14-calowy (36 cm) monitor LCD, natomiast Epson wypuścił kolorowy projektor LCD, VPJ-700, w styczniu 1989 roku. Badania nad ekranami LCD były kontynuowane i ostatecznie ekrany LCD mogły być produkowane w rozmiarach konkurencyjnych do wyświetlaczy plazmowych. Dzięki temu mogły być używane zarówno na rynku małych rozmiarów (gdzie wyświetlacze plazmowe nie nadawały się do użycia),

jak i na rynku dużych rozmiarów, gdzie wyświetlacze plazmowe dominowały.

W 1994 roku na targach w Japonii zademonstrowano ekran LCD o przekątnej 21 cali (53 cm). Pod koniec lat 90. zademonstrowano prototypowe wyświetlacze o przekątnej 40 cali/1 m.

W 1995 roku firma Hitachi Ltd opracowała technologię „in-plane switching” (IPS), zapewniającą znacznie szerszy kąt widzenia niż istniejąca technologia TN (twisted nematic) bez nadmiernych zmian kolorów lub jasności.

Następnie, w 1997 roku, firma Fujitsu Ltd wyprodukowała wyświetlacz LCD z technologią „wyrównania pionowego” (VA), która zapewniała znacznie lepszy kontrast i głębszą czerń, gdy nie było podłączone napięcie.

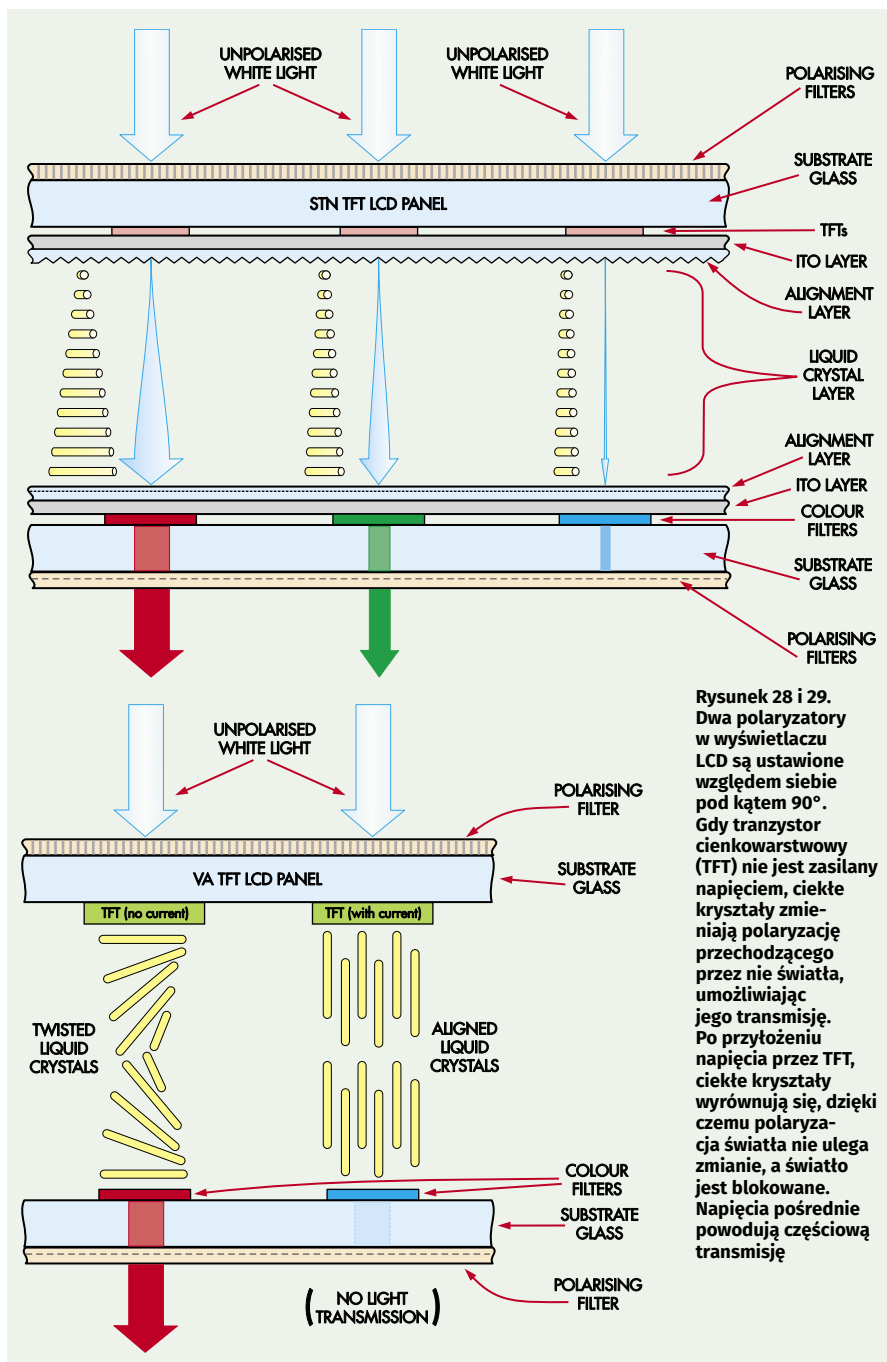
Większość dzisiejszych ekranów LCD nadal wykorzystuje technologię TN, IPS lub VA. TN jest używana głównie tam, gdzie wymagany jest bardzo szybki czas reakcji, ponieważ ma gorsze odwzorowanie kolorów i kąty widzenia. IPS zapewnia najlepsze kąty widzenia i odwzorowanie kolorów, ale jego kontrast nie jest tak wysoki jak VA, więc czernie mogą wyglądać na szare.

W 2000 roku opracowano nowe materiały ciekłokrystaliczne o znacznie skróconym czasie reakcji do 8 ms i jeszcze szerszych kątach widzenia dla wyświetlaczy VA o lepszych kolorach, jasności i kontraście. W 2006 roku Sharp opracował technologię VA stabilizowaną polimerami, która zapewnia lepszą transmisję światła, a tym samym niższe zapotrzebowanie na energię do podświetlenia.

W 2006 roku ceny ekranów LCD zaczęły drastycznie spadać i zaczęły wypierać z rynku wyświetlacze plazmowe, a telewizory LCD zaczęły wyprzedzać te plazmowe. Do 2008 roku telewizory LCD wyprzedziły również telewizory CRT.

Zasady działania wyświetlacza matrycowego LCD są dość proste, jak pokazano na rysunkach 28 i 29. Liniowe filtry polaryzacyjne, stosowane w niektórych aparatach i okularach przeciwsłonecznych, zapewniają jednolitą polaryzację światła w jednym kierunku. Światło jest przepuszczane normalnie, jeśli dwa liniowe filtry polaryzacyjne są ustawione w jednej linii. Jeśli jednak zostaną obrócone względem siebie o 90°, światło zostanie zatrzymane. Dlatego sterując polaryzacją jednej z dwóch warstw, ilość przechodzącego światła może być kontrolowana płynnie, od blisko 100% do blisko 0%.

W wyświetlaczu LCD warstwa ciekłych kryształów jest umieszczona pomiędzy dwoma skrzyżowanymi polaryzatorami. Pomiedzy polaryzatorami znajdują się również przezroczyste elektrody wykonane z tlenku indowocynowego, z warstwą wyrównującą i kolorowymi filtrami (dla kolorowych wyświetlaczy LCD) reprezentującymi kolory subpikseli. Cały zespół nazywany jest „kanapką”.



Rysunek 28 i 29. Dwa polaryzatory w wyświetlaczu LCD są ustawione względem siebie pod kątem 90°. Gdy tranzystor cienkowarstwowy (TFT) nie jest zasilany napięciem, ciekłe kryształy zmieniają polaryzację przechodzącego przez nie światła, umożliwiając jego transmisję. Po przyłożeniu napięcia przez TFT, ciekłe kryształy wyrównują się, dzięki czemu polaryzacja światła nie ulega zmianie, a światło jest blokowane. Napięcia pośrednie powodują częściową transmisję

Warstwy wyrównujące składają się z dwóch płytek poliimidowych, po jednej z każdej strony ciekłych kryształów, które zostały poddane obróbce w celu wyrównania z nimi ciekłych kryształów. Każda płytka jest ustawiona pod kątem prostym do drugiej. Co zaskakujące, jedną z metod tworzenia wzoru wyrównania jest pocieranie płytki aksamitną szmatką w pożądanym kierunku.

Gdy do ciekłego kryształu nie jest doprowadzany prąd, wyrównanie przez grubość kryształu zmienia się z kierunku jednej płytki na kierunek drugiej. Powoduje to zmianę polaryzacji światła z jednego ustawienia na drugie, a tym samym transmisję światła.

Jeśli przez ciekłe kryształy zostanie przyłożone napięcie za pośrednictwem zwykłych elektrod lub tranzystorów cienkowarstwowych (TFT) w podstawie każdego elementu piksela wyświetlacza, ciekłe kryształy wyrównują się i blokują światło. Stopień blokowania zależy od przyłożonego napięcia.

Wcześniejsze ekrany LCD były typu „pasywnej matrycy” z elektrodami po obu stronach warstwy LCD. Nowsze wyświetlacze to „aktywne matryce”, w których elektrody dla każdego elementu subpiksele są zastąpione cienkowarstwowymi (półprzezroczystymi) tranzystorami, co skutkuje szybszym czasem reakcji oraz ostrzejszym i jaśniejszym obrazem.

Źródłem światła dla paneli LCD przez długi czas były zimnokatodowe lampy fluorescencyjne (CCFL), ale obecnie są to głównie diody LED. Dodatkowe uwagi na temat tego rozróżnienia znajdują się w wydzielonej sekcji na końcu artykułu.

Nawiasem mówiąc, można stwierdzić, czy okulary przeciwsłoneczne są polaryzacyjne, czy nie, patrząc na działający ekran LCD i obracając je. Jeśli ekran przyciemni się lub zgaśnie pod pewnym kątem, okulary mają soczewki polaryzacyjne.

Ekran quantum-dot

Ekran quantum-dot (zawierające kropki kwantowe) dzieli się na dwa rodzaje: fotoemisyjne i elektroemisyjne. Są one formą nanotechnologii.

Fotoemisyjne quantum-dot są stosowane w każdej technologii wyświetlania, w której stosowane są filtry kolorów, głównie w wyświetlaczach LCD z podświetleniem LED. W wyświetlaczu LCD są one umieszczane jako folia w „kanapce” wykonanej z innych folii, polaryzatorów, szkła, TFT i elektrod. Gdy światło przechodzi przez folię z punktami kwantowymi, jest ponownie emitowane jako czysty kolor czerwony, zielony lub niebieski.

Ma to na celu zapewnienie bardziej realistycznych kolorów niż jest to możliwe



Rysunek 30. Wyświetlacz Sony Crystal LED (CLEDIS) tworzy ściany na tym zdjęciu. Wyświetlacze są modułowe, więc mogą być wykonane w zasadzie w dowolnym rozmiarze. Źródło: https://pro.sony/en_PT/products/led-video-walls/crystal-led-walls

w przypadku samego oświetlenia LED. Mówi się, że ekrany LCD stosujące punkty kwantowe są porównywalne lub lepsze od wyświetlaczy OLED (organicznych diod elektroluminescencyjnych). Wyświetlacze z punktami kwantowymi są jednak tańsze i mogą zapewnić lepsze kolory przy pełnej jasności niż OLED.

Elektroemisyjne wyświetlacze quantum-dot same emitują światło, ale na tym etapie są eksperymentalne. Są to cienkie, elastyczne wyświetlacze, które odznaczają się lepszą żywotnością niż OLED.

Wyświetlacze LED i microLED

Wyświetlacze LED to płaskie wyświetlacze panelowe składające się z pojedynczych diod LED tworzących subpiksele, które są rzeczywistymi elementami emitującymi światło. Nie należy ich mylić z wyświetlaczami LCD, które wykorzystują podświetlenie LED (patrz panel). Wyświetlacze LED są używane do dużych ekranów zewnętrznych, takich jak podczas wydarzeń sportowych, rozrywkowych lub interaktywnych znaków drogowych.

Diody MicroLED są produkowane w mniejszym rozmiarze niż standardowe diody LED i dlatego nadają się do mniejszych urządzeń wyświetlających (lub urządzeń

o wyższej rozdzielczości) niż zwykle diody LED. Wyświetlacze te są nieorganiczne i teoretycznie mają dłuższą żywotność niż diody OLED, które mają pewne ograniczenia (jak wyjaśniono poniżej).

W porównaniu do wyświetlaczy LCD, mają potencjalnie szybszy czas reakcji, niższe zużycie energii, większą jasność, lepszy współczynnik kontrastu i lepsze nasycenie kolorów.

Nie były one jeszcze produkowane masowo dla mniejszych urządzeń, takich jak telewizory konsumenckie, ale Sony opracowało CLEDIS lub Crystal LED Integrated Structure, w którym zastosowano MicroLED. Jest to modułowy system, który umożliwia zbudowanie wyświetlacza o niemal dowolnym rozmiarze. Mogłyby być on stosowany do obsługi wystaw publicznych lub jako ekran kinowy (rysunek 30).

W styczniu tego roku Samsung ogłosił plany sprzedaży telewizorów microLED w rozmiarach 89 cali (2,25 m), 101 cali (2,5 m) i 110 cali (2,75 m), ale w chwili pisania tego tekstu nie są one jeszcze dostępne na rynku.

OLED

OLED to skrót od organicznej diody elektroluminescencyjnej (Organic Light-Emitting Diode). W przeciwieństwie do tradycyjnych

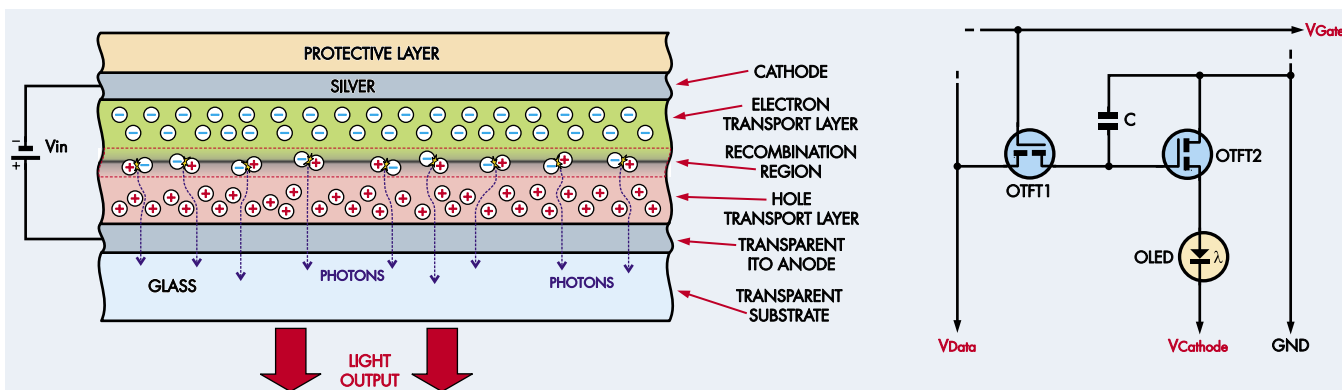
Niedziałające lub uszkodzone piksele w wyświetlaczach

W wyświetlaczach matrycowych, takich jak ekrany plazmowe, LCD i OLED, istnieje możliwość otrzymania ekranu z niedziałającymi pikselami (zwanymi również „martwymi pikselami”). Możliwe usterki obejmują piksele lub subpiksele, które są włączone lub wyłączone.

Międzynarodowa norma ISO 13406-2 została opracowana w celu kategoryzacji typów i ilości defektów pikseli, które są uważane za dopuszczalne. Liczba dopuszczalnych defektów różni się w zależności od producenta. Zależy ona od rodzaju defektów, lokalizacji wadliwych pikseli na ekranie i odległości wadliwych pikseli od siebie.

Źródło obrazu: <https://w.wiki/5JET>





Rysunek 31. Zasada działania piksela ekranu OLED. Jest on nieco podobny do zwykłej diody LED, ale wykorzystuje organiczne półprzewodniki polimerowe. Oznacza to między innymi, że ekrany OLED mogą być elastyczne

diod LED, które są wykonane z nieorganicznych półprzewodników, takich jak azotek galu, diody OLED są wykonane z organicznych półprzewodników. Są to złożone materiały organiczne oparte na małych cząsteczkach lub cząsteczkach połączonych ze sobą jako polimery (tworzywa sztuczne).

Wszystkie te materiały charakteryzują się luźno związanymi elektronami, co umożliwia im przewodzenie prądu elektrycznego w różnym stopniu. Są one znane jako przewodniki organiczne. Warstwa aktywna (obszar rekombinacji) diody OLED jest elektroluminescencyjna, co oznacza, że emituje światło w odpowiedzi na przyłożone napięcie.

Elektroluminescencję w materiałach organicznych zaobserwowano w latach 50. XX wieku, a podstawowe badania przeprowadzono w latach 60., ale Eastman Kodak opracował pierwsze praktyczne diody OLED w 1987 roku.

Białe diody OLED zostały po raz pierwszy wyprodukowane i skomercjalizowane w Japonii w 1995 roku do podświetlania wyświetlaczy i innych celów oświetleniowych.

W 1999 roku Kodak i Sanyo nawiązały współpracę i wyprodukowały 2,4-calowy (61 mm) wyświetlacz OLED, a następnie 15-calowy (38 cm) ekran HDTV w 2002 roku. W 2007 roku Sony wypuściło na rynek telewizor XLE-1, a w 2017 roku firma JOLED rozpoczęła produkcję paneli OLED drukowanych w procesie atramentowym.

Prosta struktura OLED składa się z warstwy ochronnej, katody (–), warstwy transportu elektronów, obszaru rekombinacji, warstwy transportu dziur, przezroczystej anody (+) i szklanego podłoża (**rysunek 31**). Bardziej zaawansowane diody OLED mają dodatkowe warstwy z różnymi obszarami w celu uzyskania różnych kolorów.

OLED do tego, by zacząć działać wymaga prostej różnicy potencjałów (napięcia). Katoda przyjmuje elektrony (–) ze źródła zasilania, a anoda traci dziury (brak elektronu, +). Przeciwnie ładunki są przyciągane do siebie i spotykają się w obszarze rekombinacji, w obszarze granicznym między warstwą transportu elektronów i warstwą transportu dziur.

Te elektrony i dziury stykają się, tworząc „ekscyton” i emitując foton światła. Dzieje się tak wiele razy, powodując ciągłą emisję światła.

Wadą wyświetlaczy OLED jest ich krótsza żywotność niż ma to miejsce w przypadku innych technologii. Zaletą jest to, że mogą być składane, jak w niektórych telefonach (**rysunek 32**).

AMOLED to szczególna technologia OLED, w której jest zastosowana aktywna matryca



Rysunek 32. Smartfony Samsung ze składanymi wyświetlaczami OLED. Docierają doniesienia o pękaniu tych ekranów po wielu miesiącach lub latach składania i rozkładania, więc przed zakupem warto zrobić rozeznanie, zwłaszcza że są one drogie. Źródło: Wikimedia Ka Kit Pang, licencja Apache 2.0



Rysunek 33. Przykłady wyświetlaczy elektroluminescencyjnych Lumineq z opcjonalnym ekranem dotykowym. Na górnym zdjęciu wyświetlana jest cena taksówki lub Ubera, podczas gdy na dolnym zdjęciu przedstawiono panel z kodem dostępu do samochodu

sterowana przez tranzystory cienkowarstwowe (TFT).

Wyświetlacze elektroluminescencyjne

Elektroluminescencja (EL) to zjawisko, w którym materiał, taki jak arsenek galu emituje światło po przyłożeniu do niego pola elektrycznego. Kolor światła różni się w zależności od aktywnego materiału, ale obecnie jedyne praktyczne wyświetlacze są jednokolorowe – żółte lub pomarańczowe. Wyświetlacze mogą mieć stałe segmenty lub matrycę do wyświetlania dowolnego obrazu.

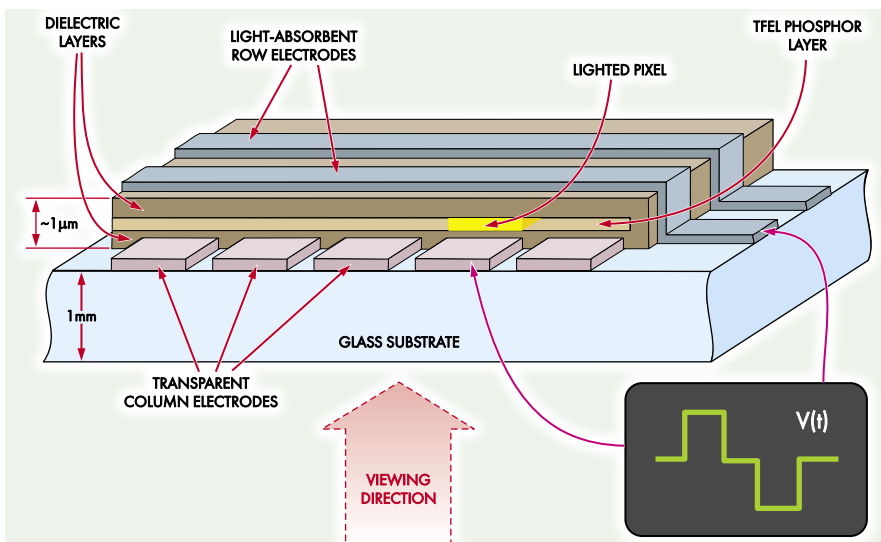
Struktura wyświetlacza jest podobna do wyświetlaczy LCD lub OLED z pasiastymi nieprzezroczystymi (lub przezroczystymi) elektrodami z tyłu, biegnącymi w jednym kierunku i przezroczystymi pasiastymi elektrodami z przodu pod kątem prostym do tych z tyłu. Jedna tylna elektroda i jedna przednia elektroda są zasilane w celu aktywacjiżądanego segmentu lub piksela (**rysunek 35**).

Istnieją dwa główne typy wyświetlaczy EL, przezroczyste lub nieprzezroczyste. Są one podobne, ale przezroczyste wyświetlacze mają przezroczyste tylne elektrody.

W przypadku wyświetlaczy przezroczystych, obszary, które nie są aktywowane, są przezroczyste w 70% w przypadku wyświetlaczy matrycowych i w 80% w przypadku wyświetlaczy segmentowych. Mogą być laminowane w szkło, takim jak szkło samochodowe, a także mogą mieć funkcję wykrywania dotyku.

Wyświetlacze elektroluminescencyjne są wytrzymałe, mogą pracować w wysokich lub niskich temperaturach, są odporne na wysokie lub niskie ciśnienie i światło słoneczne, a ich żywotność wynosi co najmniej 20 lat. Dlatego są one lepsze od wyświetlaczy LCD i OLED w niektórych zastosowaniach, takich jak praca na zewnątrz.

Beneq z Finlandii jest jedynym producentem segmentowych i matrycowych wyświetlaczy



Rysunek 35. Struktura elektroluminescencyjnej matrycy (piksela) wyświetlacza. Oryginalne źródło: Electronics Weekly – siliconchip.au/link/abfd

elektroluminescencyjnych, które są sprzedawane pod marką Lumineq (www.lumineq.com – **rysunek 33**).

Cyfrowe przetwarzanie światła (DLP)

DLP to technologia projekcji światła opracowana przez Texas Instruments (TI) w 1987 roku i skomercjalizowana w projektorze przez Digital Projection Ltd. Wykorzystuje ona chip z układem mikroluster. Można je ustawić w pozycji „włączonej”, aby odbijały światło w kierunku płaszczyzny obrazu, lub w pozycji „wyłączonej”, aby odbijały światło w innym miejscu, na przykład na radiator.

Chociaż lustro mogą znajdować się tylko w jednej z dwóch pozycji, można uzyskać pośrednie jasności poprzez szybkie włączanie lub wyłączanie luster, aby zmienić średnią ilość światła wysyłanego do płaszczyzny obrazu.

Układ ten jest znany jako cyfrowe urządzenie mikrolusterkowe lub DMD (**rysunek 34**).

Lustra są mikroskopijnie małe, o rozstawie 5,4 μm (mikronów, milionowych części metra) lub mniejszym. Liczba luster odpowiada rozdzielczości obrazu, z wyjątkiem sytuacji, gdy do zwiększenia efektywnej rozdzielczości wykorzystywany jest proces znany jako wobulacja.

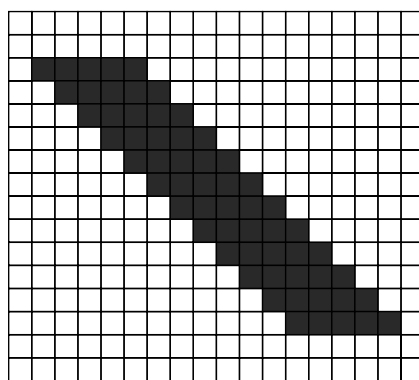
W przypadku wobulacji (**rysunek 36**) DMD jest przesuwany o niewielką wartość (w obu kierunkach X i Y), na przykład o pół piksela, w celu wyświetlenia nowej podramki. Jest ona generowana przez oprogramowanie sprzętowe projektora i w połowie nakłada się na poprzednią klatkę, zapewniając wzrost rozdzielczości bez dodatkowych kosztów związanych z DMD o wyższej rozdzielczości.

Kolory są generowane albo przez koło kolorów obracające się przed chipem, tworząc serię różnych kolorowych obrazów, które oko łączy, albo przez trzy oddzielne chipy, z których każdy wyświetla jeden kolor podstawowy.

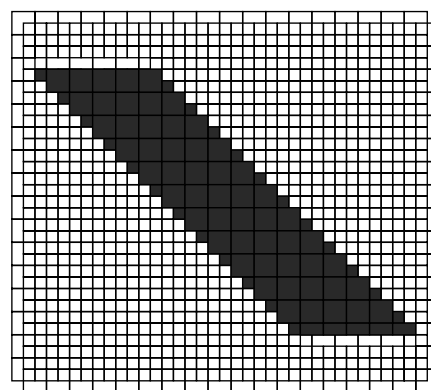
DMD to optyczny MEMS (system mikroelektromechaniczny) – w listopadowym



Rysunek 34. Przód układu DMD firmy Texas Instruments do zastosowań kinowych. Źródło: Wikimedia user Binant, CC BY-SA 4.0

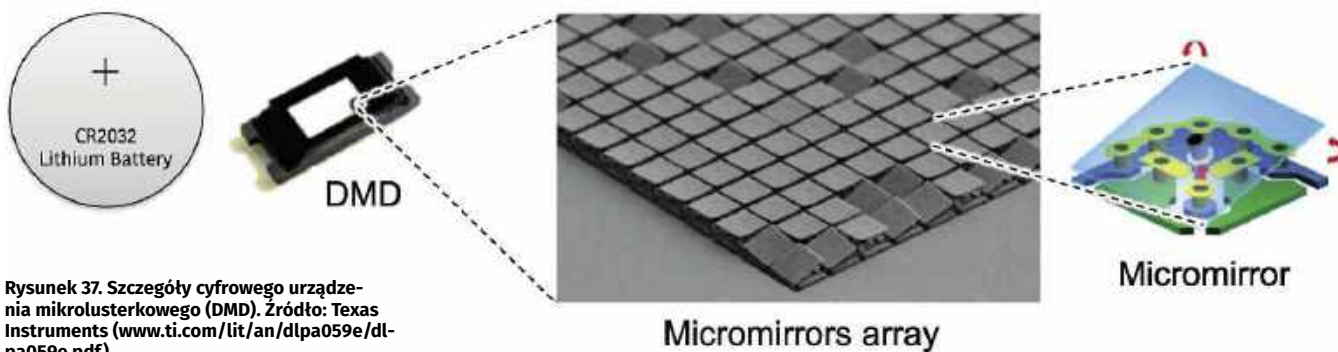


(A) NON-WOBULATED DMD PATTERN



(B) WOBULATED DMD PATTERN

Rysunek 36. Obrazy bez wobulacji i z wobulacją generowane przez DMD. Wobulacja poprawia widoczną rozdzielczość bez konieczności stosowania większej liczby luster



Rysunek 37. Szczegóły cyfrowego urządzenia mikrolusterkowego (DMD). Źródło: Texas Instruments (www.ti.com/lit/an/dlpa059e/dlpa059e.pdf)

numerze Siliconchip z roku 2020 zamieszczony był artykuł na ten temat (siliconchip.au/Article/14635).

W DMD tysiące mikroskopijnych zwierciadeł aluminiowych opiera się na jarzmie, które z kolei opiera się na zawiasie skrętnym między dwoma słupkami i obraca się o około 10° między pozycjami włączenia i wyłączenia za pomocą sił elektrostatycznych, jak pokazano na **rysunku 37**.

Podstawowa warstwa DMD zawiera komórki SRAM (statyczna pamięć o dostępie swobodnym), które przesuwają jedno lustro za pomocą ładunku elektrostatycznego zgodnie z jego aktualnym stanem. Napięcie polaryzujące jest używane do sterowania pamięcią SRAM, dzięki czemu po odłączeniu zasilania wszystkie lustra ustawiają się do tej samej pozycji początkowej, więc wszystkie lustra poruszają się razem w następnej klatce (**rysunek 38**).

Ze względu na obszerne portfolio patentów, wysokie koszty produkcji i wysoki poziom

wymaganej wiedzy technicznej, urządzenia te produkuje tylko Texas Instruments.

DMD jest produkowany zgodnie ze standardowymi procesami MEMS i litografii, które zostały opisane w trzyczęściowej serii na temat produkcji układów scalonych w Siliconchip od czerwca do sierpnia 2022 r. (siliconchip.au/Series/382). Z pewnością dokładne procesy są ściśle strzeżoną tajemnicą. Mimo to chcielibyśmy je poznać!

Technologia DLP jest stosowana w niektórych projektorach domowych i około 90% komercyjnych projektorów filmowych. TI oferuje rozdzielczości DMD do 4K UHD (3840×2160) i częstotliwości odświeżania od 60 Hz do 240 Hz z obsługą źródeł światła LED, żarowych lub laserowych.

Wideo z rozbiórki wczesnego projektora DLP jest pokazana na filmie „Extreme teardown – NEC XT5000 Projector” na stronie <https://youtu.be/RzikiKqBA1U>.

Laser TV

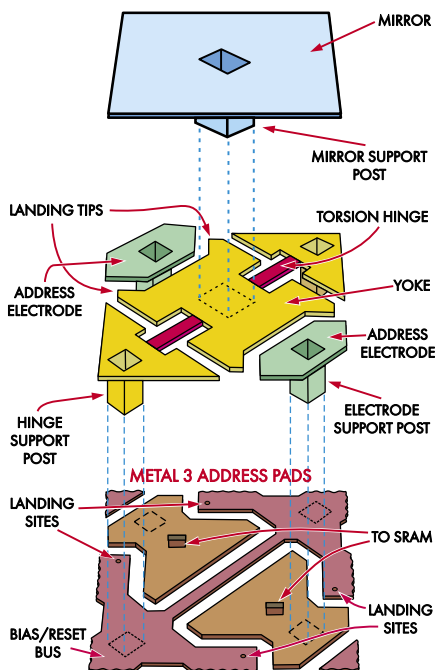
Telewizja laserowa to nowa technologia, która jest obecnie w fazie wdrażania. Aby wygenerować obraz, wiązki laserowe są skanowane

w poprzek płaszczyzny obrazu, zwykle elektromechanicznie, jak w przypadku chipu DLP. Koncepcyjnie, obraz jest tworzony podobnie jak w CRT, ale przy użyciu wiązki laserowej zamiast wiązki elektronów – **rysunek 39**.

Pomysł telewizji laserowej został po raz pierwszy zaproponowany w 1966 roku i opatentowany w 1977 roku, ale technologia laserowa była zbyt droga do czasu opracowania laserów półprzewodnikowych. System został zademonstrowany na targach elektroniki użytkowej (CES) w Las Vegas w 2006 roku przez firmę Novalux Inc. W 2008 roku Mitsubishi Electric wypuściło komercyjny 65-calowy (165 cm) model HDTV 1080p, a w 2013 roku LG wypuścił 100-calowy (2,5 m) model konsumencki 1080p.

Elektroniczny papier i atrament

Papier elektroniczny to rodzaj wyświetlacza, który naśladuje papier. Podobnie jak papier, nie wytwarza własnego światła, ale jest odczytywany przez odbite światło otoczenia. W związku z tym powoduje mniejsze zmęczenie oczu i stres. Elektroniczny papier może być odświeżany dość szybko, ale



Rysunek 38. Szczegóły poszczególnych zespołów luster w DMD. Oryginalne źródło: Texas Instruments



Rysunek 39. Dostępny w sprzedaży telewizor laserowy Hisense. Obraz jest wyświetlany z urządzenia pod ekranem



Rysunek 40. Elektroniczny papierowy rozkład jazdy w czasie rzeczywistym używany w autobusach w Sydney. Źródło: Wikimedia user MDRX, CC BY-SA 4.0

przy użyciu obecnej technologii nie wystarczająco szybko, aby wyświetlać wideo w pełnym ruchu. Może jednak wyświetlać filmy w zwolnionym tempie lub często zmieniające się liczby, takie jak na wyświetlaczu zegara.

Podobnie jak zwykły papier, papier elektroniczny zachowuje ostatni zapisany na nim obraz po wyłączeniu zasilania. Do utrzymania wyświetlacza w bieżącym stanie nie jest wymagane zasilanie.

Inne nazwy papieru elektronicznego to atrament elektroniczny i wyświetlacze elektroforetyczne. Nazwa „E Ink” jest znakiem

Ekran LCD: IPS, VA czy TN?

Jak wspomniano w tekście, są to trzy dominujące technologie LCD, choć istnieją też inne. Przy wyborze ekranu LCD jest to jedna z najważniejszych decyzji.

Podczas gdy nowoczesne panele VA (wyrównane w pionie) mają przyzwite kąty widzenia, z doświadczenia wynika, że panele IPS są nadal zauważalnie lepsze. Jest to szczególnie ważne w przypadku monitorów komputerowych, gdzie zazwyczaj siedzi się blisko ekranu. Staby kąt widzenia nie tylko oznacza, że nie można zbyt poruszać głową, ale nawet przy statycznej pozycji głowy, rogi ekranu mogą wydawać się wyblakłe lub zmieniać kolory w porównaniu do środka.

Z tego powodu używamy prawie wyłącznie paneli IPS (przełączanych w płaszczyźnie). Mają one również zwykle najlepsze odwzorowanie kolorów, choć ekrany VA przeszły długą drogę również pod tym względem.

Niektórzy preferują panele VA do odtwarzania wideo/TV lub grania w gry ze względu na wyższy współczynnik kontrastu, głębszą czerń i szybsze częstotliwości odświeżania. Jednak obecnie dostępne są ekrany IPS z odświeżaniem 144 Hz, co sprawia, że rozróżnienie częstotliwości odświeżania staje się mniej krytyczne. Panele VA mają zauważalnie lepszy kontrast niż IPS, ale kompromis nie jest opłacalny, chyba że mają znakomite kąty widzenia w swojej klasie.

Jest to sytuacja, w której naprawdę pomocne jest fizyczne wypróbowanie produktu przed jego zakupem. Upewnimy się w ten sposób, że jego odwzorowanie kolorów, jasność, kontrast i kąty widzenia są zgodne z naszymi upodobaniami.

Jedynym powodem, dla którego warto kupić ekran TN (twisted nematic), jest chęć uzyskania bardzo wysokiej częstotliwości odświeżania, np. 240 Hz lub wyższej. I w tym przypadku kompromis raczej nie jest tego wart, ponieważ obraz wygląda o wiele gorzej, ale niektórzy ludzie naprawdę lubią te wysokie częstotliwości odświeżania do gier. W takim przypadku TN jest w zasadzie jedynym wyborem.

towarowym firmy E Ink Corporation (www.eink.com). Czytnik książek elektronicznych Kindle jest popularnym zastosowaniem technologii papieru elektronicznego.

