

ELEKTRONIKA

dla wszystkich

nr 1/2024 (336) • styczeń • www.elportal.pl



Uniwersalny pełnookresowy regulator prędkości silnika Full Wave

PROJEKTY dla elektroników

- ▶ Elektroniczne dzwonki wietrzne, część 2
– wykończenie
- ▶ Dbaj o swoje akumulatory. Wysokoprądowy balanser baterii, część 2 – budowa
- ▶ 64-klawiszowa matryca MIDI sterowana modulem Arduino, część 2

DIY dla wszystkich

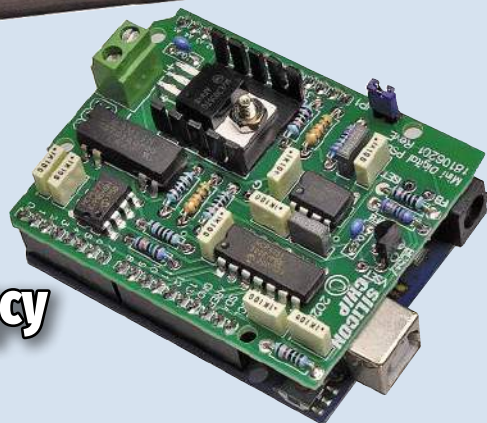
- ▶ Dokładny zegar ze wskazaniem milisekund z użyciem ESP32 RTC
- ▶ Automatyczny dozownik odmierzonej ilości cieczy
- ▶ Prosty wskaźnik temperatury wody oraz sterownik z użyciem Arduino

TUTORIALE

- ▶ Audio OUT: Wzmacniacz audio do Theremina, część 4
- ▶ Chirurgia obwodowa: Czas i metastabilność w obwodach synchronicznych, część 2
- ▶ Know-how: Komunikacja 433 MHz
- ▶ Podzespoły: ogniwa Peltiera
- ▶ Edukacja w EdW dla szkół i uczelni: Wykład 14. Multimetry. Od AVO do niemal nieskończonego wyboru
- ▶ Ekscytacje Maxa: Migające diody LED i śliniacy się inżynierowie



Regulowany
zasilacz
wykorzystujący
Arduino



EP.com.pl
Największy portal dla elektroników konstruktorów

**Król automatyki
jest w Tobie**

AutomatykaB2B.pl

FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

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

www.piekarz.pl

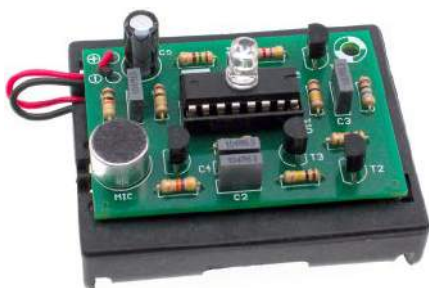
OLED

artronic
OPTOELEKTRONIKA
www.artronic.pl



Najbardziej popularne kity AVT

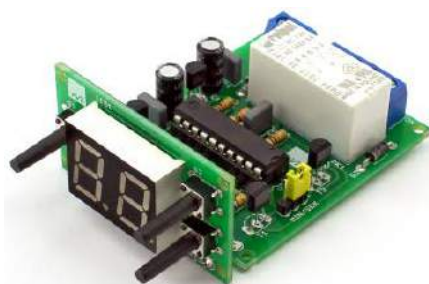
Poznaj listę **TOP 100** na www.elportal.pl/kityavt



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



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



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



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



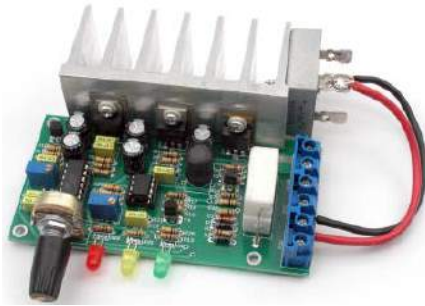
AVT723 Uniwersalna gra zręcznościowa
<https://sklep.avt.pl/avt723.html>



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



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



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



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



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



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

PRENUMERATA roczna EdW+

NA START
DO 6 WYDAŃ
GRATIS!

Cena drukowanej prenumeraty rocznej na start wynosi 185,90 zł
Przy zamówieniu prenumeraty dwuletniej za 304,20 zł
oszczędność wynosi równowartość sześciu wydań EdW

PO 5 LATACH
ZA PÓŁ CENY

Przedłuż prenumeratę drukowaną po zalogowaniu się do swojego panelu na www.UlubionyKiosk.pl/logowanie, gdzie znajdziesz atrakcyjną ofertę, która uwzględnia przysługujące Ci zniżki za lojalność. Po 5 latach nieprzerwanej prenumeraty **otrzymasz rabat 50% na drukowaną prenumeratę dwuletnią**

PRENUMERATA EdW+

Rozpocznij przygodę z elektroniką – poznaj jej podstawy, zamawiając roczną prenumeratę drukowaną EdW wraz z Praktycznym Kursem Elektroniki (PKE)

Do wysyłki prenumeraty dołączymy zestaw edukacyjny EDW A09 KPL, na który składają się:

1. projekt – układ elektroniczny samodzielnie uruchamiany przez kursanta. Wszystkie układy są montowane na dołączonej płytce stykowej, do której wkłada się „nóżki” elementów na wcisk,
2. pendrive z wykładami i materiałami multimedialnymi kursu PKE.
3. zasilacz płytek stykowych AVT3072 C
4. oraz zasilacz impulsowy 12 V, 1,4 A

Cena prenumeraty EdW+ wynosi **280,90 zł**

TYLKO prenumeratorzy* otrzymują pełny dostęp do:

ARCHIWUM

cyfrowego archiwum EdW na elportal.pl/archiwum



projektów w zbiorze DIY+ na elportal.pl/diy

* Promocja z dostępem do archiwum EdW oraz projektów DIY+ dotyczy płatnej prenumeraty drukowanej lub płatnej e-prenumeraty EdW zamawianej na www.UlubionyKiosk.pl bądź przelewem na konto Wydawnictwa AVT. Po odnotowaniu płatności wysyłamy mailowo kod dostępu, za pomocą którego zalogujesz się na elportal.pl

Zamów prenumeratę lub e-prenumeratę na www.UlubionyKiosk.pl/prenumerata

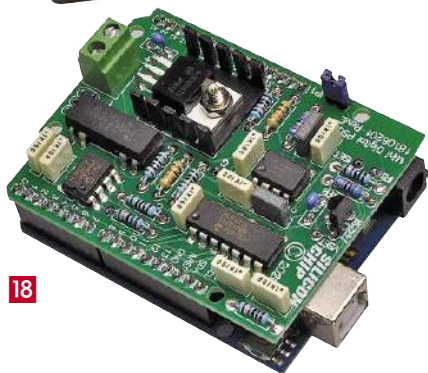
Kontakt ws. prenumeraty: 22 257 84 22 (godz. 10.00–14.00), prenumerata@avt.pl
Kontakt merytoryczny ws. kursu PKE: kity@avt.pl • Konto bankowe: AVT-Korporacja sp. z o.o.
03-197 Warszawa, ul. Leszcynowa 11, ING Bank Śląski 18 1050 1012 1000 0024 3173 1013



8

Projekty dla elektroników:

- Uniwersalny pełnookresowy regulator prędkości silnika Full Wave 8
- Wyjście 0–14 V, 0–1 A – sterowanie z komputera!
- Regulowany zasilacz wykorzystujący Arduino..... 18
- Elektroniczne dzwonki wietrzne, część 2 – wykończenie..... 26
- Dbaj o swoje akumulatory.
- Wysokoprądowy balanser baterii, część 2 – budowa..... 32
- 64-klawiszowa matryca MIDI sterowana modułem Arduino, część 2..... 40



18

Tutoriale:

- Audio OUT: Wzmacniacz audio do Theremina, część 4..... 47
- Chirurgia obwodowa:
- Czas i metastabilność w obwodach synchronicznych, część 2..... 50
- Know-how: Komunikacja 433 MHz..... 54
- Podzespoły: ogniwa Peltiera..... 60
- Edukacja w EdW dla szkół i uczelni: Wykład 14. Multimetry.
- Od AVO do niemal nieskończonego wyboru 66
- Ekscytacje Maxa:
- Migające diody LED i śliniacy się inżynierowie (4) 81
- Sprytnie porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania .. 83



26

DIY dla wszystkich:

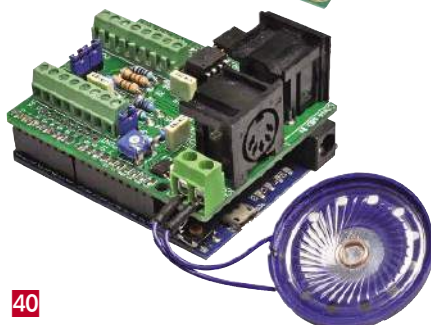
- Dokładny zegar ze wskazaniem milisekund z użyciem ESP32 RTC..... 84
- Automatyczny dozownik odmierzonej ilości cieczy 86
- Prosty wskaźnik temperatury wody oraz sterownik z użyciem Arduino..... 88



32

DIY PLUS

- 3-fazowy sterownik silnika bezszczotkowego wykorzystujący L6235 91
- Sterownik solenoidu, przekaźnika, zaworu z regulacją prądu 91



40

Rubryki stałe:

- Prenumerata 3
- Od wydawcy 5
- Poczt..... 6

A za miesiąc w lutowym EdW



* Rejestrator stanu akumulatorów, część 1

Znajomość stanu akumulatorów jest niezbędna do utrzymania ich przez długi czas w dobrej kondycji. System, który może monitorować i rejestrować istotne parametry akumulatora jest bardzo przydatny i może pomóc uniknąć konieczności kosztownych wymian. Może być również używany do rozwiązywania problemów, na przykład, gdy nie wiadomo, który moduł czy panel jest odpowiedzialny za okresowe rozładowywanie baterii.

* Pęseta do testów SMD

To sprytne małe urządzenie składa się zaledwie z 11 elementów. Mimo to może mierzyć wartości wielu rezystorów SMD i kondensatorów, a także pokazywać orientację diod i LED oraz mierzyć ich napięcia przewodzenia. Jest szybkie i łatwe w użyciu, zasilane z wbudowanego ogniwa pastylkowego, z ekranem OLED o wysokim kontraście do wyświetlania odczytów.

* Prosta, liniowa klawiatura muzyczna MIDI, część 3

Ta klawiatura Midi jest następcą naszej 64-klawiszowej matrycy MIDI, którą zaprezentowaliśmy w numerach 12/2023 i 1/2024 EdW. Jest podobnie modyfikowalna i oferuje sposób na łatwe tworzenie muzyki, chociaż można ją też wykorzystać w wielu innych zastosowaniach.

* Subwoofer tubowy „Tapped Horn”

Ta niedroga i zajmująca mało miejsca konstrukcja doskonale odtwarza dźwięki powyżej 100 dB SPL (Sound Pressure Level) w zakresie poniżej 30 Hz. Doysterowania można używać względnie mały wzmacniacz.

* Plus kolejna porcja intrygujących projektów DIY.

* Plus wiele artykułów w Twoich ulubionych cyklach Tutoriali.

**W kioskach
od 31 stycznia**



Krótką historia multimetrów

Tematem wykładu w tym numerze EdW jest przyrząd pomiarowy, bez którego nie może się obejść żaden elektronik. Równie 100 lat temu brytyjski inżynier pocztowiec Donald Macadie zrealizował pomysł zintegrowania w jednym przyrządzie pomiarów natężenia prądu (A – ampery), napięcia (V – wolty) i oporności (O – omy). Ten wielofunkcyjny miernik wskazówkowy, znany pod nazwą AVOMetr, to protoplasta współczesnego multimetru cyfrowego.

Droga wskazówkowego AVOMetru, który kosztował w latach trzydziestych ponad 30 USD (ok. 500 dzisiejszych USD), do współczesnego multimetru cyfrowego, który mieści się w dłoni i kosztuje kilka dolarów (na Allegro można kupić już za kilkanaście złotych, a oferta na Aliexpress za 2 zł jest zupełnie niepojęta), prowadziła przez kilka przełomów. Odnajmy dwa najważniejsze. Pierwszy to pojawienie się w roku 1953 multimetrów cyfrowych – dużych nastolnych przyrządów lampowych. Drugi przełom wiąże się z pierwszym multimetrem ręcznym na układach scalonych – Fluke 8020A w roku 1977. Kluczową rolę w tym multimetrze odgrywa jeden układ scalony spełniający funkcję przetwornika analogowo-cyfrowego i sterownika wyświetlacza 3½ cyfry.

Ten genialny chip, oznaczony przez Fluke jako 429100 jest powszechnie znany jako oryginalny produkt firmy Intersil pod nazwą ICL 7106 (wersja z wyświetlaczem LCD) oraz ICL 7107 (wersja z wyświetlaczem LED). Układ ICL 7106/7107 to fenomen długowieczności porównywalny ze wzmacniaczem operacyjnym 741 oraz timerem 555. Te układy scalone opracowano w latach sześćdziesiątych/siedemdziesiątych i ciągle są produkowane.

Której firmie mamy zawdzięczać początki ręcznych multimetrów cyfrowych – Fluke czy Intersil? Ten problem próbowano rozstrzygnąć w sądzie, choć zaczęło się przyjemnie, od współpracy. Firma Fluke zleciła firmie Intersil wyprodukowanie scalonego przetwornika ADC pod nomenklaturą 429100. Firma Intersil, którą założył w roku 1967 twórca technologii planarnej Jean Hoerni, dysponowała unikalną technologią i kompetencją do wykonania tego zlecenia Fluke. W istocie, ten układ scalony, po dodaniu wyświetlacza, to niemal gotowy multimetr. Po roku od wykonania zlecenia Fluke firma Intersil wypuściła na rynek układ ICL 7106, z niewielkimi zmianami klon układu 429100. Na strukturze krzemowej pozostało nawet wytrawione logo Fluke. Firma Fluke wytoczyła firmie Intersil proces sądowy, jednak pozew został wycofany, gdy zarząd Fluke zorientował się, że „wisi” na jednym jedynym dostawcy chipów, którym jest Intersil. Multimetr Fluke 8020A okazał się wielkim sukcesem handlowym Fluke, a szeroko dostępne chipy Intersila ICL 7106/7107 otworzyły drogę do uruchomienia produkcji ręcznych multimetrów w wielu firmach na świecie.

Z własnego podwórka dodajmy, że pierwszymi kitami firmy AVT były woltomierze panelowe AVT 01/02 oraz multimetry AVT 03/04 zbudowane na bazie układów ICL 7106/7107. Ponad 30 lat temu nasza oferta kitów multimetrów za 200.000 zł (na stare pieniądze) zrobiła furorę, gdyż dostępne wówczas na rynku dwa modele krajowych multimetrów 3½ cyfry kosztowały po około 900.000 zł. Można powiedzieć, że doskonały start rynkowy firma AVT zawdzięcza genialnym układom Intersila ICL 7106/7107.

Wiesław Marciniak

Niesforne LEDy

W moim domu jednorodzinny mam instalację elektryczną zakładaną 30 lat temu. Wyłączniki są podświetlane małutkimi neonówkami. Gdy ostatnio zdecydowałem się na wymianę żarówek wolframowych na lampy LED spotkała mnie niespodzianka. Niektóre LEDy nie dają się do końca wyłączyć, inne zaś generują rozbłyśki jak stroboskop. Doczytałem się w Internecie, że powinienem zrezygnować z podświetlanych wyłączników, bo neonówki stosowane do podświetlenia przewodzą niewielki prąd, wystarczający do świecenia LEDów. Są też jakieś rozwiązania z dodatkowym kondensatorem. Czy możecie poruszyć ten temat na łamach EdW?

Red. Znaleźliśmy w pocztce magazynu Silicon Chip dyskusję na ten temat, którą warto w tym miejscu przytoczyć.

Prąd upływu powoduje świecenie lub miganie lamp LED

Czy kiedykolwiek zdarzyło Ci się, że zasilane z sieci lampy sufitowe LED świeciły po wyłączeniu lub, co gorsza, powoli migały, jednocześnie świecąc? Być może nie, jeśli są one włączane za pomocą standardowego wyłącznika światła na ścianie; jeśli jednak używasz urządzenia półprzewodnikowego, takiego jak ściemniacz, prawdopodobnie tak się stało.

Wiele lat temu zbudowałem dom wakacyjny na północnym wybrzeżu Nowej Południowej Walii. Postanowiłem, że wszystkie włączniki światła będą przełączane niskim napięciem. Sygnały niskonapięciowe sterowały przekaźnikami półprzewodnikowymi wykonanymi z triaków i opto-sprzęgaczy, takich jak MOC3021. Wyobrażałem sobie, że pewnego dnia, gdy pojawi się technologia, będę mógł zdalnie włączać i wyłączać światła. Na przykład, aby dom wyglądał na zajęty podczas naszej nieobecności.

Działało dobrze, bez problemów, przez wiele lat, chociaż nigdy nie udało mi się skonfigurować zdalnie sterowanych światła.

W związku z pandemią i powodziami na północnym wybrzeżu zdecydowałem się wynająć dom, ponieważ w okolicy istniało rozpaczliwe zapotrzebowanie na wynajęte zakwaterowanie. Wszystkie światła nadal działały dobrze.

Następnie najemca poinformował mnie, że niektóre światła świeciły w nocy, gdy były wyłączone, a jedno lub dwa migały powoli.

Stało się to dopiero wtedy, gdy żarówki żarowe w końcu się przepaliły i zostały zastąpione lampami LED. Żarówki z żarnikiem są obecnie prawie niemożliwe do zdobycia.

Poszperałem trochę i szybko okazało się, że jest to duży problem, szczególnie w produkcjach teatralnych i scenicznych, gdzie specjalne efekty świetlne wymagają szybkiego przełączania światła. Używają oni półprzewodnikowego przełączania napięcia sieciowego, więc miałiby podobne problemy jak mój system.

Problem wydaje się wynikać z wysokiej sprawności lamp LED zasilanych z sieci, które pobierają bardzo mało prądu. W związku z tym, nawet bardzo małe prądy upływu przez układy przełączające triaki lub filtry i filtry wokół tych triaków, pozwalają zasilaczom



lamp LED na okresowe włączanie (miganie) lub zasilanie diod LED z niską mocą (świecenie), nawet gdy są „wyłączone”. W Internecie pojawiły się sugestie, aby po prostu zamontować prostą żarówkę pilotującą 15 W równolegle z lampą LED, aby zbocznikować ten prąd upływu i zapobiec świeceniu lampy LED. Skonfigurowałem to na próbę z lampą LED i na pewno zadziałało.

W Internecie pojawiły się historie o scenicznych konfiguracjach oświetleniowych z setkami małych żarówek żarnikowych podłączonych równolegle do lamp LED. Żarówki żarnikowe są tak nieefektywne, że nie wytwarzają żadnego widocznego strumienia świetlnego podczas przewodzenia prądów upływu. Nie wydaje się to jednak praktyczne w warunkach domowych.

Eksperymentowałem więc z alternatywnymi metodami bocznikowania lamp LED zasilanych z sieci. Jedną z prostych technik jest umieszczenie rezystora 47...100 Ω i kondensatora 100 nF szeregowo ze sobą, bezpośrednio na końcówkach lampy LED zasilanej z sieci, aby bocznikować ten prąd upływu.

Wypróbowałem to z kilkoma lampami LED zasilanymi z sieci i zadziałało to w większości przypadków, ale była jedna, w której lampa nadal świeciła po wyłączeniu. Hmm!

Miałem dwie podtynkowe lampy sufitowe LED 92 mm, identyczne pod każdym względem, z wyjątkiem tego, że jedna była oznaczona jako 7 W, a druga jako 9 W. Ta o mocy 9 W miała możliwość ściemniania, a ta o mocy 7 W nie.

Zasilałem obie lampy równolegle za pomocą przełącznika Triac. Przy wyłączonym przełączniku jedna z nich była całkowicie wyłączona, podczas gdy druga świeciła. Miałem rezystor 47 Ω i kondensator 100 nF równolegle z nimi i doszedłem do wniosku, że tworzy to dzielnik napięcia z rezystorem i kondensatorem na przełączniku Triac.

Najprostszym rozwiązaniem było usunięcie bocznika na triaku i przełączenie się na opto-sprzęgacz triaka MOC3063 z detekcją przejścia przez zero. Nawet kondensator 47 nF, zaprojektowany w celu usunięcia szumu linii z sygnału wyzwajającego triaka, musiał zostać usunięty. W końcu wszystko działało idealnie, a obie diody LED wyłączały się.

Chciałbym wiedzieć, czy inni doświadczyli podobnego problemu i czy znaleźli inny sposób na jego przezwyciężenie, poza użyciem przekaźnika zamiast triaka.

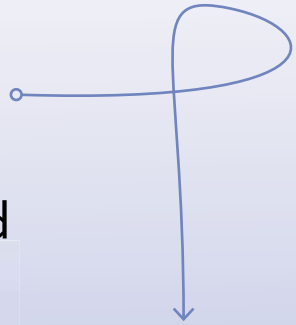
Komentarz od redakcji SC: widzieliśmy lampy fluorescencyjne świecące/migające przy wyłączonym przełączniku w budynkach ze starym okablowaniem, co uznaliśmy za spowodowane upływem przez starą izolację przewodów (która mogła w pewnym momencie ulec zawilgoceniu lub po prostu uległa degradacji z wiekiem) lub prawdopodobnie upływem w zwykłym przełączniku mechanicznym. Tak więc to, co opisujesz, może się zdarzyć nawet w przypadku standardowych przełączników.

Byłoby dobrze, gdyby lampy można było zaprojektować tak, aby pozostawały wyłączone, chyba że zostaną przyłożone do pełnego przebiegu napięcia sieci; problemem może być to, że są to konstrukcje uniwersalne (110...240 V), więc nie mogą być zbyt wrażliwe na napięcie zasilania.

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



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

Uniwersalny pełnookresowy regulator prędkości silnika Full Wave

230 V ↗ WYSOKIE NAPIĘCIE

Potrzebujesz wyjątkowo płynnej regulacji prędkości w całym zakresie dla swojego elektronarzędzia? Potrzebujesz zatem naszego nowego uniwersalnego regulatora prędkości obrotowej silnika. Jest on idealny do użytku z zasilanymi z sieci wiertarkami elektrycznymi, kosiarkami do trawników, nożycami do żywopłotów, piłami tarczowymi, polerkami, frezarkami lub innymi urządzeniami z uniwersalnymi (tj. szczotkowymi) silnikami o natężeniu prądu do 10 A.

Nasz najnowszy uniwersalny regulator prędkości obrotowej silnika Full Wave jest ulepszeniem tego, który opisaliśmy w Silicon Chip-ie w marcu 2018 roku. Tamten działał bardzo dobrze, ale znaleźliśmy kilka usprawnień i ulepszonych funkcji, które można było wprowadzić do projektu.

Jedną z głównych wad poprzedniej konstrukcji było to, że ustawianie wzmocnienia sprzężenia zwrotnego znajdowało się wewnątrz obudowy regulatora. Regulator ten ustawiał wielkość kompensacji w celu utrzymania prędkości obrotowej silnika pod obciążeniem.

Po ustawieniu kompensacja była odpowiednia tylko dla używanego urządzenia, ponieważ sterowanie sprzężeniem zwrotnym wymagałoby zmiany dla różnych silników.

Sterowanie to jest teraz ustawiane zewnętrznie za pomocą pokręteł, co ułatwia korzystanie z regulatora z wieloma różnymi elektronarzędziami i innymi urządzeniami.

Dodaliśmy także możliwość wyłączenia funkcji miękkiego startu, również za pomocą zewnętrznego przełącznika. Łagodny start jest przydatny, gdy regulator prędkości jest ustawiony na określoną prędkość,

DODATKOWE OSTRZEŻENIE!

Ten regulator prędkości obrotowej silnika działa bezpośrednio z sieci zasilającej 230 V AC, a kontakt z jakimkolwiek elementem pod napięciem jest potencjalnie śmiertelny. Nie buduj go, jeśli nie wiesz, co robisz.

NIE WOLNO DOTYKAĆ ŻADNEJ CZĘŚCI OBWODU, GDY JEST ON PODŁĄCZONY DO GNIAZDA ZASILANIA i nigdy nie używaj go poza uziemioną metalową obudową lub bez założonej pokrywy.

Obwód ten nie nadaje się do użytku z silnikami indukcyjnymi i może być używany wyłącznie z uniwersalnymi silnikami szczotkowymi (uzwojonymi szeregowo) lub silnikami ze zwartą fazą rozruchową (wentylatorowymi) o wartościach znamionowych prądu do maksymalnie 10 A. Więcej informacji można znaleźć w sekcji zatytułowanej „Jakie silniki mogą być regulowane”. Elektronarzędzia z wbudowanymi wentylatorami nie mogą pracować przy niskich prędkościach obrotowych przez dłuższy czas; w przeciwnym razie mogą się przegrzać. Przepis Red. EdW: Nie sprawdzaliśmy działania tego układu, ale jego schemat wyraźnie sugeruje, że może on poprawnie pracować tylko przy podłączeniu do sieci zasilającej w określony sposób, tzn. przewód fazowy we wtyczce zasilającej do styku fazowego w gniazdku zasilającym, analogicznie przewód neutralny. Zamiana tych przewodów może skończyć się spalaniem mP.

a silnik jest włączany i wyłączany w urządzeniu. Po włączeniu urządzenia prędkość silnika jest powoli i automatycznie zwiększana do ustawionej prędkości. Bez tej funkcji silnik przy starcie może szarpnąć.

Miękki start jest niezbędny w przypadku korzystania z regulatora przy pracy z frezarką lub piłą tarczową o dużej mocy.

W przypadku mniejszych urządzeń, gdy silnik jest często włączany i wyłączany, może okazać się, że ogranicza to szybkość pracy, ponieważ trzeba czekać, aż silnik ustabilizuje obroty.

Dzieje się tak w przypadku korzystania z nożyc do trawy i niektórych wiertarek ręcznych. Dlatego też wprowadziliśmy możliwość łatwego wyłączenia funkcji łagodnego rozruchu. Podczas wprowadzania tych zmian, skorzystaliśmy z okazji, aby poprawić zdolność regulatora do utrzymywania prędkości obrotowej silnika pod obciążeniem, szczególnie przy niskich ustawieniach prędkości i w przypadku urządzeń o niskim poborze mocy.

Uniwersalny regulator prędkości silnika Full Wave może być używany z zasilaniem sieciowym w zakresie 220...250 V AC przy 50 Hz lub 60 Hz. Oznacza to, że może być używany w wielu różnych krajach, choć nie nadaje się do użytku z zasilaniem 100...120 V AC.

Regulator jest zamontowany w stosunkowo płaskiej obudowie z odlewane aluminium, z wtyczką sieciową i gniazdem podłączenia narzędzi przymocowanymi na końcach przewodów wychodzących z jednego boku obudowy poprzez przepusty kablowe.

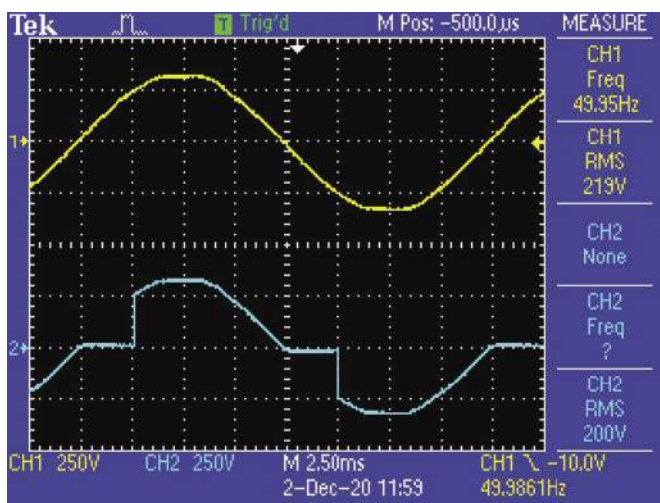
Na tym samym boku obudowy znajduje się również gniazdo bezpiecznika.

Potencjometry regulacji prędkości i wzmocnienia sprzężenia zwrotnego oraz przełącznik łagodnego rozruchu są zamontowane na pokrywie.

Dlaczego potrzebujesz sterowania prędkością?

Większość elektronarzędzi wykona lepiej pracę, jeśli posiada regulację prędkości. Na przykład, wiertarki elektryczne powinny pracować wolniej podczas używania wiertel o większej średnicy, ponieważ zapewnią to równiejsze brzozy otworu.

Podobnie przydatna jest możliwość spowolnienia obrotów frezarek, wyrzynarek, a nawet pił tarczowych, szczególnie podczas cięcia niektórych materiałów, zwłaszcza tworzyw sztucznych, ponieważ wiele z nich topi się, a nie zostaje przeciętych, jeśli prędkość narzędzia jest zbyt wysoka. Te same uwagi odnoszą się do narzędzi do szlifowania



Oscylogram 1. Przebieg wyjściowy (aktywne napięcie, w kolorze cyjan) przy ustawieniu wyższej prędkości z obciążeniem rezystancyjnym (żarówka). Widać, że napięcie wyjściowe odpowiada napięciu wejściowemu przez większość czasu, więc podłączone obciążenie otrzyma prawie pełną moc i, jeśli jest to silnik, będzie działał z dużą prędkością obrotową

Właściwości

- Do regulacji prędkości obrotowej uniwersalnych silników szczotkowych i silników ze zwartą fazą rozruchową i prądzie do 10 A
- Zasilanie 220...250 V AC przy 50 Hz lub 60 Hz
- Pełnookresowe sterowanie prędkością obrotową silnika
- Pełny zakres prędkości obrotowej (od prawie zera do blisko 100%)
- Prądowe sprzężenie zwrotne dla utrzymania prędkości obrotowej pod obciążeniem
- Regulacja wzmocnienia sprzężenia zwrotnego
- Opcjonalnie łagodny rozruch od prędkości zerowej po włączeniu zasilania
- Zoptymalizowana regulacja dla obciążeń indukcyjnych, takich jak silniki

Specyfikacja

- Zasilanie: sinusoida 230 V AC do 10 A
- Częstotliwość pracy: dowolna stała częstotliwość pomiędzy 40 Hz a 70 Hz
- Miękki start: dwie sekundy od uruchomienia do pełnej prędkości
- Prąd bramki triaka: 68 mA
- Impulsy bramki triaka, kąt fazowy <90°: impulsy bramki 40 µs powtarzane w odstępach 200 µs

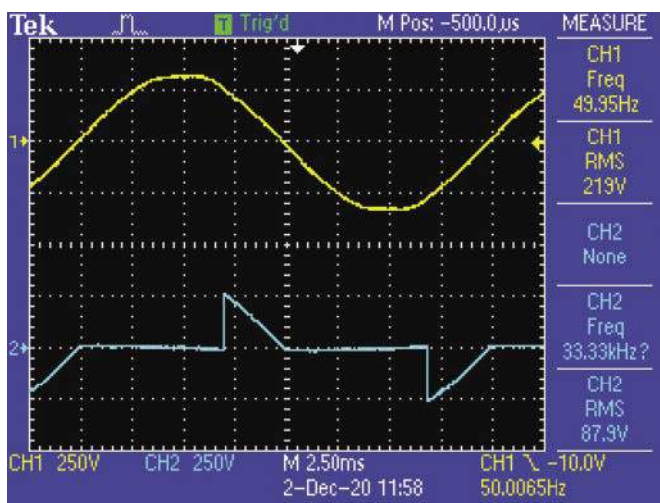
i polerowania, a nawet elektrycznych kosiarek do trawy; są one mniej podatne na zerwanie żyłki, gdy pracują wolniej.

Jakie silniki można regulować?

Regulator ten pasuje do większości elektronarzędzi i urządzeń. Zazwyczaj wykorzystują one silniki uniwersalne, które są szeregowymi silnikami szczotkowymi. Są one nazywane silnikami uniwersalnymi, ponieważ mogą działać zarówno na prąd przemienny, jak i stały.

Nie można regulować prędkości obrotowej żadnego uniwersalnego silnika, który ma już wbudowaną elektroniczną regulację prędkości, niezależnie od tego, czy jest częścią mechanizmu uruchamiania, czy oddzielnym pokrętkiem prędkości.

Nie dotyczy to narzędzi takich jak wiertarki elektryczne, które mają dwupozycyjny mechaniczny przełącznik prędkości. W takim przypadku można użyć naszego regulatora prędkości z mechanicznym



Oscylogram 2. Wyzwalając triak później w każdym półcyklu sieciowym, napięcie wyjściowe (cyjan) jest zerowe przez większość czasu, a moc na obciążeniu jest znacznie zmniejszona. Spowoduje to, że podłączony silnik będzie obracał się dość wolno, ponieważ średnie przyłożone napięcie będzie niskie

przełącznikiem ustawionym na szybkie lub wolne obroty. Wybór niższej prędkości zazwyczaj zasilają silnik napięciem półfazowym.

Silniki indukcyjne (z wyjątkiem typów ze zwartą fazą rozruchową, które często występują w wentylatorach itp.) w oczywisty sposób nie mogą być używane z tym regulatorem. Jak upewnić się, że dane elektronarzędzie lub urządzenie napędzane jest silnikiem uniwersalnym, a nie indukcyjnym? Jedną z wskazówek jest to, że większość silników uniwersalnych jest dość głośna w porównaniu do silników indukcyjnych. Jest to jednak tylko wskazówka i z pewnością nie jest niezawodna.

W wielu elektronarzędziach widać, że silnik ma szczotki i komutator (zwykle przez otwory wentylacyjne), a podczas pracy widać iskry ze szczotek. Oznacza to, że silnik jest typu uniwersalnego. Jeśli jednak nie widać szczotek, wskazówką może być tabliczka znamionowa lub instrukcja obsługi.

Większość silników indukcyjnych stosowanych w urządzeniach gospodarstwa domowego to silniki 2- lub 4-biegunowe, które pracują ze stałą prędkością, zazwyczaj 2850 obr/min dla jednostki 2-biegunowej lub 1440 obr/min dla jednostki 4-biegunowej. Prędkość obrotowa będzie podana na tabliczce znamionowej.

Szlifiarki stołowe zazwyczaj wykorzystują dwubiegunowe silniki indukcyjne.

Jeśli konieczna jest regulacja prędkości obrotowej tego typu silnika, należy użyć sterownika silnika indukcyjnego o mocy 1,5 kW opisanego w Silicon Chip-ie w kwietniu i maju 2012 r. (siliconchip.com.au/Series/25), z ważnymi modyfikacjami w wydaniu z grudnia 2012 r.

Sterowanie fazowe

Napięcie sieciowe AC jest zbliżone do sinusoidy. Zaczyna się od 0 V, wzrasta do wartości szczytowej, spada do 0 V, a następnie robi to samo w przeciwnym kierunku. Powtarza się to 50 razy na sekundę w przypadku sieci 50 Hz lub 60 razy na sekundę w przypadku sieci 60 Hz.

Silnik podłączony do sieci zasilającej w pełni wykorzystuje energię każdego cyklu, dzięki czemu pracuje z maksymalną prędkością. Jeśli więc do silnika zostanie dostarczona tylko część tej fali sinusoidalnej, przy mniejszej ilości energii dostępnej do jego zasilania, silnik nie będzie obracał się tak szybko.

Podczas każdego półcyklu zmiana czasu, w którym napięcie jest podawane do silnika, umożliwia sterowanie prędkością. Jest to podstawa sterowania fazowego: rozpocznij zasilanie bardzo wcześnie w cyklu, a silnik będzie pracował szybko; opóźnij zasilanie do znacznie późniejszej części cyklu, a silnik będzie pracował wolniej.

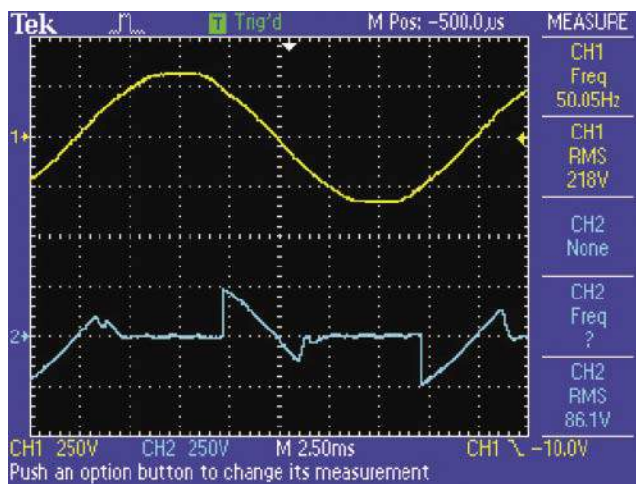
Termin „sterowanie fazowe” powstał, ponieważ czas impulsów wyzwalających jest zmieniany w odniesieniu do fazy sinusoidy sieciowej. Do przełączania napięcia sieciowego można użyć kilku podzespołów; tutaj używamy triaka. Urządzenie to może być używane do przełączania zarówno dodatnich, jak i ujemnych półokresów napięcia sieciowego.

Pokazaliśmy oscylogramy sterowania fazowego zmieniającego moc świecenia żarówki, ponieważ przedstawia to sterowanie fazowe w czystej postaci, bez dodatkowego szumu spowodowanego zasilaniem silnika.

Oscylogram 1 pokazuje pocięty przebieg z obwodu sterowania fazowego, gdy żarówka jest zasilana z dużą jasnością. Jest to odpowiednik zasilania silnika z dużą prędkością. Tutaj triak jest wyzwalany 2,5 ms po przejściu przez zero (punkt, w którym przebieg sieciowy przechodzi przez 0 V).

Napięcie przyłożone do obciążenia ma kolor cyjan i wynosi 200 V RMS. Jest to mniej niż 219 V RMS napięcia sieciowego pokazanego w postaci przebiegu żółtego.

Oscylogram 2 pokazuje przebieg sterowania fazowego zasilającego żarówkę przy niższym ustawieniu, z triakiem wyzwalanym



Oscylogram 3. To samo ustawienie prędkości, co pokazane na Oscylogramie 2, ale tym razem z podłączonym silnikiem. Indukcyjność uzwojeń silnika powoduje, że triak wyłącza się z niewielkim opóźnieniem w stosunku do przejścia napięcia przez zero, z powodu przesunięcia fazy prądu wyjściowego wynikającego z reaktancji silnika

później w cyklu. Napięcie przyłożone do obciążenia jest teraz znacznie niższe i wynosi 87,9 V RMS.