Przykłady zastosowań obejmują czytniki książek elektronicznych, aktualizowane wyświetlacze cen w sklepach, oznakowanie elektroniczne, rozkłady jazdy transportu publicznego, identyfikatory konferencyjne, niektóre smartfony i tablety – **rysunek 40**.

Papier elektroniczny został wynaleziony w Xerox Palo Alto Research Center (PARC) w latach 70. i nosił nazwę Gyricon (**rysunek 41**).

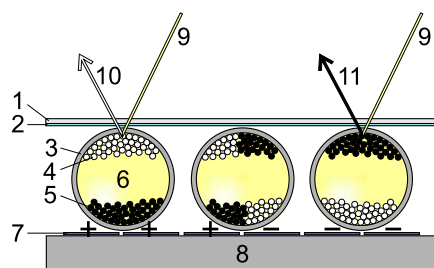


Rysunek 41. Xerox Gyricon, pierwszy papier elektroniczny. Źródło: Strona internetowa Xerox zarchiwizowana z 2005 r.

Zgodnie z pierwotnymi założeniami, papier elektroniczny nie posiadał elektrod. Obraz mógł być tworzony poprzez zastosowanie zewnętrznego pola elektrycznego we wzorze, który miał być zapisany, podobnie jak rysowanie piórem. Następnie można go było wymazać i napisać nowy wzór.

Istnieje kilka metod implementacji, ale podstawowa zasada składa się z „częstek Janusa”, pokrytych olejem lub podobnym płynem, aby umożliwić łatwy obrót. Są one osadzone w pewnego rodzaju matrycy, na przykład silikonowej (**rysunek 42**).

Cząsteczka Janusa to sferyczna nano- lub mikrocząsteczka o różnych właściwościach elektrycznych lub innych właściwościach po każdej stronie, takich jak ładunek dodatni



Rysunek 42. Technologia E Ink. 1) Górna warstwa 2) Przezroczysta warstwa elektrody 3) Przezroczyste mikrokapsułki 4) Dodatnio naładowane białe pigmenty 5) Ujemnie naładowane czarne pigmenty 6) Przezroczysty olej 7) Warstwa pikseli elektrody 8) Dolna warstwa nośna 9) Światło 10) Biały pigment 11) Czarny pigment. Grubość wyświetlacza wynosi około 0,5...1 mm. Źródło: Wikimedia user FREEscanRIP, CCA 3.0



Rysunek 43. Projektacja z wykorzystaniem ściany wodnej wykonana przez australijską firmę Laservision podczas wydarzenia w Australii. Źródło: www.laservision.com.au/galleries/photos/

lub ujemny. W przypadku papieru elektronicznego jedna strona kuli może być biała, a druga czarna. Częsteczki wyrównują się z polem, gdy pole elektryczne jest przykładane przez lub w poprzek matrycy (w zależności od orientacji elektrody).

Powoduje to ich obracanie się i wyświetlanie koloru białego, czarnego lub innych kolorów, którymi cząsteczki zostały zabarwione. Po odwróceniu pola elektrycznego cząsteczki obracają się i prezentują swoją drugą stronę. Cząsteczki Janus mają zazwyczaj rozmiar od 10 μm do 50 μm .

W celu uzyskania kolorów można stosować dodatkowe filtry barwne. Alternatywnie, pole elektryczne może kontrolować powłokę kolorowego oleju w tak zwanym procesie elektroprzędzenia. W tym drugim przypadku stosowany jest system kolorów subtraktywnych, podobnie jak w typowej drukarce kolorowej z atramentami CMYK (cyjan/magenta/żółty/czarny).

Wyświetlacze te są dostępne i odpowiednie dla eksperymentatorów. Można je kupić jako zestawy Arduino i Raspberry Pi oraz z interfejsami SPI. Artykuł o korzystaniu z wyświetlaczy e-papierowych z Micromite był opisany

Kiedy telewizor LED nie jest telewizorem LED?

Prawie wszystkie telewizory sprzedawane jako „telewizory LED” są w rzeczywistości telewizorami LCD z białym podświetleniem LED. W starszych telewizorach LCD jako podświetlenie były stosowane lampy fluorescencyjne z zimną katodą (CCFL). Telewizory opisane jako QLED to wyświetlacze LCD z kropkami kwantowymi i podświetleniem LED.

Telewizory OLED generują własne światło i nie potrzebują podświetlenia. Aby uniknąć nieporozumień, dobrze by było, aby branża przyjęła termin „telewizor LCD z podświetleniem LED” zamiast „telewizor LED”, chyba że jest to prawdziwy telewizor LED. Jednak producenci czerpią korzyści z tego zamieszania, sprawiając wrażenie, że podświetlenie LED jest większą zaletą technologiczną niż jest w rzeczywistości.

w Siliconchip z czerwca 2019 r. (siliconchip.au/Article/11668). Warto również obejrzeć filmy na temat wyświetlaczy z atramentem elektronicznym:

- „Czy kiedykolwiek widziałeś tak szybką aktualizację wyświetlacza E Ink?” – <https://youtu.be/KdrMjnYAap4>,
- „Badger 2040 – Raspberry Pi Pico z wbudowanym wyświetlaczem e-Ink” – <https://youtu.be/kI-ksiYw40>,



Rysunek 44 i 45. Dysza z sitem wodnym sprzedawana na stronie <https://fountains-decor.ie/product/water-screen-nozzle/>. Dysza ma wymiary 930 mm × 528 mm × 802 mm i zapewnia półokrągłe sito z dopływu wody 4000 l/min przy ciśnieniu 12 barów. Grubość warstwy wody wynosi 6 mm. Producent nie określił rozmiaru ekranu, jaki może wyprodukować, ale pokazano przykład. Półokrąg ma przypuszczalnie promień około 10 m



- „5 powodów, dla których warto kupić tablet z ekranem e-ink”
– <https://youtu.be/YKjXvjhe-Ss>
- „Bigme Max+ Color EINK 10.3”,
Note Taking Review”
– <https://youtu.be/RAhFzefT5DI>.

Ekran wodne

Ekran wodny to technologia wyświetlania nocnego na dużą skalę, w której obraz jest wyświetlany na ekranie wykonanym z kropelek wody za pomocą lasera lub projektorów wideo. Woda jest rozpylana w powietrzu tworząc wodosпад lub jest pompowana pod wysokim ciśnieniem w celu utworzenia ekranu lub chmury mgły (rysunek 43).

Australijska firma Laservision (www.laservision.com.au) jest liderem w tej dziedzinie. Niestety, nikt z tej firmy nie oddzwonił do autora artykułu przed publikacją, więc nie jest możliwe podanie żadnych dalszych szczegółów poza tym, co znajduje się na ich stronie internetowej. W Siliconchip z sierpnia 1990 roku opublikowany był artykuł na temat Laservision (siliconchip.au/Article/7208). Dostępne są również filmy:

- „Laservision Corporate Showreel”
– <https://youtu.be/cv04MrAJnLM>,
- <https://vimeo.com/271808280>.

Powiązane produkty innych firm można znaleźć na **rysunkach 44 i 45**. Poniższe filmy na ten temat obejmują zarówno domowe, jak i komercyjne ekrany projekcyjne:

- „Domowy ekran do projekcji wody”
– <https://youtu.be/Z7XHAKAUQuA>,
- „Test projekcji na ekranie wodnym 10’,”
– <https://youtu.be/3TPMwv2SmS8>,
- „Kurtyny wodne | Projekcja ekranu wodnego” przez Water Screen
– <https://youtu.be/27YYmowUFno>,
- „Preview 1 | Water Screen Projection”
– <https://youtu.be/tkCNHmVlQBk>.

Wnioski

Chociaż wyświetlacze LCD stanowią znaczący postęp w stosunku do wyświetlaczy plazmowych i CRT, w ciągu najbliższych kilku lat nadal będą wprowadzane ulepszenia.

Wydaje się prawdopodobne, że ostatecznie OLED i MicroLED zastąpią LCD, ale w tej chwili wszystkie są konkurencyjne na swój sposób. Konkurencja będzie napędzać rozwój wszystkich tych technologii w ciągu najbliższych kilku dekad – chyba że pojawi się coś zupełnie nowego. ■

dr David Maddison

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



Samsung ma w Sydney 14-metrowy ekran kinowy LED obsługujący treści HDR. Źródło: <https://news.samsung.com/global/samsung-unveils-the-first-onyx-cinema-led-screen-in-australia>

Wyświetlacze HDR (High Dynamic Range)

High Dynamic Range (HDR) nie jest rodzajem wyświetlacza, ale zestawem standardów zaprojektowanych w celu odzwierciedlenia możliwości nowych technologii wyświetlania. Do czasu HDR, sygnały wideo były projektowane dla monitorów CRT i nie mogły przekazywać informacji wideo, które w pełni wykorzystywałyby możliwości nowoczesnych wyświetlaczy.

Wyświetlacze obsługujące HDR mogą wyświetlać większy zakres kolorów, kontrastu, jasności, bieli i czerni, bardziej żywe kolory, wyższą częstotliwość odświeżania do 120 klatek na sekundę itp.

Jednym z krytycznych aspektów HDR jest jednak współczynnik kontrastu treści, tj. stosunek najjaśniejszych obszarów obrazu do najciemniejszych. Standardowa zawartość ma maksymalny współczynnik kontrastu do około 1000:1, podczas gdy zawartość HDR może przekraczać 5000:1. To lepiej odpowiada możliwościom ludzkiego oka w zakresie rozpoznawania jasnych i ciemnych obszarów na tym samym obrazie.

Jednym z kluczowych postępów w wyświetlaczach HDR było zastąpienie starszej technologii podświetlenia krawędziowego podświetleniem matrycowym LED. Zamiast diod LED rozmieszczonych wokół krawędzi ekranu, matryca białych diod LED znajduje się za nim. Ich jasność można indywidualnie regulować. Dzięki temu niektóre części ekranu mogą być bardzo jasne, podczas gdy inne są przyciemnione, bez „prześwitowania” związanego z podświetleniem o wysokiej jasności. Fakt, że podświetlenie nie jest równomierne, jest kompensowany przez sposób, w jaki kontroler wyświetlacza steruje samym panelem LCD.

Zazwyczaj, im więcej diod LED zastosowano w matrycy podświetlenia, tym lepsze są możliwości wyświetlacza HDR. Wyświetlacze z wieloma diodami LED w podświetleniu są czasami nazywane „mini LED”.

Wyświetlanie treści HDR

HDTV i standardowe dyski Blu-ray wykorzystują 24-bitowy kolor, co daje 16,7 miliona kolorów, ale zawartość HDR wykorzystuje 30 bitów, co daje ponad miliard kolorów.

Wymaga to większej ilości danych, które mogą być zawarte na płycie Ultra HD Blu-ray, choć takie płyty nie będą odtwarzane na standardowych odtwarzaczach. Treści HDR można również przysyłać strumieniowo, ale potrzebne jest do tego wystarczająco szybkie łącze internetowe. Jeśli może ono obsługiwać wideo 4K, powinno być wystarczająco szybkie dla HDR.

HDR ma kilka konkurencyjnych formatów: Dolby Vision (Dolby), HDR10 (UHD Alliance), HDR10+ (Samsung), Hybrid Log-Gamma/HLG (BBC i japońska NHK), Technicolor Advanced HDR i IMAX Enhanced.

Aby móc oglądać treści HDR, telewizor HDR musi obsługiwać konkretną wersję HDR. Urządzenie do strumieniowego przysyłania multimediów może być w stanie przekonwertować jeden kolorystyczny HDR na inny, który telewizor HDR może wykorzystać. Najpopularniejsze schematy to HDR10 i Dolby Vision. Należy pamiętać, że nie wszystkie telewizory 4K obsługują HDR.

Istnieją również różne standardy HDR, przy czym HDR10 jest najpopularniejszy. Inne standardy mogą być bardziej wymagające.

Fotografowie mogą również używać swoich aparatów i oprogramowania do tworzenia zdjęć HDR – siliconchip.au/link/abfe wśród wielu innych artykułów.

Fnirsi SWM-10

Inteligentna zgrzewarka punktowa

Spróbuj przylutować przewód lub metalowy pasek do akumulatora, okaże się, że jest to prawie niemożliwe. Nie wszystko można zlutować. Czego nie da się zlutować, należy zespawać. Idealnym do tego urządzeniem jest zgrzewarka Fnirsi SWM-10.

Akumulatory opanowały nasz świat. Są wszędzie – od pilotów TV do bezprzewodowych elektrycznych narzędzi ogrodowych, rowów elektrycznych, skuterów, itp. Ogniwa w akumulatorach są zgrzewane punktowo za pomocą pasków cienkiego metalu. Przy prawidłowym wykonaniu tworzy to dobre, mocne połączenia bez negatywnego wpływu na żywotność baterii, co miało by miejsce podczas długotrwałego jej nagrzewania w procesie lutowania. Niedogodnością jest jednak to, że trudno jest wymienić martwe ogniwo bez odpowiednich narzędzi. Do takiej naprawy potrzebujesz zgrzewarki punktowej, takiej jak Fnirsi SWM-10 (**fotografia 1**). To przenośna „maszyna” do zgrzewania punktowego w rozmiarze małego multimetru. Jest przeznaczona do budowy i naprawy zestawów akumulatorów i wykonywania wszelkich innych połączeń punktowych arkuszy, pasków (**fotografia 2**) lub przewodów metalowych, a także może służyć jako bank energii USB o pojemności 5000 mAh.

Rozpakowywanie Fnirsi SWM-10

Jeśli kupisz tę zgrzewarkę w sklepie AVT (www.sklep.avt.pl), to w pudełku oprócz samego urządzenia znajdziesz dwa grube kable (8 AWG, ok. 30 cm długości), dwie zapasowe końcówki spawalnicze, kabel USB-A-to-C, rolkę metalowego paska (10 mm szerokości, 0,1 mm grubości) oraz instrukcję obsługi (**fotografia 3**). Nie rozwinęliśmy metalowego paska, ale po zmierzeniu jego obwodu i oszacowaniu liczby zwojów wydaje się, że ma on około 5 m długości.

SWM-10 ma wygląd stonowany, jest cały czarny, z czarnym okienkiem i trzema przyciskami. Dwa gniazda dla kabli spawalniczych znajdują się w lewym górnym rogu. Z tyłu znajduje się składana podstawka. Przycisk zasilania znajduje się w górnej części SWM-10, obok złącza USB i otworu resetowania.

Intuicyjny wyświetlacz

Po włączeniu urządzenia pojawia się 1,8-calowy kolorowy wyświetlacz pokazujący kilka wartości (**fotografia 4**). Są one w rzeczywistości łatwe do zrozumienia, ponieważ składają się z czterech parametrów regulowanych przez użytkownika (czas podgrzewania, długość impulsu, interwał impulsu i kropki) oraz informacji o stanie (napiecie akumulatora, temperatura, dźwięk włączony lub wyłączony, prąd i licznik zgrzewania punktowego). Przyciskami lewo/prawo można przesunąć kursor w górę i w dół, a przyciski góra/dół umożliwiają regulację wybranej wartości (zaznaczonej na żółto). Długie naciśnięcie lewego/prawego przycisku otwiera (i zamyka) menu ustawień. Krótkie naciśnięcie przycisku zasilania powoduje wyświetlenie (i zamknięcie) ekranu ładowania/rozładowania. W tym miejscu można sprawdzić, ile energii pozostało w akumulatorze oraz czy i jak przebiega ładowanie.

Instrukcja jest krótka i wyjaśnia tylko, co oznaczają wartości i parametry oraz do czego służą przyciski i diody LED. Wyjaśniono krótko jak używać urządzenia do spawania – należy naładować baterię przed



1. Ogólny widok urządzenia
Fnirsi SWM-10



2. Urządzenie Fnirsi SWM-10 jest przeznaczone do spawania blach niklowych, żelaznych i ze stali nierdzewnej o grubości do 0,25 mm



5. Tak wygląda spawanie. Naciśnij lekko i odczekaj kilka sekund

3. Zawartość pudełka

spawaniem i dla uzyskania najlepszych rezultatów nie należy zbyt mocno naciskać na pasek niklowy podczas spawania.

Prostota użytkowania

Jak się okazuje, korzystanie z Fnrst SWM-10 jest niezwykle proste i prawdopodobnie dlatego twierdzą, że jest inteligentna. Wszystko, co musisz zrobić, to umieścić dwie końcówki jedna po drugiej na obrabianym przedmiocie (nie naciskaj zbyt mocno) i odczekać dwie sekundy (jak pokazano na fotografii 5). Następnie usłyszysz delikatny dźwięk

kliknięcia i to wszystko. Liczba słyszalnych kliknięć zależy od wartości Dots (od 1 do 5).

Wartości domyślne sprawdziły się w naszym przypadku podczas spawania kawałka paska niklu do ogniwa guzikowego (fotografia 6). Próba oderwania paska wymaga użycia znacznej siły. Aby uzyskać mocniejsze połączenia, można wydłużyć impulsy i zwiększyć liczbę punktów. Należy jednak uważać, ponieważ przy wysokich wartościach zaczynają pojawiać się iskry i dym, a powierzchnia pod spodem może zostać spalona lub uszkodzona w inny sposób. Podczas prac z urządzeniem zalecamy stosowanie okularów ochronnych. Należy również zachować ostrożność podczas pracy z tak cieką blachą, ponieważ bardzo łatwo jest się nią skaleczyć.

Doskonałe narzędzie

Fnrst SWM-10 to doskonałe narzędzie dla każdego, kto miał do czynienia z akumulatorem z rozładowanym ogniwem i nie wiedział, jak go naprawić. Akumulatory są zazwyczaj kosztowne. Nawet naprawa tylko jednego jest prawdopodobnie wystarczająca, aby uzasadnić zakup tego niedrogiego narzędzia. Oczywiście, gdy już je masz, możesz również montować własne zestawy akumulatorów lub używać zgrzewania punktowego do innych zastosowań. Nie zapominajmy, że jest to również power bank o pojemności 5000 mAh. ■

Redakcja EdW



4. Kolorowy wyświetlacz



6. Spawanie taśmy niklowej do ogniwa guzikowego



Migające diody LED i śliniacy się inżynierowie (13)

Ojej! Nie mogę uwierzyć, że jesteśmy już w 13 części tej wspaniałej mini-mega-serii. Hmmm, „trzyście”. Mam przyjaciela, który odmawia wychodzenia z domu w piątek trzynastego, chociaż są też i tacy, którzy twierdzą, że to ich szczęśliwa liczba.

Jesteś zdezorientowany?

Gdy mamy piksele (w postaci trójkolorowych diod LED, w naszym przypadku) ułożone w pierścienie, jedną z rzeczy, które często chcemy zrobić, jest cykliczne włączanie nowego piksela i wyłączanie starego piksela w tym samym czasie. W poprzednich odcinkach mówiliśmy o różnych technikach, których możemy użyć, aby to osiągnąć. Z przykrością muszę jednak stwierdzić, że kilku czytelników wysłało do mnie e-maile z informacją, że nadal są zaskoczeni, zdezorientowani i zakłopotani – być może dlatego, że nie mogliśmy powstrzymać się od wprowadzenia i zbadań rogu obfitości alternatywnych podejść w trakcie naszej wędrówki. Poświęćmy więc chwilę, aby zebrać wszystkie te taktiki w jednym miejscu i zrobimy wszystko, co w naszej mocy, aby je zdemistyfikować.

W przypadku mojego projektu Silnika Progностycznego pierścienie mają po 16 pikseli. Jednakże, jak zauważyłem w poprzednim odcinku (EdWo8/2024), kiedy próbuję coś ogarnąć, najpierw lubię uprościć problem, co – w tym przypadku – oznacza zmniejszenie liczby pikseli. Mając wybór, chcemy znaleźć ogólne rozwiązanie, które może działać dla dowolnej liczby pikseli, więc zazwyczaj wybieram małą liczbę pierwszą, wychodząc z założenia, że jeśli nasze rozwiązanie działa w tym przypadku, to będzie działać we wszystkim.

Tak więc, wyłącznie na potrzeby naszych eksperymentów myślowych, założmy, że mamy 5-pikselowy pierścień i że nasze piksele są ponumerowane zgodnie z ruchem wskazówek zegara, zaczynając od 0 na górze (**rysunek 1a**). Czasami możemy chcieć ustawić pojedynczy piksel ścigający się wokół pierścienia w kierunku zgodnym z ruchem wskazówek zegara (**rysunek 1b**); innym razem możemy zdecydować się na obrót w kierunku przeciwnym do ruchu wskazówek zegara (**rysunek 1c**).

Aby zachować ogólność tych dyskusji, założmy, że liczba pikseli, którymi gramy, jest zdefiniowana jako `NUM_P`, która w tym przykładzie będzie wynosić 5. Ponadto, ponieważ nasze piksele są ponumerowane od 0 do 4, założmy również, że zdefiniowaliśmy `MAX_P` jako `(NUM_P - 1)`, ponieważ pozwoli nam to uniknąć konieczności ciągłego odejmowania 1.

Jak omówiliśmy we wcześniejszych odcinkach, jeśli wykonujemy obrót zgodnie z ruchem wskazówek zegara i jeśli jesteśmy obecnie gotowi do aktywacji piksela `P`, to przez większość czasu chcemy dezaktywować

piksel `(P - 1)`. Tylko wtedy, gdy chcemy aktywować piksel 0, musimy dezaktywować piksel `MAX_P`. Podobnie, jeśli wykonujemy obrót w kierunku przeciwnym do ruchu wskazówek zegara i jeśli obecnie chcemy aktywować piksel `P`, to przez większość czasu chcemy dezaktywować piksel `(P + 1)`. Tylko wtedy, gdy chcemy aktywować piksel `MAX_P`, musimy dezaktywować piksel 0.

Czas testów (instrukcji warunkowych)

Jednym z rozwiązań jest wykonanie prostego testu w celu określenia, który piksel ma zostać dezaktywowany. Zakładając obrót zgodnie z ruchem wskazówek zegara, nasza główna funkcja `loop()` może być napisana w następujący sposób (gdzie warunek i jego sprawdzenie są pogrubione).

```
void loop ()
{
    for (int iPix = 0; iPix <= MAX_P; iPix++)
    {
        // Włącz nowy piksel
        Neos.setPixelColor(iPix, COLOR_ON);

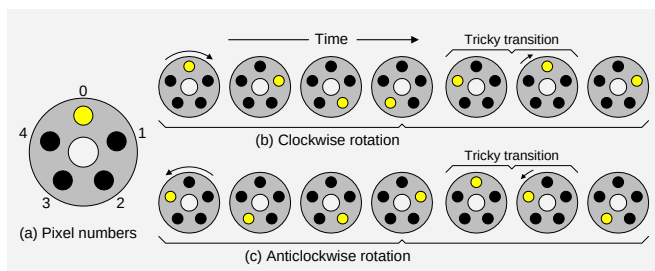
        // Wyłącz stary piksel
        if (iPix == 0)
            Neos.setPixelColor(MAX_P, COLOR_OFF);
        inny
            Neos.setPixelColor((iPix - 1), COLOR_OFF);

        // Wyświetl wynik
        Neos.show();
        delay(PadDelay);
    }
}
```

Prawdopodobnie warto zauważyć, że w poprzednich programach używaliśmy odpowiednika `iPix < NUM_P` (co odpowiada `iPix < 5` w tym przykładzie) jako warunku testowego w naszej pętli `for()`. Powodem, dla którego używamy `iPix <= MAX_P` (co odpowiada `iPix <= 4`) jako warunku testowego w tym przykładzie jest to, że będzie on lepiej uzupełniał eksperymenty, które mają nadejść. Główną kwestią, na którą należy zwrócić uwagę, jest to, że oba powodują te same działania, widziane przez osobę obserwującą, jak pierścienie wykonuje swoją magię.

Jeśli chcesz zapoznać się i rozważyć to bardziej szczegółowo, pełny szkic (program) zawarty jest w pliku `CB-Mar21-01.txt` (ten i wszystkie inne pliki powiązane z tym artykułem są dostępne na stronie PE z marca 2021 r.: <https://bit.ly/3oouhbl>).

Jeśli spojrzysz na pełny szkic, zobaczysz, że ustawiliśmy czas cyklu – czyli czas potrzebny na wykonanie pełnego obrotu pierścienia



Rysunek 1. Obroty zgodne i przeciwnie do ruchu wskazówek zegara

– na 1000 ms (milisekund) – czyli jedną sekundę. Przyjmujemy również założenie, że wykonanie wszelkich obliczeń i przesłanie nowych wartości dla każdej pozycji piksela zajmuje 1 ms. Na tej podstawie obliczamy odpowiednią wartość zmiennej `PadDelay` w naszej funkcji `setup()`.

Możemy użyć podobnego warunkowego podejścia, aby wykonać obrót w kierunku przeciwnym do ruchu wskazówek zegara, jak pokazano poniżej (plik CB-Mar21-02.txt). W tym przypadku wszelkie zmiany w naszym poprzednim kodzie zostały wyróżnione pogrubioną czcionką.

```
void loop ()
{
    for (int iPix = MAX_P; iPix >= 0; iPix--)
    {
        // Włącz nowy piksel
        Neos.setPixelColor(iPix, COLOR_ON);

        // Wyłącz stary piksel
        if (iPix == MAX_P)
            Neos.setPixelColor(0, COLOR_OFF);
        inny
            Neos.setPixelColor((iPix + 1), COLOR_OFF);

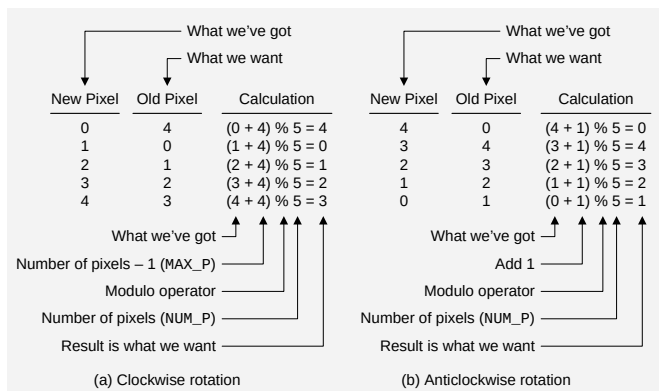
        // Wyświetl wynik
        Neos.show();
        delay(PadDelay);
    }
}
```

Zanim przejdziemy dalej, warto przypomnieć sobie, że `iPix - 1` lub `iPix + 1` w większości przypadków działa poprawnie. Jedynym powodem, dla którego musimy wykonać wspomniane testy, jest obsługa wartości granicznych występujących na początku pętli (lub na jej końcu, w zależności od punktu widzenia).

Wspaniałe moduły

Korzystanie z powyższych instrukcji warunkowych z pewnością zapewnia użyteczne rozwiązania. Z drugiej strony, wydają się one nieco „niechlujne”. Jak być może pamiętasz, operator modulo `%` zwraca resztę całkowitą z dzielenia liczb całkowitych. Cóż, jest to jeden z tych przypadków, w których operator modulo naprawdę ma szansę się wykazać.

Ponieważ nie jestem urodzonym programistą, zazwyczaj najpierw szkicuję rzeczy na papierze. Zaczniemy od obrotu zgodnie z ruchem wskazówek zegara (**rysunek 2a**). Na podstawie naszych wcześniejszych szkiców wiemy, że – dla każdego kroku wokół pierścienia – mamy numer nowego piksela, który ma zostać włączony, podczas



Rysunek 2. Użycie operatora modulo w celu uzyskania pożądanego wyniku

gdy to, czego potrzebujemy, to numer starego piksela, który ma zostać wyłączony. Jeśli dodamy `MAX_P` (czyli liczbę pikseli minus jeden) do liczby pikseli, które właśnie włączyliśmy, a następnie wykonamy dzielenie modulo z `NUM_P` (liczbą pikseli), otrzymamy żądaną wartość. Kod do tego jest pokazany poniżej z interesującymi częściami zaznaczonymi pogrubioną czcionką (plik CB-Mar21-03.txt).

```
void loop ()
{
    int tPix;

    for (int iPix = 0; iPix <= MAX_P; iPix++)
    {
        // Włącz nowy piksel
        Neos.setPixelColor(iPix, COLOR_ON);

        // Wyłącz stary piksel
        tPix = (iPix + MAX_P) % NUM_P;
        Neos.setPixelColor(tPix, COLOR_OFF);

        // Wyświetl wynik
        Neos.show();
        delay(PadDelay);
    }
}
```

Podobnie, jeśli chcemy wykonać obrót w kierunku przeciwnym do ruchu wskazówek zegara (**rysunek 2b**), wówczas dla każdego kroku wokół pierścienia znamy numer nowego piksela, który ma zostać włączony i musimy obliczyć numer starego piksela, który ma zostać wyłączony. W tym przypadku, jeśli dodamy 1 do numeru piksela, który właśnie włączyliśmy, a następnie wykonamy dzielenie modulo z `NUM_P` (liczbą pikseli), otrzymamy żądaną wartość. Poniżej znajduje się kod z interesującymi fragmentami wyróżnionymi pogrubioną czcionką (plik CB-Mar21-04.txt).

```
void loop ()
{
    int tPix;

    for (int iPix = MAX_P; iPix >= 0; iPix--)
    {
        // Włącz nowy piksel
        Neos.setPixelColor(iPix, COLOR_ON);

        // Wyłącz stary piksel
        tPix = (iPix + 1) % NUM_P;
        Neos.setPixelColor(tPix, COLOR_OFF);

        // Wyświetl wynik
        Neos.show();
        delay(PadDelay);
    }
}
```

Niegrzecznie jest pokazywać palcem

Nie wiem jak ty, ale ja, gdy byłem dzieckiem i próbowałem skierować uwagę mojej matki na coś interesującego (na przykład na starszą panią z ogromną brodawką na końcu nosa), informowała mnie

groźnym szeptem, że „niegrzecznie jest pokazywać palcem”, po czym starając się wybrnąć z niezręcznej sytuacji nagle przypominała sobie, że nasza obecność jest pilnie wymagana gdzieś indziej.

Oczywiście moja matka nie zajmowała się programowaniem, a przy najmniej tak zakładałam, ale – teraz o tym myślę – nie wiedziałam nawet, że mówi płynnie po francusku i niemiecku, dopóki nie opuściłem domu, więc może programuje jak diva. Tak czy inaczej, wskaźniki mogą być bardzo przydatne w kontekście programów.

Powinienem teraz zaznaczyć, że nie mówię o prawdziwych wskaźnikach C, które – choć są niezwykle skuteczne – są tematycznie zupełnie inną parą kaloszy. Na potrzeby tego cyklu wizualizuję prosty pseudo wskaźnik zaimplementowany przy użyciu liczby całkowitej.

Jak zwykle, zaczniemy od obrotu zgodnie z ruchem wskazówek zegara. Założmy, że zadeklarujemy globalną zmienną całkowitą o nazwie `NewP` („nowy piksel”); a także, że zainicjujemy ją tak, aby zawierała 0. W takim przypadku kod naszej funkcji `main` `loop()` mógłby wyglądać następująco (plik CB-Mar21-05.txt).

```
void loop ()
{
    int oldP;

    // Włącz nowy piksel
    Neos.setPixelColor(NewP, COLOR_ON);

    // Wyłącz stary piksel
    oldP = (NewP + MAX_P) % NUM_P;
    Neos.setPixelColor(oldP, COLOR_OFF);

    // Wyświetl wynik
    Neos.show();
    delay(PadDelay);

    // Zwiększenie wskaźnika
    NewP = (NewP + 1) % NUM_P;
}
```

Zauważ, że nie potrzebujemy już pętli `for()`. Ponadto używamy dokładnie tej samej operacji modulo do obliczenia liczby starych pikseli, które mają zostać wyłączone (zmieniliśmy tylko nazwę zmiennej). Jediną inną modyfikacją jest to, że musimy teraz zwiększyć wartość naszego wskaźnika liczby całkowitej na końcu pętli. Jak widzimy, ponownie używamy operatora modulo, aby zadbać o to, by `NewP` podążał za sekwencją 0, 1, 2, 3, 4, 0, 1, 2....

Oczywiście, jeśli mielibyśmy wykonać obrót w kierunku przeciwnym do ruchu wskazówek zegara, chcielibyśmy, aby `NewP` podążał za sekwencją 4, 3, 2, 1, 0, 4, 3, 2.... Dla śmiechu i zabawy, spróbuj sam stworzyć przeciwną do ruchu wskazówek zegara wersję tej techniki – nawet jeśli miałyby to być tylko ćwiczenie z ołówkiem i papierem – zanim przejdziemy dalej.

Co? Już skończyłeś? Jestem z ciebie dumny! Jak już pewnie zdążyłeś się zorientować, jeśli wykonujemy

obrót w kierunku przeciwnym do ruchu wskazówek zegara, to jeśli chodzi o dekrementację wskaźnika na końcu, nie możemy użyć niczego podobnego do `NewP = (NewP - 1) % NUM_P`. Dzieje się tak dlatego, że gdy dojdziemy do `NewP` równego 0, to `NewP - 1` będzie równe -1, a nie chcemy zanurzać się w gąszczu komplikacji i kłopotów, które pojawiają się, jeśli zastosujemy operator modulo w tym momencie (patrz także Sprytne porady i sztuczki z tego miesiąca).

Rozwiązanie, które wymyśliłem, jest następujące (plik CB-Mar21-06.txt). Część, w której wyłączamy stary piksel, opiera się na naszej poprzedniej procedurze przeciwnej do ruchu wskazówek zegara. Interesująca jest część, w której zmniejszamy nasz wskaźnik na końcu.

```
void loop ()
{
    int oldP;

    // Włącz nowy piksel
    Neos.setPixelColor(NewP, COLOR_ON);

    // Wyłącz stary piksel
    oldP = (NewP + 1) % NUM_P;
    Neos.setPixelColor(oldP, COLOR_OFF);

    // Wyświetl wynik
    Neos.show();
    delay(PadDelay);

    // Zmniejszenie wskaźnika
    NewP = (NewP + MAX_P) % NUM_P;
}
```

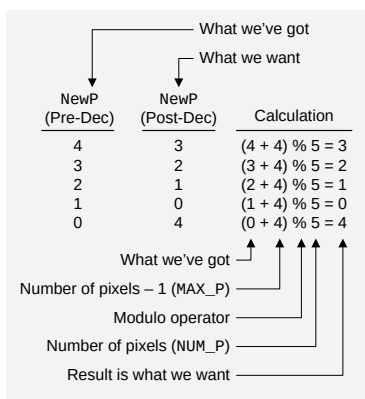
Kluczową kwestią, na którą należy zwrócić uwagę – oprócz tego, że to działa – jest fakt, że używamy `(NewP + MAX_P)`, co oznacza, że zawsze będziemy mieli do przekazania naszemu operatorowi modulo liczbę dodatnią. Przydatne może być zanotowanie twoich przypuszczeń na temat tego co się dzieje, by móc potem spojrzeć na **rysunku 3**, i skonfrontować z moją wizualizacją.

Czy zauważyłeś coś interesującego na rysunku 3? Co powiesz na fakt, że te obliczenia są identyczne z tymi pokazanymi na rysunku 2a (tylko trochę zmieniliśmy kolejność)? Oczywiście, gdy się nad tym zastanowić ma to sens. Poprzednio używaliśmy naszego równania do obliczenia starego piksela do wyłączenia podczas wykonywania obrotu zgodnie z ruchem wskazówek zegara. Teraz używamy tego samego równania do obliczenia nowego piksela do włączenia podczas wykonywania obrotu w kierunku przeciwnym do ruchu wskazówek zegara. Choć wydaje się to przerażające, wygląda na to, że mamy pojęcie, co robimy.

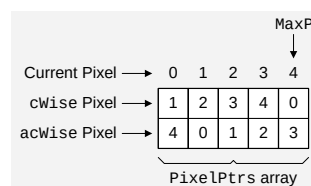
Kurczaki i jajka

Podczas pisania tego tekstu przychodzą mi do głowy różne rzeczy. Na przykład, nasze poprzednie przykłady obejmowały najpierw włączenie bieżącego piksela, a następnie obliczenie starego piksela do wyłączenia. Moglibyśmy oczywiście zrobić to na odwrót – to znaczy wyłączyć bieżący piksel, a następnie obliczyć nowy piksel, który ma zostać włączony.

Weźmy nasz poprzedni przykład obrotu zgodnie z ruchem wskazówek zegara (plik CB-Mar21-05.txt) i zmodyfikujmy go, aby



Rysunek 3. Modulo na każdym kroku



Rysunek 4. Korzystanie z tabeli wyszukiwania

odzwierciedlał ten nowy scenariusz (plik CB-Mar21-07.txt). W tym przypadku nasza funkcja `main loop()` może wyglądać następująco:

```
void loop ()
{
    // Wyłącz stary piksel
    Neos.setPixelColor(CurrentP, COLOR_OFF);

    // Zwiększenie wskaźnika
    CurrentP = (CurrentP + 1) % NUM_P;

    // Włącz nowy piksel
    Neos.setPixelColor(CurrentP, COLOR_ON);

    // Wyświetl wynik
    Neos.show();
    delay(PadDelay);
}
```

Jeśli porównasz oba rozwiązania, zauważysz, że nasza nowa wersja jest odrobinę bardziej zwięzła i odpowiednio bardziej wydajna. Poprzednio musieliśmy obliczyć numer starego piksela, a także wartość nowego wskaźnika, co wymagało dwóch operacji modulo. Dla porównania, nasze nowe wcielenie wymaga tylko jednej operacji modulo. Tra-la!

Sprawdź teraz, czy potrafisz opracować równoważne rozwiązanie dla obrotu w kierunku przeciwnym do ruchu wskazówek zegara i porównaj je z moją wersją (CB-Mar21-08.txt).

Kręte tabele

Innym podejściem opartym na wskaźnikach – nadal wykorzystującym nasze uproszczone wskaźniki liczb całkowitych w przeciwieństwie do rzeczywistych wskaźników C – jest skonstruowanie tabeli przeglądowej w postaci tablicy, której rozmiar odzwierciedla liczbę pikseli, którymi się bawimy. Następnie ładujemy naszą tabelę (tablicę) obliczonymi wartościami dla sąsiednich pikseli zgodnie z ruchem wskazówek zegara (`cWise`) i przeciwnie do ruchu wskazówek zegara (`acWise`) (rysunek 4).

Chociaż obecnie eksperymentujemy z 5-pikselowym pierścieniem, chcemy, aby nasza implementacja była łatwo skalowalna do większych pierścieni z większą liczbą pikseli. Zaczniemy od wzięcia naszej poprzedniej rotacji „Kury i jajka” zgodnie z ruchem wskazówek zegara (plik CB-Mar21-07.txt), utworzenia nowej wersji (plik CB-Mar21-09.txt) i zdefiniowania struktury o nazwie `CwAcwPair`, jak pokazano poniżej.

```
typedef struct
{
    int cWise;
    int acWise;
} CwAcwPair;
```

Pamiętaj, że wprowadziliśmy `typedef` (definicje typów), `enum` (typy wyliczeniowe) i `struct` (struktury) we wcześniejszym artykule (EdW07/2024). Następnie zadeklarujemy naszą tablicę wyszukiwania jako tablicę tych struktur w następujący sposób:

```
CwAcwPair PixelPtrs[NUM_P];
```

Na koniec dodamy pętlę `for()` do naszej funkcji `setup()`, aby załadować wartości `cWise` i `acWise` do naszej tablicy `PixelPtrs[]` w następujący sposób:

```
for (int iPix = 0; iPix < NUM_P; iPix++)
{
    PixelPtrs[iPix].cWise = (iPix + 1) % NUM_P;
    PixelPtrs[iPix].acWise = (iPix + MAX_P) % NUM_P;
}
```

Zauważ, że do załadowania tej tabeli używamy dokładnie tych samych operacji opartych na modulo, których używaliśmy w naszych poprzednich szkicach. Na tej podstawie możemy przepisać naszą główną funkcję `loop()` w następujący sposób:

```
void loop ()
{
    // Wyłącz stary piksel
    Neos.setPixelColor(CurrentP, COLOR_OFF);

    // Zaktualizuj wskaźnik
    CurrentP = PixelPtrs[CurrentP].cWise;

    // Włącz nowy piksel
    Neos.setPixelColor(CurrentP, COLOR_ON);

    // Wyświetl wynik
    Neos.show();
    delay(PadDelay);
}
```

A co jeśli chcielibyśmy zmodyfikować najnowszą wersję naszego programu tak, aby wykonywał obrót w kierunku przeciwnym do ruchu wskazówek zegara? Nic prostszego. Wszystko, co musimy zrobić, to zmodyfikować instrukcję aktualizacji wskaźnika w naszej pętli `loop()`, aby brzmiała następująco (plik CB-Mar21-10.txt):

```
CurrentP = PixelPtrs[CurrentP].acWise;
```

Jedną z wad korzystania z tabeli, tak jak robimy to tutaj, jest to, że zajmuje ona miejsce w pamięci, ale podobnie jest z każdym kodem, którego używamy do wielokrotnego wykonywania tych samych obliczeń. Zaletą podejścia tabelarycznego jest to, że obliczenia wykonujemy tylko raz, a treść programu jest łatwiejsza do zrozumienia.

Można powiedzieć, że w ramach naszej dyskusji „Kury i jajka” widzieliśmy już, jak sprowadzić wszystko do pojedynczego obliczenia, ale jest to prawdą tylko wtedy, gdy chcemy wyłączyć pojedynczy piksel i włączyć inny piksel. Jeśli zdecydowalibyśmy się na bardziej wyrafinowany efekt, taki jak końcowe zanikanie, które polega na oświetleniu nowego piksela z jasnością 100% i ustawieniu trzech końcowych pikseli odpowiednio na 75%, 50% i 25%, to przekonalibyśmy się, że podejście oparte na tabeli znacznie ułatwia nam życie.

Ach, delay(), fajnie było cię znać

W moim artykule *Is Time Truly an Illusion?* (<https://bit.ly/38mwa3k>) postawiłem pytanie: „Czy czas istniał przed Wielkim Wybuchem, czy czas był emergentną właściwością Wielkiego Wybuchu, czy czas jest tylko czymś, co powstrzymuje wszystko przed wydarzeniem się jednocześnie, czy też czas jako fundamentalna właściwość po prostu w ogóle nie istnieje?”

Powodem, dla którego wspominam o tym tutaj, jest fakt, że wszystkie programy, którym przyjrzelśmy się do tej pory, wykorzystywały funkcję `delay()`, aby wszystko nie działo się jednocześnie. Jednakże `delay()` jest funkcją blokującą, co oznacza, że całkowicie blokuje mikrokontroler, tym samym uniemożliwiając (lub blokując) cokolwiek innego.

Gdy mikrokontroler wykonuje funkcję `delay()`, nie może reagować na zmiany na żadnym ze swoich wejść, nie może wykonywać żadnych obliczeń ani podejmować żadnych decyzji, a także nie może zmieniać stanu żadnego ze swoich wyjść.

Jedną z technik zastąpienia `delay()` jest cykliczne sprawdzanie zegara systemowego, aby dowiedzieć się, kiedy nadszedł czas na działanie. Zamierzamy teraz wziąć nasz oparty na tabeli program obrotu zgodnie z ruchem wskazówek zegara (CB-Mar21-09.txt) i zmodyfikować go tak, aby korzystał z tej nowej techniki (plik CB-Mar21-11.txt). W rezultacie nasza funkcja `main loop()` będzie teraz wyglądać następująco:

```
void loop ()
{
    uint32_t currentTime = millis();

    // Czy nadszedł czas, aby coś zrobić?
    if ( (currentTime - TimeOfLastChange) >= PadDelay)
    {
        // Wyłącz stary piksel
        Neos.setPixelColor(CurrentP, COLOR_OFF);

        // Zaktualizuj wskaźnik
        CurrentP = PixelPtrs[CurrentP].cWise;

        // Włącz nowy piksel
        Neos.setPixelColor(CurrentP, COLOR_ON);

        // Wyświetl wynik
        Neos.show();

        TimeOfLastChange = currentTime;
    }
}
```

I tak jak poprzednio, jeśli chcemy zmodyfikować tę najnowszą wersję naszego programu, aby wykonywała obrót w kierunku przeciwnym

do ruchu wskazówek zegara, wszystko, co musimy zrobić, to zmodyfikować instrukcję aktualizacji wskaźnika w naszej głównej pętli `loop()`, aby brzmiała następująco (plik CB-Mar21-12.txt):

```
CurrentP = PixelPtrs[CurrentP].acWise;
```

Zwycięzcą jest...

Tak więc, spośród wszystkich technik, które pokazaliśmy powyżej – i pamiętając, że w tym miesiącu pojawi się więcej porad i wskazówek – która z nich jest najlepsza? To trochę jak pytanie „Jak długi jest kawałek sznurka?”.

W rzeczywistości wszystko zależy od tego, co chcemy osiągnąć. Jeśli chcemy tylko wyłączyć jeden piksel i włączyć inny lub odwrotnie, to możemy zdecydować się na rozwiązania przedstawione w sekcji Kury i jajka. Jeśli chcemy dodać nieco bardziej wyrafinowane efekty, to programy przedstawione w temacie Kręte tabele mogą okazać się najlepszą opcją. Alternatywnie, jeśli chcemy wykonywać dodatkowe zadania podczas migania naszych pikseli – na przykład sprawdzać stany przełączników i wykonywać złożone obliczenia – wówczas rozwiązania oferowane w sekcji Ach, `delay()`, fajnie było cię znać będą prawdopodobnie najlepszym rozwiązaniem.

Co? Nie ma GOL?

Czuje się jak stary głupiec (ale gdzie go znaleźć o tej porze dnia?). Przez ostatnie dwa artykuły obiecywałem zaimplementować wersję „Gry w życie” Hortona Conweya’a (GOL) (<https://bit.ly/pe-jan21-cgol>) na mojej tablicy 12×12 piłeczek pingpongowych.

Niemal niewiarygodne jest jednak to, że po raz kolejny dałem się zepchnąć na boczny tor (jestem tak samo zszokowany jak ty). Ale nie ma co płakać nad rozlanym mlekiem. Wszystko, co omówiliśmy w tym miesiącu w odniesieniu do użycia operatora modulo, wejdzie w grę, gdy zaimplementujemy GOL, co zrobimy w przyszłym miesiącu, albo nie nazywam się Max Wspaniały! ■

Clive „Max” Maxfield

Sprytne porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania



Jestem dość podekscytowany, gdy zaczynam pisać te słowa, ponieważ jesteśmy gotowi zebrać razem kilka różnych koncepcji i – pod koniec tego odcinka – wszyscy powiemy „Wow!” (Jeśli nie powiesz „Wow!” z wystarczającym entuzjazmem, będę musiał wysłać do ciebie moją kochaną starą matkę, aby porozmawiała z tobą o twym niewłaściwym zachowaniu).

Bity, bajty i nybble, ojej!

Najmniejszą informacją, którą można przechowywać i manipulować w maszynie cyfrowej, jest cyfra binarna (inaczej „bit”), która może być używana do ucieleśnienia dwóch wartości. Zazwyczaj myślimy o tych wartościach jako reprezentujących liczby 0 i 1. Jak omówiliśmy wcześniej, zakładając, że pracujemy z mikrokontrolerem w środowisku Arduino, te wartości 0 i 1 są równoważne odpowiednio `LOW` i `HIGH`, gdy pracujemy (odczytujemy lub zapisujemy) na cyfrowych wejściach/wyjściach (I/O). Są one również równoważne odpowiednio `false` i `true`, jeśli traktujemy je jako logiczne wartości logiczne.

Ponieważ pojedynczy bit jest ograniczony pod względem ilości informacji, które może reprezentować, zwykle wygodniej jest pracować z grupami bitów. Niektóre grupy są powszechne, więc nadaliśmy im specjalne nazwy. Na przykład termin „bajt” odnosi się do grupy 8-bitowej, podczas gdy termin „nybble” (lub „nibble” – wolę „nybble”!) odnosi się do grupy 4-bitowej. Oznacza to, że dwa nybble tworzą bajt, co pokazuje, że inżynierowie komputerowi mają poczucie humoru, choć nie jest ono zbyt wyrafinowane.

Ogólnie rzecz biorąc, ludziom trudno jest ogarnąć umysłem duże liczby przedstawione w postaci binarnej. Na przykład, 10101100 w systemie binarnym nie ma większego znaczenia dla większości ludzi, podczas gdy jego dziesiętny odpowiednik 172 jest łatwiejszy do zrozumienia.

Binarny system liczbowy ma dwie cyfry, 0 i 1. Dziesiętny system liczbowy ma 10 cyfr, 0, 1, 2, 3, 4, 5, 6, 7, 8 i 9. Z powodów, które zostaną wyjaśnione w przyszłych odcinkach, mapowanie (tłumaczenie) wartości tam i z powrotem między systemem binarnym i dziesiętnym

nie jest tak wygodne, jak można by się spodziewać. Znacznie bardziej efektywne jest użycie szesnastkowego systemu liczbowego, który składa się z 16 cyfr: 0, 1, 2, 3, 4, 5, 6, 7, 8, 9, A, B, C, D, E i F. Piękno tego systemu polega na tym, że każda cyfra szesnastkowa jest bezpośrednio odwzorowywana na 4-bitową cyfrę binarną (**rysunek 5**).

W poprzednich odcinkach omówiliśmy różnicę między liczbami całkowitymi ze znakiem i bez znaku. Również fakt, że w przypadku typu danych `int` (integer) jego rozmiar – tj. liczba bitów użytych do jego reprezentacji – zależy od komputera lub środowiska na którym pracujesz. Na przykład w przypadku Arduino Uno, `int` zajmuje dwa bajty (16 bitów) pamięci.

Załóżmy, że deklarujemy liczbę `int` o nazwie `intA`. Załóżmy ponadto, że przypiszemy mu wartość `0xCA35`, gdzie przedrostek „0x” wskazuje, że jest to wartość szesnastkowa:

```
int intA = 0xCA35;
```

Używając rysunku 5 jako odniesienia, widzimy, że faktycznie załadowaliśmy `intA` wartością binarną `1100101000110101`, jak pokazano na **rysunku 6**. Zauważ, że zamiast zapisywać wszystkie 16 bitów razem w postaci `1100101000110101`, możemy łatwiej myśleć o nich w kategoriach 4-bitowych nybbles i zapisywać je w postaci `1100 1010 0011 0101` (to właśnie zrobimy w dalszej części tego odcinka).

Bramki AND i OR

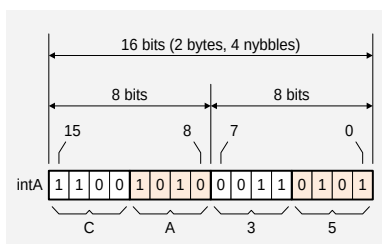
Na najniższym poziomie komputer cyfrowy składa się z ogromnej liczby kluczy przełączających. Przełączniki te mogą być implementowane przy użyciu różnych technologii, w tym mechanicznych, elektromechanicznych (przełączniki), pneumatycznych i półprzewodnikowych (tranzystory). W dzisiejszych czasach, oczywiście, tranzystory są w modzie – i to właśnie założymy tutaj – ale kto wie, co nadejdzie jutro?

Kolejnym poziomem jest gromadzenie małych grup tranzystorów i wykorzystywanie ich do implementacji prostych funkcji logicznych – zwanych prymitywnymi bramkami logicznymi – a następnie łączenie tych bramek ze sobą w sprytny sposób. Dwie bramki, na których skupimy się w tym artykule, to bramki AND i OR (**rysunek 7**).

Zauważ, że używamy znaku `&` do reprezentowania AND, podczas gdy znak `|` jest używany do reprezentowania OR. Jak widzimy, wyjście (`y`) z bramki AND ma wartość 1 tylko wtedy, gdy oba jej wejścia (`a` i `b`) mają wartość 1; w przeciwnym razie wyjście ma wartość 0. Dla porównania, wyjście bramki OR ma wartość 1, jeśli którekolwiek z jej wejść ma wartość 1.

Decimal	Hexadecimal	Binary
0	0	0000
1	1	0001
2	2	0010
3	3	0011
4	4	0100
5	5	0101
6	6	0110
7	7	0111
8	8	1000
9	9	1001
10	A	1010
11	B	1011
12	C	1100
13	D	1101
14	E	1110
15	F	1111

Rysunek 5. Mapowanie binarne i szesnastkowe



Rysunek 6. Przykładowa wartość 16-bitowa

AND			OR		
a	b	y	a	b	y
0	0	0	0	0	0
0	1	0	0	1	1
1	0	0	1	0	1
1	1	1	1	1	1

Rysunek 7. Bramki AND i OR

Operatory logiczne && i ||

Wcześniej mówiliśmy o AND i OR w kontekście sprzętu, czyli fizycznych bramek logicznych używanych do budowy komputera. Teraz rozważymy pokrewne funkcje w oprogramowaniu, czyli naszych programach.

W odcinku z ubiegłego miesiąca wprowadziliśmy pojęcie zmiennych logicznych, którym można przypisać wartości `false` i `true`. Zauważyliśmy również, że zmienne te są często traktowane podobnie do liczb całkowitych, ponieważ `false` jest równe 0, a `true` jest równe 1 lub – dokładniej – `true` jest równe dowolnej wartości niezerowej.

Operatory logiczne `&&` i `||` pozwalają nam konstruować złożone instrukcje warunkowe. Na przykład, zakładając, że zadeklarowaliśmy zmienną całkowitą o nazwach `intA` i `intB`, możemy napisać instrukcję warunkową w następujący sposób:

```
if (((intA == 6) && (intB == 4)) == true)...
```

Zauważ, że w rzeczywistości nie potrzebujemy wszystkich nawiasów, których tutaj użyłem, ponieważ dwa operatory relacyjne `==` (porównania) po lewej stronie mają wyższy priorytet niż logiczny operator AND `&&`, więc moglibyśmy napisać to w następujący sposób:

```
if ((intA == 6 && intB == 4) == true)...
```

Jeśli jednak nie masz pewności co do pierwszeństwa operatorów (<https://bit.ly/355G9rC>), zawsze najlepiej jest użyć nawiasów, aby wymusić działanie w żądanej kolejności, z dodatkową zaletą, że nawiasy zwykle ułatwiają innym osobom zrozumienie tego, co próbujesz osiągnąć.

Zastanówmy się, w jaki sposób powyższe wyrażenie zostanie obliczone przez kompilator, a zatem w jaki sposób zostanie wykonane przez komputer. Jeśli `intA` jest równe 6, to to podwyrażenie zwróci wartość `true` (1), w przeciwnym razie zwróci wartość `false` (0). Podobnie, jeśli `intB` jest równe 4, wówczas to podwyrażenie zwróci wartość `true` (1), w przeciwnym razie zwróci wartość `false` (0). Operator `&&` zwróci `true` (1) tylko wtedy, gdy oba jego wejścia są `true` (1), w przeciwnym razie zwróci `false` (0). Na koniec używamy skrajnego prawego operatora `==`, aby porównać wynik z `&&` do wartości `true` (1). Ponownie, to porównanie zwróci wartość `true` (1) tylko wtedy, gdy oba jego wejścia są `true` (1), w przeciwnym razie zwróci `false` (0).

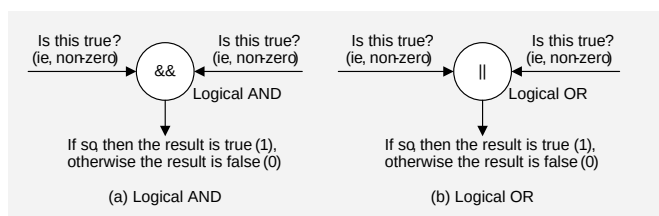
W rzeczywistości, odrobina zastanowienia pokazuje, że możemy uprościć instrukcję warunkową, pisząc ją w następujący sposób:

```
if ((intA == 6) && (intB == 4))...
```

Chociaż istnieje niebezpieczeństwo, że za chwilę pogubimy się we własnych myślach, głębsza analiza pokazuje, że – o ile zadeklarowaliśmy również zmienną całkowitą o nazwie `intY` – poniższa instrukcja jest również poprawna:

```
intY = intA && intB;
```

W tym przypadku, jeśli `intA` zawiera 0, operator `&&` uzna to za `false` (0); w przeciwnym razie, jeśli `intA` zawiera wartość niezerową, operator `&&` uzna to za `true` (1). Podobnie w przypadku `intB`. W rezultacie operator `&&` zwróci wartość `true` (1) tylko wtedy, gdy zarówno `intA`, jak i `intB` zawierają wartości niezerowe; w przeciwnym razie zwróci wartość `false` (0). Każda wartość zwrócona przez operator `&&` zostanie przypisana do `intY`.



Rysunek 8. Wizualizacja operatorów logicznych && i ||

Możemy myśleć o operatorze logicznym `&&` jako o gigantycznym programowym AND (rysunek 8a); podobnie możemy myśleć o operatorze logicznym `||` jako o gigantycznym programowym OR (rysunek 8b).

Operatory bitowe & i |

Podczas pisania programów możemy również używać operatorów `&` (bitowe AND) i `|` (bitowe OR). Bitowe AND akceptuje dwie wartości całkowite jako dane wejściowe i wykonuje AND na każdej parze odpowiadających im bitów. W komputerze jest to realizowane za pomocą kilku fizycznych bramek AND. Podobnie, operator bitowy OR wykorzystuje szereg fizycznych bramek OR, aby wykonać operację OR na każdej parze odpowiadających sobie bitów.

Założmy, że deklarujemy trzy liczby całkowite, `intA`, `intB` i `intY`. Pamiętaj, że zakładamy również, że pracujemy z Arduino Uno, więc każda z tych liczb całkowitych będzie miała szerokość dwóch bajtów (16 bitów). Założmy, że przypiszemy wartość `0xCA35` w systemie szesnastkowym (`1100 1010 0011 0101` w systemie binarnym) do `intA` i wartość `0x6DA7` w systemie szesnastkowym (`0110 1101 1010 0111` w systemie binarnym) do `intB`. Założmy teraz, że użyjemy następującej instrukcji w naszym programie, aby wykonać bitowe AND:

```
intY = intA & intB;
```

Najmniej znaczący bit `intA` to 1, podobnie jak najmniej znaczący bit `intB`, więc AND tych bitów oznacza, że najmniej znaczący bit wyniku również będzie 1. Podobne operacje są wykonywane dla wszystkich pozostałych bitów (rysunek 9a), w wyniku czego `intY` zawiera `0x4825` w systemie szesnastkowym (`0100 1000 0010 0101` w systemie binarnym).

Założmy teraz, że zaczynamy od `intA` i `intB` zawierających te same wartości co poprzednio i używamy następującej instrukcji w naszym programie do wykonania bitowego OR:

```
intY = intA | intB;
```

W tym przypadku wykonujemy operację OR na każdej parze odpowiadających sobie bitów (rysunek 9b), w wyniku czego `intY` zawiera `0xEFB7` w systemie szesnastkowym (`1110 1111 1011 0111` w systemie binarnym).

Sprawdzanie bitów

Założmy, że wiemy, że `intA` zawiera obecnie poprawną wartość, ale tak naprawdę nie wiemy, jaka jest ta wartość. Czysto dla celów tej dyskusji, możemy wizualizować to jako `0x????` w systemie szesnastkowym (`???? ???? ???? ????` w systemie binarnym). Założmy teraz, że chcemy wykonać sprawdzenie w celu ustalenia, czy bit 7 `intA` ma wartość 1, a jeśli tak, wykonać pewne czynności. Możemy to osiągnąć za pomocą operatora bitowego AND i wartości referencyjnej `0x0080` (`0000 0000 1000 0000`) w następujący sposób:

```
if ((intA & 0x0080) == true)...
```

Graficzną reprezentację tej operacji pokazano na rysunku 10a. Dla każdego bitu, który ma wartość 0 w wartości referencyjnej, odpowiedni bit wyjściowy z bitowego AND jest wymuszany na 0. Ponieważ bit 7 w wartości referencyjnej wynosi 1, odpowiedni bit wyjściowy będzie odzwierciedlał wartość bitu 7 w `intA`.

Jeśli bit 7 w `intA` wynosi 0, wartość zwrócona z `(intA & 0x0080)` wyniesie zero, co jest uważane za fałsz, więc nasze sprawdzenie sprowadzi się do `if (false == true)`. Oczywiście to się nie powiedzie, a instrukcje zawarte w `if()` nie zostaną wykonane.

Alternatywnie, jeśli bit 7 w `intA` wynosi 1, wartość zwrócona z `(intA & 0x0080)` będzie wynosić `0x0080` – czyli niezerowa – co jest uważane za prawdę, więc instrukcja warunkowa sprowadzi się do `if (true == true)`. Oczywiście test zostanie zaliczony, a instrukcje zawarte w `if()` zostaną wykonane.

Z naszych wcześniejszych dyskusji wiemy, że moglibyśmy osiągnąć ten sam efekt za pomocą następującego stwierdzenia:

```
if (intA & 0x0080)...
```

Założmy teraz, że chcemy wykonać działania objęte przez `if()` tylko wtedy, gdy bit 7 `intA` zawiera 0. W takim przypadku moglibyśmy użyć następującej instrukcji:

```
if ((intA & 0x0080) == false)...
```

Alternatywnie możemy osiągnąć ten sam efekt w następujący sposób:

```
if (!(intA & 0x0080))...
```

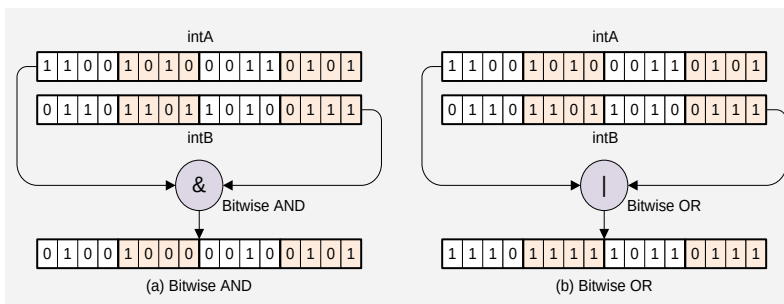
Tę drugą wersję możemy odczytać jako „If not (`intA & 0x0080`) then...” lub „If the bitwise AND of `intA` and `0x0080` doesn't return a true (non-zero) value, then...” Wiem, że na początku może to być nieco mylące, ale nie minie dużo czasu, zanim będziesz w stanie czytać i pisać tego typu rzeczy bez zastanawiania się (zbyt mocno) nad tym.

Ustawianie i kasowanie bitów

Założmy, że `intA` zawiera pewną wartość, taką jak `0xCA35` (`1100 1010 0011 0101`) i chcemy ustawić jej bit 7 na 1, pozostawiając pozostałe bity bez zmian. Możemy to osiągnąć za pomocą operatora bitowego OR i wartości referencyjnej `0x0080` (`0000 0000 1000 0000`) w następujący sposób.

```
intA = intA | 0x0080;
```

Jak pokazano na rysunku 10b, wynikiem jest pozostawienie `intA` zawierającego `0xCAB5` (`1100 1010 1011 0101`). Założmy teraz,



Rysunek 9 Wizualizacja operatorów bitowych & i |

że zmienimy zdanie i zdecydujemy, że naprawdę chcemy wyczyścić bit 7 `intA` do 0, pozostawiając pozostałe bity bez zmian. Moglibyśmy to osiągnąć za pomocą operatora bitowego AND i wartości referencyjnej `0xFF7F` (`1111 1111 0111 1111`) w następujący sposób.

```
intA = intA & 0xFF7F;
```

Jak pokazano na **rysunku 10c**, w tym przypadku wynikiem jest pozostawienie `intA` zawierającego `0xCA35` (`1100 1010 0011 0101`).

Bity maskujące

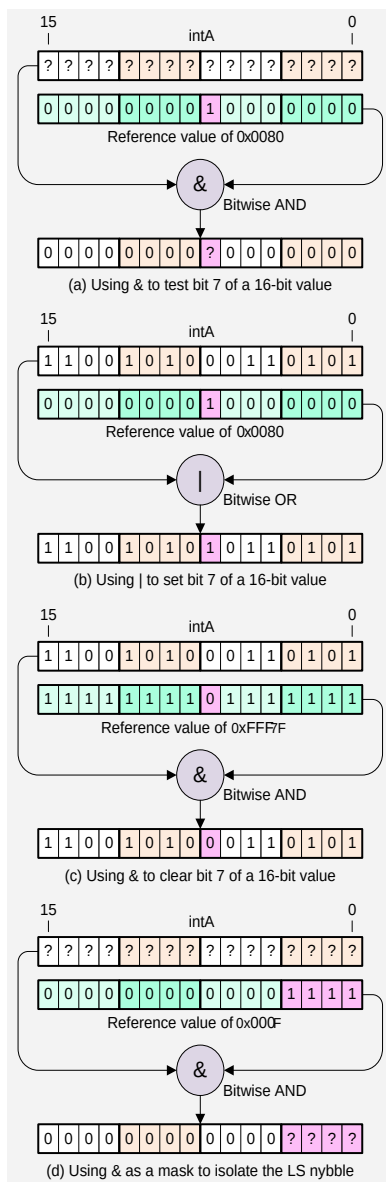
Jak widać, jest to logiczne rozszerzenie tego, co widzieliśmy wcześniej. Podobnie jak podczas sprawdzania bitów, założmy, że wiemy, że `intA` zawiera obecnie prawidłową wartość, ale nie wiemy, jaka to wartość, więc zwiualizujemy to jako `0x????` w systemie szesnastkowym (`????` `????` `????` w systemie binarnym). Załóżmy teraz, że chcemy „wyodrębnić” tylko najmniej znaczący nybble i skopiować go do `intV`. Możemy to osiągnąć w następujący sposób (**rysunek 10d**):

```
intV = intA & 0x000F;
```

Dla przypomnienia, we wcześniejszym odcinku (EdW05/2024) użyliśmy różnych kombinacji masek i przesunień, aby wyodrębnić trzy 8-bitowe kanały czerwony, zielony i niebieski z 32-bitowej wartości koloru (najbardziej znaczące 8 bitów nie jest używane). Dlaczego więc operacje maskowania są dla nas interesujące? Cóż, ciesz się, że zapytałeś...

Mów mi po prostu „Władca Pierścieni”

Wracając do pierścieni diod LED, które omawialiśmy w naszym cyklu, jeśli chodzi o obracanie piksela wokół pierścienia, istnieje sztuczka, której możemy użyć, jeśli liczba pikseli w pierścieniu jest potęgą dwóch. W przypadku mojego projektu Silnika Progностycznego, na przykład, pierścienie mają po 16 pikseli, gdzie $16=2^4$. Jak wiemy, te piksele są ponumerowane od 0 w systemie dziesiętnym (0 w systemie szesnastkowym; `0000` w systemie binarnym) do 15 w systemie



Rysunek 10. Testowanie, ustawianie, kasowanie i maskowanie bitów

dziesiętnym (F w systemie szesnastkowym; `1111` w systemie binarnym).

Wróćmy myślami do tematu „Kury i jajka” w odcinku z tego miesiąca. Założmy, że chcielibyśmy zastosować to podejście z 16-pikselowymi pierścieniami Silnika Progностycznego. Z naszą nowo odkrytą wiedzą na temat masek, jeśli chodzi o zwiększanie wskaźnika w naszym obrocie zgodnie z ruchem wskazówek zegara (plik `CB-Mar21-07.txt`), moglibyśmy zastąpić oryginalne obliczenia modulo (`CurrentP = (CurrentP + 1) % NUM_P;`) w naszej głównej funkcji `loop()` operacją maskowania w następujący sposób (zmodyfikowana część jest pogrubiona):

```
void loop ()
{
    // Wyłącz stary piksel
    Neos.setPixelColor(RingXref[CurrentP], COLOR_OFF);

    // Zwiększenie wskaźnika
    CurrentP = (CurrentP + 1) & MAX_P;

    // Włącz nowy piksel
    Neos.setPixelColor(RingXref[CurrentP], COLOR_ON);

    // Wyświetl wynik
    Neos.show();
    delay(PadDelay);
}
```

Jak to działa? Cóż, jeśli spojrzysz na mój kod źródłowy (plik `CB-Mar21-13.txt`), zobaczysz, że zdefiniowaliśmy `NUM_P` jako 16. Ponieważ `MAX_P` jest zdefiniowane jako (`NUM_P - 1`), oznacza to, że `MAX_P` wynosi 15 w systemie dziesiętnym (F w systemie szesnastkowym; `1111` w systemie binarnym).

Dla `CurrentP` równego 0 do `CurrentP` równego 14, (`CurrentP + 1`) będzie równe odpowiednio od 1 do 15. W kontekście naszych 16-bitowych liczb całkowitych odpowiada to wartościom (`CurrentP + 1`) od `0x0001` do `0x000F` w systemie szesnastkowym lub od `0000 0000 0000 0001` do `0000 0000 0000 1111` w systemie binarnym. Jeśli użyjemy operatora bitowego AND do wykonania operacji maskowania na tych wartościach z wartością referencyjną `MAX_P`, która zostanie automatycznie awansowana do `0x000F` (`0000 0000 0000 1111`), to wszystkie one przejdą bez szwanku.

Rozważmy teraz, co się stanie, gdy `CurrentP` będzie równe 15, co oznacza, że (`CurrentP + 1`) będzie równe 16 w systemie dziesiętnym (`0x0010` w systemie szesnastkowym; `0000 0000 0001 0000` w systemie binarnym). W tym przypadku nasze bitowe AND zamaskuje trzy najbardziej znaczące nybble i zwróci 0. W rezultacie nasze liczenie postępuje od 0 do 15, a następnie powraca do 0, co jest dokładnie tym, czego chcemy.

Na marginesie, być może zauważyłeś, że za każdym razem, gdy wywołujemy funkcję `Neos.setPixelColor()` w powyższym fragmencie kodu, zamiast pierwszego argumentu `CurrentP`, zmieniliśmy go na `RingXref[CurrentP]`. Dzieje się tak, ponieważ jak omówiono wcześniej, fizyczne piksele na 16-elementowym pierścieniu w moim prototypie nie są zorientowane i połączone w sposób, w jaki chcemy, aby były, więc używamy tablicy `RingXref[]`, aby przetłumaczyć to, co mamy w naszym wirtualnym świecie (w postaci `CurrentP`) na to, czego potrzebujemy w świecie rzeczywistym.

Na koniec rozważmy zastosowanie techniki maski dla obrotu w kierunku przeciwnym do ruchu wskazówek zegara. Zanim przejdziemy dalej, przypomnijmy sobie, że liczba całkowita w maszynie cyfrowej

zajmuje ograniczoną liczbę bitów (16 w przypadku Arduino Uno). Co więc się stanie, jeśli nasza liczba całkowita zawiera już 0xFFFF (1111 1111 1111 1111) i dodamy do niej 1? Rezultatem jest warunek przepełnienia, w którym liczba całkowita zawiera 0x0000 (0000 0000 0000 0000). Z drugiej strony, jeśli nasza liczba całkowita zawiera 0x0000 (0000 0000 0000 0000) i odejmiemy 1, otrzymamy stan niedopełnienia, w którym liczba całkowita zawiera 0xFFFF (1111 1111 1111 1111).

Rozważmy teraz nową implementację naszego obrotu w kierunku przeciwnym do ruchu wskazówek zegara (plik CB-Mar21-14.txt) i porównajmy ją z naszą oryginalną implementacją (plik CB-Mar21-08.txt). W tym przypadku możemy zastąpić oryginalne obliczenia modulo $(CurrentP + MAX_P) \% NUM_P$; w naszej głównej funkcji `loop()` operacją maskowania w następujący sposób (zmodyfikowana część jest pogrubiona):

```
void loop ()
{
    // Wyłącz stary piksel
    Neos.setPixelColor(RingXref[CurrentP], COLOR_OFF);

    // Zmniejszenie wskaźnika
    CurrentP = (CurrentP - 1) & MAX_P;

    // Włącz nowy piksel
    Neos.setPixelColor(RingXref[CurrentP], COLOR_ON);

    // Wyświetl wynik
    Neos.show();
}
```

```
delay(PadDelay);
}
```

Jeśli porównasz to z naszą wersją zgodną z ruchem wskazówek zegara, zobaczysz, że wszystko, co zrobiliśmy, to zmiana „+” na „-”. Więc, jak to działa?