Oscylogram 3 pokazuje przebiegi podczas zasilania silnika. Dolny niebieski przebieg to napięcie przyłożone do silnika, a wejściowe napięcie sieci jest pokazane na górnym (żółtym) przebiegu. Zwróć uwagę na dodatkowe impulsy i opóźnienie wyłączenia silnika na dolnym przebiegu, ponieważ silnik jest obciążeniem indukcyjnym.

Sterowanie prędkością

Aby silnik miał dobrą wydajność przy niskich prędkościach, sterownik musi kompensować spadek prędkości obrotowej silnika wraz ze wzrostem obciążenia.

Wiele fazowych regulatorów prędkości wykorzystuje fakt, że silnik może być uznany za generator, gdy obraca się bez zasilania. Gdy silnik jest obciążony, a prędkość obrotowa silnika spada, siła elektromotoryczna (back-EMF) wytwarzana przez silnik spada, a obwód kompensuje to, dostarczając wyższe napięcie zasilające do silnika, wyzwalając triak wcześniej w cyklu sieciowym.

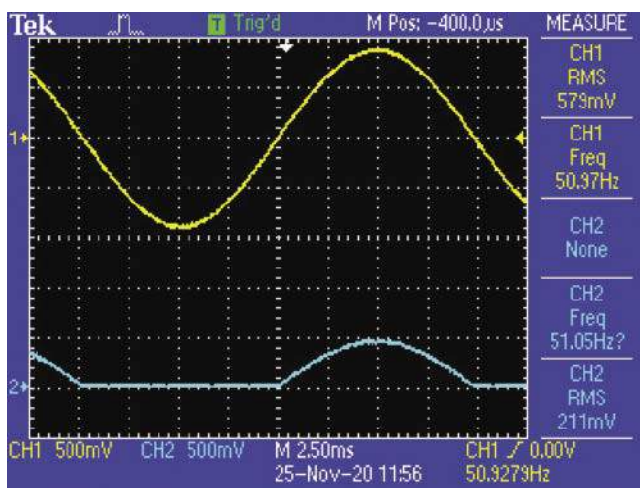
Jednak w praktyce siła elektromagnetyczna generowana przez większość silników szeregowych, gdy triak nie przewodzi, jest albo bardzo niska, albo w ogóle nie występuje. Dzieje się tak częściowo dlatego, że nie ma prądu zasilania elektromagnesów (prądu stojana), a generowanie napięcia „back-EMF” wynika jedynie ze szczątkowego magnetyzmu w stojanie silnika. Jeśli wytwarzane jest jakiekolwiek wsteczne napięcie EMF, występuje ono zbyt późno po zakończeniu każdego półcyklu, aby mieć znaczący wpływ na obwód wyzwalający w następnym półcyklu.

Używamy więc innej metody regulacji prędkości, monitorując prąd przepływający przez silnik. Gdy silnik jest nieobciążony, pobiera pewien prąd, aby utrzymać się w ruchu.

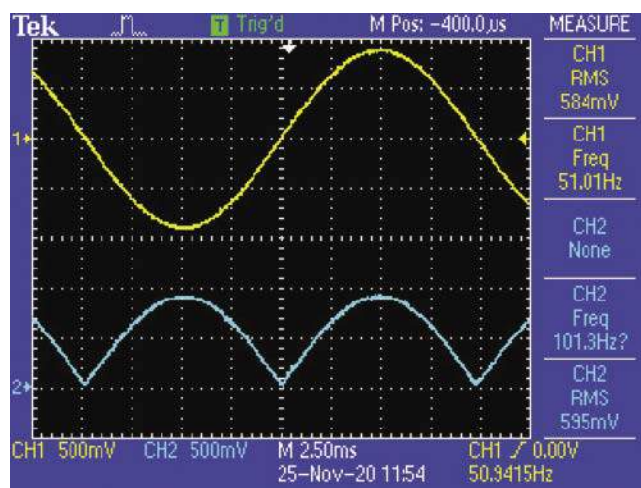
Gdy silnik jest obciążony, jego prędkość spada, a pobór prądu wzrasta. Sterownik silnika wykrywa to i kompensuje spadek prędkości poprzez zwiększenie napięcia zasilającego silnik.

Może to brzmieć jak dodatnie sprzężenie zwrotne, w którym wykrycie większego poboru prądu zwiększy napięcie, a tym samym pozwoli silnikowi pobierać więcej prądu. Prawdą jest, że może się tak zdarzyć, jeśli współczynnik kompensacji jest zbyt wysoki, dlatego dołączamy pokrętkę regulacji sprzężenia zwrotnego, aby dostosować wzmocnienie kompensacji.

Przy odpowiednim ustawieniu regulacja prędkości jest wyjątkowo zadowalająca, ale zbyt duże sprzężenie zwrotne spowoduje, że silnik



Oscylogram 4 Pierwszy stopień precyzyjnego prostownika pełnookresowego działa jako prostownik półokresowy z napięciem wyjściowym o połowę niższym od napięcia wejściowego. Oba sygnały (oryginalny i obcięty/zredukowany) są podawane do drugiego stopnia i sumowane w celu wytworzenia sygnału wyjściowego pokazanego na Oscylogramie 5



Oscylogram 5 Końcowy przebieg wyjściowy precyzyjnego prostownika pełnookresowego ma kolor cyjan. Jest on identyczny z żółtym przebiegiem, z wyjątkiem tego, że ujemne półokresy stały się dodatnimi napięciami, dzięki czemu można je podawać do przetwornika ADC typu single-ended w celu pomiaru

będzie zwiększał prędkość przy zwiększonym obciążeniu zamiast utrzymywać ustawioną prędkość.

Sterowanie triakiem przy obciążeniu indukcyjnym

Jednym z głównych problemów podczas używania triaka do pełnookresowego sterowania silnikiem jest sposób wyłączania triaka i charakter obciążenia silnika. Triak jest zwykle włączany poprzez przyłożenie prądu do jego bramki. Jeśli prąd płynący między głównymi zaciskami triaka jest większy niż jego prąd podtrzymywania, triak pozostanie włączony przez pozostałą część cyklu sieciowego.

Triak wyłącza się tylko wtedy, gdy bramka nie jest zasilana, a prąd triaka spada poniżej jego prądu podtrzymywania. Ponieważ silnik nie jest obciążeniem czysto rezystancyjnym, ale ma znaczną indukcyjność, prąd silnika jest opóźniony w fazie w stosunku do napięcia. Oznacza to, że triak zasilający silnik niekoniecznie wyłączy się przy przejściu napięcia przez zero; prąd silnika może nadal płynąć w późniejszym okresie.

Nasz obwód wykorzystuje mikroprocesor do generowania wymaganych impulsów bramki w celu prawidłowego zasilania obciążenia indukcyjnego, takiego jak silnik, za pomocą triaka. Regulator przekazuje serię impulsów do bramki triaka w celu zapewnienia pełnego zakresu sterowania fazowego.

Opis obwodu

Schemat ideowy sterownika prędkości obrotowej silnika pokazano na **rysunku 1**. Jego kluczowymi komponentami są triak Q1 i mikroprocesor PIC12F617 IC1.

Układ IC1 monitoruje na wejściu analogowym AN1 (styl 6) ustawienie potencjometru prędkości VR1 oraz na wejściu AN0 (styl 7) ustawienie potencjometru wzmacnienia sprzężenia zwrotnego VR2. Monitoruje również prąd silnika na wejściu analogowym AN3 (styl 3), przy czym sygnał ten pochodzi z przekładnika prądowego T1, po przejściu przez prostownik pełnookresowy wykorzystujący układ scalony IC2. Przebieg napięcia sieciowego jest monitorowany pod kątem przejścia przez zero na wejściu 5, po redukcji przez rezystor 330 kΩ.

W odpowiedzi na wszystkie te parametry, układ IC1 wytwarza serię impulsów na swoim wyjściu cyfrowym GP5 (styl 2), które zasilają bazę tranzystora NPN Q2, który z kolei pobiera prąd z bramki triaka Q1. Prąd bramki triaka przepływa przez rezystor 47 Ω podłączony między

zasilaniem 5,1 V a zaciskiem A1 triaka, a następnie przez bramkę i do masy obwodu przez Q2 (tj. prąd bramki jest ujemny).

Ta metoda połączenia umieszcza rezystor 47 Ω między zasilaniem sieciowym 230 V AC a zasilaniem 5,1 V, które zasila mikroprocesor PIC. Zapobiega to przedostawaniu się szumu przełączania triaka do zasilania 5,1 V, co mogłoby spowodować zatrzaśnięcie mikroprocesora.

Tłumik

Obwód tłumiący składa się z dwóch szeregowo połączonych rezystorów 220 Ω 1 W i kondensatora 220 nF 275 V AC X2, podłączonych między zaciskami A1 i A2 triaka. Tłumik zapobiega gwałtownym zmianom napięcia przyłożonego do triaka Q1, co w przeciwnym razie spowodowałoby jego włączenie (z powodu przekroczenia szybkości zmian napięcia przełączania dV/dt), gdy powinien być wyłączony.

Te gwałtowne zmiany napięcia mogą wystąpić po pierwszym podłączeniu zasilania lub mogą pochodzić ze stanów nieustalonych napięcia generowanych przez indukcyjność sterowanego silnika za każdym razem, gdy triak się wyłącza. Wymienione podzespoły tłumią stany nieustalone i zmniejsza ich amplitudę.

Zasilanie DC dla mikroprocesora pobierane jest bezpośrednio z sieci 230 V AC poprzez kondensator 470 nF/275 V AC X2 połączony szeregowo z rezystorem 1 kΩ/5 W. Impedancja kondensatora (~7 kΩ) ogranicza średni prąd pobierany z sieci, podczas gdy rezystor 1 kΩ ogranicza prąd udarowy przy pierwszym podłączeniu zasilania.

Gdy linia aktywna jest ujemna w stosunku do neutralnej, prąd płynie przez kondensator 470 nF, diodę D1 i rezystor 47 Ω do kondensatora 1000 μF, aby go naładować. Przy dodatnim półcyklu, kierunek przepływu prądu przez kondensator 470 nF jest odwrócony i przepływa on przez diodę D2, rozładowując kondensator 470 nF z powrotem do sieci. Inaczej mówiąc, kondensator zasilający jest ładowany tylko przez co drugą połowę sinusoidy sieci zasilającej, natomiast w drugiej połowie prąd płynie do masy (wg umownej konwencji z masy do przewodu neutralnego).

Dioda Zenera ZD1 ogranicza napięcie na kondensatorze 1000 μF do poziomu 5,1 V. Jest to napięcie zasilania mikroprocesora IC1, wzmacniaczy operacyjnych IC2a i IC2b oraz źródło prądu bramki triaka Q1. Zasilanie 5,1 V IC1 jest zbocznikowane kondensatorem 100 nF, podczas gdy zasilanie IC2 jest zbocznikowane kondensatorem 100 μF.

Przełącznik S1 umożliwia włączenie lub wyłączenie funkcji łagodnego rozruchu. Przełącznik ten steruje poziomem wejścia GP3 (styk 4). Gdy S1 jest otwarty, wejście GP3 jest utrzymywane w stanie wysokim na poziomie 5,1 V przez rezystor 47 kΩ, więc miękki start jest wyłączony. Gdy przełącznik S1 jest zamknięty, wejście GP3 jest spolaryzowane do stanu niskiego, a program uruchamia procedurę miękkiego startu.

S1 ustawia wejście GP3 na niskim poziomie poprzez rezystor 100 Ω, który jest dołączony w celu ochrony wejścia przed stanami nieustalonymi prądu, które mogłyby spowodować zatrzaśnięcie w układzie scalonym. Kondensator 100 nF zapewnia niską impedancję dla stanów nieustalonych, zapobiegając nieprawidłowemu wykrywaniu stanu wejścia GP3, gdy S1 jest otwarte, z powodu stanów nieustalonych lub zakłóceń.

Potencjometr regulacji prędkości obrotowej VR1 jest podłączony do zasilania 5,1 V. IC1 konwertuje napięcie z wyjścia potencjometru VR1 na wartość cyfrową za pomocą wewnętrznego przetwornika analogowo-cyfrowego (ADC). Rezystor 100 kΩ od ślizgacza do masy utrzymuje wejście AN1 na poziomie 0 V, ustawiając prędkość silnika na zero, jeśli ślizgacz VR1 straci kontakt ze ścieżką oporową.

Potencjometr VR2 jest podłączony podobnie. Jego napięcie na ślizgaczu ustawia wzmocnienie sprzężenia zwrotnego w celu utrzymania prędkości obrotowej silnika pod obciążeniem. Jest ono również konwertowane na wartość cyfrową w układzie IC1. Kondensatory na wyjściach potencjometrów VR1 i VR2 zapewniają niską impedancję źródła dla przetwornika ADC IC1 i odfiltrowują tętnienia zasilania.

Zarówno VR1, jak i VR2 są podłączone do IC1 za pomocą złączy śrubowych. CON2 zapewnia wspólne połączenia +5 V i 0 V dla VR1 i VR2, podczas gdy ślizgacz VR1 również łączy się z CON2. CON3 zapewnia połączenie ślizgacza dla VR2, a przełącznik S1 wykorzystuje pozostałe dwa połączenia w CON3.

Synchronizacja sieci

Moment przesłania impulsów wyzwalających bramkę triaka jest krytyczny dla jego prawidłowego działania. Układ IC1 monitoruje napięcie sieciowe na swoim wejściu 5, z rezystorem 330 kΩ łączącym się z przewodem neutralnym plus kondensatorem filtra dolnoprzepustowego 4,7 nF. Procedura przerwania jest wyzwalana w IC1 za każdym razem, gdy poziom napięcia na wejściu 5 zmienia się z wysokiego na niski lub odwrotnie. Przerwanie informuje układ IC1, że napięcie sieciowe właśnie przekroczyło 0 V, dzięki czemu może on zsynchronizować wyzwalanie bramki z przebiegiem sinusoidy sieci.

Opóźnienie fazowe wprowadzane przez kondensator 4,7 nF jest kompensowane przez oprogramowanie IC1, podobnie jak asymetria wyzwalania wynikająca z różnicy 5 V między poziomem niskim i wysokim.

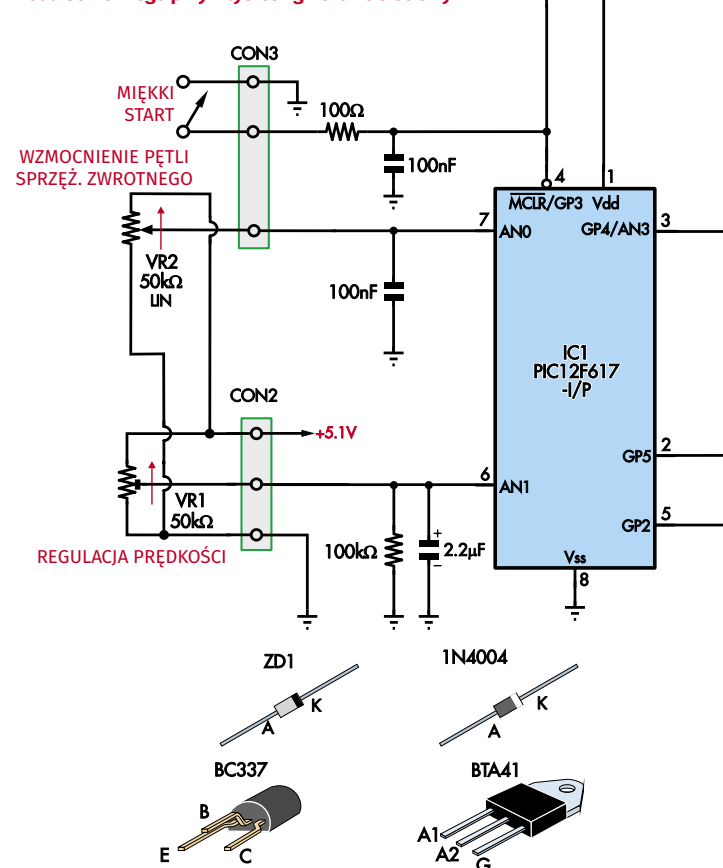
Bieżące informacje zwrotne

T1 to przekładnik prądowy składający się z ferrytowego toroidu z dwuzwojowym uzwojeniem pierwotnym połączonym szeregowo z triakiem. Uzwojenie wtórne ma 1000 zwojów i jest obciążone rezystorem 510 Ω.

Przy takim obciążeniu transformator wytwarza na wyjściu 800 mV na amper prądu obciążenia. Jest to napięcie proporcjonalne do prądu płynącego przez sterowany silnik.

Napięcie wyjściowe z transformatora T1 jest podawany na precyzyjny (od Red. EdW: precyzyjny, to znaczy taki prostownik, w którym nie występują straty napięcia na zwykłych diodach prostowniczych) prostownik pełnookresowy składający się z układów IC2a i IC2b. Ta konfiguracja jest nietypowa, ponieważ nie wykorzystuje żadnych diod. Większość precyzyjnych prostowników z diodami wymaga ujemnego zasilania

UWAGA: pod żadnym pozorem proszę nie pomylić dwóch podobnych symboli: symbolu masy w układzie elektrycznym oznaczającego po prostu wspólną szynę układu; z symbolem podłączenia do obudowy przewodu ochronnego przy wtyczce i gnieździe sieciowym



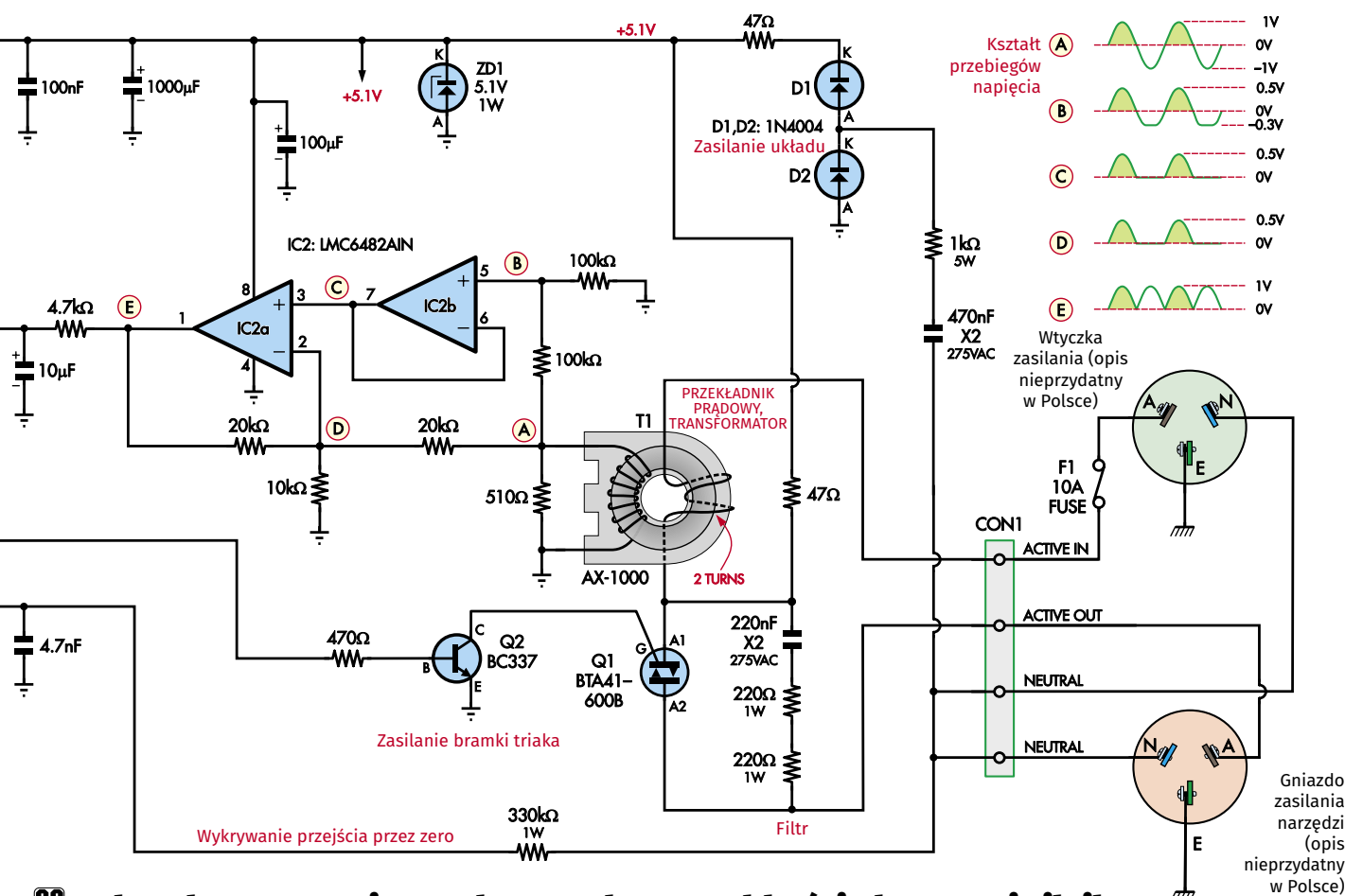
dla wzmacniaczy operacyjnych. Chociaż moglibyśmy dodać ujemne zasilanie, zwiększyłoby to złożoność obwodu i koszty.

Prostownik pełnookresowy wykorzystuje wzmacniacze operacyjne, które mają określone właściwości. Pierwszą z nich jest to, że napięcie na wyjściu wzmacniacza operacyjnego musi w pełni opadać aż do poziomu ujemnej szyny zasilania (tj. tutaj aż do 0 V). Ponadto napięcie wyjściowe 0 V musi być utrzymywane, gdy napięcie wejściowe wzmacniacza operacyjnego spadnie poniżej 0 V. Wzmacniacz operacyjny LMC6482 „rail-to-rail” (IC2a i IC2b w obwodzie) ma te cechy, a także niski prąd zasilania.

Oznaczyliśmy kilka punktów w obwodzie i pokazaliśmy oczekiwane przebiegi, aby pomóc wyjaśnić, jak działa ta sekcja. Sygnał z uzwojenia wtórnego transformatora pojawia się w punkcie A. Sygnał ten jest sinusoidą i waha się powyżej i poniżej 0 V, jak pokazano. Sygnał płynie stąd dwiema ścieżkami. Jedna prowadzi przez rezystor 20 kΩ do punktu D, a druga przez dwa szeregowo połączone rezystory 100 kΩ do 0 V.

Układ IC2b jest podłączony jako bufor o wzmocnieniu jednostkowym. Wewnętrzna dioda wzmacniacza operacyjnego zatrzaśkuje każde napięcie poniżej -0,3 V (w stosunku do potencjału masy) na wejściu nieodwracającym (styk 5). Jego wyjście (styk 7) będzie na poziomie 0 V, gdy jego wejście będzie na poziomie 0 V lub niższym.

Działanie tej części obwodu najlepiej wyjaśnić, opisując przepływ prądu osobno dla ujemnych i dodatnich połówek sinusoidy napięcia na wyjściu T1.



Pełnookresowy uniwersalny regulator szybkości obrotowej silnika

Rysunek 1. Regulator prędkości obrotowej silnika wykorzystuje transformator czujnika prądu T1 i wzmacniacze operacyjne IC2a i IC2b (działające jako precyzyjny prostownik pełnookresowy) do wykrywania prądu silnika. IC1 dostosowuje impulsy bramki z wyjścia na styku 2 kierowane do bramki triaka Q1, aby utrzymać mniej więcej stałą prędkość silnika pod obciążeniem

Ujemna część sinusoidy

Gdy napięcie w punkcie A jest ujemne, napięcie w punkcie B jest zatrzaśnięte na poziomie $-0,3\text{ V}$ przez wewnętrzną diodę zabezpieczającą na wejściu (styk 5) układu IC2b. Wyjście IC2b na styku 7 (punkt C) ma zatem potencjał 0 V , podobnie jak nieodwracające wejście IC2a na styku 3.

W rezultacie IC2a działa jako wzmacniacz odwracający ze wzmocnieniem -1 . Ustala to rezystor wejściowy $20\text{ k}\Omega$ i rezystor sprzężenia zwrotnego $20\text{ k}\Omega$ z wyjścia na styku 1 do wejścia odwracającego na styku 2.

Tak więc IC2a wytworzy dodatnie napięcie na wyjściu (styk 1), proporcjonalne do ujemnego napięcia w punkcie A.

Aby zrozumieć, jak to działa, należy wziąć pod uwagę, że wzmacniacz operacyjny działa tak, aby napięcia na jego wejściach były równe. Ponieważ wejście nieodwracające jest utrzymywane na poziomie 0 V , a rezystory o równej wartości w ścieżce sprzężenia zwrotnego tworzą dzielnik $1:1$, napięcie wyjściowe (E) musi mieć równą wielkość i przeciwną polaryzację w porównaniu do napięcia wejściowego (A), aby napięcie wejściowe odwracające (D) było na poziomie 0 V .

Na przykład, gdy punkt A jest na potencjale -1 V , punkt E będzie na potencjale $+1\text{ V}$, więc punkt D będzie na potencjale 0 V , równym C. Należy zauważyć, że rezystor $10\text{ k}\Omega$ w punkcie D nie ma w tym przypadku żadnego wpływu, ponieważ styk 2 jest na potencjale 0 V , a zatem nie ma napięcia na tym rezystorze. Pełni on funkcję tylko podczas dodatnich połówek sygnału.

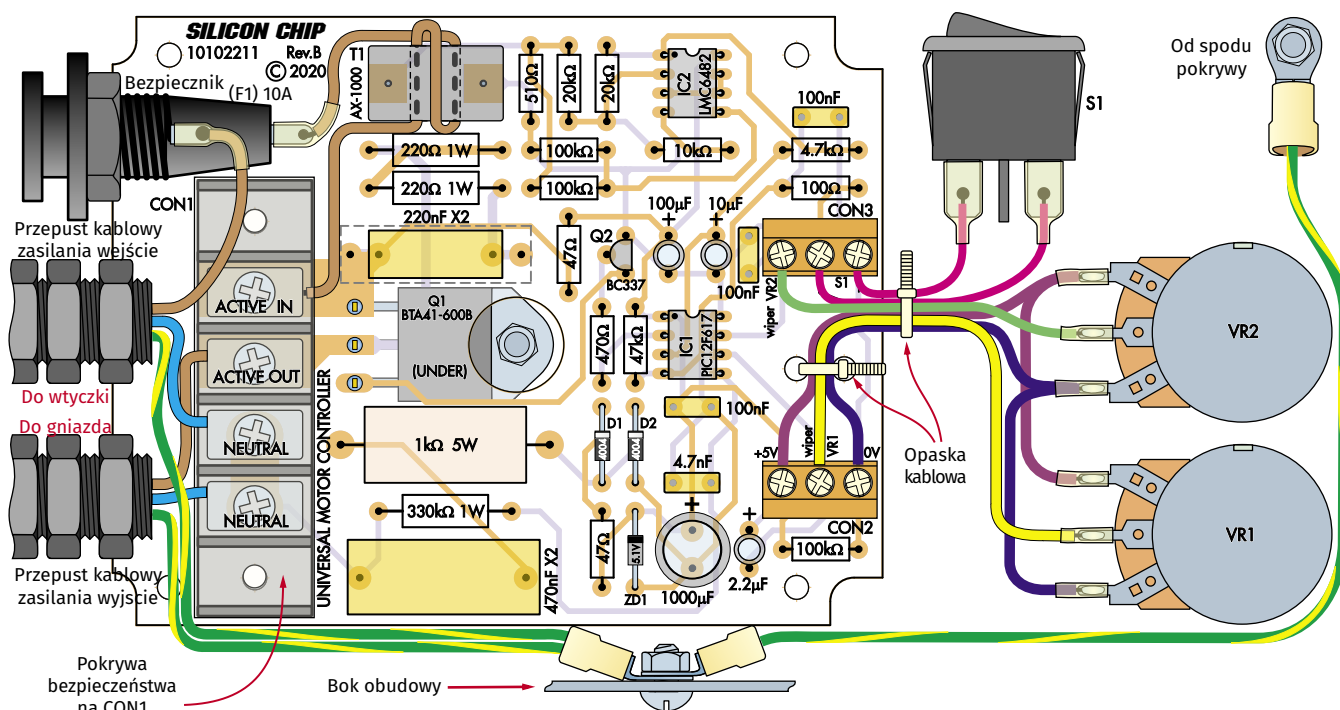
Dodatnia część sinusoidy

W przypadku dodatnich napięć w punkcie A, napięcie w punkcie B będzie o połowę niższe od napięcia w punkcie A ze względu na dzielnik rezystancyjny $100\text{ k}\Omega/100\text{ k}\Omega$. Punkt C i wejście nieodwracające układu IC2a na styku 3 również będą miały połowę napięcia wygenerowanego w punkcie A, ponieważ IC2b działa jako bufor.

Należy pamiętać, że zwykle napięcie na wejściu odwracającym będzie takie samo jak na wejściu nieodwracającym. Wzmacniacz operacyjny zapewni to, dostosowując swoje wyjście tak, aby mógł utrzymać to napięcie za pośrednictwem rezystora sprzężenia zwrotnego.

Jedynym sposobem, w jaki może się to zdarzyć w tym przypadku dla IC2a, jest sytuacja, w której wyjście wzmacniacza operacyjnego w punkcie E ma taki sam potencjał jak wejście sygnału w punkcie A. W takim przypadku to samo napięcie jest przykładane do obu rezystorów $20\text{ k}\Omega$ i są one zasadniczo równoległe, tworząc równoważny rezystor $10\text{ k}\Omega$ do punktu D. Tworzy to dzielnik $1:1$ z rezystorem $10\text{ k}\Omega$ od punktu D do masy, zmniejszając napięcie w tym punkcie o połowę w porównaniu do punktów A i E.

Podsumowując. Układ IC2a zapewnia takie samo dodatnie napięcie na wyjściu E, jak na wejściu A podczas dodatnich połówek sinusoidy generowanej w przekładniku T1. Podczas ujemnych połówek, IC2a odwraca napięcie. Tak więc wyjście IC2a jest dodatnie zarówno dla ujemnych, jak i dodatnich napięć w punkcie A i równe co do amplitudy. Mamy zatem prostownik pełnookresowy.



Rysunek 2. Większość komponentów jest zamontowana na górze płytki, z wyjątkiem triaka Q1. Jest on montowany po wewnętrznej stronie obudowy, pod płytką drukowaną. Po wykonaniu okablowania należy je dokładnie sprawdzić zgodnie z tym schematem. Śruby i końcówki oczkowe przewodów ochronnych muszą się dobrze kontaktować, a przewody sterujące należy związać opaskami kablowymi, jak pokazano na rysunku

Jego wyjście jest filtrowane dolnoprzepustowo za pomocą rezystora 4,7 kΩ i kondensatora 10 μF w celu uzyskania stabilnego napięcia DC na wyjściu, które jest następnie podawane, gotowe do digitalizacji, na wejście analogowe AN3 układu IC1, styk 3,.

Oscylogram 4 pokazuje sygnał sinusoidalny w punkcie A (na żółto), a dolny niebieski przebieg pokazuje kształt napięcia w punkcie C, dodatni przebieg wyjściowy ma połowę amplitudy przebiegu wejściowego. Gdy przebieg A spada poniżej potencjału 0 V, przebieg C pozostaje na poziomie 0 V.

Oscylogram 5 pokazuje ten sam sygnał sinusoidalny w punkcie A na żółto, a wyjście wyprostowane pełnookresowo w punkcie E na dolnym niebieskim przebiegu.

Budowa

Większość komponentów uniwersalnego sterownika prędkości obrotowej silnika Full Wave jest zamontowana na dwustronnej, galwanizowanej płytce PCB (płytką drukowaną) o kodzie 10102211 i wymiarach 103×81 mm.

Jest ona zamontowana wewnątrz odlewanej ciśnieniowej obudowy o wymiarach 119×94×34 mm.

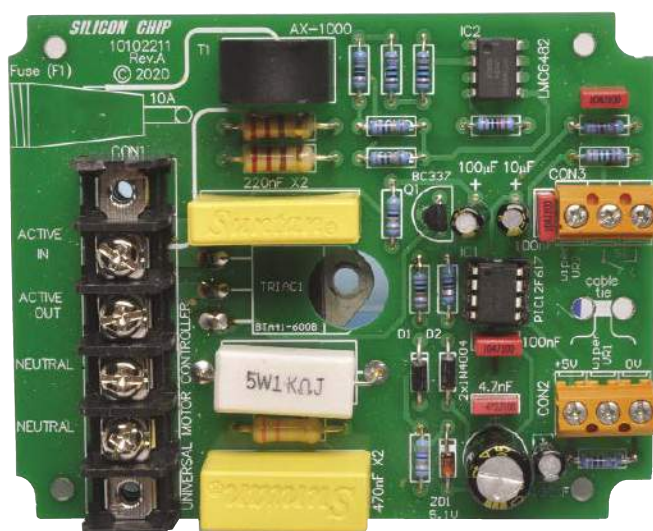
Postępuj zgodnie ze schematem montażowym PCB pokazanym na **rysunku 2**. Rozpocznij od instalacji rezystorów z wyjątkiem typu o mocy 5 W. Kolorowe kody rezystorów są znormalizowane, ale należy również dwukrotnie sprawdzić każdy rezystor za pomocą multimetru cyfrowego. Następnie należy zamontować diody, które muszą być ustawione zgodnie z ilustracją. Dostępne są dwa różne typy diod: 1N4004 dla D1 i D2 oraz dioda Zenera ZD1 typu 5,1 V 1 W (1N4733).

Układ IC1 jest zamontowany na 8-stykowej podstawce DIL-8, więc wlotuj to gniazdo teraz, zwracając uwagę na prawidłową orientację, z wycięciem skierowanym w stronę górnej części płytki drukowanej. Na razie pozostaw IC1 na zewnątrz; wcisniemy go do podstawki później. Układ IC2 można zainstalować na podstawie lub bezpośrednio na płytce drukowanej. Dodatkowo, teraz może być wlotowany tranzystor Q2.

Następnie należy umieścić kondensatory. Kondensatory typu MKT i polipropylenowe mają zwykle nadrukowany opis wskazujący ich wartość. Są one pokazane na liście części.

Natomiast kondensatory elektrolityczne są prawie zawsze oznaczone wartością w μF wraz z biegunowością. Zazwyczaj wyprowadzenie ujemne jest oznaczone paskiem. Kondensatory należy wkładać zgodnie z podaną biegunowością.

Zaciski śrubowe są następne w kolejce. 3-drożne bloki zacisków dla CON2 i CON3 są instalowane z wejściami przewodów skierowanymi do siebie, podczas gdy CON1 nie ma określonej orientacji. Następnie należy zamontować rezystor 5 W około 1 mm nad płytką drukowaną, aby zapewnić lepsze chłodzenie.



Zbliżenie płytki drukowanej regulatora prędkości obrotowej silnika w tym samym rozmiarze. Ponieważ jest to urządzenie zasilane i sterowane z sieci, konstrukcja musi być wzorowa. Nie podejmuj się tego projektu, jeśli nie masz doświadczenia z urządzeniami sieciowymi

Wykaz elementów:

1 dwustronna płytką drukowaną o kodzie 10102211 i wymiarach 103×81 mm
1 odlewana obudowa, 119×94×34 mm [Jaycar HB5067].
2 potencjometry obrotowe liniowe 50 kΩ 24 mm (VR1, VR2)
2 plastikowe pokrętki pasujące do VR1 i VR2
1 mini przełącznik wychyłny SPST (S1) [Jaycar SK0984 lub Altronics S3210].
1 przekładnik prądowy TALEMA AX-1000 10 A (T1) [RS Components 775-4928].
1 uchwyt bezpiecznika M20×5 10 A do montażu panelowego w obudowie (F1) [Altronics S5992].
1 bezpiecznik szybki M20×5 10 A
1 4-stykowy zacisk śrubowy o dużej obciążalności do montażu na płytce drukowanej (CON1) [Jaycar HM-3162].
2 3-stykowe zaciski śrubowe do montażu na płytce drukowanej, raster 5,08 mm (CON2, CON3)
2 przepusty kablowe GP9 dla kabli o średnicy 4-8 mm
1 8-stykowa podstawa IC DIL-8 (dla IC1)
1 przedłużacz sieciowy 10 A o długości 2 m
3 złącza oczkowe 4 mm
4 poliamidowe kołki dystansowe z gwintem M3 o długości 6,3 mm
3 śruby M4×10 z łbem walcowym lub stożkowym (do montażu Q1; złącza ochronne)
3 podkładki gwiazdkowe o średnicy otworu 4 mm
3 nakrętki M4
8 wkrętów M3×5 z łbem walcowym lub stożkowym
4 przyklejane gumowe nóżki
1 rurka termokurczliwa o długości 20 mm i średnicy 12 mm
1 rurka termokurczliwa o długości 80 mm i średnicy 3 mm
1 przewód sieciowy o długości 600 mm i natężeniu prądu do 7,5 A (dla VR1, VR2 i S1)
4 opaski kablowe o długości 100 mm

Półprzewodniki;

1 8-bitowy mikroprocesor PIC12F617-I/P zaprogramowany kodem 1010221A.hex, DIP-8 (IC1)
1 podwójny wzmacniacz operacyjny CMOS „rail-to-rail” LMC6482AIN, DIP-8 (IC2)
1 izolowany triak BTA41-600B 40 A 600 V TOP3 (Q1)
1 tranzystor NPN BC337 500 mA, TO-92 (Q2)
1 dioda Zenera 5,1 V 1 W (1N4733) (ZD1)
2 diody 1N4004 400 V 1 A (D1, D2)

Kondensatory:

1 kondensator elektrolityczny 1000 µF 16 V PC
1 kondensator elektrolityczny 100 µF 16 V PC
1 kondensator elektrolityczny 10 µF 16 V PC
1 kondensator elektrolityczny 2,2 µF 16 V (lub wyższe) PC
1 kondensator foliowy 470 nF 275 V AC metalizowany polipropylenowy klasy X2
1 kondensator foliowy 220 nF 275 V AC metalizowany polipropylenowy klasy X2
3 kondensatory 100 nF 63/100 V MKT poliestrowe
1 kondensator 4,7 nF 63/100 V MKT poliestrowy

Rezystory: (wszystkie 0,25 W, 1%, o ile nie podano inaczej)

1 szt. 330 kΩ 5% 1 W węglowy	3 szt. 100 kΩ	2 szt. 20 kΩ
1 szt. 1 kΩ 10% 5 W drutowy	1 szt. 10 kΩ	1 szt. 4,7 kΩ
2 szt. 220 Ω 5% 1 W węglowy	1 szt. 70 Ω	1 szt. 100 Ω
		2 szt. 47 Ω

Różne:

Super Glue (klej cyjanoakrylowy) lub Loctite, pasta termoprzewodząca, lutowie

Na koniec (na razie) zainstaluj przekładnik prądowy T1. Nie ma znaczenia, w którą stronę jest zorientowany, bo w PCB wlotujemy tylko wyprowadzenia uzwojenia wtórnego. Triak Q1 zostanie zamontowany później.