Cóż, dla `CurrentP` równego 15 do `CurrentP` równego 1, $(CurrentP - 1)$ będzie równe odpowiednio 14 do 0. Odpowiada to wartościom $(CurrentP - 1)$ od 0x000E do 0x0000 w systemie szesnastkowym lub od 0000 0000 0000 1110 do 0000 0000 0000 0000 w systemie binarnym. Tak jak poprzednio, nasza operacja maski bitowej AND przepuści wszystkie te wartości bez szwanku.

Rozważmy teraz, co się stanie, gdy `CurrentP` będzie równe 0, co oznacza, że $(CurrentP - 1)$ będzie równe 0xFFFF w systemie szesnastkowym; 1111 1111 1111 1111 w systemie binarnym. W tym przypadku, gdy nasz bitowy AND zamaskuje trzy najbardziej znaczące nibble, wynikiem będzie 0x000F (0000 0000 0000 1111), czyli 15 w systemie dziesiętnym. W rezultacie nasze liczenie postępuje od 15 do 0, a następnie powraca do 15, czego właśnie szukamy.

Zaletą podejścia maskowego jest to, że w przypadku prostego mikrokontrolera zastąpienie dzielenia modulo prostym bitowym AND może potencjalnie zaoszczędzić kilka cykli zegara. Wadą jest to, że podejście modulo działa z pierścieniami zawierającymi dowolną liczbę pikseli, podczas gdy technika maski działa tylko wtedy, gdy liczba pikseli jest potęgą dwójki. ■

Clive „Max” Maxfield

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

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>



Obwody tłumiące zakłócenia

Na rynku nie brakuje urządzeń wprowadzających zakłócenia do instalacji elektrycznej. Jednocześnie wiele innych urządzeń bywa wrażliwych na te zakłócenia. Dlatego istotnym jest, by je zwalczać i filtrować.

O nieprawidłowym działaniu i jego przyczynach

Obwody i urządzenia generujące zakłócenia. Z każdym kolejnym rokiem w gospodarstwach domowych wzrasta liczba urządzeń elektronicznych, które wprowadzają do sieci elektrycznej sporo zakłóceń. Urządzenia te można z grubsza podzielić na sześć grup:

- Fazowe regulatory mocy oparte na tyrystorach lub triakach, na przykład ściemniacze, regulatory obrotów i grzałek.
- Prostowniki dużej mocy, w których napięcie wyjściowe jest regulowane przez tyrystory.
- Generatory wysokiej częstotliwości z szybkimi tranzystorami mocy, stosowane w generatorach ultradźwiękowych w przemyśle.
- Przetwornice impulsowe.
- Urządzenia zawierające silniki, które są włączane i wyłączane w losowych momentach.
- Układy z zasilaczami beztransformatorowymi, jak żarówki CFL lub LED.

Układy elektroniczne stosowane w takich obwodach i urządzeniach wytwarzają napięcia zakłócające sieci prądu przemiennego. Te napięcia zakłócające objawiają się w postaci sygnałów o wysokiej częstotliwości, które mają częstotliwość do 20 MHz. Amplituda tych sygnałów czasami przekracza oficjalnie dozwolone wartości tysiąckrotnie. Szczególnie krytyczny jest zakres częstotliwości od 150 kHz do 20 MHz. Linie zasilające zachowują się jak idealne anteny dla tych częstotliwości, emitując w przestrzeń silne pole elektromagnetyczne o wspomnianych częstotliwościach. Te fale radiowe mogą być odbierane przez inne przewody przenoszące inne sygnały, i tą drogą wnikać do urządzeń. Może to prowadzić do niepożądanego mieszania się tych sygnałów, co w rezultacie tworzy nowe częstotliwości mogące zakłócać prawidłowe działanie sprzętu elektronicznego.

Metody eliminacji zakłóceń. Zasadniczo można obrać dwie bardzo różne drogi, aby wyeliminować wpływ sygnałów zakłócających o wysokiej częstotliwości.

Pierwszą opcją jest rozwiązanie problemu u źródła, co sprowadza się do dodania układów tłumiących przy każdym potencjalnym źródle zakłóceń. Są one nazywane filtrami EMI, gdyż filtrują zakłócenia elektromagnetyczne wytwarzane przez dany obwód.

Drugim sposobem jest zapobieganie przenikaniu sygnałów zakłócających do urządzeń na nie wrażliwych. Można to zrobić poprzez zamontowanie filtrów na wejściach zasilania i sygnałowych.

Ekranowanie urządzeń i wrażliwych części obwodów jest trzecią metodą, często stosowaną w połączeniu z pozostałymi dwoma. Przyp. tłum.

Pochodzenie sygnałów zakłócających. Zakłócenia o wysokiej częstotliwości powstają w wyniku uniwersalnego prawa fizycznego, które ma zastosowanie do wszystkich zjawisk falowych. Prawo to zostało zbadane przez Fouriera i opisane w pełni matematycznie. Prawo Fouriera stwierdza, że każde zjawisko fali okresowej składa się z pewnej liczby fal sinusoidalnych lub cosinusoidalnych, których częstotliwości są równe wielokrotności częstotliwości podstawowej fali okresowej. Fale te nazywane

są „harmonicznymi sygnału podstawowego”. Wzajemne powiązania tych fal nazywane są „widmem częstotliwości sygnału podstawowego”.

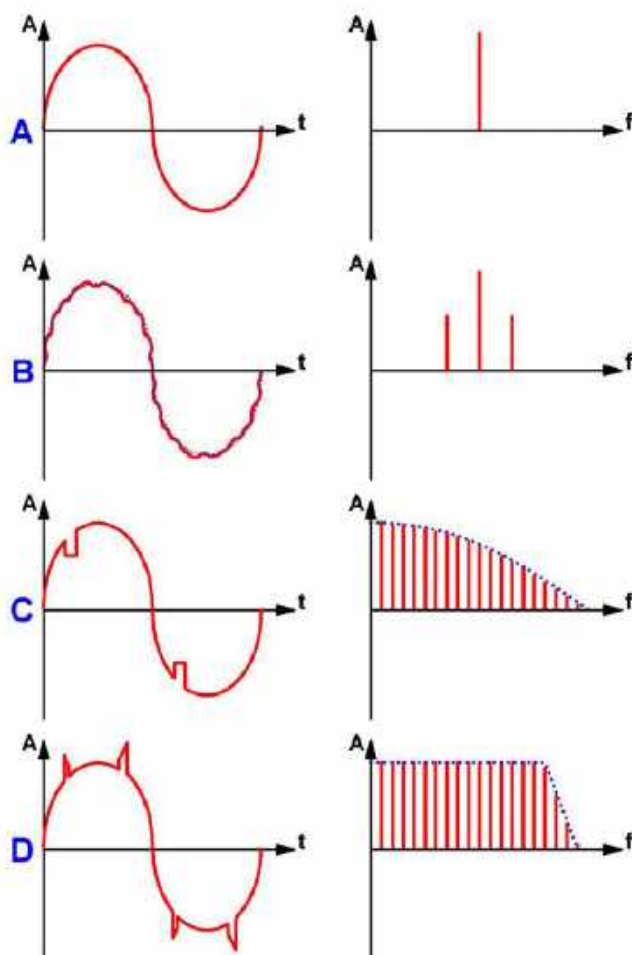
Kilka przykładów. Teorię Fouriera wyjaśniono za pomocą kilku przykładów przedstawionych na poniższym rysunku. Wykresy po lewej stronie przedstawiają szereg okresowych zjawisk falowych. Wykresy przedstawiają przebieg amplitudy A w funkcji czasu t . Wykresy po prawej stronie przedstawiają widmo częstotliwościowe sygnałów. W tym widmie częstotliwości, które składają się na sygnał, są reprezentowane przez pionowe linie na poziomej osi częstotliwości f . Wysokość linii reprezentuje amplitudę A różnych częstotliwości składowych.

- W przykładzie A narysowany jest czysty sygnał sinusoidalny. Widmo częstotliwości składa się tylko z jednej linii. Położenie tej linii na osi poziomej odpowiada oczywiście częstotliwości sygnału sinusoidalnego.
- W przykładzie B rysowany jest sygnał sinusoidalny zanieczyszczony innym sygnałem sinusoidalnym o znacznie wyższej częstotliwości. Sygnał ten objawia się jako niewielkie tętnienie na powierzchni fali sinusoidalnej. Widmo częstotliwości pokazuje, że sygnał ten składa się z trzech linii widmowych, z których dwie leżą po obu stronach linii bazowej i mają dość wysoką amplitudę.
- W przykładzie C narysowany jest sygnał sinusoidalny, zakłócony dwoma ostrymi spadkami na okres. Takie zakłócenie może wystąpić w napięciu sieciowym, na przykład w przypadku włączania dużego obciążenia co pół okresu za pomocą triaka. Ze względu na duży prąd rozruchowy, napięcie sieciowe „załamie się” na chwilę, powodując te typowe spadki. Analiza Fouriera pokazuje teraz wiele linii, z których tylko kilka jest narysowanych. W rzeczywistości linie widmowe są tak blisko siebie, że pozornie tworzą ciągły obszar. Najważniejszą rzeczą, jaką można wywnioskować z analizy Fouriera, jest to, że amplituda harmonicznych powoli maleje wraz ze wzrostem ich częstotliwości.
- W przykładzie D narysowana jest fala sinusoidalna, która jest zakłócana wąskimi szpilkami napięcia. Takie zakłócenie może wystąpić na przykład w przypadku okresowego włączania i wyłączania obciążenia indukcyjnego, takiego jak silnik. Duża indukcyjność własna uzwojeń silnika generuje wówczas za każdym razem dość wysokie napięcie wsteczne, które trafia do sieci zasilającej. Ponownie, analiza harmonicznych pokazuje, że sygnał ten zawiera wiele harmonicznych, które utrzymują dużą amplitudę w zakresie do MHz.

Rodzaje sygnałów zakłócających. Napięcia zakłócające o wysokiej częstotliwości mogą występować pod dwiema postaciami w sieci prądu przemiennego:

- jako symetryczne napięcia zakłócające,
- jako asymetryczne napięcia zakłócające.

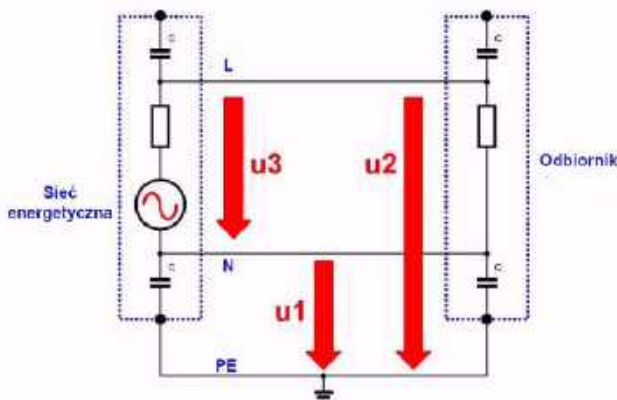
Różnicę między tymi dwoma źródłami zakłóceń wyjaśniono na poniższym rysunku. Sieć prądu przemiennego nie może i nigdy nie powinna być rozpatrywana oddzielnie od uzziemienia. W końcu jest to ważna część obwodu pierwotnego w większości urządzeń. Jak



1. Cztery przykłady widm sygnałów okresowych (© 2017 Jos Verstraten)

pokazuje schemat, uziemienie jest bezpośrednio połączone z metalową obudową uziemionego sprzętu. Obwód jest jednak zamknięty w tej obudowie. W rezultacie można zmierzyć pewną pojemność C między wszystkimi punktami obwodu a obudową. Rzeczywiście, między dowolnymi dwoma punktami przewodzącymi można z definicji założyć istnienie kondensatora!

Te pojemności pasożytnicze są ważnymi elementami w przekazywaniu sygnałów zakłócających ze źródła zakłóceń do urządzenia. Wynika to z faktu, iż pojemności pasożytnicze zapewniają bezpośrednie



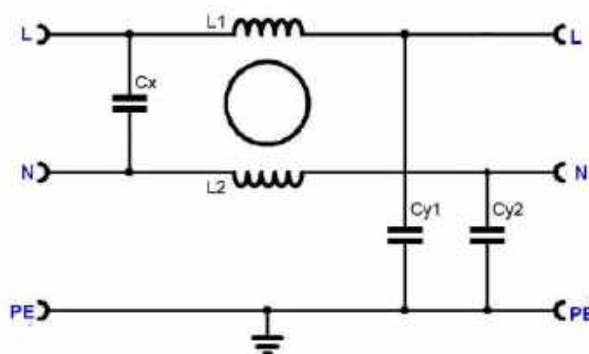
2. Różnica między symetrycznymi i niesymetrycznymi napięciami zakłócającymi (© 2017 Jos Verstraten)

połączenie między źródłem zakłóceń z jednej strony, a urządzeniem wrażliwym na zakłócenia z drugiej. Fakt, że to bezpośrednie połączenie zawiera dwa kondensatory szeregowo, nie ma większego znaczenia dla sygnałów zakłócających o wysokiej częstotliwości. W końcu dla tych wysokich częstotliwości kondensatory te mają raczej niską impedancję. Symetryczne sygnały zakłócające występują między przewodem fazowym i neutralnym sieci elektrycznej. Zakłócenia asymetryczne występują między przewodem fazowym, a uziemieniem, oraz między przewodem neutralnym, a uziemieniem.

Przewody uziemiające. Drugim ważnym czynnikiem jest to, że długi przewód uziemiający, podłączony do uziemienia na jednym końcu, może działać jak idealna antena dla przepływających przez niego prądów zakłócających o wysokiej częstotliwości. Dlatego podczas odłączania zasilania urządzeń należy nie tylko tłumić symetryczne napięcia zakłócające między fazą a przewodem neutralnym, ale także zwracać uwagę na redukcję sygnałów asymetrycznych.

Podstawowy układ filtra przeciwzakłóceniewego. Korzystając z ogólnych danych znanych do tej pory, można łatwo skonstruować ogólny schemat uniwersalnego filtra przeciwzakłóceniewego. Ten schemat został przedstawiony na poniższym rysunku. Lewa strona filtra biegnie do źródła zakłóceń, prawa strona jest podłączona do sieci prądu przemiennego, a tym samym biegnie przez tę sieć do urządzenia wrażliwego na zakłócenia. Kondensator C_x podłączony między fazą a przewodem neutralnym tłumi symetryczne sygnały zakłócające. Pomiędzy fazą a uziemieniem znajduje się filtr LC $L1/Cy1$, który pracuje jako filtr dolnoprzepustowy. Podobnie działa filtr złożony z $L2$ i $Cy2$.

Filtry tego typu spotyka się powszechnie w zasilaczach ATX, często przyłutowane bezpośrednio do gniazda IEC. Hobbysta może łatwo pozyskać taki filtr i wykorzystać go we własnym projekcie w miarę potrzeby. Przyp. tłum.



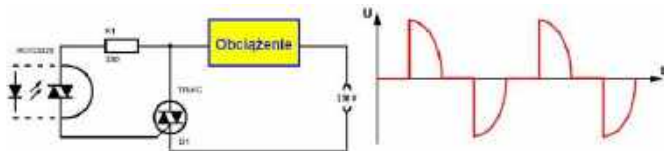
3. Ogólny schemat filtra przeciwzakłóceniewego (© 2017 Jos Verstraten)

Zakłócenia powodowane przez triaki i tyrystory

Wprowadzenie. Jednym z częściej występujących źródeł zakłóceń w konstrukcjach amatorskich są wszelkiej maści regulatory mocy zbudowane na triakach lub tyrystorach, pracujące według zasady regulacji kąta fazy przełączania. Zasada sterowania sprzężeniem fazowym została przedstawiona na poniższym rysunku. W obwodzie tym pracują triaki i optotriaki, ale zasada jest podobna dla tyrystorów. Triak włączony jest szeregowo z obciążeniem i normalnie nie przewodzi. Zostanie załączony, gdy przez bramkę zacznie przepływać niewielki prąd, ograniczony tutaj rezystorem. W tym przykładzie triak załączany jest za pomocą optotriaka zapewniającego galwaniczną izolację między obwodami pod napięciem sieci, a układem sterującym. W momencie załączenia triaka (lub tyrystora) będzie przewodził, a napięcie na nim się odkładające wyniesie 0,7 V. Triak (lub tyrystor) będzie przewodził od momentu załączenia do momentu, gdy napięcie (a zatem i prąd)

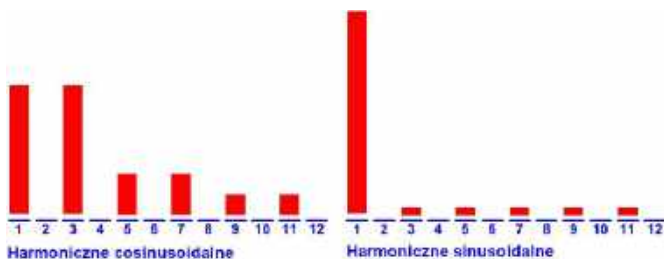
spadnie do zera, na końcu półokresu napięcia przemiennego. Należy też pamiętać, iż w momencie załączenia przepłyne maksymalny prąd, jaki przy danym napięciu pobiera urządzenie, i ten fakt jest źródłem generowanych przez fazowy regulator mocy zakłóceń.

Kąt otwarcia. Przebieg napięcia na obciążeniu jest narysowany po prawej stronie. W narysowanym przykładzie triak przewodzi dokładnie przez połowę każdego półokresu. Mówi się, że kąt otwarcia jest równy 90° . Oczywiście można zmieniać ten kąt w zakresie od 0° do 180° . Ta zmiana pozwala ustawić moc generowaną przez odbiornik w zakresie od zera do wartości maksymalnej.



4. Zasada działania obwodu z triakiem lub tyrystorem (© 2017 Jos Verstraten)

Widmo częstotliwości. Narysowany przebieg napięcia występującego na obciążeniu można, ponownie korzystając ze wzorów Fouriera, rozłożyć na składowe. Wówczas wyłania się obraz przedstawiony na poniższym rysunku. Prąd początkowo składa się z dwóch składowych, reprezentowanych na wykresach przez 1. Te dwie składowe mają taką samą częstotliwość jak napięcie sieciowe, tj. 50 Hz, ale są przesunięte względem siebie w fazie o 90° . Stąd jedna składowa nazywana jest składową sinusoidalną (w fazie z napięciem sieciowym), a druga składową cosinusoidalną (przesuniętą o 90° w fazie względem napięcia sieciowego). Jeśli pamiętasz cokolwiek z lekcji matematyki, natychmiast zrozumiesz, dlaczego te składowe nazywane są „sinusoidalną” i „cosinusoidalną”. Ponadto prąd składa się z szeregu składowych o wielokrotnościach częstotliwości sieci, nazywanych na wykresach odpowiednio 3, 5, 7 itd. Są to kolejne harmoniczne częstotliwości podstawowej, i tak 3 składowe mają częstotliwość równą trzykrotności częstotliwości napięcia sieciowego, tj. 150 Hz. Składowe sinusoidalne tych wyższych harmonicznych są dość słabe, ale amplituda prądów cosinusoidalnych jest dość duża. Na rysunku przedstawiono tylko amplitudy prądów do jedenastej harmonicznej. Nie oznacza to, że prąd nie zawiera wyższych harmonicznych! Można obliczyć (i przy okazji zmierzyć), że prądy o częstotliwościach znacznie przekraczających zakres MHz są generowane przez typowe regulatory fazowe.



5. Analiza Fouriera prądu przepływającego przez obciążenie (© 2017 Jos Verstraten)

Wpływ kąta otwarcia. Powinno być jasne, że kąt otwarcia regulatora fazowego pomaga określić wielkość harmonicznych sygnałów zakłócających. Przy kątach otwarcia 0° i 180° (pełne blokowanie i pełna moc), oczywiście nie będzie prądów harmonicznych. W pierwszym przypadku prąd w ogóle nie płynie; w drugim przypadku prąd płynie czysto sinusoidalnie. Pytanie, na które należy odpowiedzieć, dotyczy wpływu kąta otwarcia obwodu na amplitudy harmonicznych. Badania pokazują, że:

- Harmoniczne są najsilniejsze dla kątów otwarcia od 40° do 140° .
- Gradient amplitudy jest symetryczny względem kąta otwarcia 90° .
- Dla małych i dużych kątów otwarcia wszystkie harmoniczne mają w przybliżeniu taką samą wartość.
- Wraz ze wzrostem liczby harmonicznych amplituda harmonicznych staje się mniej zależna od kąta otwarcia.

Wnioski. Dane te pokazują bardzo wyraźnie, że głównie wyższe harmoniczne, które powodują najwięcej zakłóceń, można bardzo dobrze odfiltrować za pomocą pasywnego filtra dolnoprzepustowego. W rzeczywistości wielkość tych harmonicznych jest dość stała dla szerokiego zakresu kątów otwarcia, więc dość łatwo jest obliczyć zakres, w jakim dany filtr tłumi te harmoniczne.

Wyjściowe filtry przeciwzakłóceńowe

Od tego miejsca autor prezentuje ofertę producenta elementów indukcyjnych, firmy Schaffner. Produkty te są dostępne u dystrybutorów sprzedających elementy elektroniczne w Polsce, ale nie jest to jedyny producent takich komponentów. Przyp. tłum.

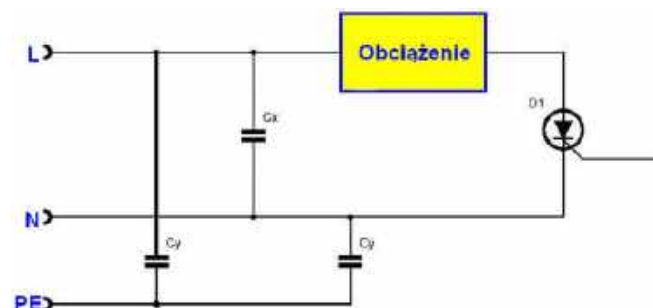
Wprowadzenie. Budowa filtra przeciwzakłóceńowego używanego jako filtr wyjściowy zależy od wielu czynników:

- Po pierwsze, oczywiście z oczekiwanej (i możliwej do obliczenia) amplitudy sygnałów zakłócających.
- Po drugie, jak bardzo chcesz tłumić sygnały zakłócające.
- Po trzecie, od rodzaju obciążenia, obciążenia czysto rezystancyjne, takie jak żarówki, wymagają innych filtrów niż obciążenia wysoce indukcyjne, takie jak silniki.

Amplituda sygnałów zakłócających jest w dużej mierze zależna od mocy przełączanej przez układ regulacji fazy. Ściemniacz oświetlenia sterujący lampą halogenową o mocy 100 W będzie generował mniej zakłóceń niż konsola sterująca oświetleniem scenicznym czy dyskotekowym o mocy kilku kilowatów.

Charakter obciążenia odgrywa ważną rolę w konstrukcji filtra. Silniki posiadają cewki, które w niektórych przypadkach mogą stanowić część filtra przeciwzakłóceńowego. Zwłaszcza jeśli mamy do czynienia z dużymi mocami, taka konstrukcja może zaoszczędzić jeden lub nawet dwa bardzo drogie i duże filtry przeciwzakłóceńowe.

Pojemnościowy filtr przeciwzakłóceńowy. Jeśli masz do czynienia z małymi, czysto rezystancyjnymi obciążeniami, możesz skonstruować filtr przeciwzakłóceńowy zgodnie z poniższym schematem z trzema kondensatorami. Kondensator C_x tłumi symetryczne sygnały zakłócające, a dwa kondensatory C_y redukują zakłócenia asymetryczne. Wartości kondensatorów C_y nie można zwiększyć do nieskończoności. Elementy te są przełączane między fazą a uziemieniem oraz między przewodem neutralnym a uziemieniem. Pomiędzy fazą a uziemieniem znajduje się pełne napięcie sieciowe. Przez kondensator popłynie więc pewien prąd, a jeśli prąd ten stanie się zbyt duży, wyłącznik RCD (różnicowoprądowy) zadziała i odetnie zasilanie. W praktyce



6. Pojemnościowy filtr przeciwzakłóceńowy może być używany z małymi, czysto rezystancyjnymi obciążeniami (© 2017 Jos Verstraten)

nie można przekroczyć wartości 2,2 nF. Ze względu na symetrię, oba kondensatory C_y powinny mieć tę samą wartość.

Wartość C_x również jest ograniczona. Jeśli kondensator ten będzie zbyt duży, po włączeniu napięcia sieciowego przez ten element popłynie bardzo duży prąd szczytowy. W rezultacie może dojść do zadziałania szybkich automatycznych bezpieczników nadprądowych lub przepalenia bezpiecznika topikowego w układzie. Ale jest jeszcze drugie niebezpieczeństwo! Kondensator jest wbudowany w urządzenie, które ma być tłumione. Urządzenie to jest w większości przypadków podłączane do gniazdka ściennego z przewodem zasilającym i wtyczką sieciową. Po wyjęciu wtyczki z gniazdka naładowany kondensator C_x znajduje się między dwoma bolcami wtyczki sieciowej. Przypadkowe dotknięcie tych styków może spowodować porażenie prądem elektrycznym. Dlatego w praktyce konstruktorzy ograniczają wartość kondensatora C_x do 220 nF.

Autor nie wspomniał o tym, na jakie napięcie powinny być te kondensatory. Zwykle stosuje się specjalizowane kondensatory przeciwzakłócenia o napięciu pracy do 400V. Druga rzecz, o której autor nie wspomniał to fakt, iż przy niesprawnej instalacji uziemiającej kondensatory C_y zachowują się jak dzielnik napięcia, przez co na obudowie będzie panować napięcie wynoszące około 115 V. Taka obudowa może „kopać”, lub sprawiać charakterystyczne „chropowate” wrażenie – nerwy czuciowe odbierają tak zmienne napięcie o częstotliwości 50Hz. Ze względu na małą pojemność kondensatorów, a przez to również wysoką impedancję dla tych częstotliwości prąd upływu jest bardzo mały. Warto jednak wezwać elektryka by sprawdził i ewentualnie naprawił instalację elektryczną. *Przyp. tłum.*

Gdy potrzebne są cewki. Do tłumienia zakłóceń dla dużych obciążeń rezystancyjnych i indukcyjnych nie można już używać filtrów czysto pojemnościowych. W takim przypadku należy przynajmniej użyć bardziej rozbudowanych filtrów przeciwzakłócenia wykorzystujących indukcyjność. Zależnie od zastosowania spotyka się:

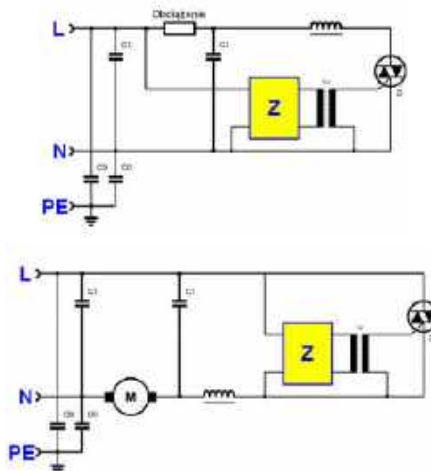
- dławiki nasycone,
- dławiki z kompensacją prądową,
- cewki z rdzeniem prętowym,
- cewki uziemiające.

Poniższe sekcje wyjaśnią obszar zastosowania tych cewek i dławików, oraz przedstawią praktyczne obwody tłumienia zakłóceń. Praktyczne obwody wykorzystują cewki tłumiące marki Schaffner. Jest to jeden z największych i najbardziej znanych producentów filtrów przeciwzakłócenia. Są to jednak tylko przykładowe komponenty, a konstruktor powinien kierować się przy zakupie części nie marką, a parametrami.

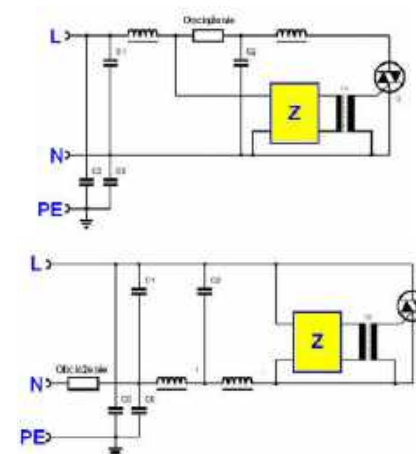
Nasycone dławiki przeciwzakłócenia

Właściwości. Nasycone dławiki są cewkami nawiniętymi na żelaznym rdzeniu. Części te mają dużą indukcyjność, gdy są włączone (prąd równy zero). Jednakże, gdy prąd wzrośnie do normalnej wartości, rdzenie cewek zostają nasycone, dzięki czemu indukcyjność staje się mała. Cewki tego typu są zawsze połączone szeregowo z tyrystorem lub triakiem. W rezultacie wysoka indukcyjność początkowa przeciwdziała nagłemu wzrostowi prądu. Wzrost prądu od zera do wartości maksymalnej trwa zatem nieco dłużej, co znacznie zmniejsza udział bardzo wysokich harmonicznych. Oczywiście już sam ten fakt skutkuje znaczną redukcją poziomu zakłóceń. Oczywiście cewki te nie są używane samodzielnie, lecz w połączeniu z kondensatorami. Utworzona w ten sposób sieć LC działa jak filtr dolnoprzepustowy, dodatkowo tłumiąc resztkowe składowe w.cz.

Obwody praktyczne. Poniższy rysunek po lewej stronie przedstawia przykład filtra z dławikiem nasyconym i kondensatorami. Obwód ten nadaje się do współpracy z obciążeniami rezystancyjnymi,



7. Praktyczne obwody z nasyconymi dławikami (© 2017 Jos Verstraten)



8. Praktyczne obwody z podwójnymi nasyconymi dławikami (© 2017 Jos Verstraten)

takimi jak reflektory teatralne lub elementy grzejne. Blok Z reprezentuje obwód wyzwalania triaka. Pojemność kondensatorów C_0 wynosi 2,2 nF. Nawiasem mówiąc, dotyczy to wszystkich dalszych praktycznych schematów. Wartość kondensatorów C_1 zależy od typu zastosowanego dławika. Schemat po prawej stronie można wykorzystać do tłumienia zakłóceń od obciążeń indukcyjnych, takich jak silniki. Indukcyjność cewki silnika jest teraz używana jako drugi dławik, tworząc dwie sieci LC.

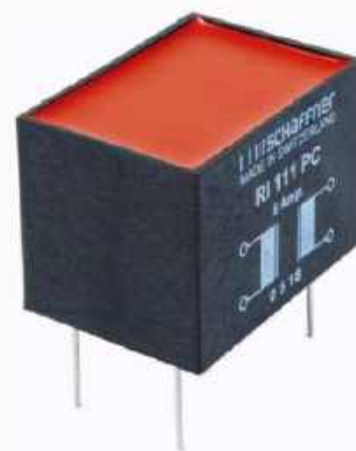
Podwójnie nasycone dławiki. Oprócz pojedynczych nasyconych dławików przeciwzakłócenia, Schaffner (i inni producenci – *przyp. tłum.*) dostarcza również podwójne dławiki o identycznej charakterystyce. Można ich używać w projektach, wobec których stawiane są bardzo wysokie wymagania dotyczące tłumienia zakłóceń. Dzięki dwóm dławikom i czterem kondensatorom, zgodnie z poniższymi schematami, można stworzyć filtry przeciwzakłócenia, które tłumią ponad 70 dB przy 150 kHz.

Dostępne typy. Poniższa tabela podsumowuje nasycone dławiki Schaffnera, z ich maksymalnym prądem i rezystancją wewnętrzną. Typy z przyrostkiem PC to wersje PCB, które można wlutować bezpośrednio w płytkę drukowaną.

Praktyczny schemat. Na zakończenie tej sekcji poświęconej dławikom nasyconym, prezentujemy poniższy rysunek przedstawiający przykład obwodu, który sprawdził się w praktyce. Obwód ten został zaprojektowany do zasilania reflektorów teatralnych o mocy 500 W.

Interesującą cechą tego schematu jest sposób sterowania bramką triaka. Otrzymuje ona prąd bramki poprzez wtórnik emiterowy T1 z transporem IC1. Napięcie +15 V, które zasilą ten obwód wyzwala, musi być całkowicie oddzielone od napięć zasilania obwodu sterującego! W skład obwodu sterującego wchodzi dioda LED transpatora oraz wszystkie obwody, które generują impulsy sterujące dla tej diody LED. To napięcie +15 V musi być zatem generowane z małego, oddzielnego transformatora zasilającego, używanego tylko do zasilania fototranzystora i wtórnika emitera tego pojedynczego stopnia. Oczywiście sterownik może zawierać więcej takich stopni, gdzie obwody zapłonowe są zasilane ze wspólnego źródła.

Choke type	Nominal current A @ 40°C	Circuit symbol	R ² mΩ/path	Weight approx. g
Ri 109 PC	2		280	65
Ri 110 PC	3		148	120
Ri 111 PC	6		42	170
Ri 113 PC	25		10	1320
Ri 207 PC	0.8		1325	50
Ri 209 PC	2		275	40
Ri 229 PC	2		265	30
Ri 230 PC	3		160	50
Ri 210 PC	3		160	65
Ri 231 PC	5		62	80
Ri 211 PC	6		43	70
Ri 221 PC	8		34	175
Ri 401 PC	1.5		620	15
Ri 403 PC	3		105	30
Ri 406 PC	6		53	55
Ri 410 PC	10		28	95
Ri 222	15		21	330
Ri 415	15		8	205
Ri 425	25		3.5	325



9. Nasycone dławiki zakłóceń Schaffnera (© Schaffner)

Dławiki z kompensacją prądową

Zasada działania. Dławiki z kompensacją prądową składają się z (zazwyczaj) toroidalnego rdzenia, na którym nawinięte są dwa identyczne uzwojenia. Są one używane do tłumienia asymetrycznych sygnałów zakłócających, które powstają między fazą a uziemieniem oraz między przewodem neutralnym a uziemieniem. Cewki te są zawsze włączone między fazę i przewód neutralny sieci zasilającej i zakończone dwoma małymi kondensatorami do uziemienia. Obie cewki są nawinięte na rdzeniu w taki sposób, że pola magnetyczne generowane przez sieć 50 Hz kompensują się wzajemnie. Cewka nie pełni więc żadnej funkcji dla napięcia sieciowego. Pełna indukcyjność działa tylko na asymetryczne napięcia zakłócające, które występują między fazą a uziemieniem oraz między przewodem neutralnym a uziemieniem.

Zastosowania. Dławiki z kompensacją prądową stosowane są w następujących sytuacjach:

- W obwodach z regulatorami fazowymi, gdzie obwody tłumiące z nasyconymi dławikami nie są wystarczająco skuteczne.
- W obwodach silnie zakłócających, takich jak generatory ultradźwiękowe i zasilacze impulsowe dużej mocy.

Praktyczne rozwiązanie układowe. Poniższy schemat przedstawia użycie dławika z kompensacją prądową w połączeniu z omawianymi wcześniej rozwiązaniami w celu lepszej filtracji zakłóceń.

Dostępne typy. Schaffner dostarcza dziesiątki dławików z kompensacją prądową z serii RN i RD. Można ich używać do prądów o natężeniu do 64 A. Równie bogatą ofertę znajdziemy wśród innych producentów elementów indukcyjnych. Zamiast więc wymieniać wszystkie dostępne typy, odsyłamy do katalogów dystrybutorów komponentów elektronicznych.

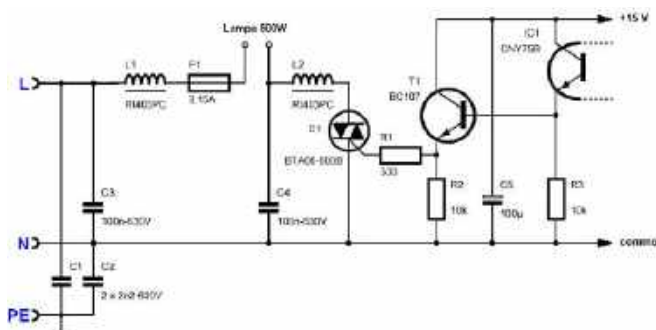
Cewki z rdzeniem prętowym

Działanie. Cewki z rdzeniem prętowym, jak sama nazwa wskazuje, składają się z pojedynczej cewki nawiniętej na rdzeń prętowy. Cewki z rdzeniem prętowym charakteryzują się stałą indukcyjnością. Rdzenie są zaprojektowane w taki sposób, aby nawet przy maksymalnym obciążeniu nie dochodziło do nasycenia materiału rdzenia.

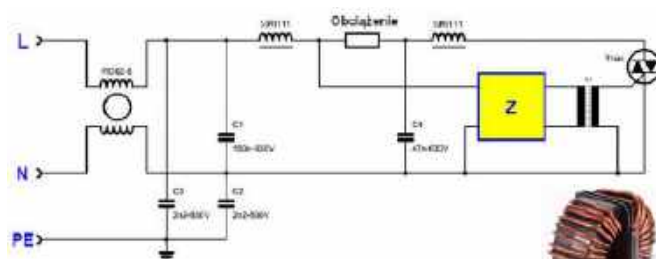
Zastosowania. Cewki te są głównie używane:

- W ciężkich warunkach przemysłowych do tłumienia odbiorników trójfazowych, takich jak silniki dużej mocy.
- W obwodach, w których zakłócenia wynikają głównie z napięć symetrycznych.

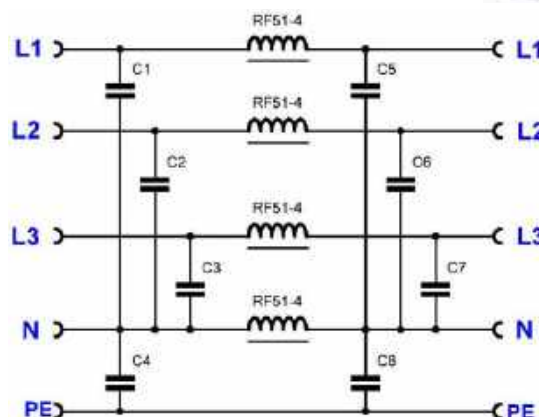
Przykładowy obwód. Na poniższym rysunku, jako przykład zastosowania cewek z rdzeniem prętowym, narysowano obwód tłumienia



10. Praktyczne zastosowanie dławików nasycenia w ściemniaczu światła (© 2017 Jos Verstraten)



11. Przykład dławika z kompensacją prądową i praktyczny schemat (© 2017 Jos Verstraten)



12. Standardowy filtr przeciwzakłóceńowy obciążenia trójfazowego z przewodem neutralnym (© 2017 Jos Verstraten)

Choke type	Nominal current A @ 40°C	Inductance L* mH	Circuit symbol	R ¹ mΩ	Weight approx. g
RF 51-4	4	2.4 (2)		310	250
RF 61-15	15	1.2 (1.2)		40	1300
RF 71-35	35	0.58 (0.35)		12	2720
RF 71-75	75	0.3 (0.06)		2	2800
RF 81-75	75	0.42 (0.3)		3.7	9060
RF 91-150	150	0.1 (0.08)		0.95	9400
RF 101-150	150	0.28 (0.22)		2.25	22000
RF 201-0.2/02	0.2	92 (90)		34000	30
RF 201-0.5/02	0.5	18.5 (18)		6300	32
RF 201-1/02	1	4.6 (4.4)		1900	35
RF 201-2/02	2	1.3 (0.84)		500	27
RF 201-0.2/07	0.2	92 (90)		34000	32
RF 201-0.5/07	0.5	18.5 (18)		6300	34
RF 201-1/07	1	4.6 (4.4)		1900	30
RF 201-2/07	2	1.3 (0.84)		520	30
RF 201-5/07	6	0.13 (0.08)		68	29
RF 211-0.5/02	0.5	50 (47)		10200	75
RF 211-1/02	1	13.6 (12.5)		3000	70
RF 211-2/02	2	3.8 (3.3)		820	70
RF 211-4/02	4	0.92 (0.68)		202	74
RF 211-6/02	6	0.39 (0.33)		100	75
RF 211-10/02	10	0.15 (0.1)		42	70
RF 211-0.5/14	0.5	50 (47)		10200	72
RF 211-1/14	1	13.6 (12.5)		3000	71
RF 211-2/14	2	3.8 (3.3)		820	74
RF 211-4/14	4	0.92 (0.68)		202	74
RF 211-6/14	6	0.39 (0.33)		90	76
RF 211-10/14	10	0.15 (0.1)		33	73

13. Przegląd cewek z rdzeniem prętowym Schaffnera (© 2017 Jos Verstraten)

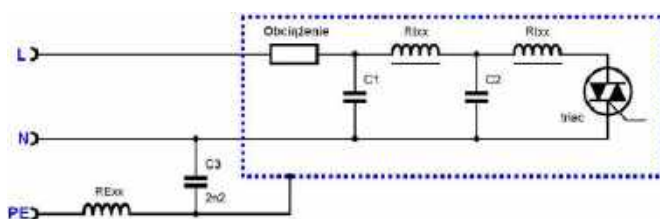
zakłóceń, który można zastosować do skutecznego tłumienia obciążenia trójfazowego z przewodem neutralnym. Cewki są połączone szeregowo z trzema przewodami fazowymi i szeregowo ze wspólnym przewodem neutralnym. Przed i za każdą cewką na przewodzie fazowym znajduje się kondensator do przewodu neutralnego. Przed i za cewką na przewodzie neutralnym znajdują się kondensatory do uziemienia.

Dostępne cewki z rdzeniem prętowym. Firma Schaffner posiada w swoim asortymencie cewki z rdzeniem prętowym, za pomocą których można wykonać wszystkie możliwe zadania związane z tłumieniem zakłóceń. Poniższa tabela pokazuje, że cewki są dostępne do maksymalnego prądu 150 A! Cewki te mają tylko dwa wyprowadzenia, więc ich użycie jest proste.

Cewki uziemiające

Działanie. Cewki uziemiające są włączone bezpośrednio w szereg z uziemieniem urządzenia i zapewniają dodatkowe tłumienie asymetrycznych prądów zakłócających, które mogą rozprasać się przez uziemienie. Istnieją dwa typy, jeden specjalnie zaprojektowany do tłumienia sygnałów zakłócających o niskiej częstotliwości, a drugi specjalnie zaprojektowany do tłumienia sygnałów zakłócających o wysokiej częstotliwości. Cechą cewek uziemiających jest to, że powinny być one nawinięte z drutu o tej samej średnicy, co zalecana w odpowiednich arkuszach norm uziemienia.

Typy o niskiej częstotliwości. Dławiki te mogą być używane w obwodach działających odwrotnie, niż normalne regulatory fazy, tj. napięcie płynie od rozpoczęcia półokresu do momentu odcięcia przez element wykonawczy. Dławiki te są umieszczane między uziemieniem obudowy urządzenia a uziemieniem sieci. Oczywiście obudowa nie może być wtedy bezpośrednio podłączona do uziemienia sieci! Charakterystyczną cechą



14. Zastosowanie cewki uziemiającej (© 2017 Jos Verstraten)

tych dławików jest częstotliwość rezonansowa wynosząca około 300 kHz.

Przykładowy obwód. Przykład zastosowania takich cewek tłumiących przedstawiono na poniższym rysunku. Oczywiście firma Schaffner (i wielu innych producentów) również dostarcza takie cewki. Dławiki te nie są dobrane na bazie prądu pracy, lecz powierzchni przekroju drutu, którym są nawinięte. W ten sposób zawsze można użyć cewki o takiej samej średnicy drutu, jaka jest wymagana przez lokalnie obowiązujące normy uziemienia.

Wejściowe filtry przeciwzakłóceńowe

Wprowadzenie. Jak już wspomniano w ogólnym wprowadzeniu, nie zawsze można wyeliminować niechciane zakłócenia u źródła. Dlatego bardzo ważne jest, aby wyposażać sprzęt wrażliwy na zakłócenia w filtry wejściowe. Filtry te tłumią sygnały zakłócające docierające przez instalację elektryczną, dzięki czemu urządzenie jest w mniejszym stopniu narażone na zakłócenia. Również w dziedzinie filtrów wejściowych istnieją oczywiście różne konfiguracje, z których każda ma określone właściwości tłumienia. Niemniej jednak filtry te nie nadają się do samodzielnego montażu. Kupuje się je gotowe, wbudowane w specjalną modułową obudowę lub nawet podłączone do sieci. W tym ostatnim przypadku bezpiecznik sieciowy jest często zintegrowany z filtrem.

Filtry wejściowe mają bardzo prostą budowę, i nic nie stoi na przeszkodzie, by je zintegrować z projektem zasilacza urządzenia. Hobbysta musi przeto pamiętać o zachowaniu bezpiecznych odstępów (jak w każdym układzie, gdzie występuje napięcie sieci), i rozważyć dodanie przekładki izolacyjnej od strony druku lub/i użycie lakieru izolacyjnego. Przyp. tłum.

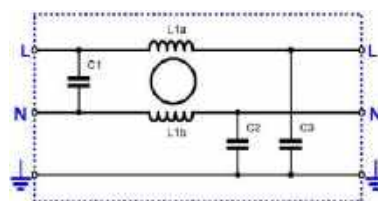
Filtry zintegrowane z gniazdem IEC. Schaffner dostarcza filtry wejściowe do urządzeń z uziemieniem, ale także do urządzeń bez uziemienia. Ponadto dostępne są typy z połączeniami lutowanymi, ale także z wtyczkami szybkorozłącznymi. Standardowy schemat takiego filtra sieciowego z uziemieniem przedstawiono na poniższym rysunku. Filtr składa się z dławika z kompensacją prądową i standardowych kondensatorów współpracujących z takim dławikiem.

Dostępne typy. Takie filtry są oferowane przez firmę Schaffner jako seria FN322 dla filtrów z uziemieniem i FN302 dla filtrów bez uziemienia. Inni producenci oferują własne filtry tych typów, dlatego warto zapoznać się z ofertami dystrybutorów.

Ważny montaż. Eliminacja zakłóceń w urządzeniu zasadniczo sprowadza się do doboru odpowiednich kondensatorów oraz dławika

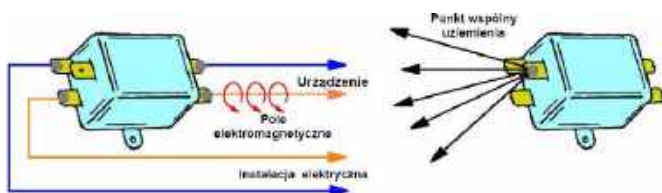
Najnowsze komentarze

Ważny montaż. Eliminacja zakłóceń w urządzeniu zasadniczo sprowadza się do doboru odpowiednich kondensatorów oraz dławika



15. Schemat filtra wejściowego i przykład takiego filtra (© 2017 Jos Verstraten)





16. Wszystko sprowadza się do przemyślanego okablowania!
(© 2017 Jos Verstraten)

bądź dławików. Należy jednak pamiętać o pewnych, istotnych detalach. Na przykład istotne jest odpowiednie prowadzenie przewodów (bądź ścieżek na płytce drukowanej). Jeśli bowiem przewody przed i za filtrem przeciwzakłóceńowym znajdują się zbyt blisko siebie, część zakłóceń może przeniknąć między nimi za sprawą sprzężenia pojemnościowego lub indukcyjnego między przewodami. Pokazuje to **rysunek 15**, po lewej.

Drugą rzeczą, o którą należy zadbać podczas instalacji filtrów przeciwzakłóceńowych, jest prawidłowe uziemienie całego urządzenia. Przewód uziemiający pochodzący z sieci zasilającej jest podłączony bezpośrednio do punktu uziemienia filtra, o ile oczywiście w urządzeniu nie ma cewki uziemiającej. Od tego centralnego punktu uziemienia, wewnętrzne połączenia uziemiające rozchodzą się gwiazdźście (w tym

miejscu należy też przyłączyć uziemienie do obudowy – przyp. tłum.), patrz **rysunek 16** po prawej.

Kondensatory. Kondensatory używane w filtrach przeciwzakłóceńowych są narażone na wysokie napięcie przemiennie i są dość delikatne. Należy używać wyłącznie części bardzo dobrej jakości! Ze względu na koszty można by pomyśleć, że kondensatory wysokonapięciowe typu ceramicznego są dobrym wyborem. Nie jest to jednak prawda! Wynika to z faktu, iż kondensatory te źle znoszą gwałtowne skoki napięcia i mogą ulec przebiciu lub innym uszkodzeniom. Jedyną użyteczną technologią są tak zwane „samonaprawiające się kondensatory z metalizowanego papieru”. Są one w stanie całkowicie zregenerować się po niewielkiej awarii spowodowanej skokiem napięcia. Wreszcie, powinno być jasne, że napięcie przebicia stosowanych kondensatorów powinno wynosić co najmniej 400 V. Zaleca się nawet stosowanie typów o napięciu przebicia 630 V.

Autor pomija zupełnie kondensatory polipropylenowe przeciwzakłóceńowe oznaczane jako X1, X2, Y1 i Y2. Kondensatory te przeznaczone są specyficznie do budowy opisanych w artykule filtrów. Tłumacz w swojej praktyce nigdy nie miał do czynienia z kondensatorami papierowymi w układach filtrów przeciwzakłóceńowych. Przyp. tłum. ■

Jos Verstraten

REKLAMA

w prezencie na każdą okazję
przejrzysz i kupisz na www.ulubionykiosk.pl

m.technik
Ciekawi świata są zawsze młodzi

eprasa.pl 4e05f27e00



Wykład 23

Filtry dolnoprzepustowe

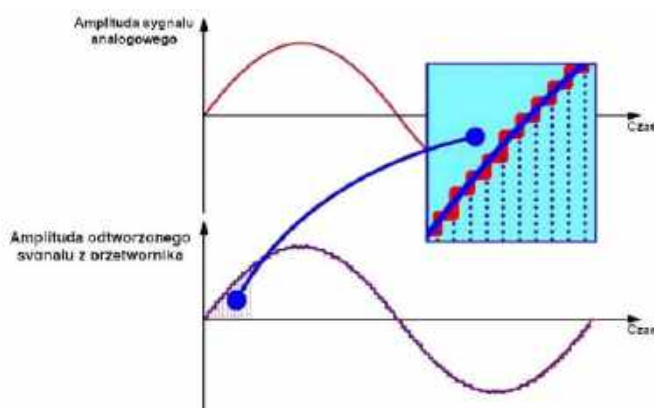
Filtry dolnoprzepustowe są bardzo użytecznymi obwodami, ponieważ są przydatne w zwalczaniu jednej z największych bolączek elektroniki: szumów.

Podstawowa zasada działania filtra dolnoprzepustowego

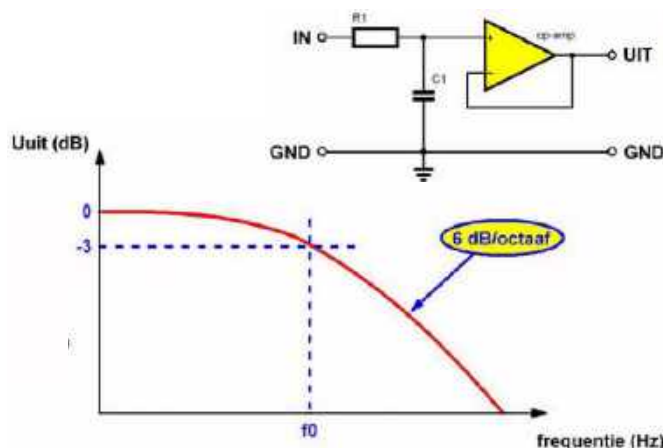
Tłumienie szumu. Zarówno w elektronice analogowej, jak i cyfrowej, szum jest jednym z głównych sygnałów zakłócających, z którymi ma do czynienia inżynier elektronik. Na poniższym rysunku przedstawiono na przykład, co się dzieje, gdy napięcie analogowe konwertuje się do postaci cyfrowej, przechowuje je w pamięci, a następnie konwertuje z powrotem na napięcie analogowe. Przetwornik cyfrowo-analogowy, który odtwarza to napięcie analogowe, z definicji zapewnia „przybliżenie krokowe” oryginalnego napięcia analogowego. Te niewielkie skoki sygnału objawiają się w postaci irytującego „szumu cyfrowego”, zwanego również „szumem kwantyzacji”. Szum ten można usunąć z odzyskanego sygnału analogowego za pomocą filtra dolnoprzepustowego.

Podczas odtwarzania starych płyt winylowych można również dobrze się bawić z filtrem dolnoprzepustowym, który w mniejszym lub większym stopniu eliminuje szumy obecne często w starych nagraniach.

Jak działa filtr dolnoprzepustowy. Wszystkie analogowe filtry dolnoprzepustowe działają na tej samej zasadzie. Filtry te składają się z kombinacji rezystorów i kondensatorów (filtry mogą też zawierać elementy indukcyjne, cewki lub dławiki – przyp. tłum.). Kondensatory mają impedancję (rezystancję napięcia przemiennego), która zależy od częstotliwości sygnału. Dla częstotliwości 0 Hz, tj. dla napięcia stałego, kondensatory mają nieskończenie wysoką impedancję. Dla nieskończenie wysokiej częstotliwości komponenty te mają impedancję 0 Ω . Dzięki tej właściwości można nadać filtrowi pożądane właściwości zależne od częstotliwości.



1. Wyjaśnienie zjawiska „szumu cyfrowego” (© 2019 Jos Verstraten)



2. Filtr dolnoprzepustowy pierwszego rzędu bez wzmacnienia (© 2019 Jos Verstraten)

Filtr pierwszego rzędu bez wzmocnienia. Podstawowy schemat dolnoprzepustowego filtra pierwszego rzędu bez wzmocnienia przedstawiono na poniższym rysunku. Obwód składa się z pasywnej sieci RC R1-C1, zakończonej wzmacniaczem operacyjnym włączonym jako wtórnik napięciowy. Wzmocnienie napięciowe w paśmie przepustowym filtra wynosi 1.

Następnie filtr będzie tłumił o 6 dB/oktawę. Krzywa charakterystyki amplitudowej wygląda zatem tak, jak pokazano na rysunku po prawej stronie. Przy częstotliwości odcięcia f_0 Hz tłumienie wynosi 3 dB.

Działanie. Działanie obwodu jest łatwe do wyjaśnienia. Przy niskich częstotliwościach kondensator ma bardzo wysoką impedancję i można go pominąć. Obwód działa wtedy jako wzmacniacz buforowy. Wraz ze wzrostem częstotliwości w grę wchodzi impedancja kondensatora. Powstaje dzielnik napięcia, utworzony przez stałą impedancję rezystora i zmienną impedancję kondensatora. Wraz ze wzrostem częstotliwości coraz większa część napięcia wejściowego będzie odkładać się na rezystorze, a coraz mniejsza na kondensatorze. Obwód będzie zatem selektywnie tłumić częstotliwość. Oczywiście taki filtr nie ma dobrych parametrów i nie będzie zbyt często używany w praktyce. **Autor się myli gdyż proste, pasywne filtry RC i LC spotyka się bardzo często w najróżniejszych obwodach. Czasami powstają zupełnie przypadkowo, na przykład z pojemności pasożytniczych tranzystora i rezystancji wypadkowej rezystorów wokół niego, co w efekcie ogranicza od góry pasmo przenoszenia obwodu. Przyp. tłum.**

Filtr pierwszego rzędu ze wzmocnieniem. Podstawowy schemat filtra dolnoprzepustowego pierwszego rzędu można rozszerzyć do schematu przedstawionego na poniższym rysunku. Wzmacniacz operacyjny jest teraz włączony jako nieodwracający wzmacniacz napięcia. Współczynnik wzmocnienia jest określany przez stosunek rezystorów R2 i R3. **Częstotliwość odcięcia f_0 .** Częstotliwość odcięcia f_0 filtra dolnoprzepustowego pierwszego rzędu jest określona wzorem:

$$f_0 = \frac{1}{2 \cdot \pi \cdot R1 \cdot C1}$$

Filtry dolnoprzepustowe drugiego rzędu

Komentarz. Filtry te noszą również nazwę filtrów Sallen-Key.

Filtr drugiego rzędu bez wzmocnienia. Znacznie bardziej interesujący jest podstawowy schemat filtra dolnoprzepustowego drugiego rzędu, przedstawiony na poniższym rysunku. Obecne są dwie sieci RC, z których druga znajduje się w obwodzie sprzężenia zwrotnego wzmacniacza operacyjnego. Działanie jest następujące. Przy niskich częstotliwościach kondensatory mają bardzo wysokie impedancje i można je pominąć. Schemat redukuje się wtedy do wzmacniacza buforowego, z dwoma rezystorami połączonymi szeregowo z wejściem. Jednakże, ponieważ wejście wzmacniacza operacyjnego ma bardzo wysoką impedancję (**nie zawsze, zależy od typu wzmacniacza operacyjnego – przyp. tłum.**), na rezystorach nie będzie żadnego spadku napięcia. Wzmocnienie obwodu jest równe 1.

Wraz ze wzrostem częstotliwości impedancja kondensatorów staje się coraz mniejsza. Sprzężenie zwrotne występuje od wyjścia do wejścia poprzez kondensator C1. Kondensator C2 z kolei tworzy dzielnik napięcia filtra pierwszego rzędu. Konsekwencją tego podwójnego działania jest to, że filtr będzie tłumić o 12 dB/oktawę po częstotliwości odcięcia.

Współczynnik dobroci Q. W większości przypadków można nadać identyczne wartości obu rezystorom. Charakterystyka filtra wokół częstotliwości odcięcia f_0 zależy od stosunku dwóch kondensatorów. Charakterystykę tę wyraża współczynnik dobroci Q filtra.

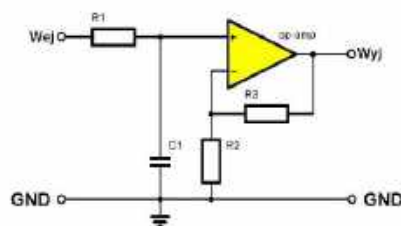
Dla opisanego obwodu można obliczyć ten współczynnik za pomocą równania:

$$Q = \frac{\sqrt{R1 \cdot R2 \cdot C1 \cdot C2}}{C2 \cdot (R1 + R2)}$$

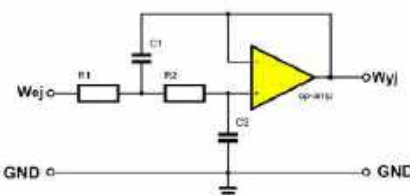
Wpływ współczynnika jakości na charakterystykę amplitudową przedstawiono na poniższym rysunku. Widoczne są trzy praktyczne krzywe narysowane dla współczynników dobroci od 1 do 0,5. Wyraźnie widać, że wszystkie trzy krzywe podążają za nachyleniem 12 dB/oktawę (czego należało się spodziewać), ale wokół częstotliwości odcięcia może występować dodatkowe tłumienie lub niewielki wzrost sygnału. Jeśli zależy Ci na możliwie płaskiej charakterystyce amplitudowej, nadaj Q wartość dokładnie 0,707!

Filtr drugiego rzędu ze wzmocnieniem. Można przekształcić schemat filtra bez wzmocnienia w obwód ze wzmocnieniem sygnału w bardzo prosty sposób. Schemat pokazano na poniższym rysunku. Wzmocnienie napięcia jest określone przez stosunek rezystorów R3 i R4.

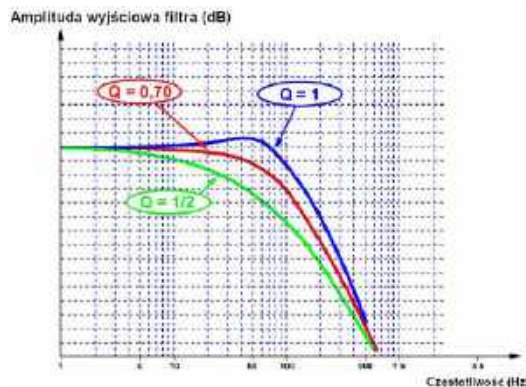
Różne koncepcje filtrów. W przypadku filtra dolnoprzepustowego drugiego rzędu, panowie Bessel, Chebyshev i Butterworth mają coś do powiedzenia. Mianowicie, jeśli nadamy dwóm rezystorom i dwóm kondensatorom sieci filtrów identyczne wartości R i C, okaże się, że stosunek między



3. Filtr dolnoprzepustowy pierwszego rzędu ze wzmocnieniem napięciowym (© 2019 Jos Verstraten)



4. Podstawowy schemat dolnoprzepustowego filtra drugiego rzędu bez wzmocnienia (© 2019 Jos Verstraten)



5. Wpływ współczynnika dobroci Q na charakterystykę amplitudową filtra drugiego rzędu (© 2019 Jos Verstraten)

rezystorami R3 i R4 ma duży wpływ na charakterystykę filtra. Można zdefiniować współczynnik sprzężenia zwrotnego α , którego wartość jest określona wzorem:

$$\alpha = \frac{R3}{R3 + R4}$$

Nadając temu współczynnikowi α nieprawidłową wartość, można nawet spowodować oscylacje obwodu! Poniższa tabela podaje wartości tych rezystorów dla trzech koncepcji filtrów. Wartość pozostałych komponentów wynosi:

$R1=R2=R=15,9 \text{ k}\Omega$

$C1=C2=C=10 \text{ nF}$

Istnieją dwie wartości dla Chebysheva, a mianowicie +1,25 dB i +3,0 dB. Są to wartości maksymalnych oscylacji w charakterystyce amplitudy, które mają być tolerowane.

Wszystkie koncepcje w jednym obwodzie. Poniższy rysunek pokazuje sposób, w jaki można przełączać się z jednej koncepcji filtra na inną, po prostu obracając potencjometr. Jest to oczywiście proces płynny i nie można powiedzieć „teraz opuszczam obszar Bessela i wchodzę w obszar Butterwortha” ani nic podobnego!

Gdy potencjometr znajduje się w górnym położeniu ($\alpha=1$), wzmacnienie wzmacniacza operacyjnego jest wyłączone i wzmacniacz pracuje jako bufor napięcia. W miarę przesuwania suwaka potencjometru bardziej w dół, wartość α maleje i pojawiają się odpowiednio charakterystyki Bessela, Butterwortha i Chebysheva. Ale jest pewna granica! Mianowicie, jeśli przesuniesz suwak potencjometru zbyt mocno w dół, obwód zacznie się coraz bardziej wzmacniać, a wartość współczynnika dobroci Q będzie coraz bardziej spadać. Przy α równym 0,333 obwód przestaje działać jako filtr, a zaczyna jako oscylator!

Filtr MFB. Omówiony do tej pory schemat przedstawia standardowy układ aktywnego filtra dolnoprzepustowego drugiego rzędu. Powstały jednak niezliczone wariacje tego schematu, z których jedna zostanie teraz pokrótce omówiona. Ten tak zwany filtr MFB, skrót od „Multiple FeedBack”, ma wiele sprzężeń zwrotnych, patrz rysunek poniżej.

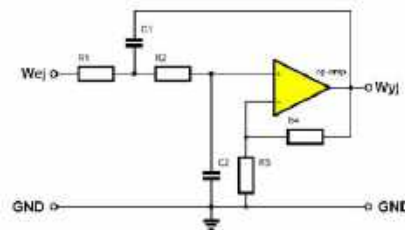
Omówienie wszystkich właściwości tego filtra byłoby zbyt obszerne w kontekście tego artykułu. Jednak zaletą dodatkowego sprzężenia zwrotnego jest to, że obwód jest bardziej stabilny i mniej prawdopodobne jest, że wejdzie w stan, który może spowodować oscylacje.

Zastosowanie filtra dolnoprzepustowego drugiego rzędu

Filtr szumów dla płyt winylowych. Obecnie płyty winylowe stały się znów modne wśród melomanów. Może to sprawić, że i Ty znów zaczniesz słuchać płyt odnalezionych gdzieś na strychu. Nostalgia! Chociaż nie jest to absolutnie konieczne, dobry filtr szumów jest przydatnym dodatkiem do analogowego systemu odtwarzania dźwięku. Dzięki takiemu filtrowi można znacznie zredukować syczenie i plumkanie starych lub często odtwarzanych płyt.

Poniższy rysunek przedstawia praktycznie użyteczny schemat filtra dolnoprzepustowego drugiego rzędu z trzema częstotliwościami odcięcia: 5, 7 i 11 kHz. Stromość tłumienia wysokich częstotliwości wynosi 12 dB/oktawę, dzięki czemu użyteczny sygnał audio jest oszczędzany w jak największym stopniu. Dodatkowy przełącznik S2 umożliwia wyłączenie filtra, bezpośrednio łącząc wejście i wyjście. Wzmocnienie nietłumionych sygnałów jest równe jeden, więc włączenie filtra do systemu reprodukcji dźwięku nie ma nieprzyjemnego wpływu na ustawienie regulacji głośności.

Ponieważ jednak większość gramofonów jest wyposażona w wkładkę elektrodynamiczną, należy również użyć przedwzmacniacza RIAA. Wyjście tego przedwzmacniacza można podłączyć do wejścia tego filtra szumów. Wyjście filtra należy podłączyć do jednego z wejść AUX wzmacniacza.



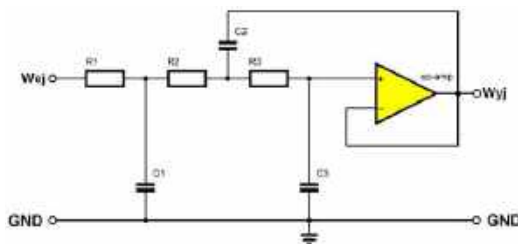
W przeciwieństwie do wszystkich obwodów opisanych do tej pory w tym artykule, ten obwód nie wykorzystuje wzmacniacza operacyjnego, ale dwa stare tranzystory typu BC109C. Są one znane z doskonałej charakterystyki szumów i dlatego zostały użyte w tym obwodzie.

Omawiany wcześniej filtr drugiego rzędu zbudowany został z wykorzystaniem tranzystora T2. Tranzystor T1 pracuje jako wtórnik emiterowy o niskiej impedancji wyjściowej, jednocześnie wstępnie polaryzując bazę T2. Takie rozwiązanie jest konieczne gdyż filtr aktywny musi być sterowany z niskiej i stałej impedancji, w przeciwnym razie obwód, do którego podłączony jest filtr, będzie miał wpływ na częstotliwość graniczną filtra i jego stromość. Rezystory R1 i R2 ustalają punkt pracy T1 tak, by na rezystorze emiterowym panowało napięcie równe połowie napięcia zasilania. Drugi tranzystor jest spolaryzowany przez rezystory szeregowo R4 i R5.

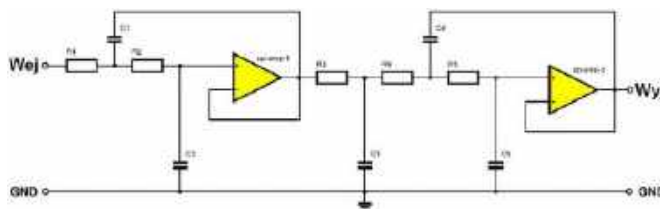
Napięcie wejściowe jest podawane na bazę pierwszego tranzystora poprzez kondensator separujący C1. Kondensator ten oddziela napięcie sygnału od napięcia polaryzacji bazy.

Filtr ma trzy częstotliwości odcięcia, zależne od wartości kondensatorów w filtrze. Zmiana częstotliwości filtra odbywa się za pomocą przełączników S1a i S1b. Napięcie wyjściowe filtra jest pobierane z emitera drugiego tranzystora za pomocą filtra RC C8-R7.

Oczywiście drugi kanał ma identyczną budowę. Zasilanie +Ub dla obu filtrów można podłączyć z odpowiednio wygładzonego zasilacza o napięciu co najmniej 12 V.



10. Filtr dolnoprzepustowy trzeciego rzędu (© 2019 Jos Verstraten)



11. Filtr dolnoprzepustowy piątego rzędu (© 2019 Jos Verstraten)

Filtry wyższego rzędu

Kaskadowanie omawianych obwodów. Filtry dolnoprzepustowe trzeciego, czwartego, piątego itd. rzędu są tworzone poprzez szeregowe połączenie kolejnych filtrów pierwszego i drugiego rzędu. Nazywa się to „kaskadowaniem” obwodów.

Filtr trzeciego rzędu. Tak więc na poniższym rysunku narysowany jest filtr trzeciego rzędu, składający się z już znanych obwodów filtra pierwszego i drugiego rzędu. Należy zauważyć, że wzmacniacz operacyjny włączony jako bufor związany z filtrem pierwszego rzędu nie jest uwzględniony na schemacie. W rezultacie filtr pierwszego rzędu ma pewien wpływ na filtr drugiego rzędu. Dzieje się tak, ponieważ rezystory są połączone szeregowo.

Filtr piątego rzędu. W identyczny sposób, patrz rysunek poniżej, można zaprojektować filtr piątego rzędu, łącząc kolejno dwa filtry drugiego i trzeciego rzędu. Taki filtr ma stromość 30 dB/oktawę, czyli całkiem sporo!

Obliczanie filtrów dolnoprzepustowych za pomocą FilterLab

FilterLab 2.0. Poniższa sekcja artykułu dotyczy starszej wersji oprogramowania Microchip FilterLab. Obecna wersja 3.0 działa w przeglądarce i nie wymaga instalacji, oferując te same możliwości. Przyp. tłum.

Czasami, w których trzeba było używać skomplikowanych formuł do obliczania filtrów analogowych, na szczęście mamy już za sobą. Istnieje darmowe oprogramowanie, które wykonuje tę pracę za użytkownika. Jednym z najprzyjemniejszych programów, ale to oczywiście bardzo osobisty gust, jest „FilterLab” udostępniony za darmo przez producenta układów scalonych Microchip. W chwili publikacji oryginalnej wersji tego artykułu dostępna jest wersja 2.0 tego programu dla systemu Windows. Program i bardzo obszerny podręcznik w języku angielskim, liczący nie mniej niż 76 stron, można pobrać ze strony: <https://www.microchip.com/developmenttools/ProductDetails/filterlabdesignsoftware>.

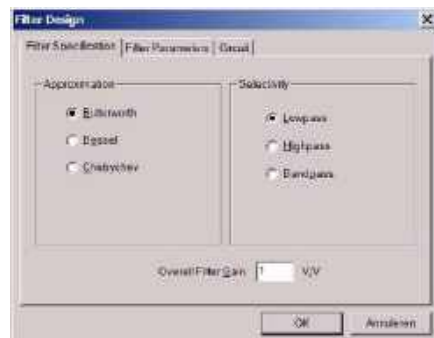
W przypadku, gdy w momencie czytania tego artykułu oprogramowanie to nie będzie już dostępne przez Internet, zawsze można poprosić o jego kopię autora tego bloga pod adresem josverstraten@live.nl.

Możliwości z FilterLab 2.0. Za pomocą tego bezpłatnego oprogramowania można obliczać filtry dolnoprzepustowe do ósmego rzędu i zgodnie z koncepcjami filtrów Czebyszewa, Bessela i Butterwortha. Zakres częstotliwości wynosi od 0,1 Hz do 1 MHz.

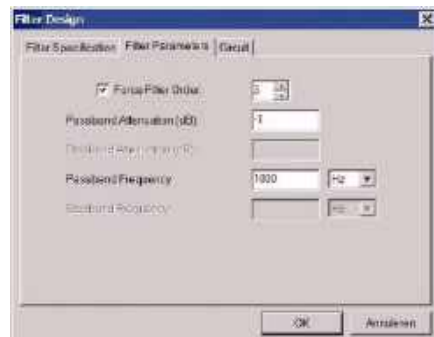
Również wersja 3.0 ogranicza wybór częstotliwości do zakresu 0,1 Hz...1 MHz. Przyp. tłum.

Instalacja FilterLab 2.0. Zainstalowaliśmy to oprogramowanie w systemie Windows 7.0 i odbywa się to całkowicie automatycznie i bez żadnych problemów. Najpierw należy utworzyć nowy folder dla tego oprogramowania, aby można je było łatwo znaleźć. Żaden skrót nie jest instalowany na pulpicie.