Odetnij wystające od spodu PCB zbyt długie dolne końcówki wszystkich komponentów, aby zapobiec zwarcia z podstawą obudowy.

Wiercenie obudowy

Rysunek 4 przedstawia szablon/przewodnik do wiercenia otworów w obudowie. Pokrywa wymaga otworów o średnicy 9,5 mm na potencjometry VR1 i VR2, prostokątnego wycięcia o wymiarach 19×10 mm na przełącznik S1 oraz otworu o średnicy 4 mm na śrubę mocowania przewodu ochronnego.

Płytkę PCB jest montowana w podstawie obudowy za pomocą gwintowanych poliamidowych kołków dystansowych M3 o długości 6,3 mm, które wymagają otworów montażowych.

Użyj płytki drukowanej jako szablonu i zwróć uwagę, że koniec zacisku śrubowego CON1 znajduje się dalej od końca obudowy w porównaniu z drugim końcem. Zapewnia to miejsce

na nakrętki przepustów kablowych. Po umieszczeniu płytki drukowanej na miejscu, zaznacz pozycje otworów, wyjmij ją i wywierć je wiertłem 3 mm.

Przymocuj poliamidowe kołki dystansowe o długości 6,3 mm do płytki drukowanej za pomocą krótkich wkrętów, a następnie wygnij do przodu przewody triaka o 90° 4 mm od jego korpusu. Włóż przewody do płytki drukowanej od spodu (patrz rysunek 3).

Przymocuj płytkę drukowaną do obudowy za pomocą śrub od spodu i zaznacz położenie otworu montażowego triaka w podstawie obudowy, wykorzystując otwór montażowy w PCB. Ponownie zdejmij płytkę PCB i wywierć w dnie obudowy otwór o średnicy 4 mm.

Usuń wszelkie metalowe opiłki i lekko szfajuj krawędzie otworu, a następnie ponownie przymocuj płytkę drukowaną i wyreguluj wysokość wyprowadzeń triaka, tak aby metalowy radiator przylegał do płaskiej powierzchni na dnie.

Przymocuj radiator triaka do obudowy za pomocą śruby M4 i nakrętki. Metalowy radiator jest wewnętrznie odizolowany elektrycznie od wyprowadzeń, więc nie wymaga dodatkowej izolacji między radiatorem a obudową.

Nie stwarza to zagrożenia nawet w przypadku uszkodzenia triaka, ponieważ obie części obudowy są połączone z przewodem ochronnym.

Przylutuj przewody triaka na górze płytki drukowanej i przytnij je. Teraz odkręć śruby, aby uzyskać dostęp do spodu płytki drukowanej i przylutuj przewody triaka również od spodu płytki drukowanej.

To dobry moment, aby przymocować gumowe nóżki do podstawy obudowy.

Przygotowanie panelu

Oprócz wywiercenia otworów w pokrywie, o których mowa powyżej, należy wywiercić otwór o średnicy 4 mm po wewnętrznej stronie pokrywy, wg poniższego opisu.

Jeśli użyjesz śruby z łbem stożkowym i odpowiednio pogłębisz od zewnątrz jej otwór, można ją zamontować pod etykietą panelu, aby uzyskać schludniejszy wygląd. W przeciwnym razie konieczne będzie wycięcie otworu w etykiecie panelu (ostrym nożem modelarskim) po przyklejeniu etykiety, tak jak na redakcyjnej fotografii.

Plik etykiety panelu można pobrać ze strony internetowej Silicon Chip-a i wydrukować. Aby wydrukować etykietę na panel przedni, można skorzystać z kilku opcji. Aby uzyskać bardziej wytrzymałą etykietę, wydrukuj ją jako odbicie lustrzane na przezroczystej folii do rzutnika (używając folii odpowiedniej dla danego typu drukarki).

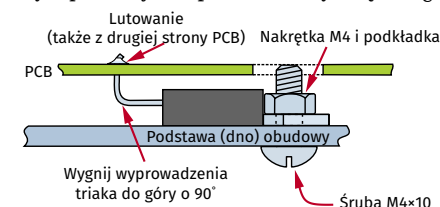
Przymocuj etykietę zadrukowaną stroną do dołu do pokrywy za pomocą jasnego lub przezroczystego uszczelnacza silikonowego.

Alternatywnie można wydrukować na syntetycznej etykiecie samoprzylepnej „Dataflex”, która nadaje się do drukarek atramentowych, lub etykiecie samoprzylepnej „Datapol” do drukarek laserowych. Następnie należy przykleić etykietę za pomocą kleju samoprzylepnego.

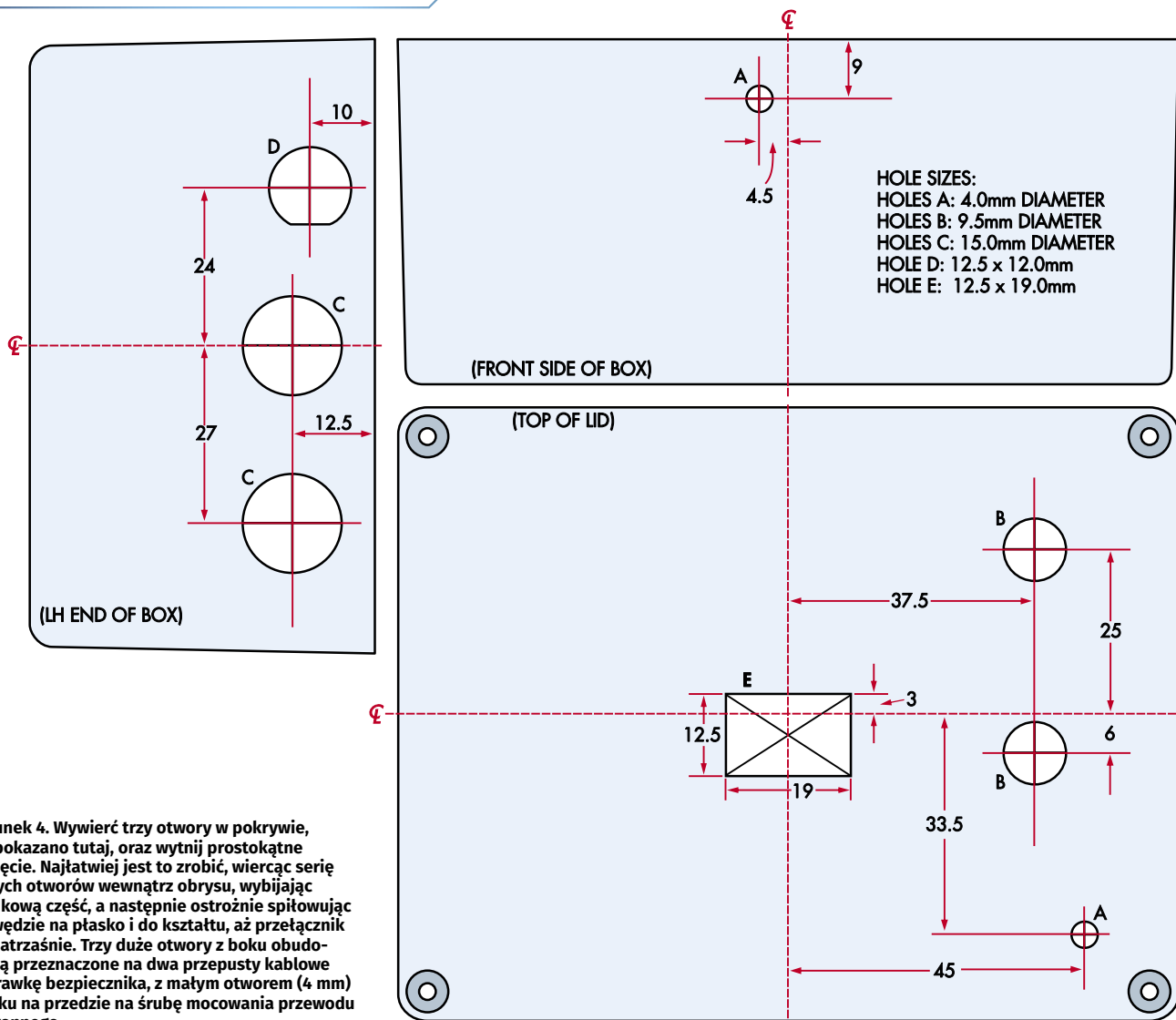
Więcej informacji na temat etykiet Dataflex można znaleźć na stronie siliconchip.com.au/link/aabw, a na temat etykiet Datapol na stronie siliconchip.com.au/link/aabw, natomiast wskazówki dotyczące tworzenia etykiet można znaleźć na stronie siliconchip.com.au/Help/FrontPanels.

Okablowanie

Przetnij przedłużacz 10 A na dwie części, aby zapewnić jeden przewód z wtyczką i drugi



Rysunek 3. Triak Q1 montuje się na dnie obudowy, wykorzystując ją jako radiator. Otwór w płycie drukowanej umożliwia przytrzymanie nakrętki podczas dokręcania śruby



Rysunek 4. Wywierć trzy otwory w pokrywie, jak pokazano tutaj, oraz wytnij prostokątne wycięcie. Najłatwiej jest to zrobić, wierząc serię małych otworów wewnątrz obrysu, wybijając środkową część, a następnie ostrożnie spłuwając krawędzie na płasko i do kształtu, aż przetacznik się zatrzaśnie. Trzy duże otwory z boku obudowy są przeznaczone na dwa przepusty kablowe i oprawkę bezpiecznika, z małym otworem (4 mm) z boku na przodzie na śrubę mocowania przewodu ochronnego

z gniazdem. Miejsce przecięcia przewodu zależy od długości każdej sekcji. Możesz preferować długi przewód z wtyczką i krótki przewód z gniazdem, aby urządzenie znajdowało się w pobliżu regulatora, lub przewód można przeciąć na dwie równe części.

Przed cięciem upewnij się, że masz wystarczający zapas długości, aby zdjąć izolację, jak opisano w następnych dwóch akapitach. Upewnij się, że oba przewody są poprowadzone przez odpowiedni przepust i podłączone, jak pokazano na schemacie okablowania, rysunek 2.

W przypadku przewodu gniazda (wyjściowego) potrzebny jest przewód ochronny o długości 100 mm (zielono-żółty pasek) do połączenia między obudową a pokrywą, więc zdejmij zewnętrzną powłokę izolacyjną na długości około 200 mm. Przytnij przewody aktywne (brązowy) i neutralny (niebieski) na długość około 50 mm, natomiast przewód ochronny na 100 mm, i zachowaj ścinki.

Odcięty zapasowy brązowy przewód o długości 150 mm może zostać użyty później do podłączenia bezpiecznika do CON1 przez transformator T1. Wymaga to dwóch zwojów przewodu aktywnego zapętłonego przez otwór transformatora.

100-milimetrowy przewód ochronny z przewodu gniazda (zielono-żółty pasek) jest poprowadzony wokół krawędzi płytki drukowanej i skręcony razem z przewodem ochronnym z przewodu wtyczki (wejściowego), w celu zaciśnięcia w jednej z końcówek oczkowych zacisku ochronnego.

Zdejmij zewnętrzną izolację przewodu wtyczki, aby odsłonić 100 mm przewodu. Wszystkie trzy przewody należy przełożyć przez dławik kablowy i podłączyć w sposób pokazany na rysunku 2. Przytnij przewód neutralny do 50 mm i zdejmij izolację przed podłączeniem go do listwy zaciskowej. Teraz zamontuj uchwyt bezpiecznika w wykonanym wcześniej otworze i przygotuj się do przylutowania do niego przewodu aktywnego (brązowego), jak pokazano na rysunku.

Ale zanim to zrobisz, nasuń rurkę termokurczliwą o średnicy 10 mm na przewód aktywny (brązowy). Po przylutowaniu tego przewodu, przesunij rurkę w górę i nad oprawkę bezpiecznika, aby zakryć boczny zacisk oprawki bezpiecznika i obkurczyć rurkę.

Podobnie, użyj rurki termokurczliwej o średnicy 3 mm, aby osłonić zacisk końcowy oprawki bezpiecznika po przylutowaniu tego przewodu.

Teraz skręć końce wejściowego przewodu ochronnego (zielony/żółty pasek) i wyjściowego przewodu ochronnego razem i zaciśnij oba w jednej z końcówek oczkowych.

Przytnij pokrętła potencjometrów VR1 i VR2 do długości 12 mm od przodu korpusów i wygładź krawędzie. Następnie przymocuj trzy 100-milimetrowe odcinki przewodu sieciowego 7,5 A do trzech zacisków VR1 oraz czwarty 100-milimetrowy przewód do środkowego zacisku VR2. Użyj krótkich odcinków tego samego przewodu, aby podłączyć dwa końce toru VR2 do tych samych zacisków na VR1. Zabezpiecz wszystkie sześć zacisków rurką termokurczliwą o średnicy 3 mm.



To „otwarte” zdjęcie pasuje do schematu PCB/okablowania na rysunku 2. Oczywiście upewniliśmy się, że regulator nie był podłączony do zasilania sieciowego przed zdjęciem pokrywki!

Następnie podłącz wolne końce tych przewodów do CON2 i CON3, upewniając się, że robisz to tak, jak pokazano na Rysunku 2. W podobny sposób należy teraz podłączyć przełącznik S1. Wystarczy podłączyć go do dwóch pozostałych zacisków w dowolnej kolejności.

Teraz przymocuj wszystkie te przewody do płytki drukowanej za pomocą opaski kablowej, która przechodzi przez otwory w płytce drukowanej. Przymocuj VR1, VR2 i S1 do pokrywki obudowy, zwracając uwagę, że potencjometry muszą być umieszczone tak, jak pokazano (tj. z ich przewodami wychodzącymi z krawędzi obudowy). Dzięki temu zmieszczą się między dwoma kondensatorami sieciowymi typu X2 na płytce drukowanej.

Dodaj również opaski kablowe wokół wiązek przewodów bliżej VR1, VR2 i S2.

Zamontuj teraz pokręta; być może będziesz musiał podnieść nasadki pokręteł nożem modelarskim i zmienić ich orientację tak, aby strzałki pasowały do oznaczeń położenia na panelu pokrywki.

Przewód ochronny o długości 100 mm, który został odcięty od przewodu wyjściowego, można teraz zacisnąć w dwóch końcówkach oczkowych, które są przykręcone do boku obudowy z przodu (razem z przewodami ochronnymi z kabli wejściowego i wyjściowego zaciśniętymi w jednej końcówce oczkowej) i pokrywki obudowy, za pomocą śrub M4, podkładek gwiazdowych i nakrętek. Upewnij się, że nakrętki są całkowicie dokręcone.

Montaż końcowy

Przed zainstalowaniem płytki drukowanej w obudowie należy nałożyć pastę termoprzewodzącą na spód radiatora triaka. Jak wspomniano, radiator triaka jest izolowany, więc może stykać się z obudową.

Ostatnimi komponentami do włożenia są: IC1 (uważając, aby był prawidłowo zorientowany), bezpiecznik 10 A w jego uchwycie i pokrywka zacisków listwy mocy (CON1). Jest ona po prostu wciskana, aby zakryć zaciski śrubowe. Na koniec obróć VR2 całkowicie w lewo, aby początkowo wyłączyć sprzężenie zwrotne.

Teraz należy dokładnie sprawdzić konstrukcję. Upewnij się, że przewody ochronne (w zielono-żółte paski) łączą obudowę, pokrywę i styki ochronne zarówno we wtyczce sieciowej, jak i gnieździe. Sprawdź to za pomocą multimetru ustawionego na odczyt przejścia/diody. Odczyty między wszystkimi punktami ochronnymi powinny wynosić poniżej 1 Ω .

Przepusty kablowe muszą być dokręcone, aby utrzymać przewody sieciowe na miejscu.

Ponieważ można je łatwo odkręcić, przed dokręceniem należy nałożyć kroplę kleju Super Glue (lepiej Loctite) na gwint dławików. W ten sposób dławiki nie będą mogły zostać odkręcone przez dzieci lub inne ciekawskie osoby.

Przymocuj pokrywę, upewniając się, że okablowanie jest poprowadzone tak, aby pasowało do wyższych elementów na płytce drukowanej. Użyj czterech śrub dostarczonych z obudową; nie ulegaj pokusie uruchomienia regulatora prędkości bez założonej pokrywki!

Testowanie

Podłącz urządzenie z silnikiem uniwersalnym (np. wiertarkę elektryczną zasilaną z sieci) do sterownika, podłącz zasilanie i sprawdź, czy silnik może być sterowany podczas regulacji potencjometrem prędkości obrotowej.

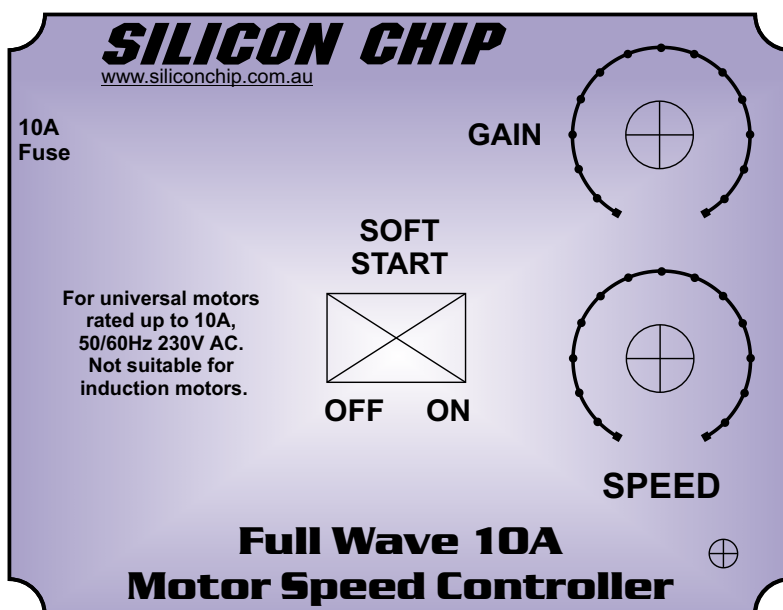
VR2 może wymagać regulacji, aby uniknąć zmian prędkości pod obciążeniem. Obróć go w prawo, jeśli prędkość spada zbyt wyraźnie pod obciążeniem, i w lewo, jeśli silnik przyspiesza pod obciążeniem.

Sprawdź, czy funkcja łagodnego startu działa, gdy jest włączona, wyłączając zasilanie, pozwalając narzędziu zwolnić obroty, a następnie włączając je ponownie, aby sprawdzić, czy uruchamia się płynnie z przełącznikiem S1 w prawidłowej pozycji. ■

John Clarke

Adaptacja do wydania polskiego – Andrzej Nowicki

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



Rysunek 5. Pełnowymiarowa grafika „panelu przedniego”, którą można skopiować lub pobrać i wydrukować (ze strony siliconchip.com.au). Jest ona przyklejona do górnej części odlewanej ciśnieniowo obudowy – może być również użyta jako szablon do wywiercenia trzech otworów w panelu i wycięcia na przełącznik łagodnego rozruchu

Wyjście 0–14 V, 0–1 A – sterowanie z komputera!

Regulowany zasilacz wykorzystujący Arduino



Przez lata publikowaliśmy w Silicon Chip różnego rodzaju wymyślne zasilacze laboratoryjne: liniowe, impulsowe, hybrydowe, wysokonapięciowe, wysokoprądowe, z podwójnym sterowaniem... Ale czasami wszystko, czego potrzebujesz, to podstawowy zasilacz z monitorowaniem i ograniczaniem napięcia i prądu; coś, co jest wygodne i łatwe w konfiguracji oraz obsłudze. To jest właśnie to – bardzo przydatny mały zasilacz zbudowany jako nakładka Arduino!

Ostatnio, podobnie jak wielu innych, pracuję głównie w domu. Niestety, mój domowy warsztat nie jest wyposażony w takim samym stopniu jak biuro/laboratorium Silicon Chip-a.

Mógłbym przynieść do domu mój prototyp zasilacza liniowego 45 V/8 A (opublikowany w październiku-grudniu 2019 r.; patrz siliconchip.com.au/Series/339).

Zrobiłby prawie wszystko, czego potrzebuję, ale moja przestrzeń jest ograniczona, a wykorzystanie jego pełnych możliwości byłoby rzadkim wydarzeniem.

Potrzebuję więc czegoś bardziej kompaktowego, ale wciąż użytecznego. Postanowiłem wykorzystać coś, co już miałem w domu, moduł Arduino Uno. Zasilacz jest w stanie dostarczyć do 14 V przy maksymalnym natężeniu 1 A.

To z pewnością skromna wartość, ale wystarczająca do większości mniejszych projektów. A jeśli potrzebujesz kilku różnych napięć (np. 5 V i 3,3 V), możesz połączyć kilka jednostek.

Uwagi dotyczące Arduino

Korzystanie ze sprzętu Arduino oznacza, że możliwe byłoby dodanie jednej z wielu nakładek i modułów, aby podłączyć niestandardowy wyświetlacz lub elementy sterujące zasilacza. Ale ponieważ mam już komputer na biurku, postanowiłem wykorzystać do sterowania istniejący ekran i klawiaturę.

Napisałem mały program komputerowy, który steruje zasilaczem, zapewniając wszystkie jego przydatne funkcje bez zajmowania cennego miejsca na biurku.

W ten sposób zasilacz może być schowany poza zasięgiem wzroku, z dwoma przewodami wyjściowymi wyprowadzonymi tam, gdzie są potrzebne. Program sterujący zajmuje tylko niewielką ilość miejsca na ekranie.



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

Połączenie mikroukładu Uno i programu sterującego pozwala na dodanie wielu funkcji bez dodatkowego sprzętu.

Na przykład, program sterujący umożliwia utworzenie pięciu wstępnie ustawionych kombinacji napięcia i prądu, które są natychmiast aktywowane. Zapobiega to spowodowaniu uszkodzeń przez nieumyślne ustawienie niewłaściwego napięcia lub limitu prądu.

Chociaż zasilacz nie ma żadnej formy wykrywania temperatury, może oszacować wpływ

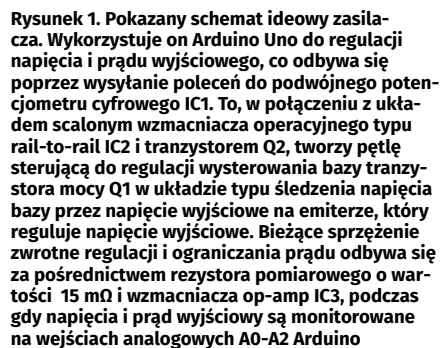
Funkcje i specyfikacje:

- Napięcie i prąd wyjściowy: 0...14 V, 0...1 A
- Regulacja i monitorowanie za pomocą komputera (stacjonarnego, laptopa, notebooka, jednokładowego itp.)
- Wszystkie funkcje pod nadzorem oprogramowania
- Rozdzielczość napięcia: około 20 mV
- Rozdzielczość prądu: około 20 mA
- Konstrukcja oparta na Arduino oznacza możliwość rozbudowy

Układ IC1 to podwójny potencjometr cyfrowy MCP4251; zawiera on dwa potencjometry po 5 k Ω z 257 krokami sterowanymi cyfrowo. Ten układ jest nadzorowany przez Uno

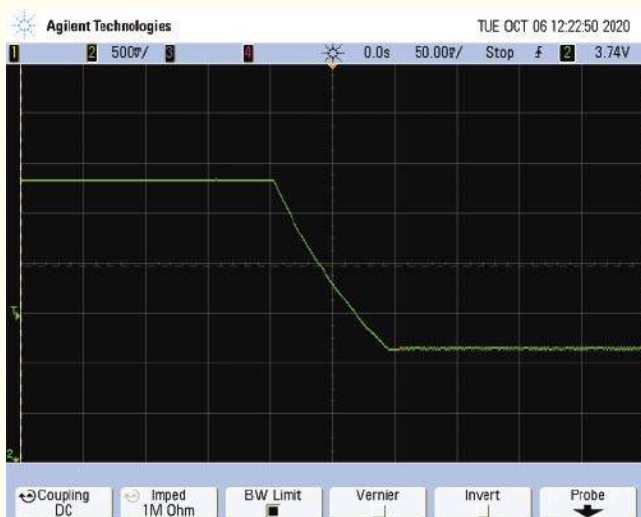
„Ślizgacz” na wyprowadzeniu 6 musi mieć napięcie będące ułamkiem pożądanego napięcia wyjściowego, ponieważ cyfrowy układ scalony ma maksymalne napięcie zasilania

W ten sposób P1 (P1W, wyprowadzenie 6) wytwarza napięcie w zakresie 0...3,3 V, które jest filtrowane dolnoprzepustowo przez obwód RC 10 kΩ/100 nF, a następnie podawane na nieodwracające wejście 3 wzmacniacza operacyjnego IC2a. Jest to połówka





Oscylogram 1. Reakcja na wzrost obciążenia, który wyzwala ograniczenie prądu. Żółty przebieg to napięcie na rezystorze pomiarowym, więc jest proporcjonalne do prądu, podczas gdy zielony przebieg jest proporcjonalny do napięcia wyjściowego. Występuje pewne przekroczenie prądu, głównie z powodu pojemności na wyjściu, po czym włącza się ograniczenie prądu, zmniejszając napięcie wyjściowe, aby osiągnąć stan ustalony w ciągu 1 ms



Oscylogram 2. Reakcja na skokową zmianę ustawionego napięcia z 5 V do 3,3 V (bez obciążenia). Trwa to nieco poniżej 100 ms ze względu na rozładowanie kondensatora 10 μ F na wyjściu przez dzielnik napięcia. Każde znaczące obciążenie znacznie przyspieszyłoby ten proces



Oscylogram 3. Skokowy wzrost ustawionego napięcia (tym razem z 3,3 V do 5 V przy obciążeniu 12 Ω) jest znacznie szybszy ze względu na niższą impedancję tranzystora wyjściowego, trwając zaledwie kilka milisekund

podwójnego wzmacniacza operacyjnego CMOS LMC6482 z wejściem/wyjściem typu rail-to-rail, który umożliwia obniżenie napięcia wyjściowego do 0 V bez szyny zasilania ujemnego, co znacznie ułatwia mierzenie prądu (jak opisano później).

Ten wzmacniacz operacyjny porównuje napięcie wyjścia potencjometru z ustaloną częścią napięcia wyjściowego, wytwarzaną przez dzielnik 51 k Ω /10 k Ω . Ułamek napięcia wyjściowego zasila wejście odwracające IC2a na wyprowadzeniu 2. Daje to 6,1-krotne wzmocnienie napięcia z wyjścia potencjometru (albo, jak kto woli, 6,1-krotne zmniejszenie napięcia wyjściowego), w celu porównania z napięciem wyjściowym. Tak więc około 20 V na wyjściu odpowiada 3,3 V w pełnej skali na wyjściu cyfrowego potencjometru IC1.

Wyjście z wyprowadzenia 1 IC2a trafia na bazę tranzystora NPN Q1, który jest skonfigurowany w układzie („emitter-follower”) naśladowania napięcia bazy przez napięcie wyjściowe w obwodzie emitera (wspólny kolektor). Jego kolektor czerpie z zasilania VIN Arduino, podczas gdy jego emiter zasila wyjście zasilania w złączu CON1 poprzez styki przekaźnika RLY1 (więcej na ten temat później).

Tranzystor ten skutecznie zwiększa wydajność prądową wyjścia wzmacniacza operacyjnego, dzięki czemu może on dostarczać do 1 A (z zasilania VIN).

Spadek napięcia na formalnej diodzie baza-emiter tranzystora Q1 jest niwelowany, ponieważ Q1 znajduje się w pętli ujemnego sprzężenia zwrotnego – od wyjścia 1 IC2a, przez Q1, a następnie przez dzielnik wyjściowy 51 k Ω /10 k Ω z powrotem do wejścia 2 IC2a. W związku z tym IC2a ustawia swoje napięcie wyjściowe nieco wyżej, aby osiągnąć požądane napięcie na wspólnym styku RLY1.

Chociaż obwód jest skonfigurowany tak, aby umożliwić napięcie wyjściowe do 20 V, w praktyce inne elementy obwodu ograniczają praktyczne napięcie wyjściowe do około 14 V. Głównym ograniczeniem jest regulator 5 V na płytce Arduino, który w przypadku klonów płytek ma napięcie znamionowe tylko 15 V (więcej informacji można znaleźć w naszym artykule na temat naprawiania Arduino z marca 2020 r. na stronie siliconchip.com.au/Article/12582).

Regulacja napięcia

Tranzystor mocy Q1 to MJE3055. Zwykle napięcie na jego emiterze (tj. wyjściowe) wynosi około 0,7 V poniżej jego napięcia na bazie (z wyjścia 1 wzmacniacza operacyjnego IC2a). Jeśli napięcie emiter/wyjście wzrasta (na przykład z powodu obciążenia pobierającego mniej prądu), wówczas napięcie baza-emiter spada, co zaczyna go wyłączać, powodując spadek napięcia na emiterze.

I odwrotnie, jeśli napięcie emiter/wyjście spada, napięcie baza-emiter wzrasta i Q1 włącza się w większym stopniu, zatrzymując spadek napięcia na emiterze. To „lokalne sprzężenie zwrotne” zapewnia bardzo szybką reakcję na stany nieustalone obciążenia.

Podczas gdy obwód naśladowania napięcia bazy przez napięcie wyjściowe na emiterze jest dość dobry w nadążaniu swojego wyjścia za wejściem, samo napięcie baza-emiter zmienia się nieco w zależności od obciążenia. Aby temu zaradzić, wzmacniacz operacyjny dostosowuje napięcie sterujące bazą Q1, aby utrzymać napięcie na wyjściowym dzielniku napięcia w pobliżu wartości odniesienia na potencjometrze cyfrowym. Wzmacniacz operacyjny reaguje jednak wolniej ze względu na ograniczoną szerokość pasma przenoszenia.

Tranzystor Q1 jest wyposażony w niewielki żebrowany radiator, ponieważ działa on jako element liniowo-przepustowy, rozpraszając nadmiar napięcia między zasilaniem a wyjściem. Został wybrany radiator niskoprofilowy, aby można było umieścić nad nim inną płytkę, jeśli konieczne będzie dodanie niestandardowego elementu sterującego lub wyświetlacza.

Zaprojektowaliśmy nakładkę zasilacza tak, aby nie kolidowała z łączami używanymi w adapterze LCD opisanym w Silicon Chip-ie w maju 2019 roku (siliconchip.com.au/Article/11629), co oznacza, że możemy przekształcić zestaw Arduino z nakładką zasilacza w urządzenie typu „wszystko w jednym”, dodając w przyszłości ekran dotykowy LCD.

Jednak obecna wersja oprogramowania nie obsługuje tej opcji.

Filtr wyjściowy 10 μ F zapewnia skromną filtrację na wyjściu, co również poprawia regulację stanów przejściowych. Wartość ta jest kompromisem, ponieważ zbyt mała pojemność wyjściowa pogorszyłaby regulację, a zbyt duża pojemność ograniczyłaby zdolność zasilacza do szybkiego ograniczenia prądu wyjściowego w warunkach zwarcia.

Pętla sprzężenia zwrotnego pomiędzy Q1 i IC2 ma duże wzmocnienie, więc należy uważać, aby nie wpadła w oscylację. Kondensator odsprężający 100 nF napięcia odniesienia na wyprowadzeniach 7 i 8 układu IC1 zapobiega widzeniu stanów nieustalonych przez wzmacniacz operacyjny. W przeciwnym razie byłyby powielane na wyjściach IC2. Podobnie, sygnał pożądanego napięcia na wejściu 3 układu IC2a jest stabilizowany za pomocą innego kondensatora 100 nF.

Istnieje również kondensator odsprężający 100 nF równolegle z rezystorem 51 k Ω dzielnika sprzężenia zwrotnego, który zmniejsza wzmocnienie zamkniętej pętli sześciokrotnie w przypadku szybkich stanów nieustalonych. Ponadto kondensator 1 nF jest podłączony między wyjściem (wyprowadzenie 1) a wejściem odwracającym (wyprowadzenie 2) IC2a, ograniczając szybkość narastania sygnału na wyjściu wzmacniacza operacyjnego. Innym tłumaczeniem obecności tego kondensatora jest zapewnienie zwiększonego ujemnego sprzężenia zwrotnego przy wysokich częstotliwościach. Zapobiega to oscylacjom.

Filtr dolnoprzepustowy utworzony przez rezystor 100 Ω i kondensator 10 μ F w bazie Q1, bocznikujący bazę do masy, również pomaga ustabilizować pętlę sprzężenia zwrotnego.

Przełącznik wyjściowy

Przełącznik przełączający na wyjściu jest przełącznikiem kontaktronowym. Jego cewka jest zasilana z wyjścia cyfrowego D5 Arduino. Jest to możliwe, ponieważ prąd cewki przełącznika kontaktronowego jest niewielki.

Niestety, potencjometry cyfrowe w IC1 uruchamiają się ze ślizgaczami w położeniu środkowym, więc na wyjściu będzie obecne napięcie, jeśli początkowo przełącznik RLY1 nie będzie odłączony. Dlatego RLY1 jest zasilany dopiero wtedy, gdy napięcie

wyjściowe regulatora ustabilizuje się na żądanym poziomie.

RLY1 działa również jako przełącznik odłączający obciążenie, umożliwiając obwodowi uzyskanie żadanego napięcia wyjściowego bez podłączonego obciążenia. Może wtedy szybko podłączyć obciążenie do już poprawnego ustawionego napięcia, bez konieczności jego zwiększania. Podobnie, może szybko odłączyć obciążenie w przypadku wystąpienia przeciążenia lub zwarcia.

Ograniczanie prądu

Ograniczenie prądu wykorzystujemy podobną pętlę sprzężenia zwrotnego jak sterowanie napięciem. Tutaj używamy najprostszego możliwego sposobu wyprowadzenia natężenia prądu. Na rezystorze pomiarowym o wartości 15 m Ω w obwodzie masy obciążenia odkłada się napięcie proporcjonalne do prądu płynącego przez obciążenie.

Jest ono podawane przez filtr dolnoprzepustowy RC 1 k Ω /100 nF do wejścia nieodwracającego (wyprowadzenie 3) układu IC3a drugiego wzmacniacza operacyjnego. Ponieważ musi on obsługiwać tylko napięcie do około 3,3 V, używamy tańszego układu scalonego MCP6272 z podwójnym wzmacniaczem operacyjnym (jego druga połowa nie jest używana); układ typu „rail-to-rail” z wejściami/wyjściami w całym zakresie zasilania jest zbyt cenny.

Układ IC3 wzmacnia napięcie na boczniku przez współczynnik 151 (150 k Ω /1 k Ω +1). Wzmocnione napięcie na boczniku jest następnie podawane na wejście 5 układu scalonego wzmacniacza IC2b (jego wejście nieodwracające). Napięcie 2,2 V na wejściu 5 wzmacniacza IC2b odpowiada w przybliżeniu prądowi wyjściowemu 1 A.

Napięcie to jest porównywane z napięciem z wyjścia drugiego potencjometru cyfrowego w układzie IC1. Jeśli prąd wyjściowy jest powyżej wartości zadanej, wyjście 7 układu IC2b przechodzi w stan wysoki, polaryzując do stanu przewodzenia złącze baza-emiter tranzystora NPN Q2.

Gdy tranzystor Q2 jest włączony, obniża napięcie na wyprowadzeniu 3 wzmacniacza IC2a, zmniejszając napięcie wyjściowe. Powinno to prowadzić do zmniejszenia prądu pobieranego przez obciążenie, aż do osiągnięcia limitu prądu, w którym to momencie

napięcie bazy Q2 jest ustalane, więc napięcie wyjściowe powinno ustabilizować się na poziomie, przy którym prąd wyjściowy jest zbliżony do ustawionego limitu prądu.

Należy tu zwrócić uwagę na kilka rzeczy. Po pierwsze, pozorne odwrócenie wejścia odwracającego i nieodwracającego na IC2b wynika z faktu, że wzmacniacz na tranzystorze Q2 ze wspólnym emiterem odwraca polaryzację sygnału w pętli sprzężenia zwrotnego. Zamieniając wejścia: odwracające i nieodwracające, ponownie odwracamy polaryzację sygnału i uzyskujemy prawidłowe sterowanie.

Podobnie jak w przypadku pętli napięciowego sprzężenia zwrotnego, stabilność poprawia kondensator 1 nF między wyjściem (wyprowadzenie 7) a wejściem odwracającym (wyprowadzenie 6), a także kondensator 100 nF stabilizujący ustawione napięcie ekwiwalentu prądu na wejściu odwracającym 6.

Sygnały sprzężenia zwrotnego napięcia i prądu trafiają również do dwóch analogowych portów wejściowych na płytce Uno. W ten sposób moduł Uno może wykrywać (za pomocą swojego urządzenia peryferyjnego, konwertera ADC) napięcie i prąd za pomocą odpowiednio wejść A1 i A0.

Napięcie zasilania VIN jest mierzone przez drugi dzielnik 51 k Ω /10 k Ω na wejściu analogowym A2. Pozwala to mikrokontrolerowi obliczyć spadek napięcia na Q1 i wnioskować o poziomie rozpraszania ciepła przez Q1.

Na płytce drukowanej znajdują się punkty testowe dla czterech napięć podawane/ustawione. Są one oznaczone jako VFB, IFB (sprzężenie zwrotne napięcia i prądu), VSET i ISET (wartości zadane napięcia i prądu) oraz jeden dla GND.

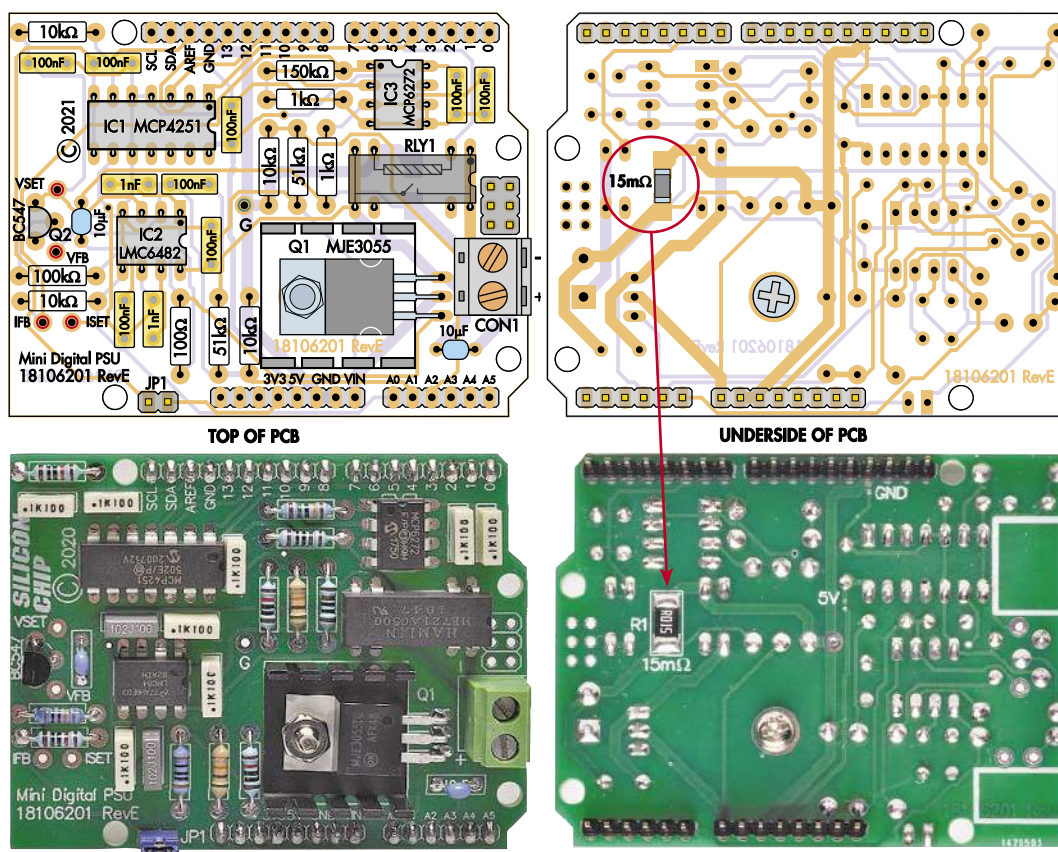
Oprogramowanie Arduino

Oprogramowanie układowe Arduino generuje dane magistrali SPI w celu ustawienia żądanych limitów napięcia i prądu, a następnie zamyka przełącznik, aby włączyć wyjście

REKLAMA

OBWODY DRUKOWANE Produkcja, Projektowanie, Montaż			
ELMAX 1988 Certyfikat Underwriters Laboratories UL 94V-0 E480148 TYPE 1 Zakład produkcyjny: 05-640 Warka ul. M. Rzepiełkowskiej 17 tel. 22 781 63 95 22 761 95 80 fax. 22 781 63 95 w. 23 www.elmax.waw.pl elmax@elmax.waw.pl	Płytki jednostronne Płytki dwustronne Płytki na podłożu aluminium	Serie dowolne Prototypy Maksymalny wymiar płytki 1m 630 mm	Dokumentacja technologiczna Dokumentacja konstrukcyjna Montaż elektroniczny ilości modelowe produkcyjne
Aktywny kalkulator prototypów na stronie internetowej	Pokrycie Sn lub SnPb line na życzenie Maski, opisy montażowe w różnych kolorach	Płyty czolowe FR4 Trawione szablony SMD	Krótkie terminy Wykonania super ekspresowe

Rysunek 2. Ta zachęcająco prosta nakładka modułu Arduino zmienia Uno w regulowany zasilacz stołowy. Pokazany został schemat montażowy obu stron płytki drukowanej oraz odpowiadające fotografie rzeczywistego zasilacza, poza lutowaniem końcówek listew kołkowych, jedynym elementem na spodzie (i jedynym elementem SMD) jest rezystor pomiarowy 15 mΩ. Transzystor mocy Q1 ma niewielki radiator, ponieważ musi rozpraszać moc kilka watów. Układy scalone, przekaznik i tranzystory są elementami kierunkowymi, więc muszą być zorientowane tak, jak pokazano na schemacie, podczas gdy inne komponenty mogą być umieszczone dowolnie. W zestawie znajduje się kilka punktów testowych, ale nie są one potrzebne do kalibracji. Postaraj się, aby Twoja nakładka wyglądała tak schludnie, jak prototyp Silicon Chip-a.



po akceptacji przez użytkownika. Nakładka modułu Uno zarządza następnie napięciem wyjściowym, zmniejszając je, jeśli limit prądu zostanie osiągnięty, jak opisano powyżej.

Po ustawieniu napięcia i prądu działanie regulatora jest automatyczne; nie zależy od oprogramowania sterującego.

Mikroprocesor mierzy napięcie zasilania i wyjściowe oraz prąd obciążenia, a następnie wysyła te dane w celu wyświetlenia do programu uruchomionego na komputerze.

Kalibracja urządzenia polega na określeniu dokładnej zależności między wartościami cyfrowymi (odczytami ADC i ustawieniami potencjometru cyfrowego) a wynikowymi napięciami analogowymi. Współczynniki te można obliczyć na podstawie zmierzonych wartości napięcia i prądu wyjściowego.

Zasilacz będzie dość dokładny „po wyjęciu z pudełka”. Jego dokładność można jednak poprawić, dokonując odczytów za pomocą multimetru, określając dokładne współczynniki i programując je wewnątrz kodu. Procedura kalibracji w programie komputerowym upraszcza ten proces, automatycznie obliczając nowe współczynniki na podstawie pomiarów.

Budowa

Główną częścią zrealizowania projektu jest montaż nakładki. Wszystkie części mieszczą się na dwustronnej płytce drukowanej o kodzie

18106201 i wymiarach 69×54 mm – patrz schemat montażowy na **rysunku 2**.

Pierwszą decyzją do podjęcia jest to, czy chcesz zbudować nakładkę ze zwykłymi slotami, czy ze slotami do układania nakładek modułu Arduino w stosy. Złącza do podłączenia rozszerzeń będą potrzebne, jeśli planujesz dodać do nakładki zasilacza jakiegokolwiek inne nakładki. W prototypie Silicon Chip-a użyliśmy jednak zwykłych listew kołkowych, ponieważ na razie nie planujemy tego robić.

Po tej wstępnej decyzji montaż jest względnie prosty. Aby upewnić się, że wszystko znajduje się na właściwym miejscu i we właściwej orientacji, sprawdź podczas montażu części schemat montażowy na **rysunku 2**, warstwę opisową (silkscreen) PCB i odpowiadające zdjęcia.

Zacznij od zamontowania rezystora SMD 15 mΩ, który znajduje się na spodzie płytki drukowanej. Niektórzy z Czytelników lubią używać (drewnianego) spinacza do białizny, aby przytrzymać element SMD na miejscu podczas lutowania.

Odwróć płytkę i przylutuj jeden styk lutownicą. Jeśli element leży płasko i mieści się w obrębie prostokątnego oznaczenia na warstwie opisowej, przylutuj drugą końcówkę. W przeciwnym razie ponownie rozgrzej pierwsze pole lutownicze doprowadzając cynę do stanu płynnego i ustaw rezystor, używając pęsety, jeśli to konieczne, aż zostanie prawidłowo

ulożony. Następnie przylutuj drugie wprowadzenie i odwróć płytkę PCB.

Zamontuj w górnej części płytki drukowanej 11 rezystorów o montażu przewlekany, zgodnie z oznaczeniami na warstwie opisowej. Przed zamontowaniem sprawdź ich wartości za pomocą multimetru, ponieważ niektóre oznaczenia mogą wyglądać dość podobnie.

Kontynuuj pracę z montażem ośmiu kondensatorów 100 nF i dwóch 1 nF, które powinny mieć oznaczenia swoich wartości na obudowach (lub kody reprezentujące je, jak odpowiednio 104 i 102). Żaden z nich nie jest spolaryzowany; podobnie jak kondensatory ceramiczne MLC 10 μF, które mogą być w obudowach do montażu przewlekane lub typu SMD. Przylutuj je teraz.

Następnie wlutuj mniejszy tranzystor, Q2. Dopasuj wyprowadzenia do pól lutowniczych PCB, upewniając się, że po zamontowaniu korpus znajduje się nisko, na wypadek gdybyś musiał dodać w przyszłości jakąś nakładkę nad nim. Upewnij się, że pasuje do konturu na warstwie opisowej PCB.

Postępuj zgodnie z opisem umiejscowienia tranzystora Q1 w obudowie TO-220. Jest on zamontowany na żebrowanym radiatorze. Najpierw zagnij wyprowadzenia tranzystora do tyłu o 90° około 7 mm od korpusu, a następnie włóż je do odpowiednich otworów metalizowanych (THT) na płytce PCB. Sprawdź, czy



większy otwór w radiatorze tranzystora jest centryczny z otworem montażowym na płytce PCB i w razie potrzeby wyreguluj wyprowadzenia.

Wymnij Q1 z płytki drukowanej i włóż śrubę M3 od spodu płytki drukowanej. Dodaj radiator po stronie montażu (większości) komponentów, następnie zamontuj tranzystor i przykręć nakrętkę. Przed dokręceniem upewnij się, że radiator równomiernie przylega, jedną stroną do powierzchni płytki PCB a drugą do tranzystora oraz czy mieszczą się w obrysie tych komponentów na warstwie opisowej płytki PCB. Ostrożnie dokręć nakrętkę (aby uniknąć uszkodzenia wyprowadzeń tranzystora), a następnie przylutuj wyprowadzenia i przytnij je od spodu.

Większość pozostałych części znajduje się w obudowach DIL. Tym razem zrezygnuj z montowania podstawek pod układy scalone, ponieważ nie tylko zapewnią one gorsze połączenie niż w przypadku bezpośredniego przylutowania układów scalonych, ale także spowodują, że komponenty będą znajdować się znacznie wyżej.

Przełącznik RLY1 ma osiem wyprowadzeń, ale jest dostarczany w obudowie DIL-14. Znajduje się powyżej Q2; wycięcie w jego obudowie jest skierowane w prawo. Delikatnie wygnij jego wyprowadzenia, aby wyrównać je z otworami w polach PCB i zamontuj przełącznik do płytki. Przylutuj dwa skrajne wyprowadzenia i sprawdź, czy obudowa przylega płasko do PCB; ewentualnie popraw, jeśli tak nie jest. Przylutuj pozostałe wyprowadzenia, a następnie wróć i odśwież lut na pierwszych dwóch wyprowadzeniach.

IC1 jest w obudowie DIL-14; jego wyprowadzenie 1 powinno znajdować się bezpośrednio przy sąsiednim kondensatorze. IC2 to LMC6482, jak zaznaczono na warstwie opisowej. Nie należy go pomylić z układem IC3, który jest oznaczony jako MCP6272, chociaż akurat zamiast niego można użyć drugiego LMC6482.

Użyj podobnej techniki jak w przypadku przełącznika RLY1, aby dopasować IC1, IC2 i IC3. Gdy to zrobisz, sprawdź, czy nie ma mostków lub „zimnych lutów” i napraw je w razie potrzeby za pomocą odsysacza cyny lub plecionki lutowniczej, aby usunąć nadmiar lutu. Stwórz nowy, poprawny lut, aby zakończyć lutowanie układów scalonych.

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

Lista części – zasilacz – nakładka Arduino

- 1 dwustronna płytka PCB o kodzie 18106201 i wymiarach 69×54 mm
- 1 moduł Arduino Uno lub płytka kompatybilna
- 1 zasilacz wtyczkowy 12–15 V 1 A z wtyczką DC 5,5/2,1 mm pasującą do gniazda Uno lub podobne źródło zasilania
- 1 2-drożny zacisk śrubowy (CON1)
- 1 6-szpilkowa jednorzędowa listwa kołkowa (lub złącze piętrowe męsko-damskie, patrz tekst)
- 2 8-szpilkowe jednorzędowe listwy kołkowe (lub złącza piętrowe męsko-damskie, patrz tekst)
- 1 10-szpilkowa jednorzędowa listwa kołkowa (lub złącze piętrowe męsko-damskie, patrz tekst)
- 1 żebrowany radiator pod obudowę TO-220 (dla Q1) [Jaycar HH8502].
- 1 2-szpilkowy odcinek listwy kołkowej i zworka (JP1)
- 1 odcinek 3×2-szpilkowy listwy kołkowej (opcjonalnie; wymagany, jeśli inna nakładka ma być podłączona powyżej)
- 1 przełącznik kontaktowy w obudowie DIL-14 z cewką 5 V (RLY1) [Altronics S4100, Jaycar SY4030]. Zasilacze zbudowane z przełącznikiem Jaycar powinny ustawiać prąd nie wyższy niż 500 mA, aby uniknąć uszkodzenia przełącznika, ponieważ ten przełącznik ma wartość znamionową prądu 500 mA

Półprzewodniki:

- 1 podwójny potencjometr cyfrowy 2×5 kΩ MCP4251-5k, DIP-14 (IC1) [Silicon Chip ONLINE SHOP SC5052; Digikey, Mouser]
- 1 podwójny wzmacniacz operacyjny LMC6482, DIP-8 (IC2) [Jaycar ZL3482].
- 1 podwójny wzmacniacz operacyjny MCP6272, DIP-8 (IC3; LMC6482 również może być użyty)
- 1 tranzystor NPN MJE3055 10 A, TO-220 (Q1) [Jaycar ZT2280]
- 1 tranzystor NPN BC547 100 mA, TO-92 (Q2) [Jaycar ZT2152]

Kondensatory:

- 2 kondensatory ceramiczne 10 μF 16 V X7R do montażu przewlekane (lub SMD 1206)
- 8 kondensatorów foliowych 100 nF MKT
- 2 kondensatory foliowe 1 nF MKT

Rezystory: (wszystkie 1/4W 1% metalizowane z wyprowadzeniami osiowymi, z wyjątkiem zaznaczonych)

- | | | | | |
|---------------|--|--------------|--------------|-------------|
| 1 szt. 150 kΩ | 1 szt. 100 kΩ | 2 szt. 51 kΩ | 4 szt. 10 kΩ | 2 szt. 1 kΩ |
| 1 szt. 100 Ω | 1 szt. 15 mΩ 1% SMD 2512 [SC ONLINE SHOP SC3943] | | | |

Listwy połączeniowe i zworka

Następnie należy przymocować złącza montażowe Arduino wzdłuż krawędzi płytki. Jeśli używasz męskich listew kołkowych, ich montaż jest prosty. Użyj Uno jako szablonu i podłącz złącza szpilkowe do Uno, a następnie umieść płytkę drukowaną na górze. Po sprawdzeniu, że wszystko jest równe i prostopadłe, przylutuj wyprowadzenia kołków od góry i odłącz zespół od Uno.

Jeśli chcesz użyć złączy, które można układać jedno na drugim, jest to nieco trudniejsze, chociaż Uno może być nadal używany jako szablon montażowy. W tym przypadku listwy przechodzą przez płytkę drukowaną od góry do modułu Uno. Odwróć zespół tak, aby płytka drukowana zasilacza znajdowała się na dole.

Teraz masz dostęp do szpilek złączy od dołu. Powinno to wystarczyć, aby przylutować skrajne kołki każdego złącza i utrzymać złącza na miejscu. Sprawdź, czy oprawki złączy przylegają płasko do płytki drukowanej i wyreguluj je w razie potrzeby.

Odłącz nakładkę od Uno, aby uzyskać lepszy dostęp do pozostałych kołków. Przylutuj je, a następnie odśwież lutowanie skrajnych kołków.

W tym przypadku prawdopodobnie konieczne będzie również przylutowanie odcinka 3×2-kołkowej listwy do styków R3 na płycie,

aby przekazać te sygnały do płytki umieszczonej powyżej.

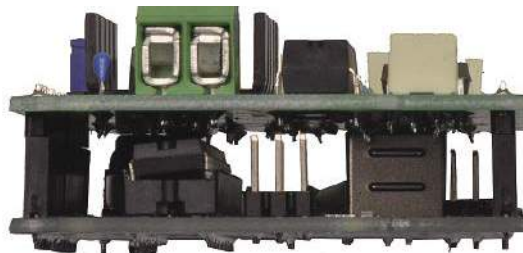
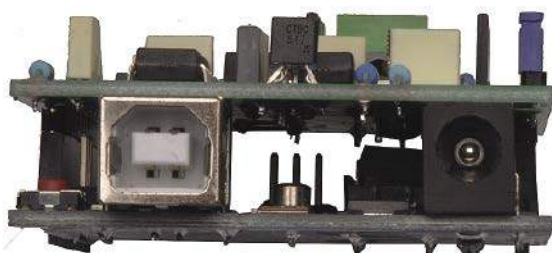
JP1 składa się z dwukołkowego odcinka listwy szpilkowej i zworki. Nałóż zworkę na kołki listwy, umieść listwę na płycie drukowanej i przylutuj jej wyprowadzenia. Zworka utrzyma kołki na miejscu, nawet jeśli plastikowa osłona trochę się stopi.

Wreszcie nadszedł czas, aby zamontować złącze wyjściowe CON1. Użyliśmy dwustykowego zaciskowego złącza śrubowego, chociaż możesz preferować coś innego, w zależności od tego, jak chcesz korzystać z zasilacza.

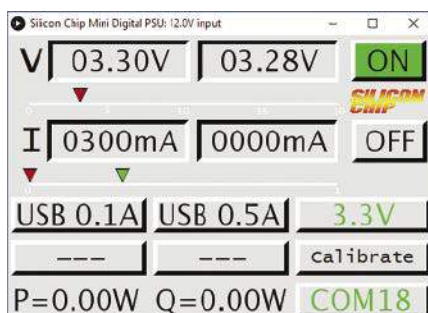
Przylutuj CON1 na miejscu, a następnie dopasuj płytkę PCB do Uno. O ile Uno nie jest nowy i niezaprogramowany, należy usunąć JP1, na wypadek gdyby istniejący szkic używał innego napięcia odniesienia, które mogłoby kolidować z zasilaniem 3,3 V i ewentualnie uszkodzić nakładkę.

Oprogramowanie

Istnieją dwa elementy oprogramowania tego projektu – pierwszy to oprogramowanie układowe, które działa na Uno. Drugim jest aplikacja komputerowa, która łączy się z modulem Uno. Szkic oprogramowania Arduino jest dostępny do pobrania ze strony Silicon Chip-a.



Widoki na płytki – nakładka zasilacza na górze; standardowe Arduino Uno (lub kompatybilne) poniżej



Ekran 1. Nasza aplikacja Processing udostępnia suwaki sterowania napięciem i prądem u góry, wraz z prostymi przełącznikami do włączania i wyłączenia wyjścia. Ustawienia wstępne są wyświetlane i wybierane poniżej, wraz z informacjami o mocy pobieranej przez zasilacz. Podłączone napięcie zasilania można obserwować na pasku tytułowym

Zakładamy, że masz pewną znajomość Arduino IDE (zintegrowane środowisko programistyczne), chociaż nie jest to zbyt trudne, nawet jeśli jesteś nowym użytkownikiem. IDE można pobrać za darmo ze strony siliconchip.com.au/link/aaatq. Używamy wersji 1.8.5, ale praktycznie każda wersja powinna działać, ponieważ szkic jest dość prosty i nie wymaga żadnych specjalnych bibliotek.

Po pobraniu szkicu, następnym krokiem jest załadowanie oprogramowania układowego do Uno. Podłącz Uno do portu USB, wybierz port szeregowy Uno z menu Narzędzia Arduino IDE, a następnie upewnij się, że płytka Uno jest wybrana jako cel (Narzędzia → Płytki → Arduino Uno). Naciśnij Upload, a gdy szkic zostanie przesłany, załóż zwórkę JP1 i otwórz Serial Monitor z prędkością 115 200 bpd (CTRL + SHIFT + M w systemie Windows).

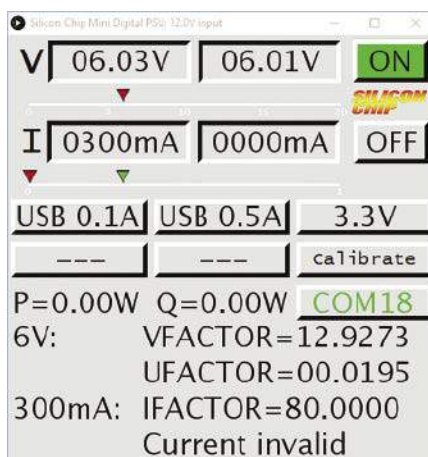
Szkic jest dość prosty; czeka na porcje szeregowy na polecenia takie jak „V100”, „I50” lub „R1”, aby ustawić odpowiednio napięcie, prąd lub stan przekaźnika. Ponieważ komunikacja do i z zasilacza odbywa się po prostu przez szynę szeregową, możemy również przetestować urządzenie, wpisując polecenia do programu terminala szeregowego, takiego jak Serial Monitor.

Tak proste rozwiązanie oznacza, że w razie potrzeby można zasilaczem sterować ręcznie. Ale oznacza to również, że zasilacz może być bardzo łatwo nadzorowany przez inne oprogramowanie; wystarczy wysłać odpowiednie polecenia i przetworzyć (proste) odpowiedzi.

Nawet jeśli nie jest dostępne zasilanie 12 V, sam moduł Uno poda do testów napięcie około 4 V na wejście VIN (a tym samym zasilacz). To wystarczy, aby przeprowadzić proste testy niskonapięciowe w celu sprawdzenia, czy urządzenie działa zgodnie z oczekiwaniami.

Testowanie

Po podłączeniu zasilacza przez port USB i otwarciu monitora szeregowego, powinieneś



Ekran 2. Procedura kalibracji jest prosta. Regulujesz elementy sterujące, aż odczyt multimetru będzie zgodny z odczytami napięcia lub prądu pokazanymi w lewym dolnym rogu, po czym po prostu kopiujesz parametry do pliku konfiguracyjnego

zobaczyć strumień linii pokazujących wartości poprzedzone nazwami J, U i S. Wartości J i U powinny być bliskie zeru, ale S będzie wynosić około 200 (wskazując około 4 V na VIN). Aby przetestować przekaźnik, wpisz „R1” lub „R0”, a następnie naciśnij Enter. Powinieneś być w stanie usłyszeć delikatne kliknięcie (po komendzie R1) i wyłączenie (po R0).

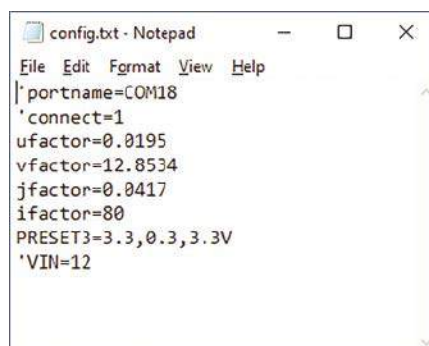
Możesz wysyłać polecenia do potencjometru cyfrowego, wpisując V lub I, a następnie liczbę z zakresu 0–256, a potem naciskając Enter.

Liczby te są surowymi cyfrowymi wartościami nastaw potencjometru, ponieważ cała kalibracja jest wykonywana w programie komputera-hosta. Po włączeniu przekaźnika i ustawieniu wartości V i I na wartości niezerowe, należy zmierzyć napięcie na zaciskach wyjściowych.

Wartości J, U i S są surowymi odczytami ADC (0–1023) napięcia i prądu wejściowego i wyjściowego, pobieranymi kilka razy na sekundę przez Uno. Litera J, U i S zostały wybrane, aby uniknąć pomyłki z poleceniami V i I. Program hosta konwertuje odczyty 0–1023 na rzeczywiste napięcia i prądy.

Aby przetestować wyjście za pomocą multimetru podłączonego do CON1, wprowadź polecenie „R1”, a następnie „V255” i „I255”. Powinno to pozwolić na uzyskanie na wyjściu napięcia o około 0,7 V niższego od napięcia zasilania VIN (ograniczonego przez nieodłączny spadek napięcia na formalnej diodzie baza-emiter wtórnika emiterowego Q1).

Wypróbuj niższe wartości V (np. V25), aby sprawdzić, czy wyjście może być regulowane do niższego poziomu. Powinno to dać około 2 V, podczas gdy V37 powinno dać około 3 V, a V13 powinno dać około 1 V. Aby sprawdzić wyższe napięcia wyjściowe, będziesz musiał



Ekran 3. Plik „config.txt” zawiera parametry kalibracji i do pięciu nazwanych ustawień wstępnych. Można również ustawić port szeregowy i to, czy aplikacja ma się z nim automatycznie łączyć podczas uruchamiania

podłączyć zasilanie 12...15 V do gniazda zasilania wtyczkowego Arduino (ale uważaj na górny limit napięcia!).

W przypadku tego testu dobrym pomysłem byłoby podłączenie Uno do komputera za pomocą zabezpieczenia portu USB, takiego jak projekt Silicon Chip-a z 2018 roku (siliconchip.com.au/Article/11065). Oznacza to, że jeśli nawet wystąpi błąd w zasilaczu, który spowoduje, że napięcie 12 V lub więcej będzie podawane z powrotem do styków sygnałowych portu USB (które działają przy napięciu 3,3 V), nie powinno to uszkodzić komputera.

Przetwarzanie aplikacji

Aplikację do sterowania komputerem napisaliśmy w języku Processing. IDE Processing jest dostępny dla systemów Windows, Mac i Linux (w tym Raspberry Pi). Za pomocą IDE można uruchomić program lub skompilować go do samodzielnego pliku wykonywalnego dla systemu. Program oparty jest na języku Java, więc do jego uruchomienia prawdopodobnie konieczne będzie zainstalowanie środowiska uruchomieniowego Java (JRE).

Processing można pobrać za darmo ze strony <https://processing.org/download/> (my korzystaliśmy z wersji 3.5.3).

Nie są wymagane żadne specjalne biblioteki ani dodatki. Otwórz szkic Processing (plik z rozszerzeniem .pde) za pomocą menu File i uruchom go za pomocą kombinacji klawiszy Ctrl-R. Samodzielny plik wykonywalny można utworzyć z menu File → Export Application.

Odnosząc się do Ekranu 1, rzeczywiste i ustawione napięcia są wyświetlane jako poziomy graf i w formie cyfrowej w górnej części okna. Podobny ekran poniżej pokazuje rzeczywiste i ustawione prądy. Dwa duże przyciski służą do włączania i wyłączania wyjścia. Poniżej znajduje się pięć przycisków ustawień wstępnych i przycisk dostępu do strony kalibracji.

Wzdłuż dolnej części znajdują się wyświetlacze mocy wyjściowej (P) i mocy traconej na tranzystorze Q1 (Q). Zmieniają one kolor wraz ze wzrostem mocy. W prawym dolnym rogu znajduje się wskaźnik portu szeregowego.

Początkowa kalibracja tego oprogramowania pochodzi z testowania prototypu zbudowanego w Redakcji Silicon Chip-a, więc powinna być z grubsza poprawna w ramach tolerancji komponentów. Można jednak mimo to łatwo dostroić zasilacz.

Korzystanie z zasilacza

Naciśnij „+” i „-” na klawiaturze, aby przełączać się między dostępnymi portami szeregowymi. Po wybraniu portu Uno naciśnij „s”, aby się połączyć – jeśli połączenie się powiedzie, symbol portu szeregowego zmieni kolor na zielony. Jeśli połączenie nie zostanie nawiązane, sprawdź, czy port nie jest używany przez inny program (na przykład Arduino Serial Monitor).

Klawisz „s” pełni funkcję przełączania, więc może być również używany do odłączania od zasilacza.

Przeciśnij suwaki na poziomych skalach za pomocą wskaźnika myszy, aby ustawić napięcie i prąd. Zielona strzałka to wartość zadana, która odpowiada najbardziej wysuniętemu na lewo wyświetlaczowi cyfrowemu. Czerwona strzałka i cyfry po prawej stronie odpowiadają rzeczywistym wartościom napięcia i prądu.

Kliknij przycisk „ON”, aby zasilić przełącznik i włączyć wyjście. Należy pamiętać, że zasilacz odczytuje napięcie przed przełącznikiem, więc pokaże wartość, nawet jeśli przełącznik jest wyłączony. Przycisk „ON” zmienia kolor na zielony, gdy przełącznik jest włączony. Użyj przycisku „OFF”, aby go wyłączyć.

Naciśnięcie dowolnego z pięciu przycisków ustawień wstępnych spowoduje załadowanie tego ustawienia wstępnego do wartości zadanych napięcia i prądu. Na Ekranie 1 załadowane jest ustawienie 3, więc jego przycisk jest podświetlony.

Kalibracja

Naciśnięcie przycisku „Kalibracja” spowoduje rozwinięcie okna w celu wyświetlenia wartości kalibracji (patrz Ekran 2). Nasza kopia Processing zatrzymuje się na kilka sekund, gdy tak się dzieje; jest to znany błąd, który, miejmy nadzieję, zostanie naprawiony w późniejszej wersji. Aby zamknąć widok kalibracji, kliknij w dolnej części okna.

Kalibracja odbywa się w dwóch etapach. Pierwszym z nich jest kalibracja napięcia, która wymaga podłączenia woltomierza do wyjścia zasilacza (CON1).

Włącz wyjście i ustaw prąd na dowolną wartość powyżej zera; ma to na celu zapewnienie, że ograniczenie prądu nie zadziała, co zmniejszyłoby napięcie wyjściowe.

Następnie wyreguluj suwak napięcia, aż multimetr wskaże wartość jak najbliższą 6 V. Do tego celu idealnie nadaje się zewnętrzny zasilacz wtyczkowy 12...15 V DC, ale nawet zasilacz 9 V DC będzie wystarczający. Należy pamiętać, że dwa wskaźniki mogą nie pokrywać się z napięciem 6 V. Jest to oczekiwane, ponieważ nadal kalibrujemy urządzenie.

Teraz zapisz wartości „VFACTOR” i „UFACTOR”, które są wyświetlane na dolnym panelu.

Abyskalibrowaćstronęprądową,wyłączwyjście i przełącz multimetr w tryb amperomierza i zakres umożliwiający odczyt do 400 mA. Konieczna będzie również zmiana sposobu podłączenia przewodów miernika.

Ponieważ multimetr skutecznie tworzy zwarcie, można dołączyć szeregowo z przewodami multimetru rezystor mocy w celu dodatkowej ochrony i zmniejszenia rozpraszania mocy w tranzystorze wyjściowym. Na przykład, dobrze sprawdzi się rezystor 10 Ω 5 W.

Włącz wyjście i przesunąć wskaźnik prądu w prawo, aż multimetr odczyta 300 mA, a następnie zanotuj dolne („IFACTOR” i „JFACTOR”) wartości kalibracji i wyłącz wyjście. Zrób to szybko, ponieważ tranzystor może się dość mocno nagrzać na tym etapie.

Konfiguracja

Współczynniki kalibracji (wraz z innymi ustawieniami) są przechowywane w pliku o nazwie „config.txt”. Musi on znajdować się w tym samym folderze/katalogu co plik .pde dla szkicu Processing. Otwórz plik „config.txt” i dodaj lub zmodyfikuj cztery zapisane współczynniki kalibracji. Wynik powinien wyglądać tak, jak pokazano na Ekranie 3. Należy pamiętać, że aplikacja nie rozróżnia w tych ustawieniach dużych lub małych liter.

Będziesz musiał ponownie uruchomić program, aby załadować nową konfigurację. Jeśli uruchamiasz go z Processing IDE (a nie z wyeksportowanej aplikacji), powinieneś zobaczyć, że kalibracje zostały załadowane i pokazane w oknie logów na dole, np. tak:

UFACTOR (ustawienie)

VFACTOR (ustawienie)

JFACTOR (ustawienie)

IFACTOR (ustawienie)

Jeśli nie są one widoczne, może to oznaczać, że wystąpił błąd i wartości nie zostały załadowane.

Plik konfiguracyjny obsługuje również inne opcje. SFACTOR jest używany do obliczania VIN; teoretycznie (w granicach tolerancji

komponentów) jest to ten sam dzielnik, co dla UFACTOR, więc można tu również użyć wartości UFACTOR. Jest on używany tylko do wyświetlania i obliczeń rozpraszania mocy, więc nie jest tak krytyczny jak inne wartości.

Jest to prosty współczynnik skalowania z przeliczenia surowego wyniku ADC (0–1023) na napięcie, więc można go również dostosować, porównując z odczytem multimetru. Na przykład, jeśli wyświetlane napięcie zasilania jest o 1% za niskie, należy zwiększyć SFACTOR o 1%.

Za pomocą parametrów PORTNAME i CONNECT można również ustawić domyślny port szeregowy i określić, czy ma się on łączyć po uruchomieniu programu. Nominalne napięcie zasilania można również podać za pomocą parametru VIN.

Parametr PORTNAME należy ustawić przed linią CONNECT, aby otworzyć właściwy port. Schemat nazewnictwa portów różni się w zależności od systemu operacyjnego.

Pięć ustawień wstępnych ustawia się za pomocą PRESET1 do PRESET5, z wartościami napięcia (w woltach), prądu (w amperach) i nazwy (skrótowej do siedmiu znaków). Parametry te są oddzielone przecinkami.

Oczywiście wszystkie zmienne konfiguracyjne mają rozsądne wartości domyślne na wypadek braku lub pustego pliku konfiguracyjnego. Pozostawiliśmy kilka potencjalnych linii w pliku poprzedzonych apostrofem; program ignoruje te linie, dopóki nie usuniesz apostrofów.

Użycie

Aplikacja do sterowania zasilaczem została zaprojektowana tak, aby korzystanie z niej było intuicyjne. Uważamy, że w ten sposób jest ona znacznie łatwiejsza w użyciu niż zasilacz z fizycznymi elementami sterującymi, takimi jak kilka potencjometrów i mały wyświetlacz.

Nie jest to w żadnym wypadku sprzęt diadnostyczny o wysokiej dokładności, ale nadal jest bardzo przydatny na biurku, zwłaszcza, że nie zajmuje dużo miejsca.

Nie opisałyśmy, jak zamontować go w jakiegokolwiek obudowie, ponieważ tak naprawdę można go po prostu używać tak, jak jest.

Jeśli chcesz go obudować, UB3 Jiffy Box jest najprostszą i najtańszą opcją, a duży rozmiar tej obudowy powinien zapewnić pewien przepływ powietrza do chłodzenia. Para otworów na każdym końcu będzie wystarczająca do przeprowadzenia wszystkich niezbędnych przewodów. ■

Tim Blythman

Adaptacja do wydania
polskiego – Andrzej Nowicki

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

Elektroniczne dzwonki wietrzne, część 2 – wykończenie

W zeszłym miesiącu opisaliśmy, jak działa nasz nowy elektroniczny dzwonek wietrzny i jak zbudować elektronikę. Teraz przechodzimy do najtrudniejszej części – modyfikacji samego dzwonka wiatrowego, aby mógł być uruchamiany przez zestaw elektromagnesów. Nie obawiaj się, ponieważ mamy szczegółowe instrukcje, jak to zrobić i zakończyć budowę, składając wszystko razem i konfigurując elektronikę.



Nasz gotowy elektroniczny dzwonek wietrzny. Wykorzystuje komercyjny zestaw dzwonków wiatrowych, ale nasz działa, gdy nie ma wiatru

Zmodyfikowaliśmy dzwonek Sonnet Wind Chime „Amazing Grace” 640 mm firmy Carson Home Accents tak, aby zawierał napęd elektromagnetyczny. Jeśli będziesz miał kłopoty z jego nabyciem, nasze wskazówki powinny wystarczyć Ci do podobnych zmian w budowie analogicznego. Zawsze możesz próbować zakupu na znanych portalach.

Jest to typ 5-dzwonkowy z rurami o średnicy zewnętrznej 31,5 mm. Najdłuższa rura ma 590 mm, a najkrótsza 450 mm.

Elektromagnesy są zamontowane na okrągłym pierścieniu wykonanym z 9-milimetrowej płyty MDF (płyta pilśniowa o średniej gęstości). Pierścień ten jest utrzymywany w miejscu za pomocą elementu w kształcie odwróconej litery „U” wykonanego z płyty MDF i kilku wsporników zamontowanych pod kątem prostym. Cała rama jest przymocowana do haka zawieszki gongu wiatrowego za pomocą śruby M5 i nakrętki.

W przypadku naszego prototypu płytka wspornika młoteczka została wykonana z kawałka blachy aluminiowej o średnicy 80 mm i grubości 1 mm. Płytkę (pokazana na **rysunku 6**) została zaprojektowana tak, aby współpracować z 5 dzwonekami rozmieszczonymi z odstępem kątowym 72° wokół osi.

Płytkę zawiera otwory na sznurki i szczylinę umożliwiającą umieszczenie młoteczka na płytce wspornika, bez konieczności demontażu żagielka (żagielek przymocowany jest na końcu tego samego sznurka).

Rama musi być tak dobrana, aby płyta podstawy mogła być umieszczona na wysokości, na której elektromagnesy i dźwignie znajdują się w jednej linii z górną krawędzią młoteczka.

Istnieją dwa otwory na sznurek wiążące każdy dzwonek (rurkę) gongu z elektromagnesem. Muszą one znajdować się na tyle daleko od siebie, aby sznurek po naciągnięciu nie dotykał rurki gongu. Płytkę wspornika młoteczka po przewleczeniu sznurka może być przyklejona na miejscu do samego młoteczka

lub przymocowana za pomocą małej śruby samogwintującej.