Specyfikacja filtra. Za pomocą menu „Filter” (Filtr) i opcji „Design” (Projekt) można przejść do poniższego okna, w którym pierwszą rzeczą jest wybranie opcji „Lowpass”



12. Wybór typu filtra w FilterLab 2.0 (© 2019 Jos Verstraten)



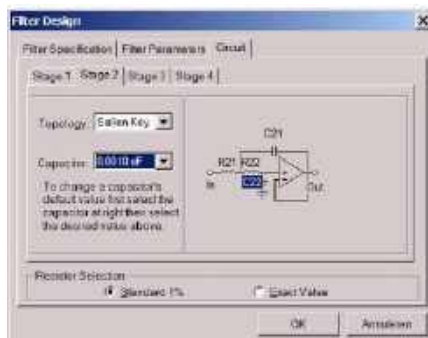
13. Wybór typu obwodu i wartości kondensatorów (© 2019 Jos Verstraten)

(Filtr dolnoprzepustowy) i żądanej koncepcji filtra w zakładce „Filter Specification” (Specyfikacja filtra). „Ogólne wzmocnienie filtra” pozwala ustawić pożądane wzmocnienie obwodu, nie w dB, ale po prostu jako współczynnik.

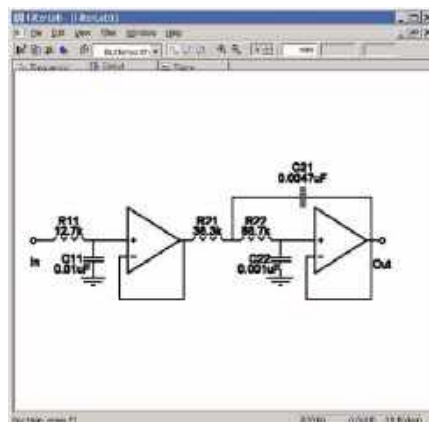
Parametry filtra. W tej zakładce należy wprowadzić kolejność filtra (Force Filter Order), żądane tłumienie w dB przy częstotliwości odcięcia f_o (Passband Attenuation) w dB oraz żadaną częstotliwość odcięcia (Passband Frequency).

Wybór żadanego obwodu. W zakładce „Obwód” można wybrać żądany obwód. Schemat Sallen & Key jest tutaj ustawiony jako domyślny, ale można również wybrać schemat MFB. Następnie należy wybrać opcję „Standard 1%”, aby program obliczył standardowe wartości rezystancji i kondensatorów. Istnieje możliwość ustalenia wartości kondensatorów, na przykład dlatego, że masz już dokładne wartości w domu i chcesz, aby program obliczył tylko rezystory pasujące do tych kondensatorów.

Obwód na ekranie. Po kliknięciu przycisku „OK”, obwód obliczony przez FilterLab natychmiast pojawi się na ekranie w własnym małym oknie. W podobnych oknach można również wyświetlić charakterystykę częstotliwościową i fazową filtra. Oczywiście można wydrukować zarówno schemat, jak i charakterystykę za pomocą drukarki systemowej Windows. ■



14. Wybór parametrów filtra (© 2019 Jos Verstraten)



15. Obliczony przez FilterLab schemat filtra (© 2019 Jos Verstraten)

Jos Verstraten

REKLAMA

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

Aplikacje prosimy kierować na adres:
redakcja@elportal.pl

Przystawka do quiz-ów dla ośmiu graczy

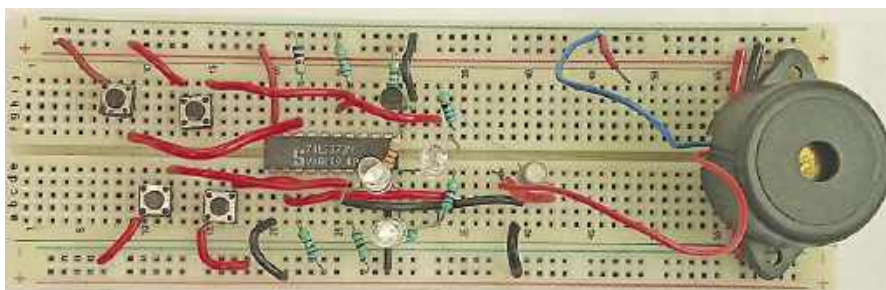
Wiele gier/quizów organizowanych w szkołach lub uczelniach wymaga szybkiej reakcji od graczy. Kto pierwszy ten odpowiada. Nawet szybka reakcja człowieka, jest na tyle wolna, że prosty obwód elektroniczny rozpozna, kto pierwszy nacisnął przycisk. Zwykle, tego rodzaju układy wykorzystują mikrokontroler, jednak nasz układ nie angażuje żadnego z takich. Sercem układu jest zatrzask typu 74LS373. Na rysunku 1 pokazano zdjęcie prototypu wykonanego przez autora na uniwersalnej płycie stykowej.

Opis układu i jego działanie

Schemat ideowy układu pokazuje rysunek 2. Głównymi podzespołami są tu: transformator obniżający napięcie sieciowe (X1), mostek prostowniczy (BR1), stabilizator 5 V typu LM7805 (IC1), osiem przycisków niestabilnych (zwierających styki po naciśnięciu; S1 do S8), zatrzask typu 74LS373 (IC2), dwa tranzystory NPN 2N2219 (T1, T2), buzzer piezoelektryczny z generatorem (PZ1), dziewięć diod LED (LED1 do LED9) i ponadto kilka dyskretnych elementów pasywnych.

Układ zasilany jest napięciem pozyskiwanym z sieci energetycznej 230 V AC za pomocą typowego układu zasilacza napięcia stałego 5 V. Jest to zasilacz liniowy o niewielkiej sprawności. Transformator obniża napięcie przemienne do poziomu 9 V, które jest następnie prostowane dwupołkowo mostkiem BR1, filtrowane pojemnością kondensatora C1 i następnie stabilizowane regulatorem liniowym LM7805. Kluczowym elementem jest zatrzask 74LS373.

Przypis redakcji: 74LS373 składa się z ośmiu niezależnych zatrzasków typu D; praca układu rozróżnia dwa stany. Albo jest „przezroczysty” powielając wejście na wyjście, albo zapamiętuje stan wejść w określonej chwili czasowej; to zależy od stanu linii LE (Latch Enable). Linia OE (Output Enable) jest cały czas aktywna (w stanie niskim); i w konstrukcji elementu jest przewidziana do wprowadzenia wyjść w stan wysokiej impedancji,



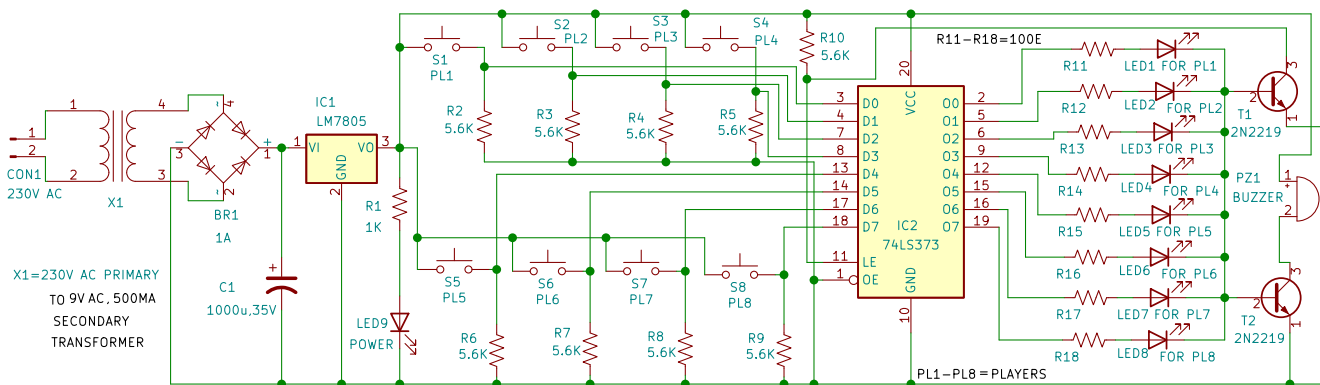
Rysunek 1. Prototyp wykonany przez autora

co jest zwykle bardzo pożądane w aplikacjach magistralowych.

Zatrzask IC2 74LS373 przenosi stan wejść D0 do D7 na wyjścia Q0 do Q7. Wszystkie wejścia (D0 do D7) wyposażono w rezystory ściągające do stanu niskiego (R2 do R9). Do wejść podłączone są także przyciski S1 do S8, których drugi koniec podłączony jest do linii zasilania +5 V. Zatem naciśnięcie dowolnego z nich wymusza stan wysoki (który nazwiemy jedynką logiczną). Idea działania układu polega na tym, że naciśnięcie dowolnego z przycisków S1 do S8 zaświeca odpowiednią diodę na wyjściu i uruchamia dźwięk buzzera. Równocześnie w tym samym momencie ulega zmianie stan wejścia LE z wysokiego na niski, co zatrzaskuje dane wejściowe w przerzutnikach (stanowiących elementarną pamięć cyfrową). Od tego momentu wyjścia nie mogą zmienić swego stanu logicznego i stan ten utrzymuje się do momentu wyzerowania układu. W obwodzie sprzężenia zwrotnego między wyjściami

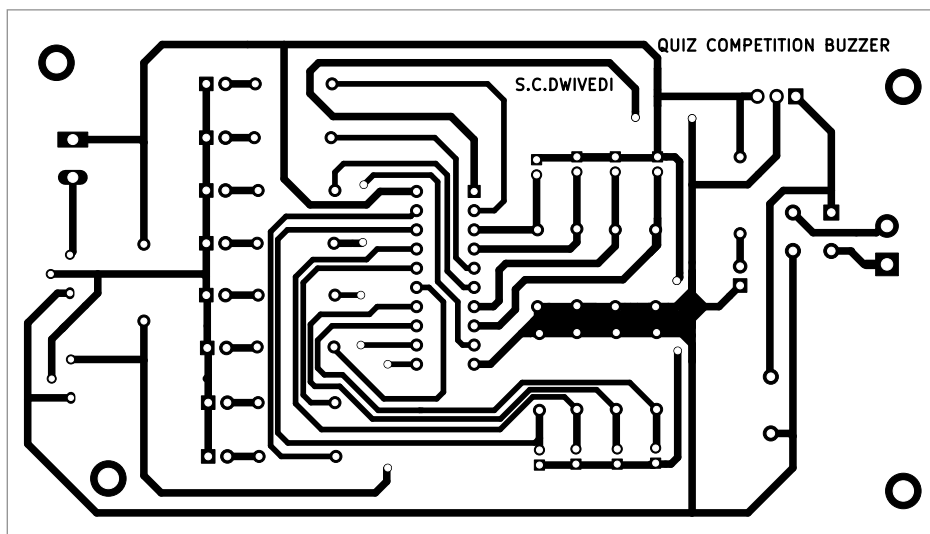
na wejściu „zatrzaskującym” pracuje tranzystor T1 wraz z rezystorem R10. Tolerancja zasilania elementu 74LS373 to wąski przedział od 4,75 V do 5,25 V. Jest to jednak napięcie typowe, a napięcia stanów logicznych zera i jedynki pozwalają na bezpośrednie łączenie z układami cyfrowymi wykonanymi w technologii CMOS, NMOS i TTL.

Przyciski S1 do S8 przewidziano dla każdego z graczy, tym samym może być ich do ośmiu. Wskazano jest, aby nie montować ich na płycie PCB, lecz przeciągnąć na przewodach tak, aby każdy z nich był w pobliżu odpowiedniego gracza. Czas reakcji powinien być szybki, a więc przycisk powinien być pod ręką. Natomiast odpowiednie LED-y (LED1 do LED8) mogą być ulokowane dalej. Na przykład mogą być na płycie PCB, lub lepiej na frontowej części obudowy „urządzenia quizowego”. Zamiast numerów nad LED-ami można umieścić nazwiska identyfikujące poszczególnych graczy.

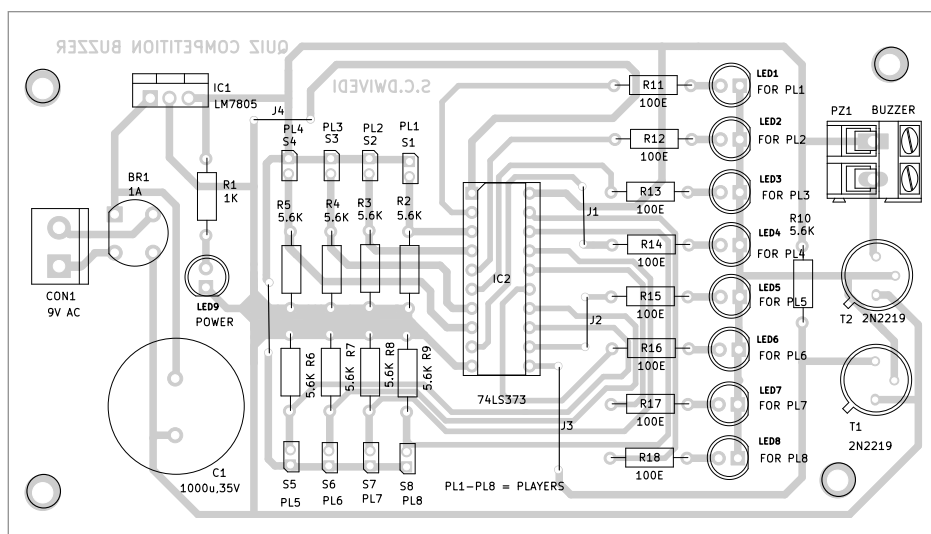


Rysunek 2. Schemat ideowy układu

Urządzenie quizowe jest gotowe do rozpoczęcia quizu natychmiast po włączeniu zasilania. Rozpoczęciu gry zawsze odpowiada stan, gdy wszystkie przyciski S1 do S8 są otwarte. Wtedy wszystkie wyjścia rejestru 373 będą w stanie niskim. Tym samym, wszystkie LED-y (od 1 do 8) pozostają wygaszone. Chwilowe naciśnięcie przycisku przez któregoś z graczy skutkuje stanem wysokim na odpowiednim wyjściu i zaświeceniem przypisanej mu diody LED. Na przykład, jeśli gracz nr 3 naciśnie swój przycisk, zaświeci dioda nr 3 identyfikując, że ten zawodnik jest pierwszy. Naciśnięcie przycisków przez kolejnych graczy jest już ignorowane, zatrask jest zablokowany. Ponieważ układ zapamiętuje stan wszystkich wejść w momencie naciśnięcia dowolnego z ośmiu przycisków jako pierwszy, system taki eliminuje jakiegokolwiek zamieszanie i niejasność kto był pierwszy. Stanowi zatem przejrzystą platformę do rozgrywek dla ośmiu (do ośmiu) uczestników zawodów lub drużyn. Oprócz diod LED wskazujących zwycięzcę, system wyposażono w buzzer z generatorem który uatrakcyjni rozgrywkę. Buzzer piezoelektryczny uruchamiany jest równocześnie z T1, który włączany jest równocześnie z T1. T1 jest elementem sprzężenia zwrotnego wymuszając stan niski na pinie 11 układu scalonego IC2. To wejście Latch Enable powodujące zatrzaśnięcie się wszystkich przerzutników rejestru. Taka konstrukcja skutkuje tym, iż kolejne zmiany stanu dowolnych przycisków S1 do S8 są już ignorowane wskazując jednoznacznie tylko zwycięzcę. System ten nie jest zatem przewidziany do gier/zawodów w których interesowałoby nas również drugie, trzecie i ew. dalsze miejsca zawodników. Zawsze zaświeci tylko jedna z ośmiu diod LED i będzie temu towarzyszył dźwięk buzzera, tak



Rysunek 3 Układ druku jednostronnej płytki PCB



Rysunek 4 Schemat montażowy ułożenia elementów na PCB

Wykaz elementów:

Półprzewodniki:

IC1: stabilizator napięcia 5 V LM7805

IC2: bufor 74LS373

T1, T2: tranzystor NPN 2N2219

BR1: mostek prostowniczy 1 A

LED1...LED9: diody LED 5 mm

Rezystory: (wszystkie 0,25 W, ±5%)

R1: 1 kΩ

R2...R10: 5,6 kΩ

R11...R18: 100 Ω

Kondensatory:

C1: 1000 µF/35 V elektrolityczny

Inne:

CON1: złącze 2-pinowe

S1...S8: przyciski niestabilne

X1: transformator sieciowy 230 V AC, uzwojenie wtórne 9 V/500 mA

długo aż obwód wyłączymy i ponownie włączymy jego zasilanie.

Konstrukcja urządzenia

Na **rysunku 3** pokazano układ ścieżek jednostronnej płytki PCB. W drukowanej wersji EdW rozmiar powinien odpowiadać rzeczywistej skali wielkości elementów. Schemat montażowy ułożenia elementów na PCB pokazano na **rysunku 4**.

Po zmontowaniu układu, należy przygotować dla niego odpowiednią obudowę. Wszystkie diody LED (LED1 do LED9) należy umieścić na froncie obudowy, natomiast piezoelektryczny buzzer można ułożyć na tylnym panelu. Dioda LED9 wskazuje obecność zasilania i należy ją wyróżnić względem ośmiu diod przypisanych poszczególnym graczom/zawodnikom. Na płycie PCB wyróżniono miejsca w których należy

podpiąć przyciski S1 do S8. Montowanie ich na płycie lub nawet obudowie urządzenia nie wydaje się wygodnym rozwiązaniem. Montaż przycisków może być zadaniem najbardziej wymagającym, tak aby quiz przebiegał sprawnie i „urządzenie quizowe” było wygodne w użyciu. Należy więc przygotować odpowiedniej długości przewody dla połączenia przycisków z płytką. Po poprawnym montażu i włączeniu zasilania, układ jest gotowy do rozpoczęcia rozgrywki. ■

S.C. Dwivedi

Film instruktażowy, można obejrzeć pod adresem:
<https://youtu.be/de6f7wEFPKU>

Ten artykuł był wcześniej opublikowany na łamach „EFY”, październik 2023 (efymag.com)

Od Redakcji EdW: Opisany tu układ jest bardzo prosty i z uruchomieniem nie powinno być problemów. Należy jednak zauważyć, iż nie jest dobrą praktyką łączenie równoległe obwodów baz tranzystorów bipolarnych. T1 i T2 są tego samego typu i powinny oba rozpoznawać ten sam stan logiczny. Tranzystor bipolarny sterowany jest prądowo i dodanie choć niewielkich rezystorów w bazach powinno być obowiązującą „szkolną praktyką”. Na dobrą sprawę układ można by uprościć rezygnując z ośmiu rezystorów na wyjściach O0 do O7, zostawiając tylko dwa, właśnie w bazach T1 i T2. W tej sytuacji oporniki mogłyby pozostać tej samej wartości. W tej aplikacji zatrasku 373, i tak tylko jedno wyjście w dowolnym czasie powinno pozostać w stanie wysokim. Zatem jasność świecenia diody LED w obu rozwiązaniach (osiem oporów R11...R18 lub dwa) powinna być taka sama.

Za nietypowe należy też uznać obwody od strony wejść rejestru. Przyciski wmuszają stan wysoki, stan zera logicznego zapewniają rezystory od strony masy. Rozwiązanie takie jest w pełni bezpieczne w technologii CMOS. W przypadku klasycznych TTL-i nie jest to rozwiązanie

dobrze, gdyż R2 do R9 są dość dużej wartości 5,6 kΩ. Odwrócenie aktywności sygnałów (z wysokiego na niski) nie stanowiłoby żadnego problemu. Modyfikacja od strony wyjścia wymagałaby jedynie odwrócenia kierunku diod LED i zamiany tranzystorów T1 i T2 na typ PNP. Natomiast, aby dostosować się do logiki wejścia Latch Enable potrzebna byłaby jeszcze jedna negacja. Najprościej, jeszcze jeden tranzystor NPN.

Zastanawiające jest także to, iż w układzie nie przewidziano żadnego przycisku resetu. W układzie z rysunku 2 wyzerowanie nastąpi dopiero po wyłączeniu zasilania. Dla zabawy (quizu) obecność przycisku resetu wydaje się niezbędna lub co najmniej bardzo pożądana. Najprościej byłoby go wpiąć na linię LE lub OE. W pierwszym przypadku w obwodzie kolektora T1 należałoby dodać jakiś rezystor. W drugim przypadku rezystor ściągający w dół też powinien być obecny, i w obu rozwiązaniach przycisk resetujący należałoby podpiąć do linii zasilania +5 V.

Gutek. Eko-Czarodziej w krainie (zużytej) elektroniki

Kasztany zaczynają spadać z drzew, pokrywając chodniki i parkowe alejki. A to może oznaczać już tylko jedno. Mamy jesień! Sezonowy fenomen spadającego kasztana to nie tylko znak, że lato ustępuje miejsca chłodniejszym dniom, ale także doskonała okazja do odkrywania kreatywnych zabaw i ciekawych zajęć pozaszkolnych. Wśród wielu jesiennych aktywności, budowanie ludzików z kasztanów to wyjątkowy sposób na rozwijanie wyobraźni i zdolności manualnych. Jednak Gutek, uczeń siódmej klasy Szkoły Podstawowej nr 257 im. prof. Mariana Falskiego w Warszawie, a zarazem bohater tej relacji, poszedł o krok dalej. Postanowił budować ludziki z... zużytych komponentów elektronicznych.



Komponenty elektroniczne w zasadzie służą do czego innego, ale kiedy się pomyśli, że wiele z nich i tak

z czasem trafia na śmietnik bo, na przykład, w partii kondensatorów wysechł już ze starości elektrolit, i nie nadają się do praktycznego użycia, lub też kiedy pomyśli się o starych zapasach wielkogabarytowych (w odniesieniu do współczesnej miniaturyzacji) rezystorów, albo o serwisach, wymieniających i kierujących do utylizacji każdego dnia dziesiątki, jeśli nie setki komponentów to pomysł jest naprawdę genialny! Z wielu takich komponentów dałoby się wiele sympatycznych figurek dziecięcymi rękoma „ulepić”, ku wielkiej frajdzie tych najbardziej zainteresowanych. Tworzenie tego typu ludzików to także praktyczny rozwój umiejętności manualnych, zarówno technicznych jak i plastycznych. W dodatku jakby nie patrzeć jest

to również recykling, tak więc w grę wchodzi również benefit ekologiczny.

W tym miejscu oddaje głos Gutekowi, wierząc, że napisze od siebie kilka zdań na temat początków tego hobby i powiązanych z nim szkolnych wydarzeń.

Moja przygoda, jak większość takich historii zaczęła się przypadkiem. Pewnego wieczoru, uciłem sobie pogawędkę z moim tatą, pytając go co robił jak był w moim wieku. Tata opowiedział, że w wieku 13 lat interesował się bardzo elektroniką, kupował w kiosku Ruchu czasopisma branżowe i wzorując się na nich samodzielnie tworzył, lutował i uruchamiał różne układy elektroniczne. Niezmiernie mnie to zaintrygowało i zaciekawiło, nigdy wcześniej nie trzymałem w ręku lutownicy, a cyna czy kalafonia... „co to takiego?” – pytałem taty.

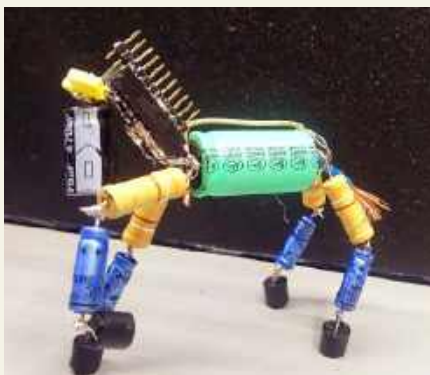
Jak mówią praktyka czyni mistrza. Tata przyniósł skrzynkę narzędziową z piwnicy i na żywo objaśniał, pokazywał i tłumaczył „co się z czym je”. Wtedy przyszedł czas na moją samodzielną próbę lutowania. Hmm... pomyślałem co tu zrobić, patrząc na kształt oporników, tranzystorów, rezystorów do głowy przyszedł mi pomysł na zrobienie... samego siebie. Ponieważ drugą moją pasją jest gra na perkusji zdecydowałem, że to będzie to! I tak powstał „Gutek” grający na perkusji. Sprawilo mi to tak dużo radości i satysfakcji, że postanowiłem kontynuować to tworzenie. Niebawem w mojej szkole miał się odbyć Piknik Rodziny, na którym za namową mojego nauczyciela od techniki pana Tadeusza Lichuty miałem pokazywać dzieciom jak lutować i przy okazji dobrze się bawić. Było to dla mnie cenne doświadczenie, dzieci były bardzo zainteresowane i jednocześnie ogromnie zdziwione, że z tak niewielkich części można zbudować coś tak kreatywnego. Jestem bardzo wdzięczny, że miałem okazję doświadczyć jak to jest być nauczycielem. Dzięki temu wszystkim, dzięki mojemu Tacie, panu Tadeuszowi, Nauczycielom i Dyrekcji mojej szkoły zaszczepiłem w sobie pasję do tworzenia elektronicznych ludzików, a jedynym ograniczeniem jest tylko moja wyobraźnia no i części... Oto cała moja historia.

Zachęcam każdego aby spróbował stworzyć samego siebie w elektronicznej wersji, to naprawdę świetna zabawa.

Powodzenia!



Gustaw Szwed (ale wszyscy mówią na mnie Gutek), lat 13



Gutek

Uniwersalny sposób zabezpieczenia bagażu przed kradzieżą

Układ będący tematem bieżącego odcinka serii „Zrób to sam” jest prostym systemem alarmowym, który może mieć szerokie zastosowanie w wielu codziennych sytuacjach. Można w ten sposób zabezpieczyć np. rower, bagaż, monitorować wejście do pomieszczenia, otwarcie drzwi, itp.

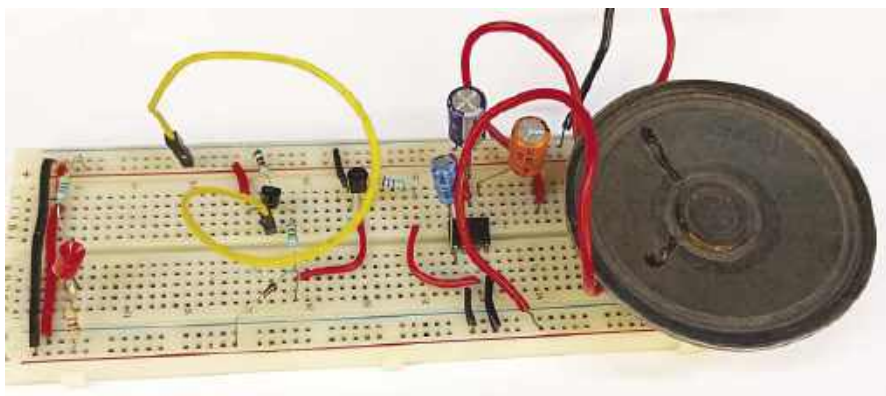
Idea zabezpieczenia wykorzystuje cienki drut którym należy owinać zabezpieczany przedmiot lub rozciągnąć go np. wzdłuż okna, drzwi lub zabezpieczanej strefy pomieszczenia. Drut ten pełni rolę czujnika, którego zerwanie wywoła alarm realizowany przez prosty generator sygnału dźwiękowego. Skuteczność tak prostego zabezpieczenia zależy przede wszystkim od tego, jakiego drutu użyjemy i jak wykonamy nim pętlę zabezpieczającą. Należy wykonać to tak, aby próba kradzieży lub wejścia na chroniony obszar nieuchronnie skutkowałą zerwaniem pętli zabezpieczającej. Można uczynić ją niewidoczną bądź widoczną jeśli uznamy, iż odstraszy to potencjalnego włamywacza. W większości sytuacji nie jest trudno przeprowadzić przewód tak, aby próba kradzieży (np. roweru) musiała skutkować zerwaniem pętli zabezpieczającej.

Opis układu i jego działanie

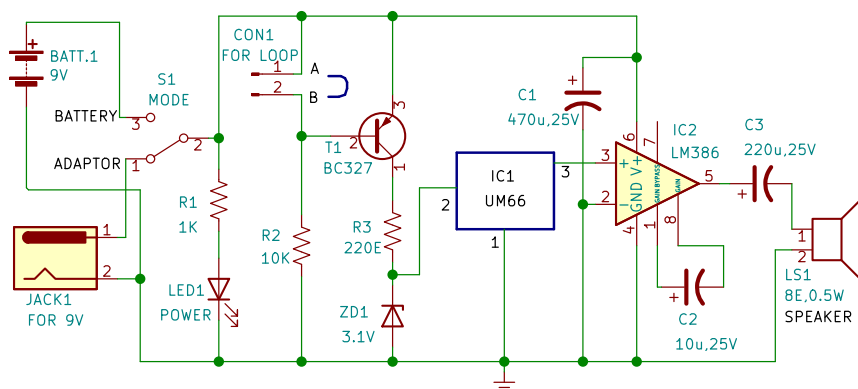
Schemat ideowy pokazano na **rysunku 2**, zaś na **rysunku 1** przedstawiono zdjęcie prototypu wykonanego przez autora na płycie uniwersalnej.

Budowa układu wykorzystuje kilka prostych i popularnych elementów. Wśród nich najistotniejszymi są: tranzystor BC327 (T1), generator melodii UM66 (IC1), prosty wzmacniacz fonii LM386 (IC2), głośnik o impedancji 8 Ω i mocy 0,5 W (LS1) oraz niewielką ilość dyskretnych elementów pasywnych. Do zasilania służy dziewięć-woltowa baterijka lub zasilacz dostarczający napięcia stałego o wartości 9 V. Wśród wykorzystanych elementów należy koniecznie wymienić wspomniany wyżej cienki, stosunkowo łatwy do zerwania drut, który jak się okazuje jest najistotniejszym elementem tego prostego systemu alarmowego. W celu podłączenia pętli utworzonej tym drutem, w układzie zamieszczono złącze opisane na schemacie ideowym desygnatorem CON1.

W prototypie wykonanym przez autora, sygnał alarmu jest mocno nietypowy. Sygnał dźwiękowy wytwarzany jest przez generator melodii UM66. Element ten wykorzystywany jest głównie przy produkcji zabawek. Dźwięk



Rysunek 1. Prototyp wykonany przez autora



Rysunek 2. Schemat ideowy układu

nie jest nadzwyczajnej jakości, a ponadto moc sygnału generowanego przez ten układ jest bardzo słaba. Choć potrafi onysterować głośniczek piezoelektryczny o wysokiej impedancji, w niniejszym zastosowaniu konieczny jest jakiś wzmacniacz sygnału audio. W tej roli wykorzystano układ LM386 będący wzmacniaczem małej mocy. Dla wyjaśnienia działania tej sekcji „systemu alarmowego”, należy zapoznać się z układem wyprowadzeń układu scalonego LM386. LM386 to aktywny element wzmacniacza wraz z kilkoma elementami sprzężenia zwrotnego. Ogranicza to możliwość sterowania głośnością, aczkolwiek taką możliwość przewidziano, i służą do tego wyprowadzenia 1 i 8. Wyprowadzenia 2 i 3 pełnią taką samą funkcję jak w przypadku typowego wzmacniacza operacyjnego. To wejścia odwracające i nieodwracające WO. Jako

wejście sygnału autor wykorzystał nóżkę 3, łącząc wejście odwracające z masą układu. Wyprowadzenia 4 i 6 to zasilanie. Nóżka 4 jest podłączona do masy, a do szóstej doprowadzono napięcie 9 V z baterii lub zasilacza wtyczkowego. Wyjście stanowi wyprowadzenie nr 5 układu scalonego. Z uwagi na obecność składowej stałej napięcia, jest ono sprzężone z głośnikiem poprzez pojemność C3. Wyprowadzenie 7 może służyć jako by-pass zasilania ale w projekcie autora nie jest ono wykorzystane.

Praca układu zestawionego według rysunku 2 jest bardzo prosta. Aktywacja generatora melodii UM66 następuje po podaniu zasilania na jego nóżkę nr 2. Sygnał audio pojawia się na nóżce nr 3, i należy go wzmocnić. UM66 ma wąski i niski zakres napięcia zasilania. Wzmacniacz mocy wymaga napięcia

Wykaz elementów:

Półprzewodniki:

IC1: generator melodii UM66
IC2: wzmacniacz małej mocy LM386
T1: tranzystor PNP BC327
ZD1: dioda Zenera 3,1 V
LED1: dioda LED 5 mm

Rezystory:

(wszystkie 0,25 W, ±5%)
R1: 1 kΩ
R2: 10 kΩ
R3: 220 Ω

Kondensatory:

C1: 470 μF/25 V elektrolityczny
C2: 10 μF/25 V elektrolityczny
C3: 220 μF/25 V elektrolityczny

Inne:

JACK1: złącze zasilania montowane na PCB
S1: przełącznik SPDT
CON1: złącze 2-pinowe
BATT1: bateria 9 V
LS1: głośnik 8 Ω/0,5 W
zasilacz wtyczkowy 9 V
elastyczny cienki przewód dla wykonania pętli zabezpieczającej (długość zależna od aktualnych wymagań aplikacji układu)

wyższego, które jest. tu na poziomie 9 V. Aby to pogodzić, autor zastosował diodę Zenera ZD1, która ograniczy napięcie na pinie 2 IC1 do poziomu 3,1 V.

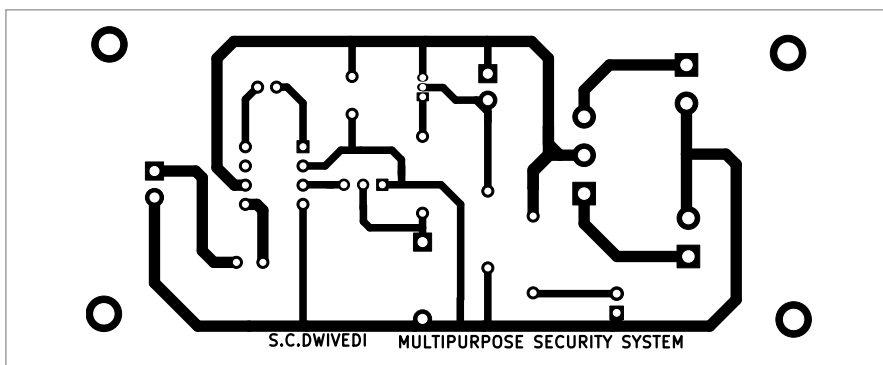
Wielkość wzmocnienia wzmacniacza LM386 zależy od aplikacji w obrębie jego nóżek 1 i 8. Jeśli pozostaną one wolne „w powietrzu” wzmocnienie wyniesie 20 dB. Jeśli podłączymy tu kondensator 10 μF wzmocnienie wzrośnie do 200 dB.

Od Redakcji EdW: O ile dolna granica jest mniej więcej zgodna z prawdą, to górna wielkość wzmocnienia jest absolutnie nierealna. 200 dB to ogromne wzmocnienie niespotykane w żadnym wzmacniaczu. Zatem należy się dopatrywać tu błędu. Wzmocnienie napięciowe mieści się w granicach 20 do 200 razy. Co w decybelach będzie od 26 dB do 46 dB.

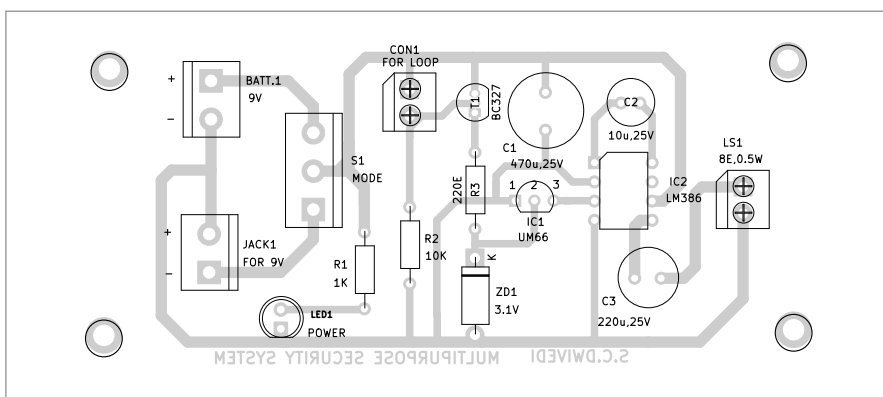
Sygnal dźwiękowy wytwarzany generatorem melodii połączono z wejściem 3 IC2, a wzmocniony sygnał jest dostępny na pinie 5. Kondensator C3 sprzęga głośnik ze wzmacniaczem i zastosowano go w celu usunięcia składowej stałej i przeniesienia do głośnika wyłącznie sygnału audio.