Prostokątne dźwignie elektromagnesu o wymiarach 100×10 mm są wykonane z blachy aluminiowej o grubości 1 mm; dwa końcowe otwory mają średnicę po 3 mm. Należy zauważyć, że dwa otwory nie są wyśrodkowane, ale umieszczone blisko siebie, aby zapewnić wystarczający zasięg ruchu obrotowego po przymocowaniu do trzpienia elektromagnesu.

Punktem obrotu jest wkręt do drewna w płycie podstawy. Śruba powinna być wystarczająco długa i wkręcona w taki sposób, aby dźwignia mogła znajdować się w pozycji poziomej, a jednocześnie nie była zbyt mocno dokręcona, patrz dalej.

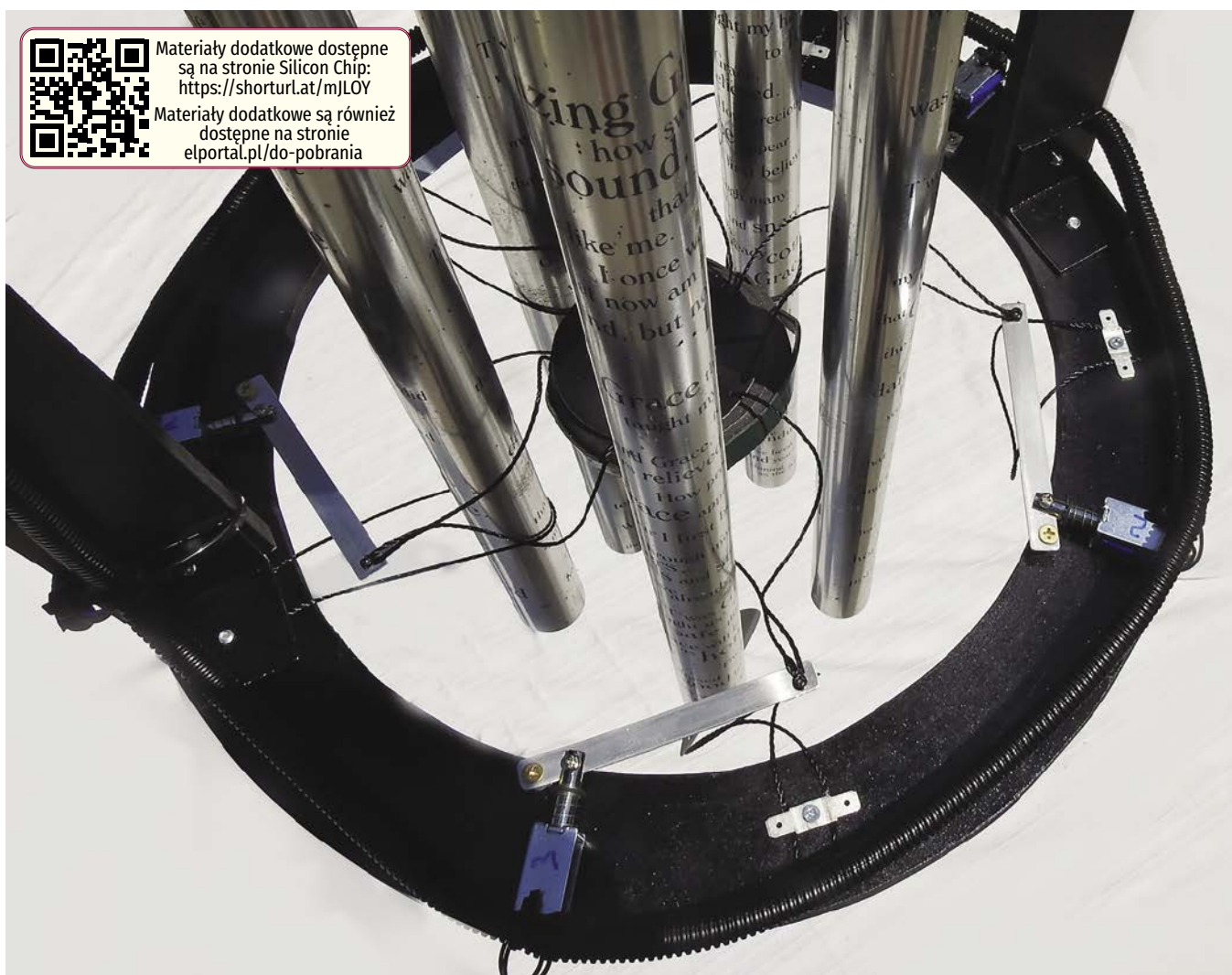
Otwór w trzpieniu elektromagnesu został wywiercony wiertłem 2,5 mm, a następnie nagwintowany pod śrubę M3. Umożliwia to zamocowanie dźwigni w punkcie podparcia za pomocą śruby M3 o długości 10 mm i bez nakrętki, przy czym śruba działa jak łożysko. Alternatywnie można wywiercić otwory o średnicy 3 mm a połączenie wykonać gwintowanymi śrubami i nakrętkami.

Otwór obrotowy na końcu dźwigni jest lekko powiększony o około 1 mm, aby dźwignia mogła się swobodnie poruszać, umożliwiając zmianę odległości między śrubami, gdy porusza się wraz z ruchem trzpienia. Tulejki bez gwintu o długości 6,3 mm utrzymują dźwignie elektromagnesów w pozycji odsuniętej od płyty bazowej – pierścienia, aby mogły poruszać się w płaszczyźnie poziomej. Całość skręcona jest wkrętami do drewna 4,5 mm × 15 mm z łbem stożkowym, wkręconymi w płytę podstawy. Jeśli użyjesz elektromagnesu o innych wymiarach obudowy, dopasuj do nich długość tulejki dystansowej.

Elektromagnesy są mocowane za pomocą śrub w obudowie elektromagnesu. Nasze mają otwory montażowe z gwintem M2,5, więc są mocowane za pomocą śrub M2,5×12. Jeśli elektromagnesy nie mają otworów, można je przykleić do podłoża bezpośrednio.



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



Zbliżenie na „zmechanizowany” koniec elektronicznego dzwonka wiatrowego, pokazujące jak są umieszczone elektromagnesy wokół pierścienia. Elektromagnesy nie uderzają w rurki gongu; a ciągną młoteczek w kierunku rurki dzwonka, która wydaje dźwięk. Na tym zdjęciu usunięto niektóre sznurki blokady i sznurki pociągające płytkę wspornika młoteczka, aby widok był klarowniejszy. Jednak uważny Czytelnik bez trudu zauważy, że sznurki przyciągające młoteczek umieszczone są, prawdę mówiąc, niepoprawnie, niezgodnie z dolną częścią Rysunku na następnej stronie. Powinny one obejmować rurę dzwonka, a nie znajdować się po jednej stronie

Inne opcje

Płytkę wspornika młoteczka i dźwignie mogą być wykonane z materiału innego niż aluminium. Dźwignie muszą być wystarczająco cienkie, aby swobodnie obracać się w szczelinie trzpienia elektromagnesu.

Łatwiejszym materiałem do obróbki jest preszpan lub podobny materiał do izolacji elektrycznej, taki jak Jaycar HG9985. Można go ciąć nożyczkami i ostrym nożem modelarskim. Możesz się tu wykazać własną inwencją i użyć np. tworzywa ze starych opakowań płyt CD, lub nawet samych płyt CD.

Podane rozmiary drewnianej ramy i podstawy są orientacyjne; zależą one od używanego dzwonka wietrznego.

Okrągła płyta podstawy w kształcie pierścienia musi mieć wystarczająco duży otwór wewnętrzny, aby rurki gongu mogły się swobodnie poruszać bez uderzania w nią.

Średnica zewnętrzna (szerokość pierścienia płyty podstawy) musi być wystarczająca do zamocowania elektromagnesów, z miejscem na dokręcenie śrub.

Do wykonania ramy w kształcie litery U i płyty podstawy użyliśmy płyty MDF, ale zamiast tego można wykonać ramę z litego drewna. Płyta podstawy nie musi być okrągła – może mieć kształt wielokąta.

Liczba prostych boków może być równa liczbie rurek dzwonek gongu; w przypadku naszego pięciorurowego gongu byłby to pięciokąt.

Należy pamiętać, że po zamontowaniu elektromagnesów i dźwigni niekoniecznie istnieje dogodny punkt do przymocowania ramy do całego zestawu dzwonek wiatrowych, w którym nie będzie ona kolidować z co najmniej jedną z dźwigni. Jest to szczególnie dokuczliwe w przypadku nieparzystej liczby elektromagnesów.

Dodatkowo jedna strona ramy powinna być bezpośrednio przymocowana do płyty bazowej. Druga strona ramy może być podparta za pomocą wspornika, który jest uniesiony nad płytą bazową za pomocą śruby i nakrętek, aby umożliwić ruch dźwigni (zobacz nasze zdjęcia i rysunki, aby uzyskać szczegółowe informacje).

Wyrównanie

Rama w kształcie litery U musi być prawidłowo wyrównana z płytą podstawy. Ma to na celu zapewnienie, że podczas swobodnego zwisu gongu zamocowanego na haku, dźwignie elektromagnesu i sznurki są prawidłowo ustawione, tak, aby płytkę wspornika młoteczka była pociągana wzdłuż promienia od środka płytki do środka rurki dzwonka dla każdego elektromagnesu.

Jeśli nie jest możliwe uzyskanie takiego wyrównania bez kolizji ramy w kształcie litery



Obracanie umieszczonego na górze „talerzyka” pozycjonującego sznurki spowoduje skręcenie linek na haku mocującym, więc nie pozostanie on w tej zmienionej, obróconej pozycji. Rozwiązaniem jest przywiązanie „talerzyka” do boku ramy. Mały otwór w boku „talerzyka” i drugi w ramie w kształcie litery U pozwolą na użycie krótkiego sznurka lub sztywnego drutu do przytrzymania „talerzyka” w jego obróconej pozycji.

Sznurki pociągowe muszą być zwykle luźne. Pociągają one młoteczek w kierunku rury gongu pod koniec skoku dźwigni.

Luźny naciąg wynika z dwóch powodów: po pierwsze, siła naciągu elektromagnesu na początku jego ruchu z pozycji spoczynkowej nie jest szczególnie duża i jest największa, gdy cewka całkowicie wciąga trzpień. Luźny naciąg pozwala cewce „nabrać siły”, zanim zacznie poruszać młoteczką.

Drugim powodem jest to, że gdy jeden elektromagnes ciągnie płytkę podstawy młoteczka w swoim kierunku, nie ma to wpływu na naprężenie sznurków po przeciwnej stronie. Poluzowanie musi być kompromisem między naciągami wystarczająco ciasnym, aby móc przyciągnąć młoteczek do rurki dzwonka gongu, a wystarczająco luźnym, aby nie wpływać na działania elektromagnesu po przeciwnej stronie. Ponadto chcielibyśmy, aby młoteczek mógł być nadal uruchamiany przez wiatr.

Sznurki przechodzą przez otwory płytki podstawy młoteczka i z powrotem do dźwigni, a następnie są zabezpieczone przez przełożenie sznurka przez otwór dźwigni. Do zamocowania sznurka w otworze można użyć śruby M3x6 i nakrętki M3. Pozwala to na łatwiejszą regulację w porównaniu do wiązania węzła.

Udoskonaleniem konstrukcji jest zastosowanie dodatkowych zaczepów, też ze sznurka. Łapią one i przytrzymują rurkę gongu, zapobiegając jej odchyleniu się do tyłu, aby zapobiec ponownemu uderzeniu w młoteczek po wcześniejszym uderzeniu w rurkę dzwonka. Ich długość jest taka, że są luźne, gdy rurka znajduje się w normalnej pozycji, ale są wystarczająco ciasne, aby zapobiec jej odchyleniu do środka i uderzeniu w młoteczek. Końce sznurka są przymocowane do płyty podstawy za pomocą zacisków.

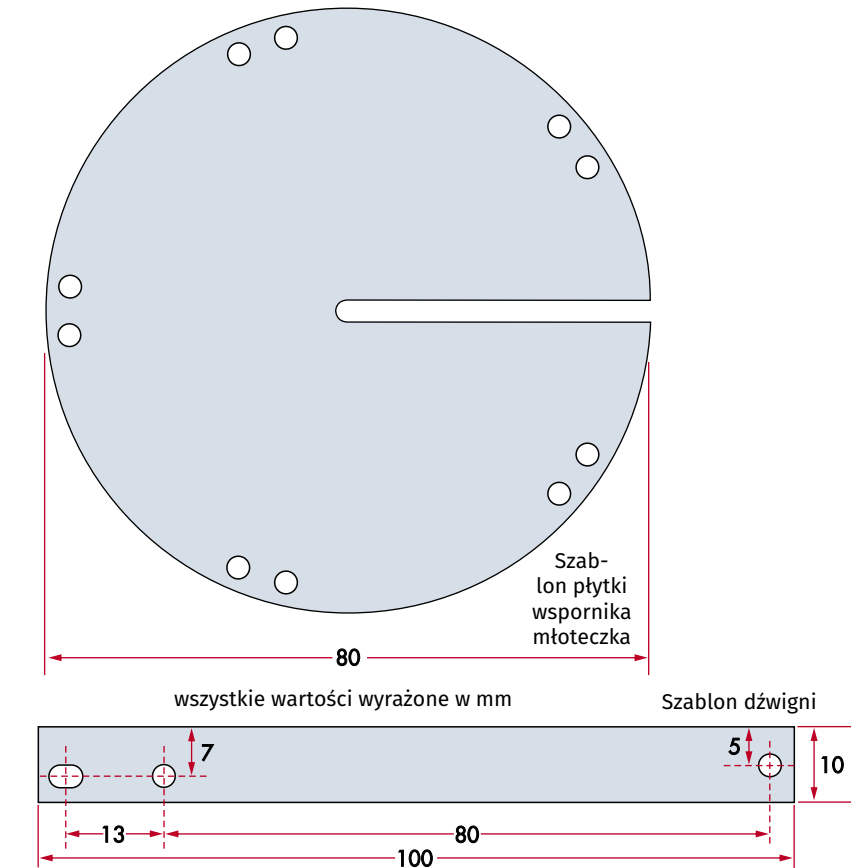
Użyliśmy sznurka poliestrowego, który rozplata się przy cięciu nożyczkami lub nożem. Zamiast tego sznurek został „przycięty” na odpowiednią długość za pomocą gorącej końcówki

WIDOK Z GÓRY

lutownicy, która zarówno przecięła, jak i zgrzała końce sznurka, aby zapobiec strzępieniu. Nie zalecamy jednak używania do tego podstawowej, wysokiej jakości grota lutownicy (szkoda byłoby go zabrudzić). Możesz zamiast tego przeciąć sznurek, a następnie użyć zapalniczki, aby zgrzać końce, zanim się rozplotą. Dawno temu takie sznurki były stosowane do sterowania żaluzjami międzyokiennymi. Może znajdziesz wystarczającą kawałek w swoich zapasach.

Okablowanie

W przypadku większych elektromagnesów należy użyć przewodu o odpowiedniej średnicy (np. linki składającej się z 19 rdzeni o średnicy 0,1 mm każdy) lub podobnego, aby spadki napięcia nie wpływały na działanie elektromagnesu. Jeśli przekrój przewodu jest zbyt mały, elektromagnesy mogą nie działać przy dłuższych przewodach idących do głównej płytki drukowanej. Aby utrzymać przewody na miejscu i zachować schludny wygląd, użyliśmy peszla o średnicy 7 mm. Przewody +12 V do każdego elektromagnesu są połączone razem i doprowadzone z powrotem do dodatniego zacisku CON1 lub CON6. Drugi przewód każdej cewki łączy się między wyjściami cewek na CON1-CON6 a odpowiednim ujemnym zaciskiem cewki – zobacz schemat elektryczny i opis układu w poprzedniej części.



Rysunek 6. Przycięliśmy blaszkę aluminiową do pokazanego kształtu i przykręciliśmy ją do górnej części naszego drewnianego młoteczka, aby umożliwić łatwe przymocowanie pięciu sznurków

Tabela 1. Działania przetącznika przy włączaniu zasilania	
S1	Losowość wyłączona
S2	Losowość włączona
S3	Opóźnienie zmienia się w 128 krokach od 10 s do 1280 s (21:20)
S4	Opóźnienie zmienia się w 64 krokach od 10 s do 640 s (10:40)
S5	Opóźnienie zmienia się w 32 krokach od 10 s do 320 s (5:20)
S6	Opóźnienie zmienia się w 16 krokach od 10 s do 160 s (2:40)
S7	Opóźnienie zmienia się w ośmiu krokach od 10 s do 80 s (1:20)
S8	Opóźnienie zmienia się w czterech krokach od 10 s do 40 s (0:40)
S9	Mnożnik opóźnienia zmienia się losowo w zakresie od jednej do pięciu wartości rzeczywistych
S10	Mnożnik opóźnienia zmienia się losowo w zakresie od jednego do trzech rzeczywistych wartości
S11	Mnożnik opóźnienia zmienia się losowo w zakresie od jednego do dwóch rzeczywistych wartości
S12	Mnożnik opóźnienia zmienia się losowo w zakresie od 1 do 1,5-krotności rzeczywistego opóźnienia

Po przylutowaniu przewodów elektromagnesu do przewodów przedłużających, zaizoluj połączenia taśmą izolacyjną lub rurką termokurczliwą. Po zakończeniu przymocowaliśmy wiązkę przewodów do górnej części każdego elektromagnesu za pomocą opasek kablowych, aby się nie poruszała. Główna obudowa zawierająca płytkę drukowaną może być umieszczona na drewnianej belce powyżej mocowania dzwonka lub dalej, poza zasięgiem wzroku.

Konfiguracja

Istnieje kilka opcji, które należy ustawić w sterowniku elektronicznych dzwonek wiatrowych, zanim będzie można z niego korzystać.

Regulacja fotorezystora

Jeśli chcesz, aby dzwonki wiatrowe działały także w ciemności, umieść zworke na JP2. W takim przypadku nie trzeba instalować fotorezystora LDR. Ale jeśli chcesz, aby zatrzymać działanie w nocy, usuń zworke z JP2 i włącz zasilanie. Dioda LED2 powinna się zaświecić, wskazując, że jest zasilanie. Potencjometr nastawny VR2 można następnie wyregulować, aby ustawić próg natężenia światła, który włącza lub wyłącza elektroniczny dzwonek wiatrowy. Przy normalnej intensywności światła dziennego umieść palec nad fotorezystorem i wyreguluj VR2 tak, aby dioda LED1 (dioda LED stanu) zaczęła migać z częstotliwością 2 Hz. Oznacza to, że odtwarzanie jest wstrzymane. Podniesienie palca z fotorezystora powinno spowodować wyłączenie tej diody LED. Im bardziej w prawo skrócony jest VR2, tym ciemniej musi być w okolicy fotorezystora, aby wstrzymać działanie.

Kalibracja

Fotorezystor LDR jest ignorowany podczas kalibracji i nagrywania. Jest on używany tylko podczas odtwarzania i tylko wtedy, gdy mostek JP2 jest otwarty. Dzięki temu kalibracja i nagrywanie nie są przerywane przez zmianę poziomu oświetlenia. Każda cewka może być skalibrowana niezależnie pod kątem napięcia zasilającego (przy użyciu PWM) i okresu włączenia. Te dwa parametry są regulowane za pomocą VR1 i JP1, jak opisano poniżej. Wypełnienie impulsów PWM 500 Hz można regulować w zakresie od około

5% do 100% w krokach po około 0,75%. Powoduje to zmianę średniego napięcia w zakresie od 600 mV do 12 V w krokach po około 90 mV. Okres włączenia można ustawić w zakresie od 2 ms do 254 ms w krokach po około 2 ms. Początkowo wszystkie elektromagnesy otrzymują pełne napięcie 12 V (100% wypełnienia) przez 254 ms.

Aby zainicjować kalibrację, naciśnij i przytrzymaj przełącznik sterowania (S13) po włączeniu zasilania.

Dioda LED stanu (LED1) zaświeci się na 200 ms, a następnie zgaśnie na 200 ms i ponownie się zaświeci. Oznacza to, że kalibracja została aktywowana.

Naciśnij przełącznik cewki (S1-S12), aby wybrać elektromagnes, który ma zostać skalibrowany. Dioda LED stanu zgaśnie, a parametry napędu elektromagnetycznego są teraz gotowe do regulacji dla wybranej cewki. Gdy JP1 jest zwarty, stopień wypełnienia PWM można regulować za pomocą VR1, a gdy JP1 jest otwarty, czas trwania zasilania elektromagnesu (okres włączenia) jest regulowany za pomocą tego samego potencjometru nastawnego VR1.

Po ustawieniu JP1 i wyregulowaniu VR1 dla ustawienia, które chcesz wprowadzić, naciśnij przełącznik sterowania (S13), aby tymczasowo zapisać ten konkretny parametr. Spowoduje to również uruchomienie odpowiedniej cewki, dzięki czemu można sprawdzić, czy ustawienie jest prawidłowe. Jeśli nie, należy ponownie wyregulować VR1 i ponownie nacisnąć S13.

Jeśli chcesz, aby inna cewka miała ten sam parametr, możesz nacisnąć przełącznik (S1-S12) dla tej cewki i ponownie nacisnąć przełącznik sterowania (S13), aby zapisać bieżącą wartość parametru dla tej cewki.

Zapewniliśmy również możliwość monitorowania bieżącego ustawienia VR1 za pomocą multimetru mierzącego napięcie między TP1 a TP GND. Ułatwia to odtworzenie odpowiednich wartości dla innych elektromagnesów.

Dioda LED stanu (LED1) świeci się po każdym naciśnięciu przełącznika sterowania przez cały czas trwania ruchu trzpienia elektromagnesu. Niższe wypełnienie PWM spowoduje wolniejszy ruch trzpienia w cewce. Wyreguluj okres włączenia elektromagnesu, aby zapewnić wystarczający czas na przyciągnięcie młoteczka do rurki dzwonka gongu, ale wystarczająco krótko, aby sznurek poluzował się, zanim rurka gongu powróci do swojego położenia po uderzeniu.

Jak wspomniano, parametry elektromagnesu są początkowo przechowywane tylko tymczasowo. Wartości zostaną utracone po wyłączeniu zasilania, chyba że zostaną zapisane w pamięci FLASH. Odbывается to również za pomocą przełącznika sterowania.

Krótkotrwałe naciśnięcie przełącznika sterującego testuje napęd elektromagnetyczny, a dłuższe naciśnięcie (jedna sekunda lub dłużej) spowoduje zapisanie wszystkich wartości parametrów elektromagnesu w pamięci stałej FLASH. Dioda LED1 zaświeci się ponownie, jeśli przełącznik zostanie przytrzymany przez jedną sekundę lub dłużej, aby wskazać, że wartości zostały zapisane w pamięci FLASH.

Aby wyjść z trybu kalibracji, wyłącz zasilanie. Po ponownym włączeniu zasilania, bez naciśnięcia przycisku S13, odtwarzacz Wind Chime Player uruchomi się w trybie odtwarzania.

Można ponownie powrócić do trybu kalibracji, powtarzając powyższą procedurę, aby ponownie

dostosować te parametry. Zmienione zostaną tylko parametry dla wybranej cewki lub cewek. Poprzednio zapisane parametry pozostaną niezmienione, chyba, że dla danego elektromagnesu zostaną ustawione nowe parametry.

Nagrywanie sekwencji

Aby wykonać nagrywanie, naciśnij przełącznik sterujący S13 po włączeniu zasilania. Dioda LED stanu, LED1, zaświeci się i pozostanie zapalona, wskazując, że rozpoczęło się nagrywanie.

Następnie można nacisnąć poszczególne przełączniki elektromagnesów, aby je aktywować, a urządzenie rejestruje podaną sekwencję i przerwy między zadziałaniem poszczególnych elektromagnesów. Można uruchomić tylko jeden elektromagnes na raz.

W obszarach białych, wypełnionych prostokątów na warstwie opisowej płytki PCB, znajdujących się nad każdym przyciskiem można opisać działanie poszczególnych przycisków (generowany dźwięk) za pomocą cienkiego markera. Mówimy o wybrzmialeń nucie (dźwięku), ponieważ dźwięk gongu zawiera wiele poddźwięków (harmonicznych), które mogą wpływać na pozorną częstotliwość. Może się również wydawać, że częstotliwość zmienia się po wcześniejszym uderzeniu.

Nie da się łatwo zmierzyć wybrzmiewającego dźwięku za pomocą analizatora widma. Prawdopodobnie najłatwiejszą metodą jest użycie stroika gitarowego lub podobnego urządzenia i dostrojenie go do porzecznej częstotliwości odpowiadającej danemu dzwonnkowi, a następnie sprawdzenie, jaka częstotliwość odpowiada brzmieniu rurki dzwonka gongu.

Oto płytkę PCB elektronicznego dzwonka umieszczoną wewnątrz obudowy, aczkolwiek bez podłączonych kabli, podczas gdy po prawej stronie znajduje się panel przedni i etykieta



Więcej informacji na temat percepcji dźwięków wydawanych przez dzwonki wietrzne można znaleźć na stronach www.leehite.org/Chimes.htm#The%20strike%20note i www.sarahtulga.com/Glock.htm. Jeśli masz wyuczucie muzyczne, podczas nagrywania możesz odtworzyć melodię, lub po prostu zapisać ładne dźwięki, które Ci się podobają. Krótkie przerwy między uderzeniami gongu można przeczekać w czasie rzeczywistym przed uruchomieniem cewki kolejnego dzwonka. Dłuższe przerwy mogą stać się uciążliwe do odczekania w czasie rzeczywistym, ale mamy na to rozwiązanie...

Sterowanie czasem spoczynku

Naciśnięcie przełącznika sterującego na dłużej niż jedną sekundę powoduje pomnożenie okresu zapisanego dla bieżącej pauzy przez dziesięć. Dioda LED stanu miga z częstotliwością 1 Hz, aby odmierzyć czas (jedno mignięcie oznacza jedną sekundę czasu rzeczywistego, ale dziesięć sekund opóźnienia). Zachowaj ostrożność podczas naciskania przycisku S13, ponieważ jeśli naciśniesz go na krócej niż jedną sekundę, zamiast aktywować zarządzanie przerwą czasową, nagrywanie zakończy się.

Po krótkim naciśnięciu przełącznika sterującego wprowadzona sekwencja zostanie zapisana w pamięci FLASH i nastąpi powrót do trybu odtwarzania. Jeśli w trybie nagrywania nie został naciśnięty żaden przełącznik elektromagnesu, poprzednie nagranie pozostanie w pamięci.

Odtwarzanie

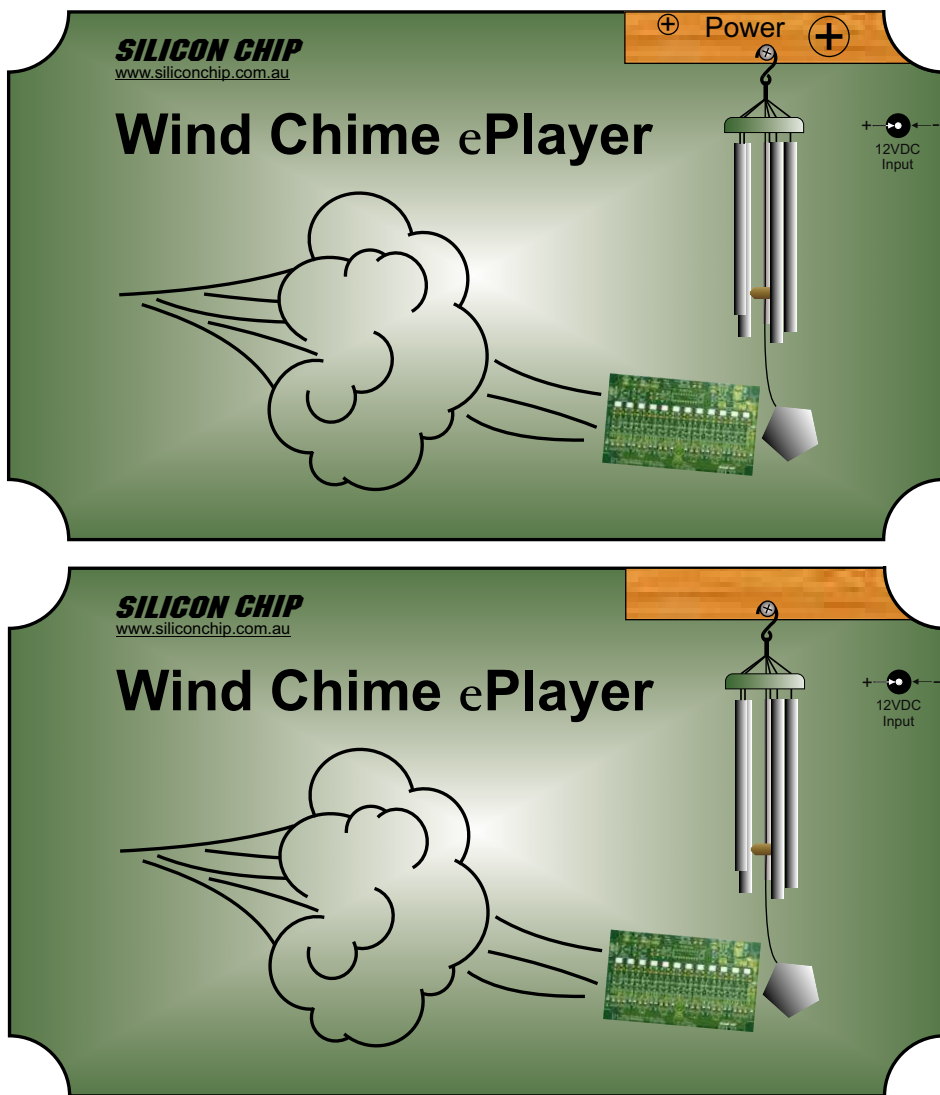
Po włączeniu zasilania elektroniczny dzwonek uruchamia się w trybie odtwarzania. Odtwarza on nagraną sekwencję, powtarzając ją w ciągłej pętli. Ustawieniem początkowym jest brak losowości w okresach opóźnienia między uderzeniami młoteczka – innymi słowy, wierne odtwarzanie nagranej sekwencji.

Dodawanie losowości

Jak wspomniano wcześniej, można dodać losowość do opóźnienia między uderzeniami młoteczka. Wyboru dokonuje się poprzez naciśnięcie przełącznika S2 podczas włączania zasilania. Przed zwolnieniem S2 poczekaj, aż dioda LED stanu (LED1) zacznie migać po około jednej sekundzie nacisku, wskazując, że funkcja losowości została włączona.

Ustawienie jest zapisywane w pamięci stałej. Jeśli chcesz wyłączyć losowość, przytrzymaj przełącznik S1 po włączeniu zasilania i poczekaj, aż dioda LED stanu zaświeci się, przed jego zwolnieniem.

Istnieją dwa parametry losowości, które można ustawiać. Jednym z nich jest szybkość, czyli jak często zmienia się wartość losowa. Można ustawić sześć różnych wartości.



Dostępne są dwa projekty paneli przednich – jeden z nich posiada przelotowy włącznik zasilania i diodę LED, podczas gdy drugi panel ich nie posiada. Szablony można również pobrać ze strony [siliconchip.com.au](http://www.siliconchip.com.au)

Losowość zmienia się w odstępach od dziesięciu sekund do wybranej wartości maksymalnej. Dostępne opcje to 1280s (21:20), 640s (10:40), 320s (5:20), 160s (2:40), 80s (1:20) i 40s (0:40).

Opcje te są wybierane poprzez przytrzymanie jednego z przełączników S3, S4, S5, S6, S7 i S8 podczas włączania zasilania – patrz tabela 1.

Można również zmienić pożądaną zmienność opóźnień. Dostępne są cztery opcje, wybierane przez przytrzymanie jednego z przełączników S9, S10, S11 lub S12 podczas włączania zasilania.

Mnożnik opóźnienia zmienia się losowo w zakresie od jednego do wybranej wartości maksymalnej. S9 wybiera zakres 1-5 razy, S10 1-3 razy, S11 1-2 razy i S12 1-1,5 razy (patrz również tabela 1).

Jeśli żaden z tych przełączników nie został jeszcze naciśnięty podczas włączania zasilania, to początkowym ustawieniem jest wyłączenie

losowości. Jeśli losowość jest włączona (za pomocą S2), wybierana jest szybkość zmiany losowości od 10 s do 1280 s (21:20) oraz zakres opóźnień od 1 do 5 razy.

Należy pamiętać, że można nacisnąć i przytrzymać więcej niż jeden przełącznik po włączeniu zasilania, aby wybrać więcej niż jedną opcję jednocześnie.

Na przykład, możesz włączyć losowość (za pomocą S2), ustawić szybkość zmiany losowości na maksymalnie 320 sekund za pomocą S5, a zmienność losowości od jednego do trzech razy za pomocą S10, wszystko w tym samym czasie. ■

John Clarke

Adaptacja do wydania polskiego – Andrzej Nowicki

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



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://shorturl.at/ptGJO>

Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania



Dbaj o swoje akumulatory. Wysokoprądowy balanser baterii, część 2 – budowa

Nasz nowy wysokoprądowy balanser baterii, którego opis rozpoczęliśmy w zeszłym miesiącu, to zaawansowana konstrukcja, która zapewnia wysoką wydajność i szybkie równoważenie poprzez wydajne przenoszenie ładunku między podłączonymi ogniwami lub akumulatorami. Może obsługiwać ogniwa lub akumulatory o napięciu do 16 V na jedną sekcję, a w przypadku większych instalacji można połączyć ze sobą dwie jednostki. Ten drugi i ostatni artykuł opisuje etapy montażu i testowania oraz sposób korzystania z urządzenia.

Włożyliśmy wiele wysiłku w to, aby projekt był jak najprostszy, a jednocześnie zapewniał doskonałą wydajność i wiele przydatnych funkcji.

W rezultacie liczba podzespołów nie jest szczególnie duża. Musieliśmy jednak użyć głównie części SMD, aby zachować rozsądny wymiar PCB, a także dlatego, że wiele najlepszych podzespołów nie było w ogóle dostępnych w wersji do montażu przewlekane.

Chociaż montaż płytki nie jest zbyt trudny, nie jest to jednak zadanie odpowiednie dla początkujących. Pożądane jest pewne doświadczenie w lutowaniu elementów SMD, zwłaszcza, że elementy są montowane po obu stronach płytki.

Potrzebna będzie przyzwoita stacja lutownicza z kontrolowaną temperaturą (a najlepiej piec rozplwowy lub stacja lutownicza na gorące

powietrze), strzykawka z pastą topnikową, miedziana plecionka lutownicza, pęseta z cienką końcówką, lupa i silne źródło światła. **Od Red. EdW:** Autor artykułu delikatnie pomija fakt, że w amatorskim piecu rozplwowym nie uda się polutować opisanej płytki balansera, choćby ze względu na ilość podzespołów SMD i ich rozmieszczenie po obu stronach płytki. Co więcej, Autor nie udostępnia danych, które pozwoliłyby na wykonanie montażu płytki w profesjonalnej firmie, w tym kompletu plików Gerbera. Jeśli więc zdecydujesz się na wykonanie opisanego balansera, pozostaje Ci tylko żmudne lutowanie pojedynczych elementów SMD, w tym naprawdę miniaturowych rezystorów o rozmiarze 0603.

Żadna z części SMD nie jest szczególnie trudna do montażu, chociaż mniejsze sześciostykowe elementy w obudowach SOT-363

należą do tych trudniejszych, wraz z układami scalonymi w obudowach QSOP-16, które mają styki położone dość blisko siebie. Wreszcie, transformatory mogą stanowić pewne wyzwanie w tworzeniu dobrych połączeń lutowanych ze względu na ich wysoką masę termiczną. Jednak przy odrobinie ostrożności płytkę PCB można zmontować ręcznie.

Budowa

Wysokowydajny balanser akumulatorów jest zbudowany na czterowarstwowej płytce drukowanej o kodzie 14102211 i wymiarach 108×80 mm.

Szczegółowe informacje na temat tego, które części gdzie się znajdują, znajdziesz na schematach montażowych PCB, **rysunek 4a i 4b** na odwrocie. Sugerujemy rozpoczęcie montażu od lutowania elementów SMD na spodzie

plytki, następnie elementów SMD na górnej stronie, a na końcu elementów przewlekanych.

Jak wspomniano wcześniej, możesz użyć różnych metod montażu, w tym lutowania rozpliwowego (patrz wzmiankę powyżej) lub lutowania ręcznego. Opiszemy metodę lutowania ręcznego, ponieważ wymaga ona najmniejszej liczby specjalistycznych narzędzi z listy wymienionych powyżej.

Generalnie, procedura polega na umieszczeniu każdej części (z prawidłową orientacją dla części spolarizowanych, czyli prawie wszystkich układów scalonych, diod i MOSFET-ów) i przylutowaniu jednego wyprowadzenia. Następnie sprawdzasz wyrównanie pozostałych styków i ponownie pozycjonujesz część, topiąc lut i delikatnie poruszając elementem, jeśli nie jest idealnie wyrównany z polami lutowniczymi. Po wyrównaniu dobrym pomysłem jest dodanie pasty topnikowej do wszystkich styków, ponieważ znacznie zmniejsza to ryzyko pojawienia się „zimnych lutów”.

Następnie lutujesz pozostałe końcówki, odświeżasz początkowo przylutowany styk (jeśli dodałeś pastę topnikową, wystarczy dotknąć go grotem lutownicy), a następnie używasz miedzianej plecionki lutowniczej i topnika, aby usunąć wszelkie mostki, które mogły się utworzyć

Kolejność umieszczania komponentów nie jest krytyczna, ale uważamy, że najlepiej jest umieścić najtrudniejsze do montażu części po każdej stronie, aby nie musieć mieć do czynienia z sąsiednimi komponentami, które mogłyby przeszkadzać przy tej odpowiedzialnej czynności. Poniższa procedura wykorzystuje tę metodę.

Należy pamiętać, że rezystory SMD są zwykle oznaczone małym kodem na górze, który wskazuje wartość (np. 47 k Ω =473 [47 $\times$ 10³] lub 4702 [470 $\times$ 10²]). Aby zobaczyć te oznaczenia, prawdopodobnie będziesz musiał użyć lupy. Kondensatory ceramiczne SMD są zazwyczaj nieoznaczone, dlatego bardzo uważaj, aby ich nie pomylić.

Na koniec pamiętaj, że większość używanych elementów półprzewodnikowych jest wrażliwa na wyładowania elektrostatyczne (ESD) – szczególnie te w mniejszych obudowach. Dlatego podczas manipulowania nimi unikaj dotykania ich styków. Uziemiony antystatyczny pasek na nadgarstek zwykle zapewnia bezpieczną pracę z takimi elementami, istnieją też inne możliwości zapewnienia bezpieczeństwa ESD.

Szczegóły montażu na spodniej stronie

Zacznij od zamontowania ośmiu rezystorów 1 Ω sterujących bramkami, ponieważ są to najmniejsze elementy (rozmiar 0603)

pasywne na płytce i zasadniczo nie przeszkadzają przy montażu innych części. Następnie wlutuj osiem par komplementarnych tranzystorów N-MOSFET/P-MOSFET z bramkami sterowanymi tymi rezystorami: Q27, Q28, Q22, Q23, Q16, Q17, Q11 i Q12.

Są one stosunkowo duże jak na sześciostykowe układy SMD, więc nie powinny sprawiać większych problemów, ale musisz uważać na ich orientację! Możesz potrzebować lupy, aby znaleźć kropkę pinu 1 na górze każdego podzespołu, która po lutowaniu w każdym przypadku znajduje się w prawym dolnym rogu, jak pokazano na rysunku 4b i odpowiadającej mu fotografii.

Następnie zamontuj osiem kondensatorów 4,7 μ F, które sąsiadują z tymi podwójnymi MOSFET-ami. Potem zamontuj po tej stronie płytki pięć rezystorów 330 Ω oraz cztery rezystory 20 Ω , a następnie osiem kondensatorów 10 μ F obok pól montażowych dla MOSFET-ów Q1-Q4.

Elementy oznaczone jako „Rsnub” i „Csnub” są wymagane w przypadku stosowania akumulatorów 12 V, ale nie są potrzebne do równoważenia akumulatorów o niższym napięciu, takich jak ogniwa Liion/LiPo/LiFePO₄. Jeśli ich potrzebujesz, zamontuj je teraz, używając wartości sugerowanych na liście części opublikowanej w zeszłym miesiącu (30 Ω i 470 pF).

Teraz zainstaluj MOSFET-y Q1...Q5. Są one w obudowach SMD LFPAK56, które są podobne do 8-stykowych podzespołów SOIC, ale z radiatorem zastępującym cztery wyprowadzenia po jednej stronie. W związku z tym powinno być oczywiste, w którą stronę mają być skierowane, ale nie pomył BUK9Y8R5-80E użytego jako Q5 z podobnymi BUK9Y4R8-60E używanymi jako Q1...Q4.

W każdym przypadku rozprowadź niewielką ilość topnika na dużym polu lutowniczym przed przylutowaniem jednego z małych wyprowadzeń, a następnie przylutuj pozostałe trzy małe styki przed lutowaniem radiatora. Może być konieczne zwiększenie temperatury lutownicy, aby przylutować spore radiatory, ponieważ mają one dużą masę termiczną. Dodana wcześniej pasta topnikowa powinna pomóc we wpłynięciu lutowia pod radiator, zapewniając dobre połączenie termiczne i elektryczne.

Po ich przylutowaniu, zamontuj osiem pozostałych MOSFET-ów po tej stronie płytki, używając tej samej techniki. Wszystkie są typu BUK9Y14-80E (inny typ niż Q1...Q4 i Q5) i są oznaczone Q9, Q10, Q14, Q15, Q20, Q21, Q25 i Q26.

Teraz zamontuj cztery diody SMB TVS, ZD1-ZD4, upewniając się, że ich paski wskazujące katodę są zorientowane tak, jak pokazano na rysunku 4b. Pamiętaj, że napięcie

znamionowe tych części różni się w zależności od typu ogniwa lub akumulatorów (patrz wykaz elementów opublikowany w zeszłym miesiącu).

Przylutuj je podobnie jak elementy pasywne, ale ponieważ są większe, wymagają nieco więcej ciepła. Ich wyprowadzenia owijają się wokół boków, więc upewnij się, że lut przylega zarówno do płytki drukowanej, jak i do wyprowadzeń diod (topnik znacznie ułatwia osiągnięcie tego celu).

Kolejnym zadaniem jest wlutowanie czterech bezpieczników SMD 3 A, które zamontuj podobnie jak rezystory (nie są spolarizowane). Pozostają tylko dwa małe rezystory: jeden rezystor 100 k Ω 0,1% i drugi rezystor 0,1%, którego wartość zależy od zastosowania. Znajdują się one po prawej stronie tranzystora Q14. Upewnij się, że ich nie pomylisz.

Podzespoły SMD na górnej stronie

Odwróć płytkę i kontynuuj montaż, przylutowując cztery większe kondensatory ceramiczne o wymiarach 3,2 $\times$ 2,6 mm (SMD 1210) w pobliżu pól transformatorów T1-T4. Użyliśmy kondensatorów 4,7 μ F/75 V (TDK CNA6P1X7R2A475K250AE), ale można użyć większej pojemności ponad dwukrotnie, używając kondensatorów 10 μ F/75 V, które kosztują tylko trochę więcej (TDK CGA6P1X7R1N106K250AC).

Następnie zamontuj pięć małych podwójnych MOSFET-ów, Q8, Q18, Q13, Q19 i Q24. W każdym przypadku należy najpierw upewnić się, że kropka na styku 1 jest prawidłowo ustawiona.

Są one w mniejszych obudowach niż te, które zamontowałeś na spodzie PCB, z bliżej rozmieszczonymi stykami, więc mogą być nieco trudniejsze do lutowania. Ale nie jest to niewykonalne, o ile pamiętasz, aby dokładnie sprawdzić ew. nadmiarowe mostki między końcówkami za pomocą lupy i usunąć wszelkie znalezione zwarcia za pomocą topnika i miedzianej plecionki lutowniczej.

MOSFET Q7, w prawym dolnym rogu, posiada taką samą obudowę, co wspomniane pięć, ale jest to nieco inny podzespół, więc nie należy go pomylić. Ponownie powinieneś dokładnie sprawdzić jego orientację przed wlutowaniem go na miejsce.

Teraz jest dobry moment na zamontowanie mikrokontrolera IC2. Powinno to być stosunkowo łatwe w porównaniu do elementów, które już przylutowałeś, ale upewnij się przed przylutowaniem, że styki ze wszystkich czterech stron są wyrównane przed lutowaniem więcej niż jednego wyprowadzenia i jak zwykle uważaj, aby pin 1 znajdował się we właściwym miejscu.

Postępuj podobnie z czterema scalonymi izolatorami galwanicznymi: IC4, IC6, IC8 i IC10. W każdym przypadku pin 1 znajduje się

w lewym górnym rogu. Mają one podobny rozstaw końcówek jak małe podwójne MOSFET-y, które już zamontowałeś, więc ich montaż nie powinien być trudniejszy.

Następnie zamontuj osiem kondensatorów 470 nF, a następnie pięć regulatorów napięcia. W przypadku regulatorów rozprowadź niewielką ilość pasty topnikowej na dużym polu lutowniczym przed przyłutowaniem jednego z mniejszych styków, a następnie przyłutuj pozostałe małe wyprowadzenia przed przyłutowaniem dużych końcówek radiatorów.

Możesz potrzebować zwiększenia temperatury lutownicy podczas ich lutowania.

Po ich zamontowaniu możesz zamontować sześć kondensatorów SMD 1 μ F, dwa koraliki ferrytowe oraz rezystory SMD 680 Ω i 100 Ω . Następnie zamontuj dwa układy zabezpieczające przed wyładowaniami elektrostatycznymi, D6 i D10, które znajdują się w czterostykowych obudowach prawie przy samej dolnej krawędzi PCB. Jeden ze styków jest większy od pozostałych.

Sprawdź rysunek 4a i zweryfikuj ich orientację, jeśli masz wątpliwości.

Teraz zainstaluj osiem rezystorów 10 k Ω , a następnie pięć kondensatorów 1 nF, trzy 100 nF i trzy 10 μ F. Następnie zamontuj pozostały TVS (ten o wyższym napięciu, ZD5). Upewnij się, że jest prawidłowo ustawiony, paskiem katodowym do góry. Następnie zamontuj dwa bezpieczniki, przy czym bezpiecznik o niższym prądzie (0,75 A) to F7, w pobliżu 8-stykowego złącza CON15, a bezpiecznik F5 o wyższym prądzie (3 A) w pobliżu CON2 w lewym górnym rogu.

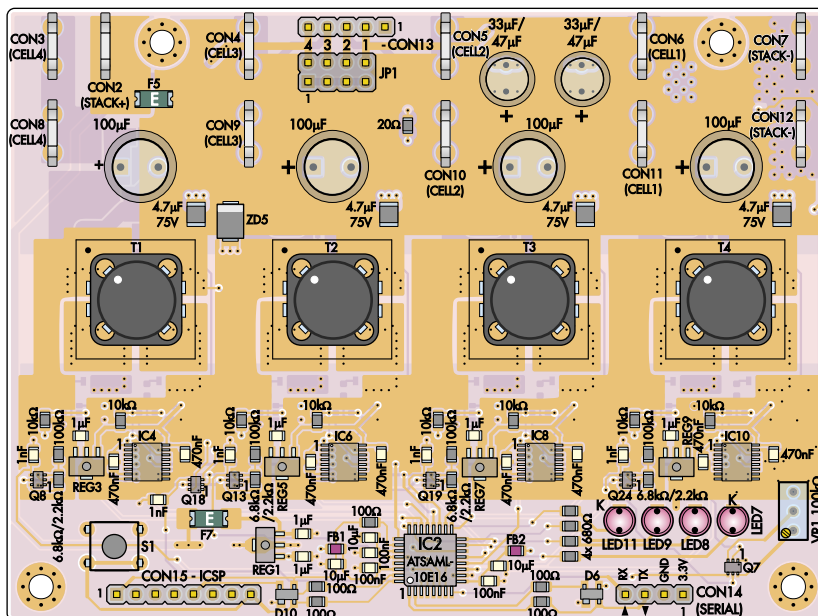
Jeśli chodzi o elementy pasywne, pozostaje tylko jeden rezystor 20 Ω w pobliżu CON10 oraz osiem rezystorów wysokiej precyzji 0,1%.

Jak wspomniano w zeszłym miesiącu, niższe wartości rezystorów 0,1% należy dobrać w zależności od napięcia akumulatora.

Górny rezystor w każdej parze ma wartość 100 k Ω . Upewnij się, że dolny rezystor ma wartość 6,8 k Ω dla całkowitego napięcia zestawu ogniwo o około 24 V lub 2,2 k Ω dla wyższych napięć zestawu.

Montaż transformatorów

Ze względu na znaczną masę termiczną transformatorów i duże płaszczyzny lutownicze, z którymi się łączą, zalecamy unikanie stosowania pasty lutowniczej do montażu tych części, chyba że posiadasz bardzo wysokiej jakości piec rozpliwowy (od Red. EdW: nie uda Ci się zamontować transformatorów w piecu rozpliwowym, jeśli przedtem zastosowałeś montaż ręczny, gdyż zniszczy to całą pracę. Lutowanie transformatorów w piecu rozpliwowym ma sens tylko wtedy, gdy cała płytka drukowana jest



TOP OF PCB

Rysunek 4a. Schemat montażowy górnej strony PCB, z pasującym zdjęciem poniżej



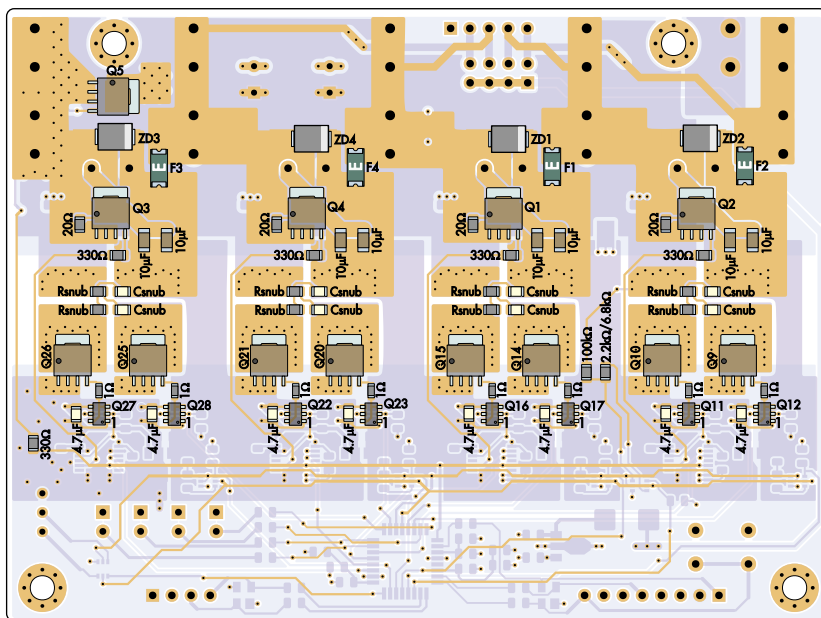
montowana jednocześnie i profesjonalnie w ten sposób lub w piecu rozpliwowym lutujesz tylko transformator, na samym początku montażu).

Zamiast tego sugerujemy umieszczenie ich tak dokładnie, jak to możliwe, przytrzymanie ich na miejscu taśmą Kapton, a następnie przyłutowanie ich czterech końcówek gorącą lutownicą i cyną z topnikiem. Po zamontowaniu transformatorów upewnij się, że wszystkie pozostałości topnika zostały usunięte. Może to być trudne, ponieważ większość pozostałości będzie ukryta między spodem transformatorów a płytką drukowaną.

Jeśli pozostałości topnika nie będą usunięte, prąd jałowy może wzrosnąć o rzędy wielkości

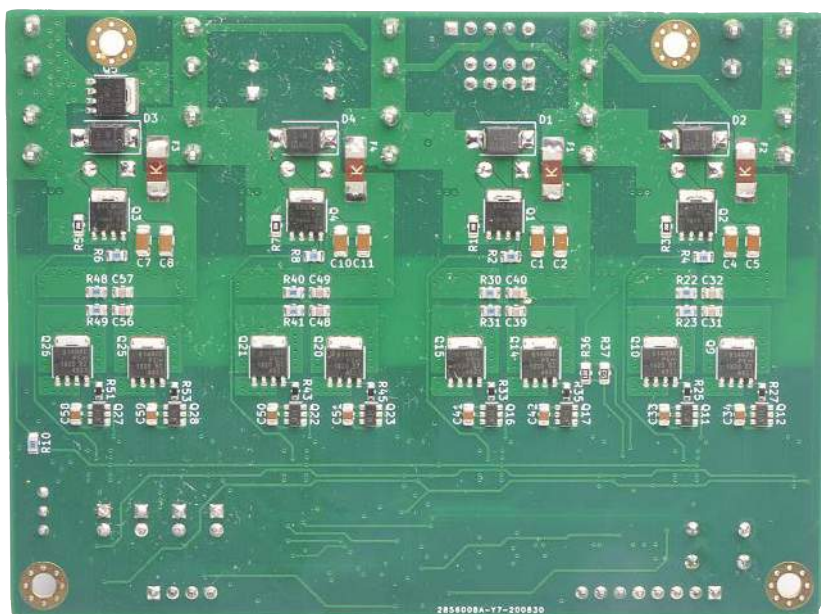
(powyżej 1 mA). Pozostałości topnika mogą ulec zwęgleniu przy wyższych napięciach, powodując nieregularne zachowanie transformatora, a nawet wyładowania łukowe.

W tym przypadku pamiętaj o starym powiedzeniu: „uncja zapobiegania jest warta funta leczenia”, więc staraj się ograniczyć gromadzenie się pozostałości topnika, nie pozwalając na zebranie się zbyt dużej jego ilości. Jeśli masz wybór, spróbuj użyć topnika nie wymagającego czyszczenia (odparowującego). Jeśli jednak używany topnik wymaga zmywania, musisz dokładnie umyć transformatory wysokiej jakości środkiem do usuwania topnika i zetrzeć wszelkie widoczne pozostałości. Od Red. EdW: Z podanych



UNDERSIDE OF PCB

Rysunek 4b. A oto spód płytki drukowanej ze schematem montażowym elementów, z odpowiadającym mu zdjęciem poniżej



względów sugerowalibyśmy, abyś transformatory przylutował na samym początku pracy. Dookoła nich jest sporo pustego miejsca i nie przeszkadzają one w lutowaniu pozostałych elementów. Po ich przylutowaniu możesz wtedy umyć PCB w myjce ultradźwiękowej stosując czysty etanol (nie denaturat!) lub izopropanol, a następnie delikatnie i dokładnie wysuszyć płytkę.

Komponenty do montażu przewlekanego

Zamontuj teraz przełącznik dotykowy S1. Ma on standardowe wymiary, ponadto dostępne są przełączniki o różnych wysokościach przycisku. Jeśli będziesz często regulować

urządzenie, możesz rozważyć zamontowanie przełącznika w obudowie i podłączenie go przewodami do pól lutowniczych. W takim przypadku należy upewnić się, że przewody łączą się z jedną z górnych par i jedną z dolnych par (która jest masą).

Następnie możesz zamontować męskie konektory płaskie 6,3 mm do podłączenia akumulatora/ogniwa zgodnie z wymaganiami. Większość typów konektorów ma dwa styki lutownicze o module 5,08 mm. Będą one pasować do PCB, ale powinieneś sprawdzić, czy do konektorów pasują planowane złącza żeńskie.

Przykładowym konektorem jest Wurth Elektronik 7471286, ewentualnie użyj

części Altronics sugerowanych na liście w zeszłym miesiącu.

W niektórych instalacjach może być korzystniejsze lutowanie przewodów zamiast złączy konektorowych, aby umieścić złącza zamontowane na panelu.

Balanser nie powinien być jednak bezpośrednio lutowany do akumulatorów. Awaria balansera lub baterii będzie trudniejsza do usunięcia, jeśli oba te elementy są ze sobą trwale połączone.

Nie ma potrzeby stosowania przewodu o szczególnie dużej obciążalności, ponieważ prądy równoważące przy ładowaniu są skromne, ale przewody w przybliżeniu powinny być w stanie przenosić prąd o natężeniu 2 A przy pomijalnym wzroście temperatury. Przewód miedziany o średnicy 0,8 mm (20AWG) jest rozsądnym wyborem.

W przypadku korzystania z końcówek konektorowych upewnij się, że żadna część ostrza, końcówki lub przewodu nie styka się z innymi pobliskimi elementami. Dostępne są izolowane końcówki konektorowe i warto z nich korzystać.

W przypadku niektórych instalacji może być konieczne zamontowanie płytki drukowanej w pozycji odwróconej, plecami do góry, i wyprowadzenie zacisków lub przewodów z tyłu płyty.

Możesz również zamontować 5-stykowe jednorzędowe złącze kołkowe 2,54 mm (pionowe lub kątowe) w pozycji CON13 do zastosowań o niższym poborze mocy, takich jak równoważenie ładowania mniejszych akumulatorów litowopolimerowych (LiPo) i podobnych.

Złącze CON13 jest dogodnie zlokalizowane na krawędzi płytki. Jeśli płytka jest zamontowana bezpośrednio na rogu obudowy o odciecinym boku, do listwy kołkowej można podłączyć standardowe złącze płaskie.

Uważaj jednak na polaryzację!

Teraz jest dobry moment, abyś zamontował listwę kołkową 2×4 jako złącza JP1. W przypadku niektórych instalacji, w których baterie są niewymienne, w razie potrzeby można ją zastąpić lutowaną zworką. Zgodnie ze schematem zamontuj wieloobrotowy potencjometr nastawny VR1, upewniając się, że jego śruba regulacyjna znajduje się tak, jak pokazano na rysunku 4a.

Jeśli będziesz często regulować napięcie równoważenia, możesz zamiast tego użyć typowego potencjometru 100 kΩ montowanego w obudowie i poprowadzić przewody z powrotem do pól VR1, ewentualnie podłączając je do listwy kołkowej.

Ani dokładność potencjometru, ani rozpraszanie mocy nie są krytyczne, ale sugerujemy użycie uszczelnionej konstrukcji dla większej żywotności i niezawodności.

Uwagi dotyczące bezpieczeństwa

Praca z akumulatorami wiąże się z pewnymi zagrożeniami. Najważniejszą rzeczą do zrobienia jest dokładne zapoznanie się z wymaganiami dotyczącymi bezpiecznego używania poszczególnych akumulatorów. Ogólnie rzecz biorąc, posiadanie bezpieczników w pobliżu zacisków wszystkich większych akumulatorów jest dobrym pomysłem, aby zapobiec zwarciu i zapaleniu się kabli czy nawet akumulatorów litowych.

Można kupić bezpieczniki, które podłącza się bezpośrednio do zacisków, z możliwością podłączenia grubych przewodów na drugim końcu. Można również użyć bezpieczników wbudowanych w przewody, ale najlepiej byłoby, gdyby odcinek przewodu między zaciskiem a bezpiecznikiem był krótki.

Jest jeszcze kilka rzeczy, o których należy pamiętać podczas korzystania z balansera:

- Przed podłączeniem urządzenia do akumulatorów lub innych źródeł zasilania należy zawsze sprawdzić, czy działa ono zgodnie z przeznaczeniem. Najlepiej zrobić to z zasilaczami o ograniczonym prądzie, jak opisano w tekście głównym.
- Nie pozostawiaj balansera bez nadzoru, dopóki nie upewnisz się, że działa on niezawodnie w Twoim konkretnym zastosowaniu. Należy zachować szczególną ostrożność przy ustawianiu długiego czasu oczekiwania.
- Balanser powinien być fizycznie oddzielony od akumulatorów. Jeśli będą one zbyt blisko siebie, ciepło pochodzące z balansera może uszkodzić akumulatory lub doprowadzić do niebezpiecznej sytuacji.
- Należy upewnić się, że balanser jest zawsze czysty i suchy.
- Balansera nie należy podłączać na stałe do akumulatorów lub innych źródeł zasilania; w razie wystąpienia niebezpiecznej sytuacji dobrze jest mieć możliwość szybkiego odłączenia balansera.
- Okresowo sprawdzaj, czy ogniwa są sprawne: jeśli balanser stale pracuje z jednym ogniwem lub jeśli zauważysz, że baterie tracą pojemność, sprawdź i wymień wszystkie uszkodzone ogniwa.
- Należy pamiętać, że balanser nie jest w stanie powstrzymać ładowania lub rozładowywania akumulatora przez zewnętrzne obwody: nadmierne ładowanie i rozładowywanie ogniw może nie tylko je uszkodzić, ale także prowadzić do niebezpiecznych sytuacji.
- Należy pamiętać, że w zastosowaniach o wyższym napięciu, niektóre z napięć obecnych na balanserze mogą być niebezpieczne (choć jego maksymalna wartość znamionowa wynosząca 60 V mieści się w zakresie bardzo niskiego napięcia lub ELV), dlatego nie należy dotykać balansera. Ponadto, niektóre elementy balansera mogą się nagrzewać podczas pracy.

Postępuj analogicznie ze czterema diodami LED, upewniając się, że katody są skierowane do góry, jak pokazano na rysunku 4a. Jeśli montujesz złącza LED na płycie drukowanej, musisz użyć typów o średnicy 3 mm.

Można jednak zamiast tego zamontować złącza kołkowe lub luźne przewody i zamontować diody w miejscu, które będzie widoczne z zewnątrz (np. zamontowane na panelu lub z boku obudowy za pomocą ramek), w którym to przypadku można użyć diod LED o średnicy 5 mm lub praktycznie dowolnego innego typu.

W przypadku niektórych kolorów LED-ów, może być pożądana inna niż 680 Ω wartość rezystora ograniczającego prąd w celu zwiększenia jasności lub zmniejszenia zużycia energii.

Ponieważ napięcie zasilania wynosi 3,3 V, nie są zalecane niebieskie lub białe diody LED, chociaż może się okazać, że takie typy dają wystarczającą ilość światła, biorąc pod uwagę ich wysoką wydajność.

Złącza CON14 i CON15 są opcjonalne. CON14 jest wymagane tylko wtedy, gdy potrzebny jest dostęp do portu szeregowego, na przykład do debugowania lub podłączenia dwóch balancerów do pracy razem (przez izolator) z 8-sekcyjnym zestawem akumulatorów.

CON15 jest potrzebny tylko wtedy, gdy zamontowano nie zaprogramowany mikroprocesor i trzeba go zaprogramować na płycie lub przeprogramować później.

Pozostaje jeszcze tylko sześć kondensatorów elektrolitycznych. Nie pomył dwóch różnych wartości i upewnij się, że wkładasz

dłuższe przewody do otworów oznaczonych znakami +.

Wybraliśmy organiczne kondensatory polimerowe, a nie zwykłe kondensatory elektrolityczne, ze względu na ich znacznie lepszą charakterystykę działania.

Programowanie

Co do programowania IC2, jak wspomniano w zeszłym miesiącu, możesz to zrobić za pomocą PICKIT 4 podłączonego do CON15 (styl 1 do styku 1). Można to zrobić za pomocą oprogramowania MPLAB X IDE, które jest dostarczane z bezpłatnym oprogramowaniem MPLAB X IDE firmy Microchip (lub po prostu użyć wstępnie zaprogramowanego układu ze sklepu internetowego Silicon Chip Shop). Na temat oprogramowania MPLAB X IDE miesięcznik EdW zamieścił artykuł w numerze 11/2023 r.

Bardzo ważne jest, abyś nie programował urządzenia, gdy jest ono podłączone do jakiegokolwiek źródła zasilania (ogniwa/baterii lub innego), więc włącz opcję „power target from PICKIT”.

Po zaprogramowaniu koniecznie przetestuj urządzenie w warunkach niskiego poboru mocy/prądu, jak opisano poniżej, na wypadek wystąpienia błędu w nowo zaprogramowanym mP.

Testowanie

Przed podłączeniem balansera do akumulatorów musisz go przetestować, aby upewnić się, że nic nie uległo uszkodzeniu, co mogłoby wpłynąć na bezpieczeństwo lub niezawodność.

Najłatwiej to zrobić za pomocą pary izolowanych zasilaczy z ograniczeniem prądowym. Ustaw ich napięcia wyjściowe na takie same (np. 4 V każde), a ich ograniczenia prądowe na około 500 mA.

Podłącz jedno zasilanie między STACK (CON7) i CELL1 (CON6), z dodatnim zaciskiem do CON6. Podłącz drugie zasilanie między CELL1 (CON6) i CELL2 (CON5), z dodatnim zaciskiem do CON5.

Upewnij się, że zworka jest zainstalowana tak, aby blok sterowania był zasilany z jednego z tych dwóch punktów, tj. w pozycjach oznaczonych 1 lub 2 dla JP1 (między szpilkami 1 i 2 lub szpilkami 3 i 4).

Za pomocą oscyloskopu sprawdź okresowe impulsy na liniach SENSE_EN i SAMPLE (odpowiednio styki 19 i 20 układu IC2). Ich brak oznacza, że mikroprocesor jest uszkodzony lub nie otrzymuje zasilania.

Jeśli nie masz oscyloskopu, możesz być w stanie wychwycić impulsy za pomocą ustawienia trybu licznika częstotliwości w multimetrze cyfrowym, a nawet woltomierza analogowego.

Jeśli mikroprocesor jest sprawny, podłącz wyższy zasilacz do szyny napięcia stosu (łącznie CON5 [CELL2] z CON2 [STACK+]) i powoli dokonaj niewielkiej zmiany napięcia jednego z zasilaczy.

Powinieneś zauważyć, że napięcie na zasilaczu o niższym napięciu wyjściowym wzrasta.

Jeśli jest to trudne do zaobserwowania, można użyć oscyloskopu do sprawdzenia linii CSPWM/SSPWM na wyprowadzeniach

mikroprocesora (styki 11 i 17 dla najniższej komórki lub styki 12 i 18 dla drugiej najniższej).

Na tych liniach powinny być widoczne wąskie, prostokątne impulsy.

Jeśli ten test zakończy się pomyślnie, sprawdź sekcje trzeciego i czwartego ogniwa, ale pamiętaj, że ogniwa/akumulatory muszą być zawsze podłączane w kolejności od dolnego; żadna sekcja nie może pozostać pusta, z wyjątkiem górnej. Czyli stosując zasilacze laboratoryjne dla sekcji trzeciej i czwartej, musisz najpierw podłączyć standardowe ogniwa w sekcji pierwszej i drugiej.

Jeśli rozważasz aplikację zestawów o wyższym napięciu, przetestuj je dokładnie, zwracając szczególną uwagę na zastosowanie odpowiednich limitów prądu i upewniając się, że sekcja logiki sterowania jest zasilana tylko z najniższego dostępnego ogniwa.

Pozwala to uniknąć marnowania energii w regulatorze napięcia (REG1) i potencjalnych uszkodzeń w przypadku przekroczenia jego maksymalnego napięcia wejściowego.

Ogólnie rzecz biorąc, jeśli najniższe oczekiwane napięcie ogniwa lub akumulatora wynosi powyżej 3,6 V, należy zawsze pozostawić JP1 w pozycji 1, aby obwód sterujący działał zasilany z najniższego ogniwa.

Jeśli najniższe oczekiwane napięcie ogniwa jest niższe niż 3,6 V, bliskie minimalnego obciążanego 2,5 V, to zawsze powinno być możliwe bezpieczne uruchomienie obwodu sterującego przy zasilaniu z pierwszego i drugiego ogniwa (pozycja 2 na JP1).

Wyższe pozycje są przydatne tylko wtedy, gdy trzeba upewnić się, że niewielki prąd zasilający sekcję sterowania pochodzi z całego stosu ogniw, co byłoby niefortunne.

Montaż końcowy

Po przetestowaniu płyty balansera należy ją zamknąć w obudowie, aby chronić ją przed kurzem i innymi zanieczyszczeniami.

Możesz użyć dowolnego pudełka, które jest wystarczająco duże, aby zmieścić moduł PCB i które umożliwia przeprowadzenie kabli. Idealnie byłoby, gdyby oferowało ono jakąś metodę wizualizacji diod LED (np. przezroczystą pokrywę), dostęp do potencjometru nastawnego VR1 i przycisku S1 (ewentualnie za pomocą śrubokręta przez małe otwory w pokrywie).

Zamontuj płytkę drukowaną do dolnej części obudowy za pomocą kołków dystansowych, aby płytka nie wyginała się, i upewnij się, że wszystkie komponenty mają odpowiedni odstęp od ścianek obudowy, ponieważ może ona rozpraszać trochę ciepła.

Cztery otwory montażowe pasują do śrub M3 i można użyć plastikowych lub metalowych elementów dystansowych. W przypadku

stosowania metalowych elementów dystansowych należy uważać, aby zmieściły się one w miedzianych obszarach wokół otworów.

W przypadku żadnego z komponentów radiatorów nie są zwykle wymagane, ale umożliwienie nawet niewielkiego przepływu powietrza wokół PCB znacznie przyczyni się do utrzymania niskiej temperatury balansera, przedłużając jego żywotność.

W trudnych warunkach można zastosować mały wentylator z termostatem (np. z czujnikiem temperatury przyklejonym do transformatora T1).

W większości przypadków wystarczy jednak pasywny przepływ powietrza, z kilkoma otworami wentylacyjnymi lub otworami wywierconymi w dolnej i górnej pokrywie lub po bokach obudowy, aby konwekcyjnie odprowadzać ciepło.

Korzystanie z urządzenia

Teraz, gdy zbudowałeś i przetestowałeś swój balanser, zapytasz zapewne, jak możesz go używać?

Przed podłączeniem urządzenia do akumulatora należy zapoznać się z poniższą listą kontrolną, aby upewnić się, że urządzenie jest prawidłowo skonfigurowane:

1. Skonfiguruj źródło zasilania urządzenia. Jak opisano powyżej, w przypadku balansowania akumulatorów 12 V należy upewnić się, że zworka wyboru źródła zasilania sterowania jest bezpiecznie ustawiona w skrajnym prawym położeniu (oznaczonym jako 1), tak aby to najniższe ogniwo zapewniało zasilanie sterowania.

Jeśli balansujesz ogniwa ~3,6 V (np. Li-ion, LiPo lub LiFePO₄), prawdopodobnie będziesz chciał ustawić zworkę wyboru źródła zasilania w pozycji drugiej od pozycji najbardziej wysuniętej na prawo, tak aby to najniższa para ogniw zapewniała zasilanie.

2. Podłącz przewody akumulatora do odpowiednich zacisków. Sugerujemy sekwencyjnie

ich podłączenie (CELL1, CELL2...) lub jednocześnie (w przypadku korzystania z zewnętrznego złącza).

Zastosuj płaskie złącza konektorowe do połączenia z CON8-CON12 dla akumulatorów o wyższym natężeniu prądu lub wtyczkę zaprojektowaną do łączenia ze szpilkami listwy kołkowej o rozstawie 2,54 mm na CON13 do równoważenia prądu o natężeniu do 1A.

Jeśli używasz CON13, upewnij się, że orientacja wtyczki jest prawidłowa, z najbardziej ujemnym zaciskiem do styku 1 po prawej stronie! Podczas podłączania przewodów akumulatora mogą pojawić się niewielkie iskry, ale powinny wystąpić tylko chwilowo.

3. Na koniec podłącz przewody stosu (STACK- do CON7 i STACK+ do CON2). W przypadku balansowania można po prostu zmostkować dodatni zacisk napięcia stosu z najwyższym zaciskiem ogniwa.

W celu ładowania podłącz ujemny zacisk stosu do ujemnego złącza źródła zasilania, a dodatni zacisk stosu do dodatniego wyjścia zasilacza

Jeśli to możliwe, zalecamy ustawienie rozsądnie niskiego limitu prądu na źródle zasilania, aby zapobiec uszkodzeniu baterii/akumulatora w przypadku awarii.

Wprowadzanie korekt

Działanie jest zasadniczo automatyczne, a balanser po prostu przenosi ładunek w zależności od różnic, które wykrywa w napięciu ogniw lub ich sekcji. Istnieją jednak pewne opcje, które można ustawić za pomocą potencjometru nastawnego VR1 i przycisku S1 lub za pośrednictwem interfejsu szeregowego.

Opcje obejmują minimalną różnicę między napięciami ogniwa/baterii, aby rozpocząć równoważenie, maksymalny prąd równoważenia oraz minimalne i maksymalne napięcia akumulatora/ogniwa, poza którymi równoważenie zostanie przerwane.

```
[156012.181] Resuming balancing
[156012.639] Stack=> 16.46V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.08V
[156022.669] Stack=> 16.46V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.08V
[156032.692] Stack=> 16.46V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.08V
[156042.720] Stack=> 16.46V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.08V
[156052.750] Stack=> 16.46V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.08V
[156062.793] Stack=> 16.46V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.08V
[156072.822] Stack=> 16.46V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156082.851] Stack=> 16.46V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156092.884] Stack=> 16.46V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156102.912] Stack=> 16.46V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156112.944] Stack=> 16.46V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156122.974] Stack=> 16.45V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156133.006] Stack=> 16.46V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156143.031] Stack=> 16.45V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156153.060] Stack=> 16.45V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156163.100] Stack=> 16.45V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156173.131] Stack=> 16.45V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156183.162] Stack=> 16.45V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156193.198] Stack=> 16.45V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156203.224] Stack=> 16.45V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.09V
[156213.252] Stack=> 16.45V | C4 4.12V | C3 4.12V | C2 4.11V | =>C1 4.10V
[156213.373] Balance reached, sleeping
```

Ekran 1. Przykładowe odczyty na wyjściu portu szeregowego

Tabela 1. Funkcje dostępne po naciśnięciu przycisku S1

Funkcja	Liczba naciśnieć przycisku S1
Sprawdź, czy urządzenie jest włączone	Jedno krótkie (<500 ms)
Wstrzymanie/wznowienie równoważenia	Jedno średniej długości (1...2 s)
Przełączanie między ustawieniami wstępnymi dla akumulatorów litowo-jonowych i kwasowo-ołowiowych	Jedno długie (5s+)
Ustawienie dopuszczalnego napięcia delta (niezrównoważenia (0...300 mV/0...1 V)	Dwa krótkie
Ustawienie maksymalnego prądu równoważenia (0...2,5 A)	Trzy krótkie
Ustawienie minimalnego napięcia akumulatora/ogniwa (0...5/0...15 V)	Cztery krótkie
Ustawienie maksymalnego napięcia akumulatora/ogniwa (0...5/0...15 V)	Pięć krótkich

Domyślnie ustawiony jest maksymalny możliwy prąd równoważenia (około 2,5 A), aby rozpocząć równoważenie przy niezrównoważeniu napięciowym 50 mV dla akumulatorów kwasowo-ołowiowych 12 V lub niezrównoważeniu napięciowym 10 mV dla ogniw litowo-jonowych oraz dla zakresu napięcia roboczego ogniwa 2,5...4,3 V dla zastosowań litowo-jonowych i 10...14,8 V dla akumulatorów o nominale 12 V.

Większość z tych ustawień można zmienić za pomocą potencjometru nastawczego VR1 i przełącznika przyciskowego S1, chociaż większy zakres ustawień konfiguracji i kalibracji jest dostępny za pośrednictwem interfejsu szeregowego/USB.

Tabela 1 przedstawia różne polecenia, które mogą być wydawane przez naciśnięcie przycisku S1 na różne sposoby – pojedynczym, długim naciśnięciem lub kilkoma krótkimi naciśnięciami po kolei z rzędu. Niektóre z nich sterują urządzeniem, podczas gdy inne dostosowują ustawienia w połączeniu z bieżącą pozycją potencjometru VR1.

Niestety, wprowadzanie zmian ustawień w ten sposób jest nieco nieprecyzyjne. Możesz zmierzyć napięcie na wyjściu VR1, sondując jego środkowy styk na spodzie płytki za pomocą DVM lub sondując wyprowadzenie 3 pobliskiego MOSFET-a Q7 względem GND. Następnie należy podzielić ten odczyt przez 1,65 V (lub lepiej, przez rzeczywiste zmierzone napięcie odniesienia ADC 1,65 V), a następnie pomnożyć przez zakres podany w tabeli 1.