Po zmontowaniu układu, czynnością którą należy starannie wykonać jest poprowadzenie pętli zabezpieczającej wspomnianym wyżej cienkim drutem. Po owinięciu zabezpieczanego przedmiotu, oba końce należy podłączyć do złącza CON1. Na płytce PCB przewidziano zasilanie z baterii lub zasilacza napięcia stałego 9 V do opcjonalnego wykorzystania w zależności od aplikacji. W każdej sytuacji należy odpowiednio ustawić położenie

Film pokazujący działanie systemu, który można obejrzeć pod adresem:
<https://youtu.be/1UlhF1lOZx8>



Rysunek 3. Projekt druku PCB



Rysunek 4. Schemat montażowy elementów na PCB

przełącznika S1. Obecność poprawnego zasilania powinna być sygnalizowana świeceniem diody LED1.

W stanie czuwania tranzystor T1 jest zatkany. Przerwanie pętli zabezpieczającej skutkuje załączeniem się tego tranzystora, co aktywuje generator melodii przez podanie napięcia 3,1 V na jego nóżkę nr 2. W tym prostym systemie elastyczny cienki drut pełni funkcję czujnika który aktywuje generator melodii w układzie scalonym IC1. Dioda LED1 sygnalizuje obecność zasilania systemem. Zwarcie na złączu CON1 blokuje tranzystor T1. Przerwa w tym obwodzie skutkuje włączeniem tranzystora dzięki obecności rezystora R2. Dalszą reakcją jest aktywacja generatora melodii IC1 oraz wzmocnienie sygnału za pośrednictwem IC2.

Konstrukcja i testowanie układu

Rysunek 3 pokazuje projekt jednostronnej płytki drukowanej PCB, która w drukowanej wersji naszego pisma powinna odpowiadać skali 1:1. Ułożenie elementów na PCB pokazuje **rysunek 4**.

Po zmontowaniu elementów na płytce PCB, należy przygotować odpowiednią obudowę. Diodę LED1 i złącze CON1 należy umieścić z przodu obudowy, zaś złącze zasilania JACK1 i baterię od tyłu. Oba końce

pętli zabezpieczającej należy umieścić w punktach A i B (zgodnie z oznaczeniami na rysunku 2). Długość pętli jest zależna od wielkości zabezpieczanego przedmiotu lub obszaru. Istotnym jest jedynie, aby uległ on zerwaniu gdy ktoś niepowołany będzie manipulował w celach kradzieży przedmiotu i/lub wniknięcia w chroniony obszar. ■

Ten artykuł był wcześniej opublikowany na łamach „EFY”, kwiecień 2024 (efymag.com)

Suresh Chandra Dwivedi

Od Redakcji EdW: Autor wykorzystał dwa nietypowe układy scalone, przynajmniej jak na zastosowanie w systemie alarmowym. I oba zasługują na co najmniej kilka słów komentarza. Argumentem za takim wyborem jest niski pobór mocy w przypadku zasilania bateryjnego. Należy zminimalizować moc pobieraną w stanie czuwania. Nie należy bowiem oczekiwać długotrwałego alarmowania, natomiast wymogi czuwania aczkolwiek zależą od konkretnej aplikacji, mogą być wielogodzinne, wielodniowe lub jeszcze dłuższe. Prąd spoczynkowy generatora melodii UM66 jest na poziomie 1 μA zaś moc tracona w LM386 (bez podania sygnału na jego wejście) jest na poziomie 20...30 mW. Jeśli alarm ten ma długo czuć nie jest pomijalna też moc w diodzie sygnalizującej obecność zasilania

(i rezystorze R1), a także prąd który płynie przez rezystor R2. Zatem, wtedy wskazana jest opcja z zasilaczem sieciowym.

UM66 to wbrew pozorom skomplikowany układ scalony. Melodia zaszyta jest w wewnętrznej pamięci ROM, i jedynie od jej pojemności zależy długość sekwencji generowanej melodii. Dostępne są UM-y 66 z wieloma popularnymi melodyjkami, a ograniczenia aplikacji tego układu scalonego wynikają głównie ze skrajnego ograniczenia ilości jego wyprowadzeń. To jedynie 3 nóżki (!) i obudowa taka jak sąsiedniego tranzystora T1 (BC327). Należy przede wszystkim zwrócić uwagę od której strony nóżki liczymy. Bowiem katalog UM66 podaje inną kolejność niż autor opisał na schemacie z rysunku 2. Funkcja nóżki 2 jest taka sama, natomiast masa to nóżka 3, a wyjściem jest wyprowadzenie nr 1. Pomyłka w tym zakresie nie jest jedynym potencjalnym źródłem rozczarowania, że tak prosty układ nie zadziała. Głośnik z wyjściem wzmacniacza LM386 jest sprzężony pojemnościowo, gdyż na nóżce 5 LM-a występuje składowa stała. Ale podobny konflikt składowych stałych występuje w miejscu sprzężenia generatora melodii ze wzmacniaczem akustycznym. UM66 może (i musi) być zasilany niskim napięciem z przedziału 1,3 V do 3,3 V. Dioda ZD1 zapewnia, że w tym przedziale się mieścimy. Ale na nóżce 3 (jak na schemacie, lub nóżki pierwszej wg katalogu) występuje oprócz zmiennego sygnału (melodii) jakaś składowa stała. Jest ona bezpośrednio podana na wejście nieodwracające wzmacniacza LM386. Drugie wejście (odwracające) jest natomiast podłączone do masy. LM386 jest rodzajem wzmacniacza operacyjnego, jedynie z wbudowanymi rezystorami ujemnego sprzężenia zwrotnego. Klasyczny WO w takiej konfiguracji uległby nasyceniu, i nie należałoby oczekiwać jakiegokolwiek wzmocnienia. Mniejsze wątpliwości budzi praca z wejściami na poziomie

dolnego zasilania 0 V (wejście nieodwracające połączono z masą). Wiele wzmacniaczy potrafi pracować w takich warunkach, i LM386 także. Zapewnia to konstrukcja w której na wejściu różnicowym pracują tranzystory PNP. W LM386 jest nawet „coś w rodzaju Darlington-a” dwu takich tranzystorów, co skutkuje bardzo małym prądem polaryzacji wejść oraz faktem, iż prąd ten wypływa z obu wejść wzmacniacza. Nie mniej, stwarza to konieczność zwrócenia uwagi na impedancję źródła widzianą z wejść wzmacniacza. Na wejściu odwracającym jest to zero, bo wejście to zwarte jest z masą. Na drugim wejściu (nieodwracającym) jest to impedancja wyjścia generatora melodii UM66, która jest prawdopodobnie wysoka. Sytuacja nie jest jednak zła, jak można by w tym miejscu domniemywać. Wejścia LM386 są bowiem zabezpieczone rezystorami o wartości 50 kΩ. Ogranicza to impedancję wejściową wzmacniacza, co może mieć znaczenie jedynie w niewielu sytuacjach nietypowych aplikacji. Tutaj niesymetria impedancji widzianych z obu wejść wzmacniacza może jedynie skutkować niewielkim przesunięciem składowej stałej na jego wyjściu. Powinno to być połowa napięcia zasilania układu scalonego. Obecność kondensatora między wyjściem IC1 i wejściem IC2 powinna uwolnić od sugerowanych wyżej wątpliwości. Powieździeliśmy, że klasyczny WO w takiej konfiguracji by nie zadziałał. Mimo wszystko, LM386 powinien działać i należy wierzyć, że autor to sprawdził. LM386 to prosty układ scalony (o wiele prostszy od trzy-nóżkowego UM66), a przede wszystkim nietypowy. Nietypowa jest też korekta wartości jego wzmocnienia. Wewnętrzne rezystory sprzężenia zwrotnego ustalają nie tylko wartość wzmocnienia, ale też balans składowej stałej tak, aby napięcie na wyjściu (noga 5) utrzymywało się w okolicy połowy zasilania. Nóżki 1 i 8 stanowią dostęp do wewnętrznych rezystorów feedbacku. Zwarcie

nóżki pierwszej z ósmą, zwiera jeden z rezystorów feedbacku. Można go zwierać „całkowicie” lub dołączać zewnętrznie dobrany rezystor. Jak autor podaje w opisie swojego projektu, dostępny zakres regulacji wzmocnienia jest dziesięcio-krotny (od 20 do 200, 26 dB do 46 dB). Jednak, aby nie ucierpieła składowa stała (wyjścia) owo „zwieranie” musi odbywać się przez kondensator (tu C2=10 μF). Obecność wyprowadzeń 1, 8 jak i 5 i 2 stwarza praktycznie szerszą możliwość manipulacji wzmocnieniem. Należy jednak pamiętać, że ten wzmacniacz nie jest skompensowany do pracy ze wzmocnieniem mniejszym od około 10. Zatem, nie należy być zbyt zaskoczonym jeśli układ się wzбудzi.

Pokazany tu układ jest bardzo prosty. Powyższe uwagi wskazują jednak, gdzie można się spodziewać niespodzianek.

Generowany dźwięk alarmu będzie melodyjką o niezbyt dużej jakości dźwięku. W przypadku alarmu nie ma to pewnie istotnego znaczenia. Gorzej, że mimo zastosowania wzmacniacza, głos wciąż będzie raczej cichy.

Tak samo, idea zerwania cienkiego drutu do wykrywania próby kradzieży nie wydaje się zbyt wyszukana. Mimo to, w wielu zastosowaniach może być rozsądna.

Na koniec jeszcze dwa zdania od tłumacza tego tekstu DIY.

Około 42...43 lata temu (rok 1981 lub 1982) podobny alarm wykonałem na potrzeby wyjazdu urlopowego. Na dachu malucha (Fiata 126p) miałem bagażnik na którym pozostawiłem sporo „gratów”. Owinąłem je cienkim drutem, a cały alarm był wykonany na jednym tranzystorze i bez żadnego układu scalonego. Wystarczył prosty przekładnik, a elementem wykonawczym był klakson samochodu. Pewnego dnia a raczej nocy, mój alarm obudził całe pole namiotowe. Zaproponowany tu układ DIY będzie przyjemniejszy „w odsłuchu”, ale pewnie mniej skuteczny.

REKLAMA

numery archiwalne
• prenumerata • książki
www.UlubionyKiosk.pl



Jak upłynął Ci pierwszy miesiąc roku szkolnego? Mam nadzieję, że świetnie, i nie musisz nadrabiać żadnych zaległości, bo chciałbym Cię teraz na chwilę odciągnąć od szkolnych zeszytów i książek. Co powiesz na zbudowanie w ramach relaksu... statku kosmicznego? Niezależnie od tego czy wierzysz w UFO (Unidentified Flying Object) czyli niezidentyfikowane obiekty latające, chciałbym, żebyś uwierzył dzisiaj w coś naprawdę niezwykłego. Otóż wyobraź sobie, że zagłębiając się w coraz to głębsze meandry elektroniki i sięgając po coraz to bardziej zaawansowane technologie, może się okazać, że materiał wcale nie staje się trudniejszy, a wręcz przeciwnie, zaczyna być coraz prostszy! Jak to możliwe? Na własnej skórze przekonasz się o tym już za chwilę!

Jak być może jeszcze pamiętasz, w lipcu poskładaliśmy **zmierniczkową lampkę LED**. Nie zawierała ona ani jednego układu scalonego, a co najwyżej dwa tranzystory okraszane kilkorgiem elementów dyskretnych. Układ był bardzo prosty, a mimo to, poznanie zasady jego działania wymagało lekkiej gimnastyki umysłu. Wszak „zatkanie tranzystora innym tranzystorem”, które jest w zasadzie „załączeniem wyłączenia” może w pierwszej chwili wydawać się nieco dziwne. Miesiąc później, budując **konsole audiochaos** poznałeś twór zwany układem scalonym, który zawierał w swojej strukturze bramki logiczne. Budowana zabawka zawierała aż dwa takie układy. „Nowość” jednak nie zbiła Cię z tropu i zakładam, że całkiem nieźle poradziłeś sobie, zarówno ze zrozumieniem z grubsza zasady działania budowanego urządzenia, jak i z poprawnym jego zmontowaniem. A wiesz co stanie się dzisiaj? Dzisiaj dla odmiany użyjesz mikrokontroler, czyli mikroprocesor z zaimplementowanymi w jednej obudowie peryferiami, takimi jak pamięć czy układy wejścia/wyjścia. Można by rzec, że zbudujesz dzisiaj mały komputer. I wiesz co? Będzie

to w zasadzie najprostszy układ, spośród tych, które dotychczas zbudowałeś podczas wcześniejszych spotkań w ramach cyklu EdW Junior!

Zerknij na **rysunek 1**, na którym przedstawiony został schemat budowanego statku kosmicznego.

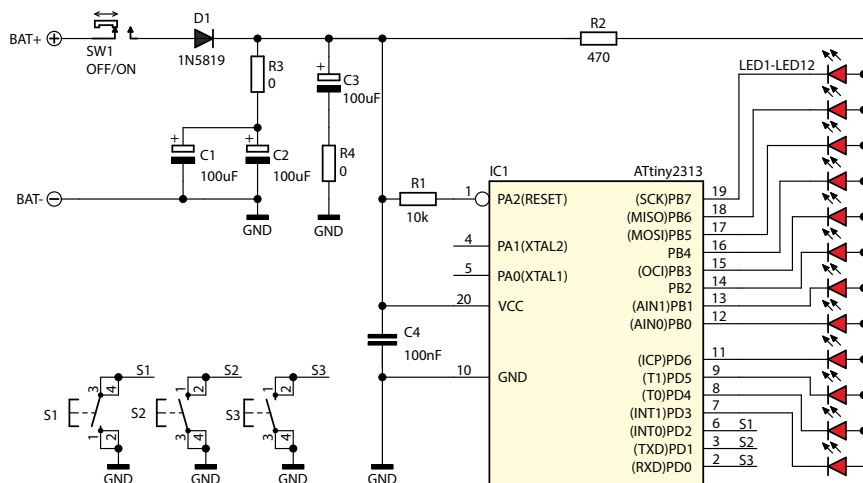
Z powyższego schematu aż „bije” prosta. Spoglądając na schemat nie trudno dostrzec aż dwunastu znajdujących się tam **diod LED** (Light Emitting Diode) czyli diod emitujących światło, a krócej, po prostu diod świecących. Szybki rzut okiem na listę elementów widoczną w instrukcji dołączonej do zestawu **AVTEDU632**, czyli tzw. BOM (Bill Of Materials) pozwoli zorientować się, że elementy o desygnatorach **LED1...LED12** to diody LED o średnicy 3 mm w kolorze czerwonym. Widzimy również, że wszystkie te diody LED swoimi katodami (wyprowadzeniami do których podłącza się ujemny biegun zasilania) podłączone zostały do „skrzynki” podpisanej desygnatorem **IC1** oraz opatrzonej opisem **ATTiny2313**. Cóż to za magiczne pudełko? Jeśli widząc ów żółty prostokąt wyciągasz



Fotografia 1. Zmontowany AVTEDU632

dłoń, rwąc się czym prędzej do odpowiedzi, że jest to układ scalony, odpowiem: bingo! Masz rację. To układ scalony. Ale nie byle jaki! To mikrokontroler! Drobiazg, który czyni budowaną zabawkę układem mikroprocesorowym. W zasadzie czyni z niej mikrokomputer.

Oprócz diod LED (LED1...LED12) i mikrokontrolera (IC1) na schemacie znajdują się trzy przyciski (**S1...S3**) za pomocą których będziesz mógł dostrajać generowany przez mikrokontroler i pokazywany na diodach LED efekt świetlny, cztery rezystory (**R1, R2, R3 i R4**), z których tylko dwa wprowadzają do obwodu jakąś rezystancję stawiając opór przepływającemu prądowi. Jeden z tych dwóch, które „coś robią” to **R1**, który odpowiedzialny jest za poprawny start programu zaszytego w pamięci mikrokontrolera, a drugi to **R2**, który ogranicza sumaryczny prąd płynący przez wszystkie diody LED, z jednej strony zapobiegając przed uszkodzeniem diod LED na skutek przepływającego przez nie zbyt dużego prądu, z drugiej strony, zabezpieczając sam mikrokontroler przed ewentualnym uszkodzeniem na skutek obciążenia zbyt dużym prądem jego wyjść. Na schemacie znajdziemy dodatkowo cztery kondensatory (**C1...C3** oraz **C4**), wszystkie filtrują napięcie wejściowe. Te elektrolityczne (**C1...C3**)



Rysunek 1. Schemat ideowy układu

o pojemności 100 μF (sto mikrofaradów) każdy odfiltrowują ewentualne tętnienia o małych częstotliwościach lub pojedyncze większe skoki napięć, które mogłyby się ewentualnie pojawić w chwilach zmniejszającego się i zwiększającego poboru prądu przez układ, na przykład na skutek pracy mikrokontrolera, załączającego mniejszą bądź większą liczbę diod LED. Kondensator elektrolityczny działa jak magazyn energii, z którego układ może czerpać prąd w momencie, gdy zapotrzebowanie na energię nagle wzrośnie. Bateria, w przeciwieństwie do alternatywnych źródeł zasilania, takich jak zasilacz sieciowy, nie wprowadzi do układu żadnych dodatkowych zakłóceń. Kondensator ceramiczny **C4** o pojemności 100 nF (sto nanofaradów) to standardowy filtr przeciwzakłóceńowy. Niweluje on wysokoczęstotliwościowe szумы i zakłócenia, które mogą powstawać na skutek pracy samego mikrokontrolera (przełączanie wewnętrznych tranzystorów, szумы cyfrowe) lub które mogą trafić do układu z zewnątrz, na przykład na skutek pracujących w pobliżu innych urządzeń generujących zakłócenia elektromagnetyczne (takich jak silniki, przekładniki, elektromagnesy). Jak wspomniałem w poprzednim odcinku kondensator ten powinien zostać dołączony możliwie najbliżej nóżek zasilających układu scalonego (w naszym przypadku – mikrokontrolera). Kolejnym elementem na schemacie jest dioda **D1**, która, jak każda dioda, przewodzi prąd elektryczny w jednym kierunku, od anody w stronę katody. Dołożenie jej w szereg z zasilaniem (patrz **D1** na schemacie) sprawi, że gdy pomylimy się i podłączymy baterię na odwrót (plusem baterii do minusa układu i minusem baterii do plusa układu) prąd elektryczny nie popłynie i mikrokontroler (najdroższy element w zestawie) nie ulegnie zniszczeniu. Przełącznik **SW1** oczywiście zamyka lub przerywa obwód zasilania, podłączając (lub odłączając) tym samym baterię do (lub od) elektroniki znajdującej się na płytce drukowanej.

To w zasadzie byłoby na tyle w temacie opisu działania układu elektronicznego. Biorąc pod uwagę Twoje doświadczenia z montażem elektroniki podczas budowania układów w poprzednich miesiącach, montaż tej płytki będzie wręcz banalny. Zbierzmy zatem wnioski i podsumujmy temat.

Wyrobiliśmy się w niecałych dwóch stronach naszego czasopisma. Kiedy patrzemy na schemat, ten wydaje się fenomenalnie prosty. Zasilanie, przyciski, mikrokontroler i diody. Sam montaż jawi się niczym drobnostka.

Wiesz do czego zmierzam? Jeśli nie, zerknij na ostatnie zdanie wstępu do dzisiejszego spotkania. Przekonałem?

Zastosowanie w układzie elektronicznym mikrokontrolera zawsze minimalizuje komplikację schematu do niezbędnego minimum. Nie trzeba wówczas implementować wielu tranzystorów i układów scalonych i okraszać ich dziesiątkami a czasem setkami elementów dyskretnych, ponieważ większość zadań (pomiarów, porównań, obliczeń, konwersji) zadzieje się wewnątrz samego mikrokontrolera zgodnie z wcześniej napisanym przez programistę i wgranym do nieulotnej pamięci mikrokontrolera programem, czyli przepisem na rozwiązanie jakiegoś zadania (na przykład cyklicznego załączania i wyłączania diod LED wedle jakiegoś pomysłu). W świecie mikrokontrolerów ten program nazywa się **oprogramowaniem układowym**. Często spotkasz się również z nazwami **firmware** lub **wsad**. Dla ścisłości należałoby dodać, że dwie ostatnie nazwy dotyczą programu w postaci wynikowej, a więc już skompilowanego, czyli przełożonego z kodu w miarę czytelnej dla człowieka (pisanego zazwyczaj w języku wysokiego poziomu, na przykład w języku C) na kod wynikowy (w zasadzie ciąg zer i jedynek), który można zapisać wprost do pamięci mikrokontrolera. Niemniej wszystkie te nazwy (program, oprogramowanie układowe, firmware, wsad) powinny Ci się kojarzyć z przepisem na wykonanie jakiegoś zadania przez mikrokontroler, który można trwale „wgrać” do mikrokontrolera. I na razie zapamiętaj tylko tyle. Zapamiętaj też (bo to ważne), że przed wgraniem programu do mikrokontrolera, nowy mikrokontroler, taki prosto ze sklepu, sam z siebie nie robi nic i nic nie potrafi. Włożenie do zbudowanego urządzenia nowo zakupionego mikrokontrolera i podłączenie takiego urządzenia do zasilania może rozczarować, gdy nie masz tej wiedzy. Urządzenie po prostu nie zadziała. Program trzeba albo samodzielnie napisać i skompilować, albo użyć gotowego wsadu (pobrać z Internetu, na przykład ze strony projektu, jeśli autor taki wsad udostępnił i z użyciem odpowiedniego programatora, zaprogramować tym wsadem mikrokontroler). Na szczęście Ty nie musisz zaprzątać sobie w tej chwili tym głowy, ponieważ dołączony do UFOledka mikrokontroler został wcześniej zaprogramowany przez producenta zestawu. Wystarczy, że zamontujesz go w swoim urządzeniu.

Zanim przystąpisz do budowania układu mam dla Ciebie jeszcze kilka zagadek i niespodzianek.

Zwory na płytkach drukowanych (jumper wire)

Nieco wyżej wspomniałem, że dwa rezystory, które mają rezystancję opisaną jako 0 Ω , nie wprowadzają do obwodu żadnej rezystancji i równie dobrze można by je było zastąpić kawałkiem drutu, na przykład pozostałością wyprowadzenia innego komponentu, obciętą po jego przylutowaniu. Mówiąc jeszcze inaczej, rezystory o wartości 0 Ω pełnią funkcję zwory. Jeśli przez cały ten czas, od momentu w którym o tym wspomniałem zastanawiasz się po co w urządzeniu zwory skoro dałoby się taką zworę zrealizować bezpośrednio na płytce drukowanej w postaci miedzianej ścieżki, podpowiem Ci, że całkiem nieźle kombinujesz. I znowu bingo! Masz całkowitą rację! Rysując taką płytkę samemu, niezbędne połączenie mógłbyś dodać w postaci ścieżki. Zwłaszcza gdybyś miał do czynienia z bardzo prostym układem elektronicznym o niewielkiej liczbie połączeń między komponentami albo gdybyś miał do dyspozycji dwie lub więcej warstwy miedzi do wykorzystania. Mogłoby się jednak okazać, że realizacja projektu na płytce dwustronnej będzie droższa (zwłaszcza, jeśli produkcję płytek będziesz chciał zlecić któremuś z rodzimych dostawców) i w takich okolicznościach mimo wszystko wolałbyś sięgnąć po laminat z miedzią po jednej stronie. Wówczas przy układach z nieco bardziej skomplikowanym układem połączeń może się okazać, że zwora skutecznie uratuje całą sytuację. Gdy na płytce zacznie się robić ciasno i ciężko już będzie znaleźć wolną trasę dla wykonania kolejnego połączenia bez kolizji z innymi ścieżkami, wówczas z jednego miejsca na płytce można „przeskoczyć” nad istniejącą ścieżką (bądź wieloma istniejącymi ścieżkami) w inne miejsce. I tutaj właśnie z pomocą przychodzi zwora. Na przykład w postaci rezystora o wartość 0 Ω , który potrafi kosztować tyle co nic a może być nieco zgrabniejszy tudzież efektywniejszy w montażu niż drut czy pozostałość wyprowadzenia po innym komponentcie. W roli zwor świetnie sprawdza się cienki posrebrzany drut o nazwie kynar, niemniej nie zawsze jest on pod ręką. Kynar z reguły sprzedawany jest w rolkach (podobnie jak cyna) a nie zawsze tak duża ilość jest nam potrzebna. I w takich sytuacjach świetnie sprawdzi się dołączony do zestawu rezystor 0 Ω . W zestawie **AVTEDU632** znajdują się dwa takie rezystory – zwory: **R3** i **R4**.

Istnieją też inne powody, dla których konstruktor może zaplanować w układzie obecność zwor. Wymontowując już zamontowaną, lub też świadomie nie montując

wybranych zwor, mamy możliwości rozłączenia poszczególnych sekcji obwodu, co może być cenne podczas diagnostyki i napraw lub też przy rozwoju projektu. Celem rozłączenia wybranych sekcji obwodu nie trzeba wówczas przecinać ścieżek, a zamiast tego wystarczy zdemontować istniejącą zworę (wystarczy wówczas odlutować i wyciągnąć płytki jeden jej koniec). Jeszcze innym powodem montażu zwor może być chęć bezpiecznego i skutecznego przeniesienia dużych prądów w wybranym obszarze PCB, których za pomocą bardzo cienkich miedzianych ścieżek (nawet obficie pocynowanych) nie sposób byłoby przenieść. W takich sytuacjach jednak dużo lepiej sprawdzi się zwora z odpowiednio grubego drutu (0 Ω rezystor z uwagi na ograniczoną moc może się tam nie sprawdzić).

Układ resetu mikrokontrolera

W pobieżnym opisie wspomniałem, że rezystor R1 odpowiada za poprawne uruchomienie się mikrokontrolera. Ale co to właściwie oznacza? Otóż po włączeniu zasilania (po włożeniu do koszyeczka baterii i włączeniu przełącznika SW1) na liniach zasilania mikrokontrolera, a więc pomiędzy pinem 20 (VCC) oraz pinem 10 (GND) niemal natychmiast pojawi się stabilne napięcie o wartości około 4,5 V (3 baterie AA $\times$ 1,5 V) pomniejszone o spadek napięcia na diodzie D1. Ponieważ jest to dioda Schottky'ego (1N5819) spadek napięcia będzie nieco mniejszy, niż w przypadku krzemowych diod prostowniczych i wyniesie on około 0,3 V. Wystąpi też spadek na rezystancji wewnętrznej baterii, ale przy (zmierzonym) prądzie pobieranym przez nasz układ wynoszącym mniej niż 10 mA (dziesięć miliamperów) jest on do pominięcia. Innymi słowy pomiędzy nogą 20 i 10 mikrokontrolera można się spodziewać napięcia ok. 4,2 V. Zanim to (zasugerowane kilka zdań wyżej) „niemal natychmiast” nastąpi na liniach zasilania będzie panowało napięcie nieustalone, na przykład pojawią się skoki napięcia na skutek drgania styków przełącznika SW1 podczas załączania. Mikrokontroler sam w sobie jest dosyć skomplikowanym elementem, wrażliwym na anomalie pojawiające się na liniach zasilania. Jeśli napięcie będzie niestabilne, może źle wykonywać swoje zadania a nawet „zawiesić się” i przerwać wykonywanie jakichkolwiek zadań do czasu ponownego restartu (wyłączenia i ponownego włączenia) zasilania. Istnieją całkiem zaawansowane sensory badające stabilność napięcia i odpowiednio ustawiające wejście RESET mikrokontrolera w stan wysoki lub niski (bliski VCC lub GND)

w zależności od parametrów jakościowych monitorowanego napięcia, adekwatnie załączając lub blokując pracę mikrokontrolera, lub też restartując jego pracę. Istnieją też wewnętrzne rozwiązania i procedury charakterystyczne dla danego producenta i modelu mikrokontrolera gwarantujące skuteczne wyjście ze stanu zawieszania jeśli do takiego już dojdzie (watchdog). Jednak nie każda aplikacja będzie krytyczna. Inaczej podchodzi się do tematu w przypadku konstruowania zabawki, a inaczej w przypadku budowania sprzętu medycznego ratującego ludzkie życie. Tym samym nie wszędzie zaawansowane mechanizmy programowe czy też rozwiązania sprzętowe będą stosowane. Najpopularniejszym układem resetu stosowanym w większości układów z mikrokontrolerami będzie para odpowiednio połączonych ze sobą rezystora i kondensatora, generująca stałą czasową, odliczającą czas od podania napięcia zasilającego, po którym sygnał utrzymujący mikrokontroler w stanie resetu zostanie zwolniony, tym samym umożliwiając wystartowanie mikrokontrolera i rozpoczęcie wykonywania zaprogramowanego w nim kodu gdy już napięcie będzie stabilne. W budowanym UFOledku, rolę układu resetu pełni pojedynczy rezystor R1 i jest to w tej aplikacji w zupełności wystarczające. Zwróć uwagę na kółeczko narysowane na schemacie przy pierwszej nóżce mikrokontrolera. Na naszym spotkaniu w ubiegłym miesiącu, gdy opowiadałem Ci o bramkach logicznych NOT i NAND wspomniałem co to kółeczko oznacza. Pamiętasz może? Wierzę, że tak, ale jeśli jakimś sposobem wypadło Ci to z głowy, przypomnę, że to symbol negacji, który możesz sobie raz na zawsze skojarzyć z buntem młodszego rodzeństwa, czasem dla zasady torpedującego każdy najlepszy Twój pomysł. Kiedy powiesz „tak”, w odpowiedzi natychmiast usłyszysz „nie”. Gdy powiesz „nie” usłyszysz „tak”. I tak bez końca. Takie bywają uroki posiadania lub bycia rodzeństwem, choć zdarzają się też rodzeństwa zaskakująco zgodne i wzajemnie rozumiejące się niemal bez słów (czego z całego serducha Ci zresztą życzę). Wracając jednak do elektroniki, kółeczko oznacza, jak wspomniano, negację. Gdy jest narysowane na wejściu układu scalonego, oznacza to, że sygnałem aktywnym (inaczej niż sugerowałaby to logika dodatnia) jest stan niski (bliski napięciu panującemu na linii GND). Znowu takie trochę „załączenie (mikrokontrolera) przez wyłączenie (stanu reset)”. Z drugiej strony „wyłączenie sygnału reset, poprzez podanie stanu wysokiego” brzmi całkiem rozsądnie. W każdym razie,

dopóki na nogę mikrokontrolera IC1 nie zostanie podany stan wysoki (napięcie z linii zasilania VCC na odpowiednim poziomie) mikrokontroler będzie utrzymywany w stanie resetu. Pojawienie się stabilnego napięcia wyzeruje mikrokontroler (mikrokontroler zostanie uruchomiony/odblokowany i znacznie wykonywać swój „przepis na wykonanie zadania” od początku).

Jeśli przyjrzyj się na schemacie opisowi pierwszego wyprowadzenia zauważysz PA2(RESET). Opis ten sugeruje, że docelowo noga ta może pełnić alternatywnie dwie funkcje: wejście/wyjście portu A oraz RESET. Dla nas na razie ciekawsza i ważniejsza jest funkcja RESET i do niej póki co ograniczę niniejszy opis.

Linie łączeniowe od przycisków porwało UFO?

Niemal wszystkie komponenty na schemacie są ze sobą wzajemnie połączone liniami, a w przypadku przycisków linii tych brak. Co tu się właściwie podziało, i czy przypadkiem nasze UFO nie ma z tym czegoś wspólnego? A może to znów jakieś konstruktorskie działanie czytelne wyłącznie dla wtajemniczonych? Jedno jest pewne. Jeśli to było UFO, zostawiło nam bardzo konkretną wskazówkę.

Otóż każdy z przycisków S1...S3 podłączony jest z jednej strony do wspólnej masy, czyli ujemnego bieguna zasilania (etykieta GND) a po drugiej stronie tych przycisków znajdujemy etykiety tekstowe odpowiadające nazwom desygatora danego przycisku (etykieta S1 dla przycisku S1, etykieta S2 dla przycisku S2 oraz etykieta S3 dla przycisku S3). Przyciski te podłączone są do wyprowadzeń mikrokontrolera skonfigurowanych programowo jako wejścia. Jeśli zastanawiasz się, dlaczego napisałem, że przyciski zostały podłączone do wyprowadzeń mikrokontrolera, skoro nie widzisz linii połączeń między przyciskami i mikrokontrolerem a w złośliwą działalność UFO w tym wypadku postanowiłeś mi nie uwierzyć, spiesz się z wyjaśnieniem. Otóż UFO chyba rzeczywiście miało ciekawsze zajęcia, bo same przyciski, owszem zostały do mikrokontrolera podłączone tyle, że nie za pomocą linii, a z użyciem etykiet. Jeśli spojrzysz na wyprowadzenia mikrokontrolera o numerach 6, 3 i 2, to zobaczysz, że przy tych wyprowadzeniach widnieją etykiety odpowiednio S1, S2 i S3. Dzięki identycznym etykietom program do tworzenia płytek drukowanych „wie”, że pomiędzy tymi konkretnymi dwoma wyprowadzeniami istnieje tzw. net (czyli połączenie) i ma między nimi zostać narysowana ścieżka. Przykładowo odnajdując

etykiętę S1 na jednym z wyprowadzeń przycisku S1 oraz przy wyprowadzeniu numer 6 mikrokontrolera, program wspierający projektowanie płytek drukowanych „wie”, że przycisk S1 ma być podłączony do wyprowadzenia numer 6 mikrokontrolera. Tak samo odczyta to każdy elektronik, który miał okazję trochę pracować z różnymi schematami. Witaj w klubie bo całkiem sporymi krokami, z miesiąca na miesiąc wkraczasz do owego zacnego grona.

Można by się zastanawiać, jaka forma łączenia elementów na schemacie jest bardziej czytelna: linie czy etykiety? W przypadku prostych schematów z kilkoma połączeniami czytelniejsze mogą okazać się linie łączeniowe. Jednak uwierz mi, że przy większych schematach etykiety ratują życie (a jeśli nie życie to co najmniej inżynierskie zdrowie). Bez nich gąszcz połączeń na schemacie byłby tak wielki, że niczego nie dałoby się z nich w sposób wygodny i sensowny odczytać. Dodatkowo, jeśli schemat będzie narysowany na większej liczbie arkuszy to o łączeniu wszystkiego ciągłą linią można zwyczajnie zapomnieć.

Istnieje jeszcze jeden sposób łączenia różnych elementów na schemacie. To magistrale gromadzące w pojedynczych (nieco grubszych) liniach łączeniowych wiele oddzielnych netów (sygnałów). Stosuje się je w przypadku schematów do dużo bardziej zaawansowanych systemów mikroprocesorowych. Niemniej osobiście mam bardzo mieszane uczucia co do wykorzystywania magistral, nawet przy większych projektach. Zamiast tego gdzie tylko się da stosuję pojedyncze etykiety z rozsądnymi brzmiącymi nazwami.