Jeśli możesz podłączyć interfejs szeregowy, znacznie lepiej jest wprowadzać zmiany w ten sposób, ponieważ będą one dokładne, a także można w ten sposób prawidłowo skalibrować urządzenie. Czytaj dalej, aby uzyskać więcej informacji na temat interfejsu szeregowego.



Z Silicon Chip-a nr 3/2021 – to izolowane łącze szeregowe jest idealne do połączenia ze sobą dwóch płytek balansujących.

Monitorowanie działania

Najprostszym sposobem na monitoring jest wizualizacja. Jedną z czterech diod LED na płycie będzie migać, aby wskazać, kiedy następuje równoważenie, przy czym najbardziej wysunięta na prawo dioda LED (LED7) odpowiada najniższej komórce, dioda LED8 następnej komórce w stosie itd.

Migają powoli, jeśli bateria/ogniwo jest ładowane, lub szybko, jeśli bateria/ogniwo jest rozładowywane. Jeśli nie odbywa się równoważenie/ładowanie, dioda LED7 będzie od czasu do czasu migotać bardzo lekko, aby dać znać, że obwód „żyje”, zużywając przy tym jak najmniej energii.

Jeśli wystąpi błąd przepięcia, wszystkie cztery diody LED będą migać jednocześnie z częstotliwością 1 Hz, z 50% wypełnieniem.

W przypadku wykrycia błędu zbyt niskiego napięcia, urządzenie po prostu wyłączy się i nie miga diodami LED (nawet w rytm „bicia serca”).

Jeśli jesteś uważny, brak „bicia serca” powie ci, że coś jest nie tak, a pozostawiając diody

LED wyłączone, nie ryzykujemy rozładowania już nadmiernie rozładowanego ogniwa lub baterii.

Jeśli chcesz uzyskać więcej szczegółów na temat działania urządzenia i upewnić się, że wykonuje ono swoją pracę, możesz monitorować dane z portu szeregowy na CON14. W idealnym przypadku port ten powinien być podłączony do komputera za pośrednictwem interfejsu izolującego (dobry interfejs opisano poniżej).

Następnie można podłączyć wyjście tego interfejsu izolującego do adaptera USB/szeregowego.

Ustaw emulator terminala na 38 400 bodów N,8,1 i powinieneś zobaczyć strumień informacji, jak pokazano na Ekranie 1. Pokazuje on zmierzone napięcie na każdym wejściu, plus napięcie całego stosu, oraz informuje, czy aktualnie jest przesyłany ładunek do lub z baterii/ogniwa i jak szybko to następuje (0...100%).

Dane są czytelne zarówno dla człowieka, jak i dla maszyny, więc dość łatwo byłoby stworzyć oprogramowanie do analizowania informacji i wyświetlania ich w różny

Tabela 2. Polecenia szeregowo

Przykład	Przykładowy wynik
p	Wstrzymuje automatyczne równoważenie
r	Wznawia automatyczne równoważenie
t 600	Ustawienie limitu czasu równoważenia na 600 s; jeśli równowaga nie zostanie osiągnięta w tym czasie, wyłączenie
l 3000	Ustawienie progu niskiego napięcia baterii/ogniwa na 3 V (3000 mV); poniżej tego progu urządzenie wyłącza się
h 4300	Ustawienie wysokiego progu akumulatora/ogniwa na 4,3 V (4300 mV); powyżej tego progu wyłącza się
d 50	Baterie/ogniwa mogą różnić się do 50 mV przed rozpoczęciem równoważenia
i2 50	Przenosi ładunek do akumulatora/ogniwa nr 2 (1...4) z prędkością 50% wartości maksymalnej (1...100)
o3 25	Wyprowadzenie ładunku z akumulatora/ogniwa nr 3 (1...4) przy 25% maksymalnego poziomu (1...100)
c2 100000 6790	Kalibracja – ustaw dzielnik akumulatora/ogniwa nr 2 tak, aby współczynnik podziału napięcia wynosił 100 kΩ:6,79 kΩ
st 100000 6812	Kalibracja – ustaw dzielnik stosu tak, aby współczynnik podziału napięcia wynosił 100 kΩ:6,812 kΩ
v 3280	Kalibracja – ustaw typowe napięcie zasilania na 3,28 V (3280 mV)

sposób lub podejmowania działań w zależności od wyników.

Jak pokazano w **tabeli 2**, można również wysyłać polecenia, aby wstrzymać lub wznowić równoważenie, zmienić ustawienia, a nawet wymusić przeniesienie ładunku do lub z danego akumulatora/ogniwa. Oznacza to, że można w przypadku korzystania z kilku urządzeń balansera scentralizować sterowanie za pomocą programu komputerowego.

Łączenie wielu balanserów

Możesz użyć dwóch urządzeń balansera, aby zrównoważyć do ośmiu akumulatorów lub ogni, o ile całkowite napięcie stosu nadal mieści się w maksymalnej wartości 60 V DC. Jedynym dodatkowym sprzętem potrzebnym do tego celu jest izolowane łącze szeregowo.

Na szczęście Silicon Chip opublikował taki projekt w marcu 2021 roku (siliconchip.com.au/Article/14785), a płytki PCB są dostępne.

Zbuduj tę płytkę, ale pomiń listwy złączy i ustaw obie zworki (JP1 i JP2) w pozycji 5 V (w rzeczywistości będą one zasilane napięciem 3,3 V, ponieważ jest to jedyna szyna niskiego napięcia dostępna na płytkach balansera).

Następnie można przylutować styki 3...6 CON1 lub CON2 bezpośrednio do CON14 na jednej z płytek balansera ogni, ponieważ układ styków jest dokładnie taki sam.

Poprowadź kabel taśmowy lub podobny z drugiego końca płytki do CON14 na drugiej płytce balansera. Okablowanie będzie takie samo jak na drugim końcu, a styk TX na balanserze powinien być podłączony do styku TX na płytce izolatora.

Podobnie, styk RX na balanserze łączy się ze stykiem RX na izolatorze. Zamiana linii sygnałowych jest wykonywana wewnątrz izolatora.

Następnie wystarczy podłączyć od jednego do czterech sąsiadujących ogni/akumulatorów w stosie do jednej płytki balansera,

zaczynając od połączenia CELL1, a resztę podłączyć do drugiej.

Podłącz oba pełne stosy do zacisków STACK i STACK+ na obu płytkach.

Oba urządzenia włączą się i będą negocjować przez łącze szeregowo, automatycznie wykrywając, że się ze sobą komunikują.

Będą one wtedy balansować ogniwa/akumulatory tak, jakby były jednym 8-wejściowym balanserem zamiast dwóch 4-wejściowych balanserów. ■

Duraïd Madina

Adaptacja do wydania
polskiego – Andrzej Nowicki

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

Patronat AVT

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

- | | | | |
|---|---|---|--|
| • Centrum Edukacji Zawodowej,
82-200 Malbork, De Gaulle'a 75a | • Kościelna 42 | • Techniczne Zakłady Naukowe w Dąbrowie
Górnicej, 41-300 Dąbrowa Górnicza,
Zawidzkiej 10 | • Zespół Szkół Spożywczych i Hotelarskich
w Radomiu, 26-600 Radom, Św. Brata Alberta 1 |
| • Centrum Edukacji Zawodowej i Biznesu,
66-400 Gorzów Wielkopolski, Pomorska 67 | • Zespół Szkół Budowlano-Elektrycznych
im. Jana III Sobieskiego w Świdnicy, 58-
100 Świdnica Śląska, Wałbrzyska 35-37 | • Zespół Szkół nr 2 im. Eugeniusza Kwiatkowskiego w Dębicy, 39-200 Dębica, Lisa 2 | • Zespół Szkół Techniczno-Informatycznych
w Elblągu, 82-300 Elbląg, Rycka 2 |
| • Gminny Zespół Szkolno-Przedszkolny nr 4
w Więkach, 42-110 Popów, Więcki,
Szkołna 1 | • Zespół Szkół Centrum Kształcenia
Ustawicznego w Gronowie, 87-162 Lubicz
Dolny, Gronowo 128 | • Zespół Szkół nr 2 im. Gen. Józefa Bema,
05-822 Milanówek, Wójtowska 3 | • Zespół Szkół Technicznych i Licealnych
w Piechowicach, 58-573 Piechowice,
Przemysłowa 21 |
| • Górnośląskie Centrum Edukacyjne im.
Marii Skłodowskiej-Curie w Gliwicach,
44-100 Gliwice, Okrzei 20 | • Zespół Szkół Elektronicznych i Telekomunikacyjnych w Olsztynie, 10-144 Olsztyn,
Bałtycka 37a | • Zespół Szkół nr 2 im. Ks. Prof. Józefa Tischnera w Żorach, 44-240 Żory, Boryńska 2 | • Zespół Szkół Technicznych i Ogólnokształcących nr 3 im. E. Abramowskiego, 40-659
Katowice, Harcerzy Września 1939 2 |
| • Noworudzka Szkoła Techniczna w Nowej
Rudzie, 57-401 Nowa Ruda, Stara Droga 4 | • Zespół Szkół Elektronicznych im. I. Domeyki w Bolesławcu, 59-700 Bolesławiec,
Tyrankiewiczów 2 | • Zespół Szkół nr 2 w Pabianicach im.
prof. Janusza Groszkowskiego, 95-200 Pabianice, Św. Jana 27 | • Zespół Szkół Technicznych im. Armii
Krajowej w Skarżysku-Kamiennej, 26-110
Skarżysko-Kamienna, Tysiąclecia 22 |
| • Regionalne Centrum Edukacji Zawodowej
w Biłgoraju, 23-400 Biłgoraj, Kościuszki 98 | • Zespół Szkół Elektronicznych w Rzeszowie,
35-078 Rzeszów, Hetmańska 120 | • Zespół Szkół nr 40 im. Stefana Starzyńskiego, 03-771 Warszawa, Objazdowa 3 | • Zespół Szkół Technicznych im. Ignacego
Mościckiego w Tarnowie, 33-101 Tarnów,
E. Kwiatkowskiego 17 |
| • Regionalne Centrum Edukacji Zawodowej
w Lubartowie, 21-100 Lubartów, 1 Maja 82 | • Zespół Szkół Elektronicznych, Elektrycznych i Mechanicznych, 43-300 Bielsko-Biała,
Słowackiego 24 | • Zespół Szkół Politechnicznych im. Bohaterów Monte Cassino we Wrześni, 62-300
Września, Wojska Polskiego 1 | • Zespół Szkół Technicznych w Kolbuszowej,
36-100 Kolbuszowa, Bytnara 2 |
| • Technikum nr 4 im. Marii Skłodowskiej-Curie, 41-902 Bytom, Katowicka 35 | • Zespół Szkół Elektrycznych nr 2 w Krakowie, 31-977 Kraków, Os. Szkolne 26 | • Zespół Szkół Ponadgimnazjalnych nr
1 w Jarocinie, 63-200 Jarocin, Franciszkańska 1 | • Zespół Szkół w Białowej, 36-030 Białowa,
Kowala 3 |
| • Zespół Placówek Edukacyjno-Wychowawczych w Gołdapi, 19-500 Gołdap, Wojska
Polskiego 18 | • Zespół Szkół Elektrycznych w Kielcach,
25-317 Kielce, Kaczorowskiego 8 | • Zespół Szkół Ponadpodstawowych nr 2
im. E. Kwiatkowskiego w Jarocinie, 63-200
Jarocin, Franciszkańska 2 | • Zespół Szkół w Gościnie, 78-120 Gościno,
Kościuszki 5 |
| • Zespół Placówek Oświatowych w Rudniku,
32-440 Sułkowice, Rudnik, Szkołna 55 | • Zespół Szkół im. Bolesława Prusa,
42-207 Częstochowa, Prusa 20 | • Zespół Szkół Ponadpodstawowych nr 3
im. Armii Krajowej w Zamościu, 22-400
Zamość, Zamoyskiego 62 | • Zespół Szkół w Zarzeczu, 37-205 Zarzecze,
Św. Jana Pawła II 7 |
| • Zespół Szkolno-Przedszkolny nr 2 w Wiśle,
43-460 Wisła, Malinka 53 | • Zespół Szkół im. ks. dra Jana Zwieryza w Ropczycach, 39-100 Ropczyce,
Mickiewicza 14 | • Zespół Szkół Powiatowych im. Stanisława
Staszica w Opocznie, 26-300 Opoczno,
Kossaka 1a | • Zespół Szkół Zawodowych nr 1 im. gen.
F. Kleeberga w Dęblinie, 08-530 Dęblin,
Tysiąclecia 3 |
| • Zespół Szkolno-Przedszkolny nr 3 w Gliwicach, 44-122 Gliwice, Żwirki i Wigury 85 | • Zespół Szkół im. Ks. Stanisława Staszica,
39-400 Tarnobrzeg, Kopernika 1 | • Zespół Szkół Publicznych w Szewnie,
27-400 Ostrowiec Świętokrzyski, Szewna,
Langiewicza 3 | • Zespół Szkół Samochodowych im. inż.
Tadeusza Tańskiego, 33-300 Nowy Sącz,
Rejtana 18a |
| • Zespół Szkolno-Przedszkolny nr 4 w Rybniku, 44-207 Rybnik, Komisji Edukacji
Narodowej 29 | • Zespół Szkół nr 1 w Przysietnicy, 36-200
Brzozów, Przysietnica 198 | | • Szkoła Podstawowa im. Rodziców Bohaterów II Wojny Światowej w Żałakowie,
83-342 Kamienica Królewska, Żałakowo 6 |
| • Zespół Szkolno-Przedszkolny w Choceniu,
87-850 Chocień, Sikorskiego 12 | • Zespół Szkół nr 10 im. Prof. Janusza Groszkowskiego w Zabrzu, 41-807 Zabrze,
Chopina 26 | | |



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://shorturl.at/nCH18>

Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania



64-klawiszowa matryca MIDI sterowana modułem Arduino, część 2

W poprzednim numerze opisaliśmy niedrogi sprzęt, który można zbudować do pracy z MIDI, w tym wszechstronną nakładkę kodera MIDI i matrycę klawiszy MIDI do jej obsługi. Opracowaliśmy więcej oprogramowania, aby jeszcze lepiej wykorzystać ten sprzęt i pokażemy, jak można go używać nawet ze smartfonami i tabletami z systemem Android.

Pierwsza część tego artykułu pokazała, jak zbudować prostą matrycę klawiszy MIDI, składającą się z siatki 64 przełączników monostabilnych (tact-switch-y). Są one skanowane przez kompatybilną z systemem Arduino płytkę Leonardo, która działa jako koder MIDI. Wykorzystuje ona swoje peryferia USB do generowania komunikatów MIDI, które mogą być odbierane na komputerze osobistym z oprogramowaniem syntezatora lub cyfrowej stacji roboczej audio (DAW).

Nakładka Leonardo jest również wyposażona w parę 5-stykowych gniazd DIN zgodnych ze standardem MIDI. Jedno ze złączy DIN zostało skonfigurowane jako port MIDI-out a drugie jako port MIDI-in. Port MIDI-out generuje „sprzętowe” komunikaty MIDI, które mogą być przesyłane do portu MIDI-in w urządzeniu takim jak syntezator MIDI.

W ten sposób matryca klawiszy MIDI i koder MIDI tworzą razem ogólne urządzenie

wyjściowe MIDI, odpowiednie do wyzwalania dźwięków i nut na dowolnej liczbie urządzeń obsługujących standard MIDI, które mają port USB lub 5-stykowe gniazdo DIN MIDI-in.

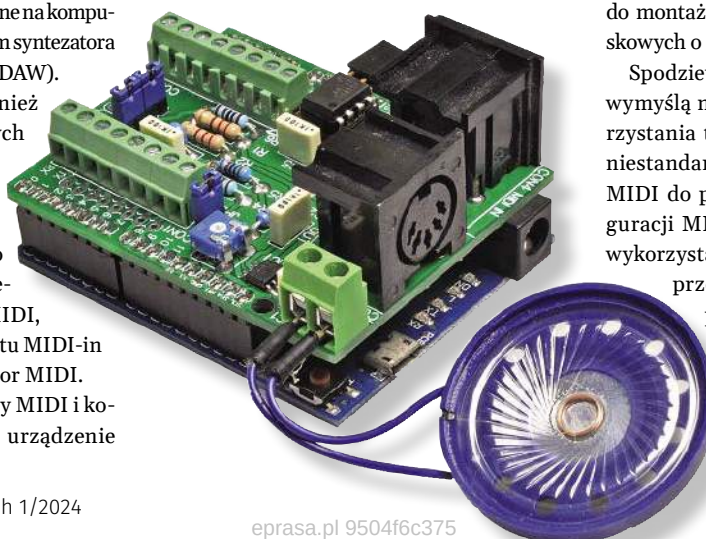
Ma również wbudowany bardzo podstawowy syntezator, który po wystąpieniu kluczowych zdarzeń przesyła dźwięki do małego głośnika

o mocy 1 W. Sam w sobie tworzy bardzo prosty instrument muzyczny.

Chociaż możliwe jest ręczne podłączenie 64 klawiszy, pokazaliśmy również, jak zbudować matrycę przełączników wykorzystującą płytkę drukowaną, co znacznie upraszcza to zadanie. Matryca przełączników przystosowana jest do montażu kilku typów przełączników naciśkowych o różnych rozmiarach.

Spodziewamy się, że niektórzy Czytelnicy wymyślą nowe i interesujące sposoby wykorzystania tego sprzętu. Może to obejmować niestandardowe oprogramowanie enkodera MIDI do pełnienia określonej roli w konfiguracji MIDI. Może to również obejmować wykorzystanie w nietypowy sposób matrycy przełączników w celu uproszczenia połączenia ze sprzętem.

Jeśli chcesz dowiedzieć się więcej o genezie i działaniu MIDI, zobacz panel „Co to jest MIDI?” pod koniec tego tekstu.



Matryca LED

Chociaż naszym zamiarem było stworzenie taniego i użytecznego urządzenia wejściowego MIDI, wykorzystaliśmy dużo miejsca na płytce drukowanej matrycy przycisków, aby dodać dodatkowe funkcje. W szczególności znajdują się tam pola umożliwiające zamontowanie podświetlanych przełączników.

Są one również podłączone w matrycy do pary 8-drożnych złączy (CON3 i CON4), z rezystorem szeregowym do ograniczania prądu LED w każdym rzędzie.

W pierwszej części pokazaliśmy podstawowy kod do sterowania diodami LED, ale oryginalny sprzęt nie mógł jednocześnie wykrywać naciśnięć klawiszy i sterować diodami LED. Teraz zajmijmy się tą kwestią.

Ograniczenia Leonardo

Problem polega na tym, że w module Arduino Leonardo nie ma zbyt wielu wolnych wejść/wyjść; z pewnością nie na tyle, aby jednocześnie sterować diodami LED i skanować przełączniki.

W rzeczywistości nie ma wielu kompatybilnych płytek Arduino, które by na to pozwalały i nadal zapewniały obsługę peryferii USB.

Arduino Mega ma wystarczającą liczbę portów, ale niestety jego obsługa magistrali USB jest ograniczona do danych szeregowych za pośrednictwem oddzielnego układu scalonego USB-serial.

W ostatnim numerze zauważyliśmy, że jeśli chcesz po prostu zapalić wszystkie diody LED, możesz zwyczajnie podłączyć szyny zasilania do CON3 i CON4 na płytce drukowanej matrycy przycisków.

Jeśli jednak chcesz niezależnie sterować diodami LED, najprostszym podejściem jest dodanie drugiej płytki mikroprocesora.

Możesz chcieć podświetlić każdy klawisz jako odpowiedź, aby wskazać, który z nich należy nacisnąć jako następny. Może to być przydatne jako narzędzie do nauki, pomagające w opanowaniu utworu muzycznego na pamięć.

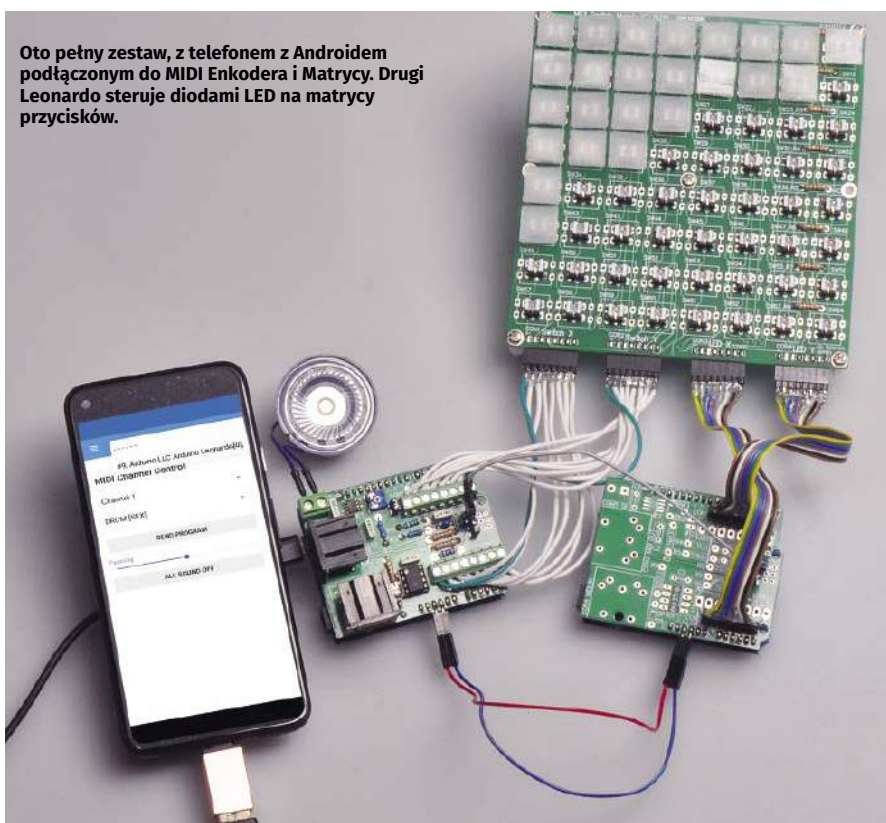
Innym przykładem może być zapalanie diod LED w rytm naciskanych klawiszy. To właśnie zrobiliśmy w naszym przykładowym kodzie.

Opowieść o dwóch mikrokontrolerach

Aby zapalić klawisze w odpowiednim czasie, wymagana jest komunikacja między dwoma mikroprocesorami, mimo że ustaliliśmy już, że pojedynczy moduł nie ma zbyt wielu portów do dyspozycji. Nasza sztuczka polega na pożyczaniu jednego, który jest już używany.

Specyfikacja MIDI obsługuje tak zwane komunikaty SYSEX (System Exclusive).

Oto pełny zestaw, z telefonem z Androidem podłączonym do MIDI Enkodera i Matrycy. Drugi Leonardo steruje diodami LED na matrycy przycisków.



Komunikaty te mają na celu umożliwienie producentom sprzętu MIDI wysyłanie niestandardowych danych, które nie pasują do standardowych komunikatów. Mogą one być używane (na przykład) do wysyłania danych próbek audio między urządzeniami.

Wielu producentów ma określone identyfikatory, ale identyfikator 0x7D może być używany do celów programistycznych i właśnie tutaj to robimy. Użycie tego identyfikatora oznacza, że nasze dane nie zostaną pomylone z sygnałami innego producenta.

Wysyłany przez nas komunikat System Exclusive składa się z bajtu stanu 0xF0, identyfikatora 0x7D, po którym następują kody ASCII dla „SC”. Te dodatkowe znaki zmniejszają prawdopodobieństwo niezrozumienia danych przez inne urządzenia.

Po tym następuje dowolna liczba bajtów danych w zakresie od 0x00 do 0x7F, co daje nam 128 kodów. Kody 0-63 wyłączają odpowiednio diody LED od 0 do 63, podczas gdy kody 64-127 włączają jedną z nich.

Każdy bajt danych większy niż 0x7F (tj. z ustawionym najbardziej znaczącym bitem) kończy komunikat SYSEX, chociaż wysyłamy 0xF7, ponieważ jest to zdefiniowane polecenie „End SYSEX”. Wszystkie urządzenia MIDI to rozumieją, a zatem wszystko pozostaje zsynchronizowane, ignorując dane wewnątrz tych pakietów.

Wysyłamy te dane na port MIDI-out. Większość urządzeń MIDI zignoruje te bajty,

wiec nie będzie zakłócać naszych komunikatów dźwiękowych.

Następnie wystarczy odebrać te dane i wyświetlić je na diodach LED przycisków. Używamy do tego drugiej płytki Leonardo.

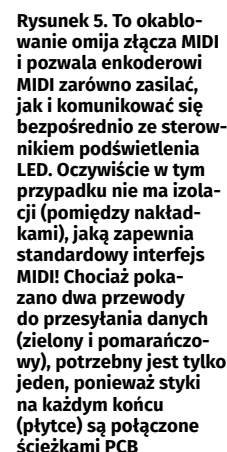
Ponieważ dane MIDI to nic innego jak szeregowy strumień bitów przesyłany z prędkością 31 250 bodów, ustawiliśmy urządzenie w tryb nasłuchiwania przychodzących danych.

Po otrzymaniu bajtów 0xF0, 0x7D, „S” i „C”, zakłada ono, że wszystkie kolejne bajty danych są poleceniami wyłączenia i włączenia diod LED, jak opisano powyżej, dopóki nie otrzyma bajtu powyżej 0x7F, który ponownie ustawia układ w stan oczekiwania na opisaną wyżej sekwencję danych.

Diody LED są multipleksowane przez przerwanie zegarowe, co zapewnia, że każda kolumna otrzymuje równą ilość czasu, a tym samym diody LED są sterowane z równomierną jasnością. Kod skanuje po kolei każdą kolumnę, zapalając diody LED zgodnie z wcześniej otrzymanymi poleceniami.

Sprzęt do sterowania diodami LED

Ponieważ nasz sterownik LED wykorzystuje standardowe pakiety MIDI, luksusowym sposobem montażu jest użycie dwóch płytek Leonardo, z których każda jest zwieńczona w pełni wyposażonym enkoderem MIDI.



Owo minimum połączeń pokazaliśmy również na zamieszczonych wcześniej zdjęciach. Przylutowaliśmy 2-stykowy odcinek gniazda żeńskiego do każdej płytki drukowanej

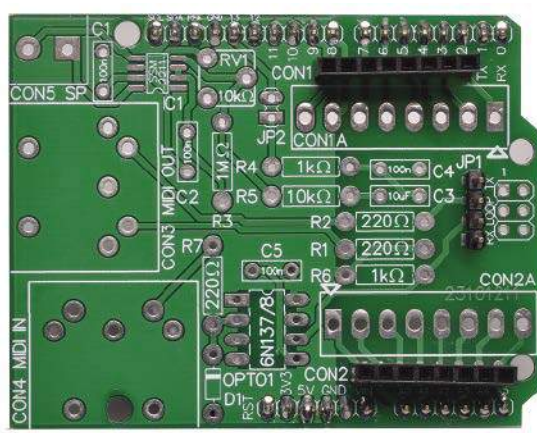
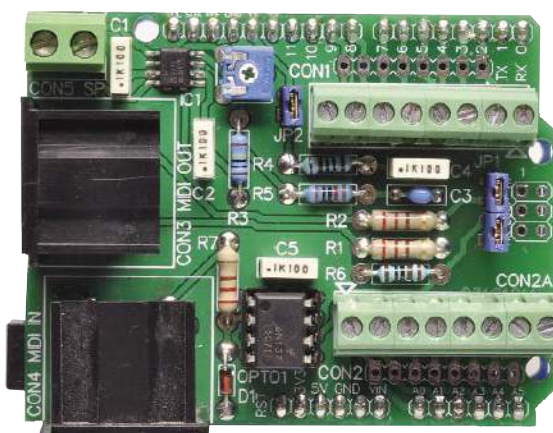
Najłatwiejszym sposobem ich przylutowania jest włożenie listew do gniazd Leonardo,

W przypadku kabli łączących naszą nową mini-nakładkę z matrycą przycisków upewnij się, że styk 1 CON2 (nakładka) jest podłączony



- 1 dwustronna płytka drukowana o kodzie 231012111 w wymiarach 69x54 mm
- 1 moduł Arduino Leonardo
- 1 odcinek 10-szpilkowej prostej listwy kotkowej (złącze nakładki do Leonardo)
- 1 odcinek 8-szpilkowej prostej listwy kotkowej (złącze nakładki do Leonardo)
- 2 odcinki 6-szpilkowej prostej listwy kotkowej (złącze nakładki do Leonardo)
- 2 odcinki 8-stykowego gniazda żeńskiego (złącze CON1 i CON2 do PCB), opcjonalnie:
- 4 szt. zaciskowych złączy śrubowych 3-kontaktowych plus 2 szt. zaciskowych złączy śrubowych 2-kontaktowych (CON1a i CON2a)
- 2 8-przewodowe odcinki kabla połączeniowego, z listwami kotkowymi na obu końcach (do podłączenia płytki drukowanej nakładki do matrycy klawiszy)
- 3 przewody połączeniowe zakończone kołkami (do podłączenia zasilania i szyny danych z enkodera MIDI)

Oto zdjęcia PCB obu płytek nakładek, które pasują do schematów połączenia nakładek na sąsiedniej stronie. Oczywiście na tylnej stronie PCB zamontowanej teraz nakładki sterownika LED-ów nie ma prawie nic poza listwami kotkowymi, których nie widać, poza ich lutowaniem do PCB. Płytkę podłącza się „na kanapkę” bezpośrednio do modułu Leonardo.



do styku 1 CON3 (matryca przycisków), a styk 1 CON1 (nakładka) do styku 1 CON4 (matryca przycisków).

Linia danych

Linia danych jest podłączona na jednym końcu do płytki enkodera MIDI poprzez dołączenie jej do styku TX na JP1, podczas gdy do sterownika LED jest podłączona przez dołączenie do styku RX na JP1.

Jeśli masz gołe moduły Leonardo na obu końcach, możesz podłączyć przewód od styku TX (styk 1) Leonardo – enkodera MIDI do styku RX (styk 0) Leonardo – sterownika LED.

Oczywiście potrzebny będzie wariant matrycy przycisków z zamontowanymi podświetlanymi stykami; anody ich LED-ów powinny być skierowane do górnej części płytki drukowanej.

Oprogramowanie

Zaktualizowaliśmy również oprogramowanie enkodera MIDI w celu wysyłania danych SYSEX LED, więc załaduj

„MIDI_ENCODER_LED_SERIAL_OUT” do Leonardo, który działa jako enkoder MIDI.

Leonardo – sterownik LED-ów powinien być podobnie zaprogramowany za pomocą szkicu „MATRIX_LED_DRIVER_SERIAL”.

Po wykonaniu tych czynności i włączeniu zasilania obu płytek, naciśnięcie dowolnego klawisza powinno spowodować zaświecenie się odpowiedniej diody LED. Jeśli podświetlenia nie zgadzają się z naciśniętymi klawiszami, spróbuj zamienić lub obrócić połączenia z CON3 lub CON4 matrycy przycisków.

Jeśli chcesz zrobić coś wymyślnego, na przykład, aby diody LED „promieniowały” z każdego naciśniętego klawisza, wystarczy wprowadzić modyfikacje do szkicu MATRIX_LED_DRIVER_SERIAL.

Nakładki na klawisze wydrukowane w 3D

Odstępy 16 mm na płytce matrycy przycisków to dobry kompromis między kompaktowością a użytecznością. Gdyby klawisze znajdowały się znacznie bliżej siebie,

istniałaby większa szansa na naciśnięcie więcej niż jednego klawisza na raz, podczas gdy szersze odstępy szybko spowodowałyby wzrost rozmiaru płytki drukowanej.

W poprzednim numerze zauważyliśmy, że drukowane w 3D nakładki na klawisze byłyby ekonomicznym sposobem na dodanie wykończenia do podświetlanej matrycy. Podczas gdy przyciski monostabilne 12 mm można znaleźć z dużymi klawiszami, które są łatwe do naciskania palcami, istnieje mniej opcji dla mniejszych, podświetlanych części.

Niestety, wiele małych, podświetlanych przycisków obsługuje tylko małe klawisze (około 10 mm), które wyglądałyby bardzo dziwnie przy zastosowaniu przez nas rozstawie 16 mm. Zaprojektowaliśmy więc w Redakcji Silicon Chip-a i wydrukowaliśmy w 3D nakładki na przyciski pasujące do przełączników z serii ILS, których użyliśmy.

Wydrukowaliśmy kilka z nich z półprzezroczystego filamentu PLA (Jaycar's Cat TL4274), który pomógł rozproszyć światło z diod LED,



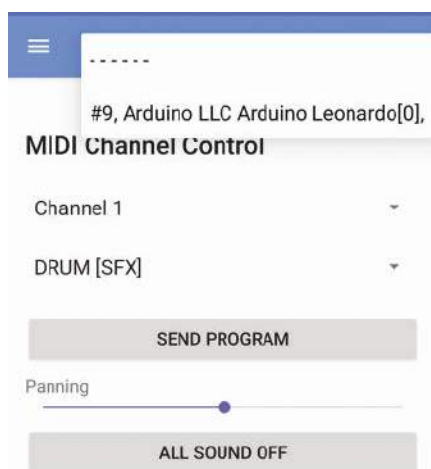
Use with your computer

- Plug your device with your computer.
- Select MIDI in the USB options dialog.
- Select Android USB Peripheral Port in the above drop-down list.
- Use your device as MIDI OUT with your favourite MIDI player/sequencer.

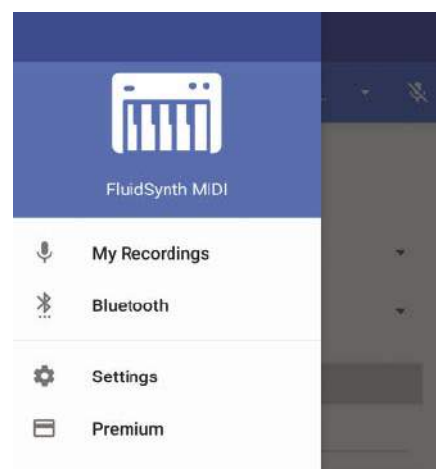
Use with a MIDI USB Keyboard

- Connect your MIDI Keyboard via an USB OTG cable.
- Select your MIDI Keyboard in the above drop-down list.

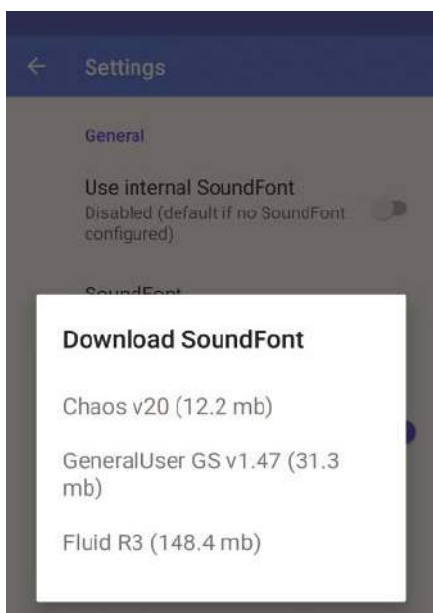
Ekran 1. Aplikacja FluidSynth MIDI Synthesiser App ma niewiele elementów sterujących, ale dzięki temu jest łatwa w użyciu. Po podłączeniu enkodera MIDI (lub innego urządzenia MIDI), staje się on dostępny na górze ekranu



Ekran 2. Po wybraniu koder MIDI pojawia się jako urządzenie Arduino Leonardo. Wyświetlany numer (tu: #9) jest inny za każdym razem, gdy enkoder jest ponownie podłączony, ale nie wydaje się to powodować żadnych problemów.



Ekran 3. Użyliśmy tylko opcji Ustawienia (Settings) z menu. Nagrywanie to płatna funkcja premium, której nie testowaliśmy, ponieważ uznaliśmy, że wystarczy podłączyć 3,5-milimetrowy stereo-foniczny przewód aux do gniazda telefonu, aby nagrywać dźwięk



Ekran 4. Konieczne jest jedynie pobranie pliku SoundFont podczas początkowej konfiguracji. Redakcja Silicon Chip-a zdecydowała się pobrać Chaos v20. Następnie wystarczy sprawdzić, przed jego użyciem, czy koder MIDI jest wybrany na ekranie głównym

choć ze względu, że są one tak małe, były dość wymagające pod kątem wykonania.

Udostępniliśmy więc do pobrania pliki 3D, jeśli chcesz wydrukować własne nakładki na przyciski i wypróbować je samodzielnie. Załączone zdjęcia pokazują, jak wygląda rezultat.

MIDI na Androidzie

Podczas testowania różnych naszych wariantów sprzętowych MIDI, w tym MIDI Enkodera PCB i Matrycy Przycisków PCB, zastanawialiśmy się, czy istnieje sposób na podłączenie płytki MIDI Enkodera do smartfona. Ponieważ wiele osób może mieć stary telefon komórkowy, który można zmienić, takie rozwiązanie byłoby szybkim, tanim i łatwym sposobem na stworzenie użytecznego instrumentu muzycznego.

Przyjrzelśmy się tylko urządzeniom z Androidem, ponieważ to właśnie z nich

korzysta większość z nas (w Redakcji Silicon Chip-a). Nasze minimalne badania sugerują, że może to być możliwe na urządzeniach Apple (takich jak iPhone'y), ale nie jest to coś, czym się zajmowaliśmy.

Należy również pamiętać, że istnieje wiele różnych typów telefonów z systemem Android i nie możemy twierdzić, że rozwiązanie będzie działać na wszystkich.

Jeśli możesz sprawdzić, czy twój telefon obsługuje USB OTG (on-the-go) i ma dość aktualną wersję Androida, to masz dużą szansę na sukces.

Android nie ogranicza się do telefonów komórkowych; niektóre tablety działają również pod kontrolą Androida, „podobnie jak” niektóre inne urządzenia (np. inteligentne telewizory).