Wracając do schematu naszego statku UFO, nie trudno jest się domyślić, że przyciski służą do „ściągnięcia” poziomów napięć panujących na odpowiednich wyprowadzeniach mikrokontrolera do masy (ujemnego bieguna zasilania). Na przykład napięcie 4,5 V (pomniejszonego o gdzieś wcześniej wspomniane spadki) panujące chwilę wcześniej na nodze numer 6 mikrokontrolera w momencie naciśnięcia przycisku S1 zostanie zwarte do masy. Dzięki temu mikrokontroler „zorientuje się”, że naciśnięty został przycisk S1 i należy wykonać jakieś zadanie (zgodnie z przepisem na rozwiązanie jakiegoś zadania, przewidzianym przez autora programu mikrokontrolera). A skąd na nodze numer 6 mikrokontrolera wzięło się napięcie o wartości bliskiej 4,5 V? Otóż zostało ono podane na tę nogę z wejścia zasilania, a więc z nogi numer 20 mikrokontrolera, za pomocą rezystora pull-up. Rezystora tego nie widać na schemacie ponieważ rezystory

pull-up znajdują się wewnątrz struktury mikrokontrolera. Rezystory pull-up łączące są programowo na etapie konfiguracji wyprowadzeń mikrokontrolera zaraz po podłączeniu zasilania (np. włożenia baterii do koszyka i/lub założenia przełącznika SW1). Wspomniana konfiguracja wewnętrznych zasobów mikrokontrolera również odbywa się zgodnie ze scenariuszem, który dostarczył programista na samym początku przygotowanego programu. W zasadzie nie musisz zaprzętać tym sobie głowy. Ważne natomiast jest to, żebyś wiedział, że część wyprowadzeń każdego mikrokontrolera stanowią tzw. porty wejścia/wyjścia (input/output, w skrócie I/O), które da się programowo skonfigurować zgodnie z zaplanowanym przeznaczeniem (wejście lub wyjście). Do portów I/O skonfigurowanych jako wyjścia podłączone zostały z kolei diody LED.

Gdzie podziały się szeregowo rezystory diod LED?

Nie czyńmy tu ponownie teorii spiskowych na temat UFO, ale rzeczywiście, patrząc na sposób w jaki diody LED zostały połączone z mikrokontrolerem mógłbyś (a nawet powinienieś) krzyknąć: „zaraz, a gdzie są rezystory?”. Na potrzeby poniższych rozważań, założmy, że nie zauważyłeś szeregowego rezystora R2.

Rzeczywiście. Schemat sprawia wrażenie jakoby brakowało tu rezystorów. Po pierwsze napięcie wystawiane przez mikrokontroler na cyfrowym wyjściu może wynosić mniej więcej 0 V dla stanu niskiego (logicznego zera) oraz około 5 V dla stanu wysokiego (logicznej jedynki).

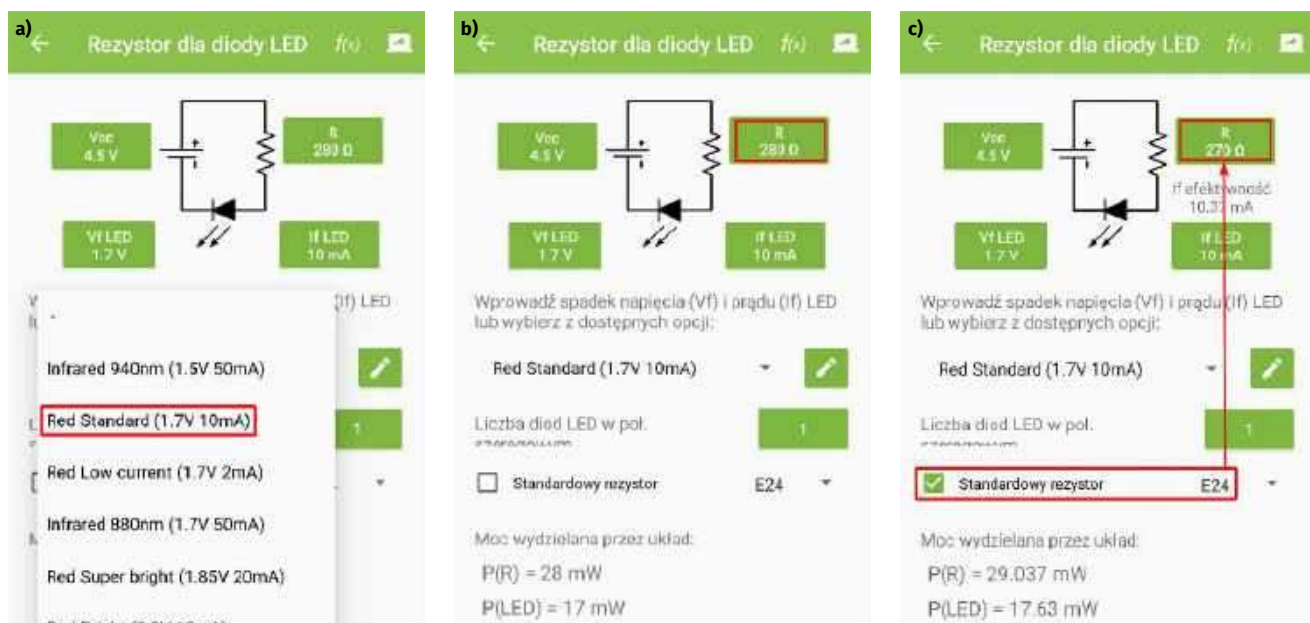
Ponieważ jednak układ zasilany jest z trzech baterii o napięciu 1,5 V, napięcie na wyjściach mikrokontrolera ustawionych w stan wysoki, nie powinno przekroczyć wartości 4,5 V. Powinno być nieco niższe z uwagi na wspomniane wcześniej spadki napięć.

Z pewnych względów, którymi nie chciałbym Ci w tym momencie zaprzętać głowy, mikrokontrolerom łatwiej jest przyjąć większy prąd, niż podać go na jakieś urządzenie (na przykład diodę). Z tego też powodu plusy diod podłączone są do plusa zasilania (przez rezystor R2) a ujemne wyprowadzenia diod podłączone są do mikrokontrolera. Tym samym prąd płynie od plusa zasilania, przez diodę D1 i rezystor R2, wpływa do diody LED i wypływa do mikrokontrolera, a ten odprowadza prąd do GND, dzięki czemu obwód elektryczny zamyka się i dioda LED zaczyna świecić. Innymi słowy dioda LED zaczyna świecić, gdy na wyprowadzeniu mikrokontrolera pojawia się logiczne 0,

czyli potencjał bliski 0 V. Kierunek przepływu prądu nie zmieni natomiast niczego w naszych dalszych rozważaniach.

Napięcie 4,5 V jest zbyt duże, by zasilać nim bezpośrednio diodę LED, ponieważ zalecane przez producenta napięcie przewodzenia dla typowej czerwonej diody LED wynosi około 1,7 V przy maksymalnym dopuszczalnym przez producenta diody prądzie przewodzenia o wartości około 10 mA. Taki prąd zapewni największą jasność świecenia diody bez szkodliwego wpływu na jej wewnętrzną strukturę. Skąd znam te wszystkie wartości? Bawiąc się spory kawałek życia różnej maści komponentami elektronicznymi, mógłbym powiedzieć, że po prostu te wartości pamiętam. Nie skłamałbym w ten sposób za nadto, jednak rozsądniej byłoby dać żądanemu wiedzy Młodszemu Koledze dobry przykład i powiedzieć, że sprawdziłem tę informację w nocy katalogowej, którą odnalazłem wpisując w wyszukiwarce internetowej tzw. partnumber (unikalny dla danego wyrobu kod producenta) użytego elementu. Problem w tym, że na liście komponentów dołączonej do zestawu AVTEDU632 znajduje się wyłącznie informacja o tym, że jest to „(jakaś) czerwona dioda LED o średnicy 3 mm”. W końcu to lista elementów dla hobbysty, a nie przemysłowy BOM produkcyjny dopuszczający jeden konkretny model od wybranego dostawcy, z ewentualnymi kilkoma dopuszczonymi alternatywami. Skąd zatem wzięłem parametry napięcia i prądu typowe dla normalnej pracy czerwonej diody LED? Nie będę Cię dłużej zwodził, i powiem jak było. Otóż nie mam w zwyczaju za nadto „walczyć” z pędzącym światem, więc jak przystało na współczesne trendy, pisząc ten tekst również nieco w biegu, sięgnąłem do kieszeni po smartfon, włączyłem zainstalowaną wieki temu aplikację **Elektrodoc**, stanowiącą świetny darmowy przybornik narzędzi użytecznych dla każdego elektronika. Następnie skorzystałem z narzędzia o nazwie **Rezystor dla diody LED** i rozwinąłem listę wyboru diody, dla której chciałem przeprowadzić obliczenia. Na liście odnalazłem pozycję Red Standard (standardowa czerwona) i odczytałem zalecaną wartość napięcia (1,7 V) i tworzącego jej prądu przewodzenia (10 mA) co można zobaczyć na **rysunku 2a**.

Znając parametry pracy diody (napięcie i prąd przewodzenia) oraz napięcie, które poda na diodę LED mikrokontroler z prawa Ohma można szybko wyliczyć wartość potrzebnego rezystora szeregowego dla pojedynczej diody LED. Skoro wiemy, że na diodzie ma się odłożyć napięcie 1,7 V to wiemy też, że pozostała wartość napięcia ($4,5\text{ V} - 1,7\text{ V} = 2,8\text{ V}$)



Rysunek 2. a) wybór diody LED o właściwym kolorze, b) wyliczona wartość rezystora przy podanym napięciu zasilania 4,5 V, c) odnalezienie najbliższej wartości rezystora dostępnej w szeregu E24

musi odłożyć się na rezystorze szeregowym. Przez gałąź rezystora ma popłynąć prąd przewodzenia diody, czyli 10 mA (dziesięć miliamperów) czyli 0,01 A (jedna setna ampera). Skoro wiemy, jaki prąd ma płynąć przez rezystor szeregowy oraz jakie ma odłożyć się na nim napięcie, wówczas posługując się prawem Ohma otrzymasz wartość potrzebnego rezystora szeregowego:

$$R_{LED} = \frac{4,5V - 1,7V}{10mA} = \frac{2,8V}{0,01A} = 280\Omega$$

Oczywiście nie musiałem tego wszystkiego wyliczać, bo zrobiła to dla mnie w międzyczasie aplikacja Elektrodoc na smartfonie (rysunek 2b). Nie mając pewności, czy da się zakupić w sklepie wyliczoną wartość rezystora, mogą też kliknąć w „standardowy rezystor”, który ma zostać wybrany na przykład z najpopularniejszego szeregu E24. Kalkulator szybko podpowie o jaką wartość rezystora zapytać w sklepie (rysunek 2c). Wyszło na to, że najbliższa wyliczonej wartości 280 Ω dostępna w szeregu E24 to 270 Ω i taką należałoby zastosować w projekcie.

Załóżmy, że właśnie dostrzegłeś na schemacie rezystor R2 o wartości 470 Ω . A to oznacza tylko tyle, że autor przewidział i zastosował rezystor szeregowy, tyle, że nie dla każdej diody z osobna, ale dla wszystkich diod łącznie. W dodatku zastosował wartość dużo większą niż ta z naszych wyliczeń. Cóż, miał takie prawo. Diody mogą świecić nieco słabiej (popłynie przez nie nieco mniejszy prąd i odłoży się na nich nieco niższe napięcie), ale za to mamy pewność, że nic im

się nie stanie, a i układ zasilany z baterii być może podziała nieco dłużej na jednym ich komplecie.

Mankamentem zastosowania wspólnego, pojedynczego rezystora szeregowego dla wszystkich diod LED jest to, że w zależności od tego, ile z tych diod LED zostanie załączonych przez mikrokontroler jednocześnie, tyle razy zwielokrotniony zostanie prąd jaki popłynie przez wspólny rezystor szeregowy a tym samym wpłynie to na jasność świecenia każdej z diod LED. Im więcej diod zostanie załączonych w tym samym czasie tym słabsze będzie światło emitowane przez każdą z nich. Niemniej ile z tych diod będzie uruchamiane jednocześnie i jak finalnie zaprezentuje się UFOledek przekonamy się już za chwilę, bo przechodzimy

właśnie do części najprzyjemniejszej naszego spotkania, czyli do montażu statku naszego zielonego gościa z przestworzy!

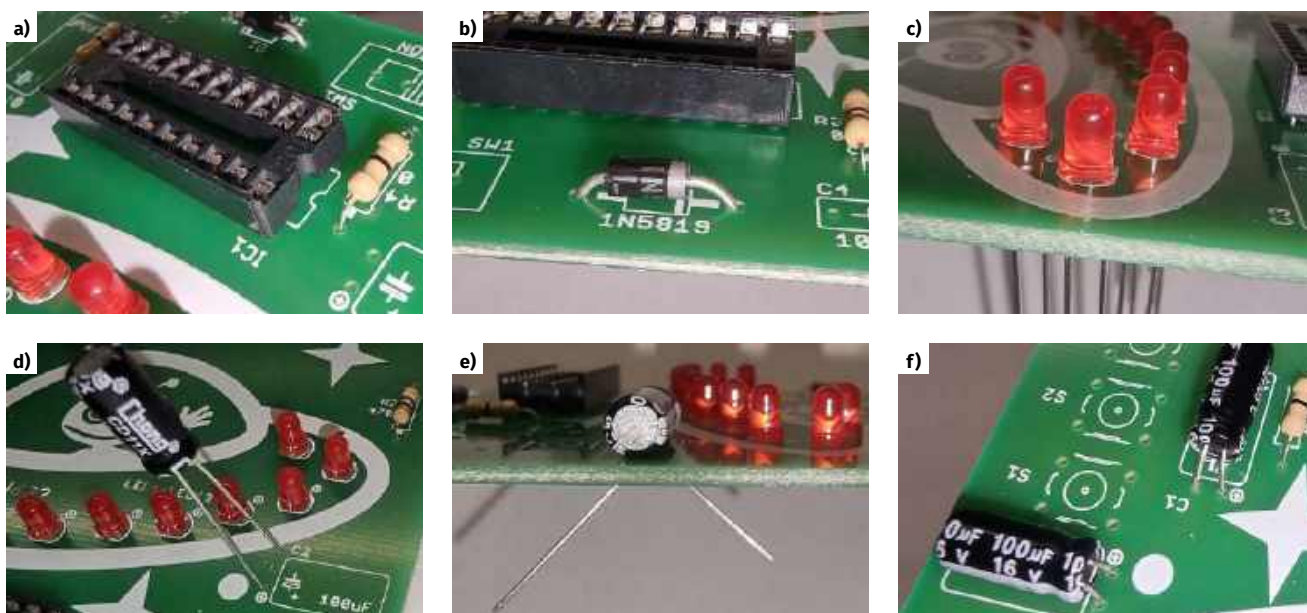
Montaż PCB

Z uwagi na zastosowanie mikrokontrolera część sprzętowa UFOledka jest banalnie prosta. Chłopaki podczas zajęć stacjonarnych z montażem poradzili sobie w mgnieniu oka (fotografia 2).

Dlatego zrobimy dziś eksperyment. Podziałaj dziś bez mojej pomocy, ale z pomocy opiekuna, którego mam nadzieję, masz teraz w pobliżu, jak najbardziej korzystaj. Nie będę dziś opowiadał o poprawnym lutowaniu, o zaginaniu wyprowadzeń przed odwróceniem płytki do góry stroną lutowania, o konieczności zachowania właściwej polaryzacji

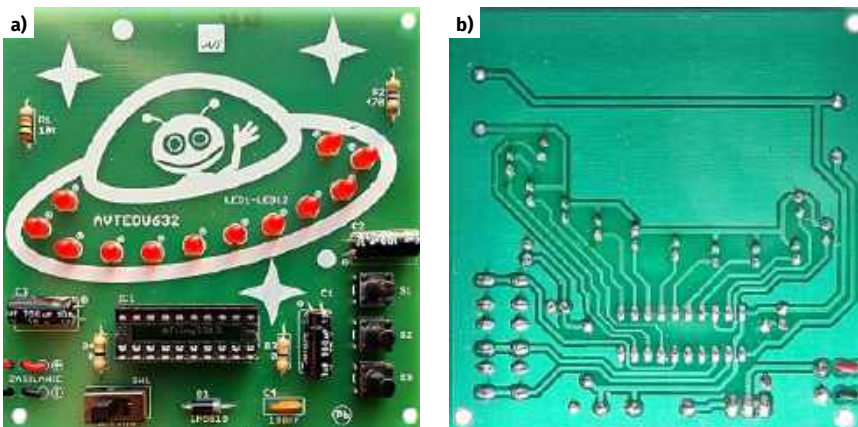
Fotografia 2. Sebastian i Tymek podczas montażu UFOledków (zajęcia koła Młodych Entuzjastów Elektroniki, Wrocław). Montaż poszedł jak z płatki. Układy zadziały od razu





▲ Fotografia 3. a) podstawka włożona do płytki we właściwym kierunku, zgodnie z opisem na PCB, b) montaż diody D1 zgodnie z polaryzacją na opisie PCB, c) montaż diod LED we właściwym kierunku, dłuższym wyprowadzeniem do otworów oznaczonych znakiem + na płycie PCB, d) poprawny montaż kondensatorów elektrolitycznych, wyprowadzeniem dłuższym do znaku + na płycie PCB, e) oraz f) poziomy montaż kondensatorów elektrolitycznych

► Fotografia 4. Zmontowany UFOledek: a) strona górna: mikrokontroler nie jest jeszcze zamontowany, zostanie on włożony do podstawki (w odpowiednim kierunku) dopiero po sprawdzeniu właściwego napięcia 5 V pomiędzy pinami 10 (GND) oraz 20 (VCC), b) strona dolna: czyste i błyszczące luty, brak zwarc



dla komponentów takich jak diody, kondensatory elektrolityczne i układy scalone. Nie będę... A, no tak, już to zrobiłem. Niech będzie moja strata. Ale teraz już milknę, i daję Ci okazję, byś mógł sam siebie sprawdzić. Zmontuj wszystko sam ale zatrzymaj się na chwilę przez włożeniem mikrokontrolera w podstawkę (ostatnie czynności przed podłączeniem baterii wykonamy razem).

Teraz załóż okulary ochronne na nos, grzej lutownicę i do dzieła! Wykorzystując wiedzę zdobytą na poprzednich spotkaniach zamontuj w płytce drukowanej w następującej kolejności:

- rezystory R1 (10 kΩ), R2 (470 Ω), R3 i R4 (0 Ω)
- podstawka pod mikrokontroler IC1 – pamiętaj o właściwym kierunku montażu – **fotografia 3a**).
- dioda D1 (1N5819) – pamiętaj o właściwym kierunku montażu – **fotografia 3b**.
- diody LED1...LED12 (dowolne czerwone diody LED o średnicy 3 mm – pamiętaj o właściwym kierunku montażu – **fotografia 3c**).

- kondensatory elektrolityczne C1...C3 (100 μF/16 V) – pamiętaj o właściwym kierunku montażu – **fotografie 3d, e, f**).
- kondensator ceramiczny C4 (100 nF)
- przyciski S1...S3 (popularne „tact switch” 6×6 mm o dowolnej wysokości)
- mikrokontrolera do podstawki jeszcze nie wkładaj, zamiast tego postępuj zgodnie z opisem poniżej.

Zakładam, że zmontowałeś już UFOledka. Efekt końcowy powinien być taki jak na **fotografii 4a i b**.

Weryfikacja poprawności napięcia zasilania w podstawie mikrokontrolera

Jeśli płytka po stronie lutowania wygląda w porządku i nie widać zwarc ani niedolotów, to świetnie, ale zanim odważysz się włożyć mikrokontrolera (najdroższy element w zestawie) do podstawki (pamiętaj o właściwym kierunku, bo przy pomyłce jego struktura wewnętrzna spłonie i nie będzie się już nadawał do użytku!) sprawdź najpierw, czy

na wyprowadzeniach podstawki po włożeniu baterii do koszyeczka i ustawieniu przełącznika SW1 na pozycji „ON” w podstawie pod mikrokontroler IC1 pomiędzy pinami 20 (VCC) i 10 (GND) pojawi się napięcie około 4,5 V.

W tym celu wykorzystaj proszę multimetr ustawiony na funkcję woltomierza napięcia



Fotografia 5. Multimetr ustawiony w tryb woltomierza w zakresie pomiarowym „do 20 V” napięcia stałego (DCV) w celu przeprowadzenia pomiaru napięcia między nogą 20 (VCC) i 10 (GND) podstawki pod mikrokontroler IC1



Fotografia 6. Sposób przyłożenia sond pomiarowych multimetru do odpowiednich pinów podstawki mikrokontrolera (sonda czerwona do pinu 20, sonda czarna do pinu 10). Przycisk SW1 ustawiony w pozycji „ON” (załączony)

stałego w zakresie pomiarowych „do 20 V” (fotografia 5).

Upewnij się, że czarny przewód jest wpięty w złącze COM multimetru, a przewód czerwony w gniazdo służące do pomiaru napięcia (fotografia 5).

Teraz ustaw przełącznik SW1 w pozycji „ON” oraz przyłóż czerwoną sondę pomiarową multimetru do wyprowadzenia numer 20 podstawki a czarną sondę do pinu 10 podstawki (fotografia 6). Z woltomierza powinno dać się w tym momencie odczytać wartość około 4,5 V. Na **fotografii 7a** można zobaczyć, że wyświetlacz multimetru wskazał podczas zajęć napięcie 4,67 V). Dzięki temu Sebastian przekonał się, że może ze spokojną głową przystąpić do montażu mikrokontrolera w podstawce (fotografia 7b).

Dopóki napięcie nie przekracza napięcia 5 V, albo na pierwszej pozycji wyświetlacza nie pojawia się znak minusa, wszystko jest w porządku. O ile w przypadku użycia trzech szeregowo połączonych baterii (rola koszyeczka) o napięciu 1,5 V każda co prawda nie ma ryzyka o to, że podamy na mikrokontroler zbyt wysokie napięcie (w przypadku tego konkretnego, większe niż 5 V), o tyle może się zdarzyć, że źle podłączysz kabelki baterii i tym samym odwrócisz polaryzację



Fotografia 8. Podłączenie kabelków baterii do płytki PCB



Fotografia 7. a) wynik pomiaru napięcia dokonanego na odpowiednich zaciskach podstawki mikrokontrolera: 4,67 V, b) Sebastian daje znak, że wynik pomiaru jest poprawny i można zamontować mikrokontroler w podstawce



zasilania, co skutecznie zabije mikrokontroler. Nawet jeśli podłączyłeś czerwony kabelek koszyeczka baterii do pola lutowniczego oznaczonego znakiem „+” a czarny do pola oznaczonego znakiem „-” (**fotografia 8**) nigdy nie możesz być do końca pewien, czy na przykład koszyczek nie został wyprodukowany wadliwie.

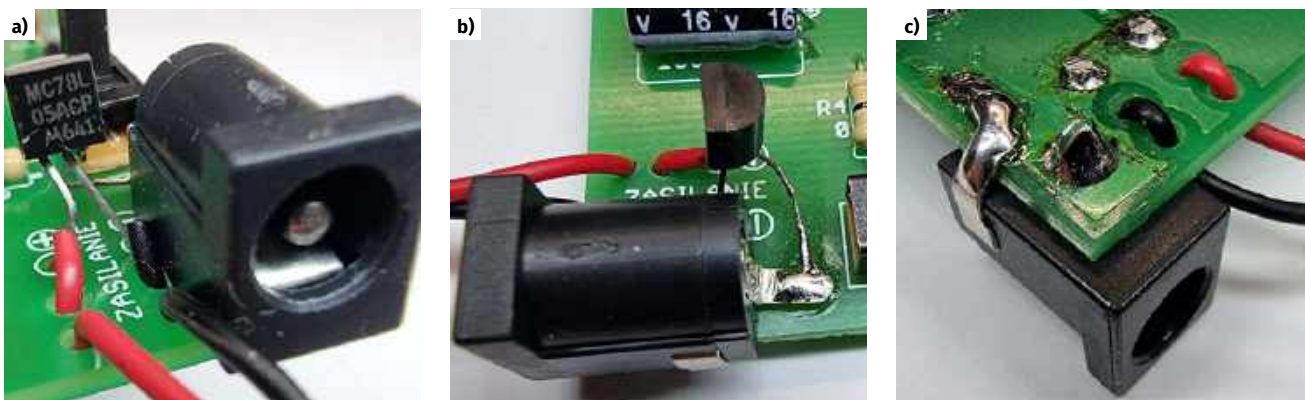
Trafić na wadliwie wyprodukowany koszyczek to trochę jakby wygrać milion w totka, albo znaleźć igłę w stogu siana, ale mi taka sytuacja przytrafiła się podczas eventu, który relacjonowałem w ramach pierwszego odcinka cyklu EdW Junior. Koszyczek zakupiony u renomowanego polskiego dostawcy wystawiał „plus” na kabelku czarnym, a „minus” na kabelku czerwonym! Ależ było moje zdziwienie gdy to odkryłem podczas namierzania przyczyny usterki złożonego układu, który nie zadziałał po włożeniu baterii! Dlatego bacznie przyjrzyj się wynikowi pomiaru. Jeśli sondę czerwoną przykładasz do „plusa” a sondę czarną do „minusa” na pierwszej pozycji wyświetlacza multimetru nie może się pojawić znak minus! Gdy się pojawi oznacza to, że „plus” jakimś sposobem został zamieniony z „minusem” i zanim włożysz mikrokontroler w podstawkę i załączysz urządzenie, najpierw musisz odnaleźć przyczynę usterki. W ten sposób uratujesz najdroższy element w zestawie przed nieodwracalnym unicestwieniem, którego w takim przypadku pewnie nawet UFO do życia nie wskrzesi.

Jeśli przykładając czerwony przewód woltomierza do pinu 20 podstawki mikrokontrolera

oraz czarny przewód do pinu 10 podstawki mikrokontrolera i odczytałeś napięcie około 4,5 V na plusie (bez minusa z przodu) to już niewiele dzieli Cię od sukcesu (fotografia 7a i 7b). Ustaw przełącznik SW1 w pozycji „OFF” (i najlepiej wyjmij jedną z baterii z koszyeczka), następnie pamiętając o właściwej polaryzacji włóż mikrokontroler w podstawkę. W razie potrzeby sięgnij po moje uwagi dotyczące montażu układów scalonych w podstawce z poprzedniego odcinka. Gdy mikrokontroler został prawidłowo włożony w podstawkę i docisnięty, możesz włożyć baterie do koszyeczka i włączyć zasilanie ustawiając przełącznik SW1 w pozycji „ON”. Twoim oczom ukaże się jeden z siedmiu możliwych do ustawienia efektów świetlnych. Zdradzę Ci tajemnicę pozyskaną od UFO, że efekty możesz zmieniać za pomocą krótkiego naciśnięcia przycisku S1. Ponadto naciskając przycisk S2 będziesz miał wpływ na szybkość odtwarzania wybranego efektu, a za pomocą przycisku S3 będziesz mógł regulować czas „smugi” pozostawianej przez UFO.

Dodanie stabilizatora 78L05 celem zasilania UFOledka z typowego zasilacza wtyczkowego o napięciu stałym 12 V

W poprzednich częściach wspominałem, że nie chcąc inwestować każdorazowo w koszyczki baterii i nowe baterie, każdy uruchamiany z chłopakami układ staramy się zasilić napięciem 12 V z zasilacza sieciowego.



Fotografia 9. Montaż dodatkowego stabilizatora 78L05 i gniazda zasilania: a) napięcie wyjściowe stabilizatora, czyli noga 1 trafia do otworu oznaczonego na płytce znakiem (+), noga 2 stabilizatora została przylutowana do otworu oznaczonego na płytce znakiem (-), b) noga 3 stabilizatora (napięcie wejściowe stabilizatora) została przylutowana do wyjścia +12 V gniazda zasilającego, c) gniazdo zasilania zostało przylutowane do PCB za pomocą styku dla napięcia ujemnego w otworze montażowym otoczonym miedzią podłączoną do potencjału GND na płytce PCB

Ponieważ podanie napięcia 12 V na mikrokontroler wymagający zasilania napięciem 5 V skończyłoby się natychmiastowym jego spalaniem, zastosowaliśmy na wejściu miniaturowy stabilizator napięcia 78L05 w obudowie TO-92, którego rolą jest obniżenie napięcia wejściowego 12 V do poziomu 5 V na wyjściu stabilizatora. Dwa z jego wyprowadzeń wlutowaliśmy w miejsce kabelków koszyeczka baterii. W miejsce otworu dla kabelka czarnego (-) umieściliśmy nogę numer 2 (GND) stabilizatora 78L05 a nogę numer 1 (V-OUT) tego stabilizatora zamontowaliśmy w miejsce kabelka czerwonego. Na wejście stabilizatora (noga 3 układu 78L05) podłączyliśmy odgiętą nogę gniazda zasilającego, którego drugą nogę (GND) zamontowaliśmy w istniejącym na płytce PCB otworze montażowym, umożliwiającym przylutowanie „ujemnego” styku gniazda do „wylania miedzi” na płytce PCB wokół tego otworu. Należałoby jeszcze przylutować dodatkowe kondensatory filtrujące na wejściu i wyjściu stabilizatora 78L05, oczywiście nie ma dla nich miejsca na płytce

i ponieważ praktyka pokazała, że układ działa świetnie i bez nich, świadomie zrezygnowaliśmy z ich dokładania. Sposób pokazujący montaż dodatkowego stabilizatora 78L05 oraz gniazda zasilania pokazano na fotografiach **9a, b, c**).

Na fotografiach 9a, b, c widzimy zamontowany stabilizator wraz z kabelkami od baterii. Sytuacja miała miejsce na prototypie, po czym koszyczek baterii został trwale zdemontowany, by uniknąć sytuacji, gdzie oba źródła zasilania (zasilacz i baterie) zostaną podłączone jednocześnie, co po dłuższym czasie mogłoby sprawić, że baterie wyleją lub w najgorszym przypadku eksplodują. Na zajęciach (wcześniejsze fotografie) koszyczek baterii w ogóle nie był montowany. Należy podjąć decyzję i zamontować albo koszyczek z bateriami albo stabilizator wraz z gniazdem zasilania i nigdy nie używać obu tych rozwiązań jednocześnie.

Podsumowanie

Elektronika to piękne hobby. Czasem o prostych rzeczach można by opowiadać całymi godzinami! Nie wiem, czy zauważyłeś, ale

diody LED UFOledka nie są jedynie załączane i wyłączane przez mikrokontroler. Odbyna się tu sterowanie jasnością, dzięki której masz okazję podziwiać efekt „smugi” pozostawianej przez „kręcący się obiekt”. Z pewnością wykorzystano tu sterowanie jasnością diod LED w oparciu o mechanizm PWM (Pulse Width Modulation), czyli modulacji szerokości impulsu. Jest to technika używana w elektronice i sterowaniu, w której szerokość impulsów sygnału jest modulowana, aby kontrolować moc dostarczaną do urządzeń lub elementów takich jak diody LED. Ale to już historia na inną opowieść.

Tymczasem w tym miesiącu już żegnam się z Tobą. Myślę, że pozostawiłem Ci na ten miesiąc wystarczająco dużo treści do pochłonięcia. Może się okazać, że na jednym „tchu” się nie skończy. Za miesiąc znów się spotkamy i będziesz miał okazję zbudować kolejną, być może tym razem nieco bardziej przyjemną (a już na pewno mniej kosmiczną) zabawkę. Do zobaczenia! ■

Mariusz Ciszewski

REKLAMA

Publikujemy dla projektantów i programistów elektroniki

ELPORTAL.pl

Znajdziesz nas również na Facebooku: facebook.com/ElportalPL

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**.

2-kanalowy moduł przekaźnikowy Wi-Fi wykorzystujący ESP8266 NodeMCU

Opisany projekt to dwukanałowy moduł przekaźnikowy ESP8266 NodeMCU. Płytką tą umożliwia użytkownikom zdalne sterowanie dwoma przekaźnikami. 2 kanałami można sterować zdalnie za pomocą serwera WWW. Projekt zasilany jest napięciem 5 VDC. Dwie diody LED wskazują działanie przekaźnika. Zaciski śrubowe pomagają użytkownikowi podłączyć urządzenia i zasilanie. Wyjście przekaźnikowe ma 3 styki: zamknięty, otwarty i wspólny. Przekaźnik może sterować obciążeniem do 10 A/250 VAC lub 10 A/30 VDC. Może być używany w aplikacjach takich jak IoT czy automatyka domowa.



8-kanalowy pilot zdalnego sterowania ESP32

Jest to 8-kanalowy pilot w kompaktowej obudowie, odpowiedni do zdalnego włączania/wyłączania za pomocą protokołu Bluetooth lub Wi-Fi. Do płytki przylutowano moduł ESP32 i 8 przełączników dotykowych podłączonych do różnych GPIO. Obwód działa z zasilaniem 3,3 VDC, a do zasilania płytki można użyć baterii litowo-jonowej. Wbudowana dioda LED wskazuje zasilanie modułu. Na pokładzie znajduje się złącze, które pomaga użytkownikom programować moduł ESP32 za pomocą Arduino IDE. Użytkownicy muszą napisać własny kod, aby korzystać z modułu zgodnie z własnymi wymaganiami.

Niektóre projekty aktualnie dostępne tylko dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl:

- Pojemnościowy czujnik wilgotności do konwertera wyjścia analogowego
- Mostek H dla wysokiej mocy szczotkowego silnika prądu stałego z czujnikiem prądu
- Przetwornica DC-DC buck 12...75 V na 10 V na wyjściu
- Czujnik prądu low-side 10 μ A...10 mA
- Kontroler ramienia robota z bezprzewodowym pilotem PS3
- Termiczny czujnik masowego przepływu powietrza – anemometr stałotemperaturowy
- Precyzyjny wzmacniacz transimpedancyjny z przetwarzaniem integratorem
- Kontroler pełnego mostka z przesunięciem fazowym i prostowaniem synchronicznym wykorzystujący UCC28950
- Wysokowydajny monofoniczny wzmacniacz audio klasy D o mocy 20 W
- Monitorowanie poziomu cieczy za pomocą czujnika ciśnienia – wyświetlacz słupkowy
- Sterowanie silnikiem DC za pomocą joysticka
- 16-kanalowy sterownik serwomechanizmów RC z interfejsem I²C
- Programowalny kondycjoner sygnału z czujnika rezystancyjnego mostkowego
- Choiinka z Arduino i pikselowymi diodami
- 20-segmentowy wyświetlacz słupkowy w rozmiarze jumbo
- Stacja pogodowa Lilygo ttgo t5-4.7 z wyświetlaczem typu e-papier
- Półprzewodnikowy przekaźnik mocy DC z prądowym sprzężeniem zwrotnym
- Wyłącznik nadprądowy – przekaźnik wyłączający nadprądowy
- Uniwersalny konwerter napięcia AC – wyjście 18 V DC z wejścia 85...265 V AC
- Moduł procesora echa głosu – urządzenie opóźniające do efektów dźwiękowych, echo, reverb
- Sterownik silnika krokowego z joystickiem
- RPI – stacja pogodowa IoT
- Niskobudżetowy monitor jakości powietrza IoT oparty o RaspberryPi 4
- Automatyczny system ogrodniczy z NodeMCU i Blynk, ArduFarmBot 2
- TinyML – Rozpoznawanie ruchu przy pomocy RPI Pico
- Wzmacniacz piezoelektryczny do gitary i skrzypiec
- Wysokowydajny i niezawodny sterownik bipolarnego silnika krokowego
- Sonarowy theremin MIDI
- Sterownik silnika prądu stałego z wykorzystaniem przekaźnika i mosfetu – interfejs Arduino
- Przedwzmacniacz do mikrofonu MEMS
- Super prosty czuły wykrywacz metali
- 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:
AVT-Korporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

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 Sobarski
jakub.sobanski@elportal.pl

Sekretarz redakcji:
Dariusz Welik
dariusz.welik@elportal.pl

Copyright AVT-Korporacja Sp. z o.o., Warszawa, ul. Leszczyńska 11. Projekty publikowane w „Elektronice dla Wszystkich” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki dla Wszystkich”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice dla Wszystkich” jest dozwolone wyłącznie po uzyskaniu pisemnej zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice dla Wszystkich”.

DTP, okładka,
Redakcja strony internetowej www.elportal.pl:
MAD Sp. z o.o.

Prenumerata:
W Wydawnictwie AVT, e-mail: prenumerata@avt.pl
tel. 22 257 84 22, (godz. 10:00–14:00)
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