Czego potrzebujesz

Jedną z rzeczy, których prawie na pewno będziesz potrzebować, jest adapter USB on-the-go (OTG). Dzięki niemu gniazdo USB w telefonie zachowuje się jak urządzenie hosta, a nie urządzenie peryferyjne. Nie wszystkie telefony obsługują OTG, ale obecnie jest to dość powszechne rozwiązanie.

Adapter OTG będzie miał wtyczkę pasującą do gniazda USB w telefonie (micro-USB lub USB-C) i gniazdo USB-A (takie jak w komputerze). Jaycar Cat WC7725 (przewód micro-USB) lub Cat WC7709 (przewód USB-C) powinny działać. Ewentualnie użyj Altronics Cat P1921 dla micro-USB lub Cat P1924 dla USB-C.

Dostępne są mniejsze adaptory, które nie mają żadnych przewodów; łączą one porty bezpośrednio. Lubimy je, ponieważ są wystarczająco małe, aby nosić je w kieszeni lub torbie. Możesz zobaczyć jeden z nich na naszym zdjęciu telefonu z Androidem.

Rzeczywiście, adapter OTG to przydatna rzecz w dzisiejszych czasach, ponieważ wiele telefonów obsługuje klawiatury USB, myszy i dyski flash. Widzieliśmy nawet aplikację, która pozwala używać go do programowania płytek Arduino!

Prawdę mówiąc, to pisanie szkiców Arduino na tak małym ekranie nie jest najłatwiejszą rzeczą na świecie.

Należy pamiętać, że telefon i aplikacje muszą obsługiwać urządzenia, które chcemy podłączyć. Na szczęście MIDI jest obsługiwane jako Device Class Definition przez standard USB, co oznacza, że każde zgodne urządzenie USB MIDI powinno działać, bez konieczności posiadania specjalistycznych sterowników.

Aplikacje MIDI

Nie mamy żadnych powiązań z następującymi aplikacjami MIDI na Androida (uprzedzamy w ten sposób ew. zarzuty Czytelników o kryptoreklamę); znaleźliśmy je, przeszukując Sklep Google Play.

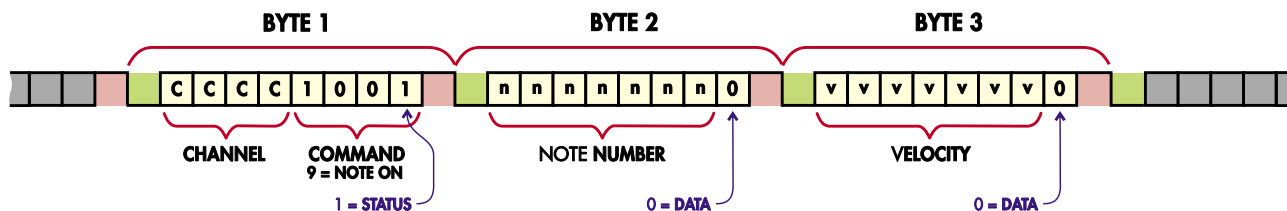
Pierwsza, którą wypróbowaliśmy, nosiła nazwę MIDI Keyboard. Była w stanie rozpoznać podłączone urządzenie Leonardo, ale była ograniczona do jednego instrumentu klawiszowego. Szukaliśmy więc dalej, aby zobaczyć, jakie inne opcje są dostępne.

Druga aplikacja, o nazwie FluidSynth MIDI Synthesizer, oferowała kilka innych opcji. W chwili pisania tego tekstu miała ona ogólnie pozytywne recenzje, a także obsługiwała pobieranie plików SoundFont (.sf2).

Dostępna jest płatna aktualizacja umożliwiająca nagrywanie na urządzeniu, ale dobrze bawiliśmy się po prostu odtwarzając dźwięki przez głośnik telefonu. Połączenie za pomocą kabla stereo aux 3,5 mm powinno wystarczyć do podłączenia dźwięku do innego sprzętu w celu wzmocnienia lub nagrywania.

Instalacja aplikacji była dość prosta. Działa ona z systemem Android 6 i nowszymi. Okazało się, że potrzebuje ona pozwolenia na dostęp do pamięci urządzenia; jest to wymagane, aby uzyskać dostęp do pobranych plików .sf2.

Ekran 1 pokazuje ekran początkowy. Gdy urządzenie MIDI (takie jak MIDI Enkoder) jest podłączone za pomocą adaptera OTG, jego nazwa pojawia się w rozwijanym menu u góry ekranu.



Struktura typowego komunikatu „Note-on”

■ = START BIT ■ = STOP BIT

UWAGA: Dane szeregowe są przesyłane z najmniej znaczącym bitem na początku!

Rysunek 6. Struktura typowego komunikatu „Note-on” (czyli startu). Pierwszy bajt ma ustawiony najbardziej znaczący bit na wysokim poziomie, aby zaznaczyć, że jest to początek komunikatu. Wysoki poziom ostatniego bitu sygnalizuje, że jest to komunikat Note-on (wysoki poziom przedostatniego bitu oznaczałby Note-off), podczas gdy 4 najmniej znaczące bity odpowiadają numerowi kanału. Kolejne bajty zawierają 7-bitowe wartości dla numeru nuty i prędkości

Co to jest MIDI?

MIDI to skrót od **M**usical **I**nstrument **D**igital **I**nterface.

Jest to znormalizowany system komunikacji między elektronicznymi instrumentami muzycznymi, kartami dźwiękowymi, klawiaturami, sterownikami i sekwencerami (w tym sekwencerami wykorzystującymi komputery PC). Oryginalny standard MIDI został uzgodniony w 1983 roku przez grupę producentów instrumentów muzycznych i od tego czasu jest używany i rozszerzany.

Ostatnim razem przyjrzelśmy się MIDI, zanim zjawisko korzystania z modułów Arduino w pełni się rozwinęło. Arduino sprawiło, że bardzo łatwo jest połączyć elektronikę ze sprzętem MIDI. Elektrycy i muzycy stworzyli za pomocą Arduino szereg modułów sterujących, instrumentów, a nawet syntezytorów z różnymi dźwiękami.

Standard MIDI 2.0 został wypuszczony w styczniu 2020 roku i ma być kompatybilny wstecz z oryginalną specyfikacją MIDI. Nowy standard nie jest jeszcze powszechnie używany i wciąż przechodzi testy.

Interfejs sprzętowy?

Standard określa występujące zdarzenia (takie jak początek lub koniec odtwarzanych nut), które są zawarte w komunikatach przesyłanych między urządzeniami. „Oryginalne” MIDI wykorzystuje protokół szeregowej transmisji danych z prędkością 31,25 kb/s przy użyciu asynchronicznej sygnałowej pętli prądowej 5 mA, z prądem dostarczanym przez stronę nadawczą.

Oznacza to, że transmisja każdego bajtu komunikatu MIDI zajmuje tylko 320 µs (wliczając bity startu i stopu). Ponieważ komunikaty MIDI mają długość jednego, dwóch lub trzech bajtów, oznacza to, że co sekundę przez pojedynczy kabel MIDI może być wysyłanych ponad 1000 takich komunikatów.

Każdy kabel MIDI przenosi tylko jeden sygnał, więc do komunikacji dwukierunkowej należy użyć dwóch kabli. Same kable wykorzystują ekranowany przewód dwużyłowy.

Wszystkie kable MIDI są wyposażone na obu końcach w standardowe 5-kołkowe wtyczki DIN 180°. Jednak tylko kołki 4 i 5 są używane do rzeczywistej pętli sygnalizacji prądowej (okablowane 4-4 i 5-5). Kołki 1 i 3 pozostają niepodłączone, podczas gdy oplot ekranujący w każdej wtyczce jest podłączony do kołka 2.

Wewnątrz urządzeń MIDI, styk 2 jest podłączony do uziemienia tylko w gniazdach MIDI-out. Zapewnia to ekranowanie przez uziemione oploty ekranujące kabli bez tworzenia problemów z pętlą masy.

W przeciwieństwie do większości innych protokołów sygnalizacyjnych pętli prądowej, prąd płynie w łańcuchu MIDI tylko wtedy, gdy przesyłane są dane. Pozwala to na bezproblemowe podłączanie i odłączanie kabli MIDI, o ile nie są one aktywnie używane.

Wszystkie wejścia MIDI mają przy pomocy transoptora izolację galwaniczną i elektrostatyczną na napięcie 3 kV, aby zapobiec uszkodzeniu sprzętu spowodowanym błędami okablowania lub usterkami komponentów. Aby zapewnić prawidłową komunikację MIDI między urządzeniami, gniazdo MIDI-out lub MIDI-thru na jednym końcu musi być podłączone do gniazda MIDI-in na drugim końcu.

To właśnie nazywamy „sprzętowo” implementacją MIDI. Istnieje również implementacja USB, która pozwala na znacznie szybszą komunikację. Nie jest to jedynie tłumaczenie typu USB-serial (choć istnieje oprogramowanie do korzystania w tym celu z konwerterów USB-serial).

Interfejs USB MIDI Cat XC4934 firmy Jaycar jest jednym z przykładów dostępnego sprzętu do tłumaczenia między tymi dwoma protokołami w celu połączenia urządzeń.

Ponieważ MIDI to niewiele więcej niż seria bajtów danych, zastoso-
wano różne inne „transporty sprzętowe”. Należą do nich FireWire, LAN, a nawet metody transmisji bezprzewodowej.

Protokół transmisji danych

Każdy komunikat zaczyna się od bajtu, który ma ustawiony na wysokim poziomie najbardziej znaczący bit, a pozostałe bajty (dla praktycznie wszystkich komunikatów) mają najbardziej znaczący bit ustawiony na niskim poziomie. Oznacza to, że bardzo łatwo jest przesyłać dane w seriach i pakietach w celu konwersji na inne techniki przesyłania.

W większości systemów MIDI istnieje jeden główny sterownik lub sekwencer, z którego pochodzi większość komunikatów MIDI (często jest to komputer, klawiatura lub DAW). Gdy komunikaty te muszą zostać wysłane do więcej niż jednego instrumentu, mogą być one dystrybuowane w sposób „gwiazdasty” lub „łańcuchowy”, w zależności od potrzeb.

Możliwe jest połączenie dwóch strumieni MIDI. Jednak sprzęt do tego celu nie jest trywialny w implementacji, ponieważ musi obsługiwać przypadek, gdy dwa komunikaty docierają w tym samym czasie i kolejować je do następnych wyjść bez nadmiernego ich opóźniania.

Nie ma potrzeby martwić się o faktyczne komunikaty kodowe wysyłane przez łańcucha MIDI. W dzisiejszych czasach jest to obsługiwane przez sekwencer lub inne oprogramowanie działające na komputerze oraz przez oprogramowanie układowe działające w innych instrumentach i klawiaturach.

Prawdopodobnie wystarczy wiedzieć, że większość komunikatów MIDI to krótkie polecenia przypisujące konkretny instrument do określonego kanału, nakazujące mu rozpoczęcie lub zatrzymanie odtwarzania określonej nuty, zmianę narastania/zanikania dźwięku instrumentu lub innych parametrów wydajności i tak dalej.

Jak wspomniano wcześniej, polecenia te mają zazwyczaj postać trybajtowych komunikatów (patrz Rysunek 6). Jednak niektóre komunikaty konfiguracyjne i/lub zarządzające systemem mają tylko jeden lub dwa bajty długości.

Istnieją również dłuższe, specyficzne dla sprzętu komunikaty konfiguracyjne, na przykład w celu załadowania cyfrowych próbek audio do samplera.

Format pliku

Korzystając z programu do edycji i sekwencjonowania muzyki na komputerze i być może z klawiatury muzycznej MIDI do wprowadzania rzeczywistych nut, można złożyć pełną sekwencję poleceń MIDI, aby odtworzyć utwór muzyczny – np. na „instrumentach” w syntezytorze.

Można następnie zmusić syntezytor do „wykonania” tego utworu muzycznego, wysyłając do niego sekwencję za pośrednictwem łańcucha MIDI. Gdy jesteś zadowolony z wyniku, możesz zapisać sekwencję na dysku jako plik muzyczny MIDI. Mają one standardowy format i są oznaczone rozszerzeniem „.mid”.

Format pliku .mid jest zasadniczo serią komunikatów MIDI po nagłówku, przechowywanych w kawałkach (w podobny sposób jak kawałki plików .WAV) i przeplatanych danymi czasowymi, aby zapewnić prawidłowe odtwarzanie komunikatów.

Ważne jest, aby zdać sobie sprawę, że chociaż plik muzyczny MIDI może wyglądać powierzchownie podobnie do pliku .wav cyfrowego nagrania dźwiękowego, to tak naprawdę jest zupełnie inny. Jest to raczej elektroniczny odpowiednik nut – po prostu sekwencja szczegółowych instrukcji opisujących sposób odtwarzania muzyki.

W tym przypadku są to instrukcje dla instrumentów elektronicznych, a nie dla muzyków. W zależności od instrumentu, który zapewnia odtwarzanie, dźwięk może się znacznie różnić.

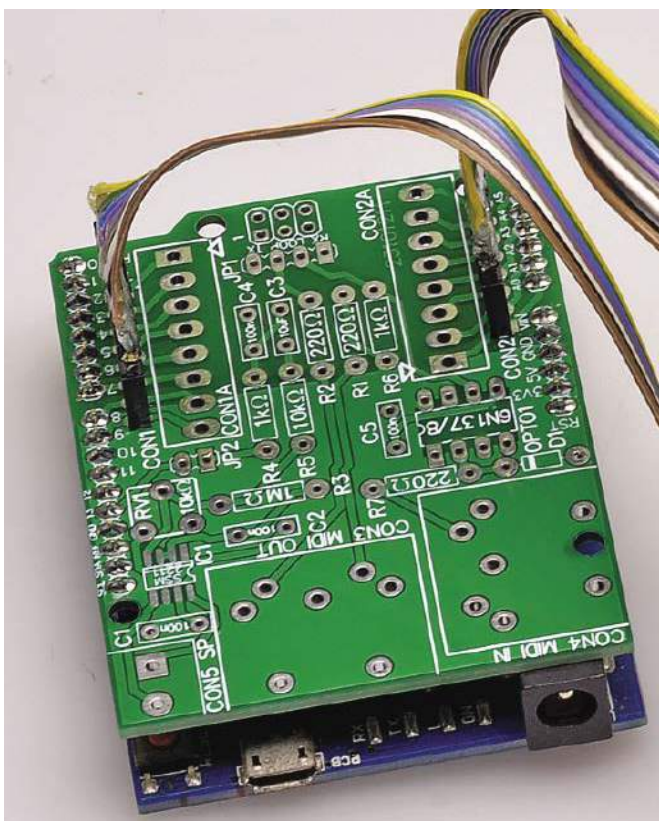
Oprogramowanie

Pojawienie się płytek kompatybilnych z modułami Arduino ze zintegrowanym peryferyjnym portem USB, takich jak Leonardo (wykorzystujący mikroprocesor ATmega32u4), ułatwia implementację interfejsu USB MIDI. Istnieje również wiele programów komputerowych, które mogą współpracować z urządzeniami USB MIDI.

Podczas testowania naszego projektu eksperymentowaliśmy z MuseScore (<https://musescore.org/en>), programem o otwartym kodzie źródłowym do pisania partytur oraz Anvil Studio (www.anvilstudio.com), programem do cyfrowej stacji roboczej audio.

Oba mogą działać jako syntezytor, generując dźwięki w oparciu o przychodzące komunikaty MIDI.

Jeśli chcesz zagłębić się w szczegóły techniczne MIDI, zajrzyj na stronę MIDI Manufacturers Association pod adresem www.midi.org/specifications.



Widok nakładki do sterowania podświetleniem klawiszy diodami LED, podłączonej do modułu Arduino Leonardo „na kanapkę”. Płytką drukowaną jest tutaj wyposażona tylko w jednorzędowe gniazda żeńskie do połączenia z matrycą przełączników

Ekran 2 jest wyświetlany po wybraniu takiego urządzenia, dając sterowanie urządzeniem MIDI. Z ikony menu w lewym górnym rogu wybierz Ustawienia (Ekran 3) i wybierz „Pobierz SoundFont”. Opcja Chaos V20 (Ekran 4) jest najmniejszą z opcji do załadowania i daje wiele różnych dźwięków.

Powrót do ekranu głównego (Ekran 2) umożliwia skonfigurowanie różnych kanałów i instrumentów. Po wykonaniu tej czynności można użyć kodera MIDI do odtwarzania przez głośnik telefonu.

Rozwijane menu Kanał i Instrument umożliwiają przypisanie określonych instrumentów do różnych kanałów. Po ustawieniu każdej kombinacji kanału/instrumentu należy nacisnąć przycisk „SEND PROGRAM”, aby ją aktywować.

Para 2-stykowych gniazd żeńskich przylutowanych do siebie, ze zmostkowanymi wyprowadzeniami, może być użyta jako zworka, która zapewnia dodatkowe gniazda do podłączenia. Użyliśmy tego układu, aby oddzielić sygnał MIDI z płytki enkodera MIDI od JP1, z dodatkowym przewodem potężniejszym prowadzącym do płytki sterownika podświetlenia LED



Domyślnie enkoder MIDI dostarcza dane tylko na kanale 1, ale stworzyliśmy wariant oprogramowania, który zamiast tego wysyła dane na czterech kanałach.

Szkic Arduino nosi nazwę „MIDI_ENCODER_4_CHANNEL”, a po przesłaniu do sprzętu MIDI Enkodera i Matrycy Przycisków będzie generował zdarzenia na kanale 1 po naciśnięciu S1-S16, kanale 2 dla S17-S32, na kanale 3 dla S33-S48 i na kanale 4 dla S49-S64.

Może się okazać, że ta kombinacja działa bardzo dobrze w przypadku generowania efektów dźwiękowych. Szukaj w dolnej części listy instrumentów, wśród instrumentów perkusyjnych.

Wnioski

Dodając drugą płytkę Leonardo i niewiele więcej, możemy dodać sterowanie diodami LED do naszego enkodera MIDI i matrycy przycisków. A ponieważ nasz system wykorzystuje istniejący sprzęt MIDI, nasze oprogramowanie można zmodyfikować, aby zapewnić sterowanie MIDI również w przypadku innych rzeczy.

Korzystanie z naszego enkodera MIDI i matrycy przełączników ze starym telefonem z systemem Android to bardzo ekonomiczny sposób na uzyskanie dostępu do pełnej gamy odtwarzanych dźwięków.

Oczywiście nie jesteś ograniczony w korzystaniu z tej aplikacji do naszego kodera MIDI. Możesz zmodyfikować szkic Arduino, aby zapewnić własny interfejs, a nawet podłączyć inne urządzenie USB MIDI. ■

Tim Blythman

Adaptacja do wydania polskiego – Andrzej Nowicki

Odnosiniki:

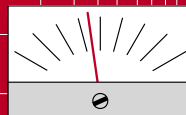
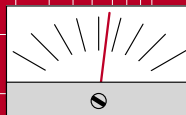
- Pliki SoundFont: <https://musescore.org/en/handbook/3/soundfonts-and-sfz-files> (najnowsza wersja)
- Aplikacja syntezatora FluidSynth MIDI: <https://play.google.com/store/apps/details?id=net.volcanomobile.fluidsynthmidi>

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

Słowniczek

- **Kanał:** Każdy komunikat MIDI może określać jeden z 16 kanałów, umożliwiając kierowanie danych do lub z różnych instrumentów, zachowując ich źródło lub identyfikując, że powinny być odtwarzane na określonym instrumencie. Nasze oprogramowanie używa jednego, stałego kanału, który można zmienić w kodzie.
- **Komunikat:** Najmniejszą jednostką danych MIDI jest komunikat i odpowiada on pojedynczemu zdarzeniu (np. naciśnięciu lub zwolnieniu klawisza fortepianu) lub ustawieniu, które ma zostać zmienione.
- **Numer nuty:** Każda nuta muzyczna jest powiązana z numerem w zakresie 0-127. Środkowe C jest przypisane do numeru 60. Odpowiada to mniej więcej klawiaturze fortepianu, z uwzględnieniem półtonów. Projekt Redakcji Silicon Chip-a implementuje 64 z nich (w matrycy 8×8), z początkiem i końcem zakresu zdefiniowanym w kodzie.
- **Komunikaty Note-on i Note-off:** Komunikaty Note-on i Note-off odpowiadają zdarzeniom, występującym podczas odtwarzania muzyki i są jedynymi komunikatami generowanymi przez enkoder. Zawierają one informacje o kanale, prędkości (patrz poniżej) i numerze nuty.
- **Status:** Pierwszy bajt komunikatu nazywany jest bajtem statusu i jest oznaczony jako taki poprzez ustawienie jego najbardziej znaczącego bitu. Zazwyczaj dolny „nybble” (cztery bity) bajtu statusu zawiera czterobitową wartość wskazującą numer kanału.
- **Prędkość:** Velocity to szybkość naciskania klawisza na instrumencie muzycznym, ale zwykle jest ona manifestowana jako głośność dźwięku nuty (co jest związane z szybkością naciskania np. klawisza fortepianu). Nasze oprogramowanie używa prędkości 64, która jest domyślna dla urządzeń, które nie mogą wykręcić prędkości naciskania klawiszy. Podobnie jak numer Note, może mieć wartość 0-127, przy czym 0 odpowiada również wyłączeniu Note na niektórych urządzeniach.

AUDIO OUT



Wzmacniacz audio do Theremina, część 4

W zeszłym miesiącu, w części 3, przyjrzelśmy się niektórym wersjom wzmacniacza do Theremina zbudowanym na tranzystorach germanowych. W tym odcinku szczegółowo opiszemy opcje komponentów i zobaczymy, jak sprawdziły się różne strategie projektowania.

Radiatory

Do wzmacniacza krzemowego, wystarczy standardowy radiator przypinany. Tranzystory germanowe wymagają jednak znacznie większych radiatorów. Aby zapewnić im długą żywotność, tranzystory powinny być ledwo ciepłe. Mocowanie tranzystorów do konstrukcji metalowej obudowy lub chassis jest rozsądne. Gotowe płytki i układy radiatorów pokazano na rysunkach 12 i 13.

Dane techniczne

Wzmocnienie obu wzmacniaczy wyniosło 16,7 (24 dB). Jest to w zupełności wystarczające do obsługi tunera radiowego. W przypadku korzystania z odtwarzacza CD wartość R1 należy zwiększyć do 27 kΩ. Wzmacniacz germanowy pracował do napięcia 5 V, podczas gdy krzemowy do około 6,5 V.

Moc wyjściowa

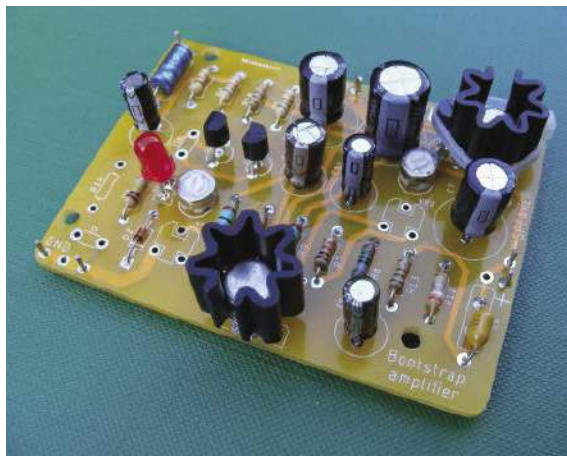
Zgodnie z oczekiwaniami, wzmacniacz na tranzystorach germanowych dał większą moc, dostarczając 880 mW RMS (7,5 V_{pk-pk}).

Obwód na tranzystorach krzemowych dał 720 mW (6,5 V_{pk-pk}). Było to mniej niż w przypadku rozwiązań Siemens, ze względu na dodanie wtórnika emiterowego (TR2). Stopień ten umożliwia znaczne niedopasowanie tranzystorów wyjściowych ze względu naysterowanie niskiej impedancji obciążenia. Wtórnik emiterowy, pracujący w klasie A, zasilany jest wysokim prądem 17 mA, co pozwala na wykorzystanie skromnego h_{FE} tranzystorów germanowych i starych krzemowych w obudowach TO5, które często może być rzędu 40. Zatem całkowity prąd spoczynkowy wynosi około 20 do 30 mA, co wyklucza praktycznie pracę z baterii. Maksymalny prąd obcinania wynosił 190 mA dla obwodu krzemowego i 220 mA dla obwodu germanowego.

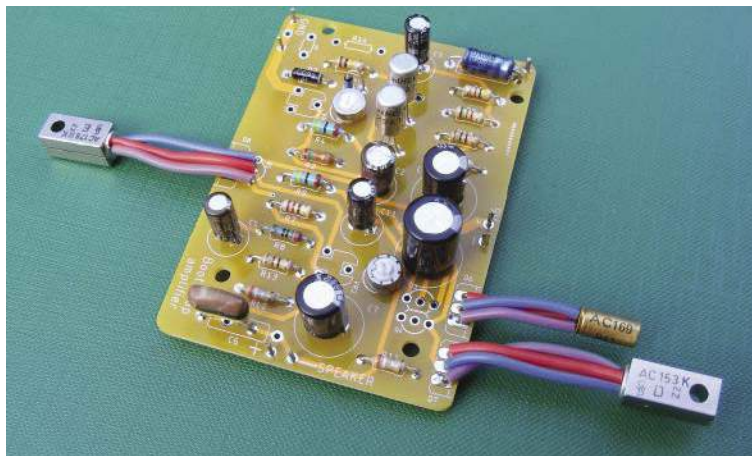
Użycie germanowego stopnia wyjściowego, ale z krzemowymi TR1 i TR2, dało 810 mW (7,2 V_{pk-pk}). W ramach eksperymentu stwierdzono, że wzmacniacz germanowy zapewnia 1,32 W (6,5 V_{pk-pk}) przy 4 Ω, ale tylko przy dobrych radiatorach.

Zniekształcenia

Zniekształcenia skrośne były bardziej miękkie i płytsze z germanowym stopniem wyjściowym oraz potrzebny był mniejszy prąd spoczynkowy (3 mA), aby je wygładzić. Wzmacniacz z tranzystorami krzemowymi potrzebował I_q około 10 mA. Bez ujemnego sprzężenia zwrotnego germanowy stopień wyjściowy teoretycznie wytwarzałby mniej zniekształceń. Jednakże niższe wzmocnienie w otwartej pętli obwodu germanowego oznaczało, że współczynnik ujemnego sprzężenia zwrotnego był mniejszy, co zmniejszało redukcję zniekształceń. Dało to również efekt bardziej miękkiego przycinania sygnału. Wzmocnienie pętli zostało zmniejszone głównie przez niższą wartość R2 zastosowaną we wzmacniaczu germanowym w celu zmniejszenia wpływu prądów upływowych. W rezultacie całkowite zniekształcenia harmoniczne były we wzmacniaczu germanowym 2,5-krotnie wyższe w środkowym paśmie (500 Hz) i 10-krotnie wyższe przy wysokich częstotliwościach (8 kHz).



Rysunek 12. Gotowy wzmacniacz krzemowy. Należy zwrócić uwagę, że tranzystor polaryzacji TR3 jest przymocowany do radiatora tranzystora wyjściowego TR4, zapewniając pewien stopień stabilizacji termicznej prądu spoczynkowego



Rysunek 13. Gotowy wzmacniacz germanowy. Tranzystory wyjściowe muszą być przykręcone do wystarczającej ilości metalu, aby były chłodne, prawie w temperaturze otoczenia. Jeśli tak się stanie, tranzystor polaryzacji można pozostawić nieprzykręcony do radiatora, aby śledził temperaturę otoczenia

Słyszalny efekt tych różnic można opisać jako brzmienie wzmacniacza germanowego „miększe i bardziej rozmyte”, podczas gdy krzemowy dźwięk jest „czystszy, ale ostrzejszy” po przycięciu sygnału. Wzmacniacz germanowy był lepszy do ćwiczeń na gitarze, a wersja krzemowa do Hi-Fi. Krzywe zniekształceń pokazano na **rysunkach 14 i 15**.

Na **rysunku 16** widać efekt zastosowania szybkich tranzystorów 2SA12 jako TR1 i TR2, zmniejszających zniekształcenia przy 10 kHz z 1,7% do 0,35%. Technicznie rzecz biorąc, wzmacniacz germanowy był „mniej liniowy” niż krzemowy. Jednakże wzmacniacz

krzemowy o tym samym poziomie całkowitych zniekształceń harmonicznych brzmiałby gorzej ze względu na większy udział wysokich harmonicznych. Ostatecznie fizyka działania tranzystora bipolarnego jest taka sama, wykładniczy proces konwersji napięcia na prąd jest taki sam, niezależnie od tego, czy używasz germanu, czy krzemu. To drugorzędne różnice wpływają na ogólne wyniki. Kiedy uruchomię analizator widma, przyjrzymy się temu głębiej. Co ciekawe, gitarzyści są gotowi zapłacić 700 funtów za małe wzmacniacze germanowe, takie jak wzmacniacze Deacy (od Red. EdW: konstrukcja basisty grupy Queen,

Johna Deacona). Podejrzewam, że dzieje się tak z powodu braku klarowności jaką charakteryzują się te niedoskonałe wzmacniacze, maskującej błędy techniczne grających!

Połącz oba rozwiązania

Najlepszy wynik uzyskano z krzemowym tranzystorem wejściowym TR1, typu BC557B, resztę pozostawiając elementom germanowym (TR2, 2SA12). Krzywą zniekształceń dla tego układu pokazano na **rysunku 17**. Wykorzystanie tranzystora krzemowego jako TR2 spowodowało, że zniekształcenia środkowego pasma były nieco gorsze i wzmacniacz dawał mniejszą moc wyjściową. Brytyjska firma Leak postąpiła słusznie, wymieniając tranzystor wejściowy OC44 na BC153 w wzmacniaczu Stereo 30.

Odtwarzanie basów

Kondensatory ograniczają od dołu (−3 dB) częstotliwość przenoszenia wzmacniacza do około 65 Hz. Jest to dobre rozwiązanie w przypadku małych głośników z papierowymi membranami, które są używane w zastosowaniach o małej mocy (wzmacniacze do ćwiczeń gitarowych), które zazwyczaj obcinają pasmo przy około 100 Hz. Dla uzyskania prawdziwie zabytkowego brzmienia gitary/radia sugeruję użycie głośników Celestion i Philips 4072 pokazanych na **rysunku 18**. Pasma przenoszenia wzmacniaczy krzemowych i germanowych pokazano na **rysunkach 19 i 20**.

Szumy

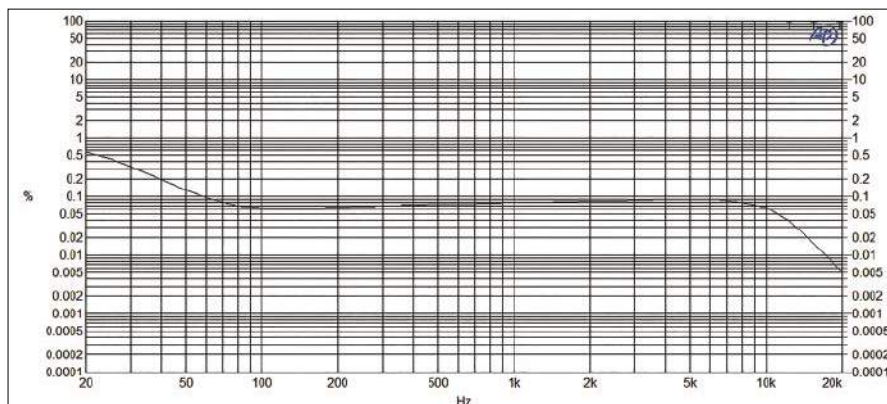
Poziom szumów w obu obwodach wynosił 1 mV_{pk-pk}, mierzony przy wejściu zwartym do masy. W konstrukcji krzemowej było więcej szumów o wysokiej częstotliwości, a w wersji germanowej więcej szumów o niskiej częstotliwości. Tranzystor germanowy był bardzo szumiący i musiał być wymieniony w stopniu wejściowym. Obydwie wersje układu sprawdziły się jako wzmacniacze słuchawkowe.

Zabezpieczenie przed przeciążeniem

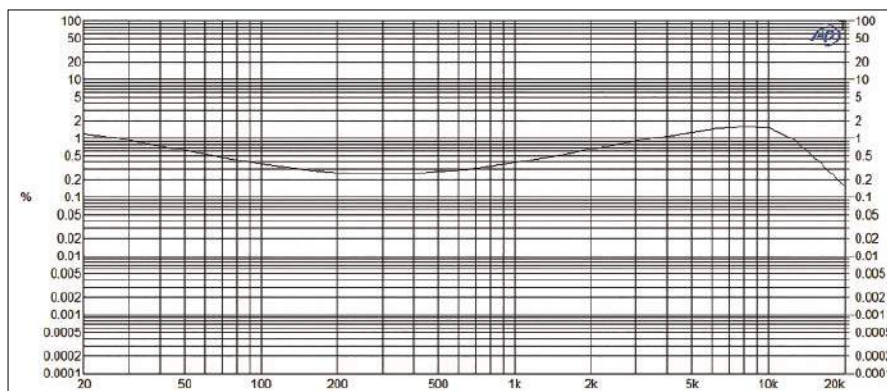
Całkowity spadek napięcia, w kierunku przewodzenia, sieci diod przeciwzwarciowych (D1 i D2) musiał wynosić 1 V. Uzyskano to poprzez zsumowanie spadku napięcia na diodzie Zenera w kierunku przewodzenia (0,75 V) i diodzie germanowej (0,25 V).

Lista elementów

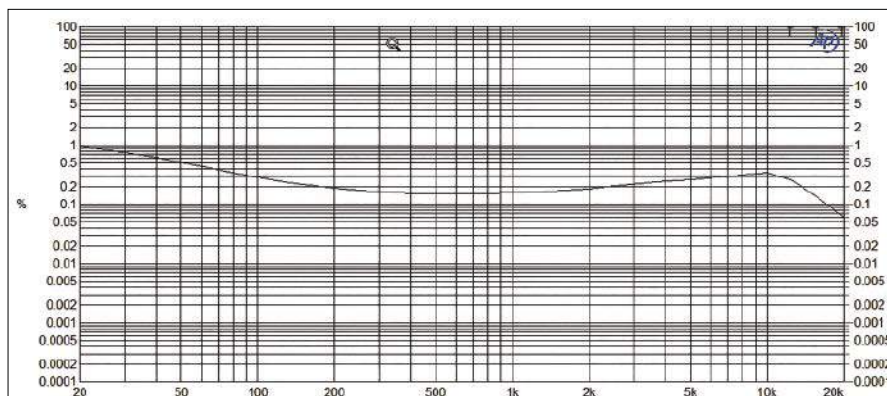
PCB – jest to ta sama płytka, która jest używana w artykułach w numerach listopadowym i grudniowym 2020 *Wyjścia audio*, część nr AO-1220-01, dostępna w serwisie PCB PE na stronie: www.electronpublishing.com.



Rysunek 14. Zniekształcenia wzmacniacza krzemowego: przy 3 Vpk-pk na 8 Ω



Rysunek 15. Zniekształcenia wzmacniacza germanowego w tych samych warunkach. Widać, że jest ich znacznie więcej niż w przypadku wzmacniacza krzemowego. Wykorzystuję tu wolne tranzystory typu NKT214 jako TR1 i TR2



Rysunek 16. Zastosowanie szybkich tranzystorów RF typu 2SA12 jako TR1 i TR2 poprawiło zniekształcenia wzmacniacza germanowego. To był trik, którego Leak użył przy projektowaniu wzmacniacza Stereo 30

Wykaz elementów:

Rezystory:

Wszystkie, rezystory są węglowe 0,25 W, 5% (dla wersji krzemowej, wartości podano w nawiasach)

R1 = 4,7 kΩ

R2 = 15 kΩ (47 kΩ dla krzemowego TR1, takiego jak BC557B)

R3 = 100 kΩ

R4, R9 = 47 Ω

R5, R6 = 1,5 kΩ

R7 = 220 Ω (1 kΩ)

R8 = 160 Ω

R10, R13 = 10 Ω

R11, R12 = 0,39 Ω

R14 nieużywany

VR1 = 500 Ω, pot. montażowy, obrys T05

VR2 = 5 kΩ, pot. montażowy, obrys T05

Kondensatory:

Wszystkie elektrolityczne są radialne, o napięciu 10 V lub wyższym, chyba że podano inaczej.

C1 = 10 μF

C2 = 100 μF

C3 = 22 μF

C4 = 470 μF

C5 = 100 μF

C6 = 100 nF, ceramiczny lub poliestrowy

C7 = 470 μF

C8 = 1000 μF

C9 nie używany

C10 = 33 pF ceramiczny (zwykle nie używany, zobacz tekst.)

C11 = 10 μF

Półprzewodniki:

TR1, TR2 – w przypadku wzmacniacza Ge użyj mało-sygnałowego pnp wysokiej częstotliwości OC44, 2SA12, GET872 lub tranz. audio takie jak OC71, OC75, NKT214; dla wzmacniacza Si użyj BC549.

TR3 – w przypadku wzmacniacza Ge użyj AC153, AC128 lub dowolnego germanowego średniej i małej częstotliwości. Można również zastosować tranzystor polaryzujący (Brimar/Thorn/Mazda) typu AC169; w przypadku wzmacniacza Si użyj BC337.

TR4 – w przypadku wzmacniacza Ge użyj germanowego tranzystora pnp mocy typu AC153K, AC188; w przypadku wzmacniacza Si użyj krzemowych tranzystorów npn średniej mocy BC138, BC140, BFY51 w obudowie TO5.

TR5 – W przypadku wzmacniacza Ge użyj germanowego tranzystora npn mocy typu AC176K lub AC187; w przypadku wzmacniacza Si należy zastosować tranzystory pnp średniej mocy BC143, 2N2905 w obudowie TO5.

D1 Dla wzmacniacza Ge użyj niskonapięciowego Zenera 3,3 V do 12 V, 400 mW; np. BZY88C5V6; dla wzmacniacza Si użyj standardowej czerwonej diody LED o dużej jasności.

D2 Dla wzmacniacza Ge użyj germanowej diody mało-sygnałowej; np. CG92, OA91; dla wzmacniacza Si użyj małej diody Schottky'ego BAT86

Elementy germanowe są dostępne u autora

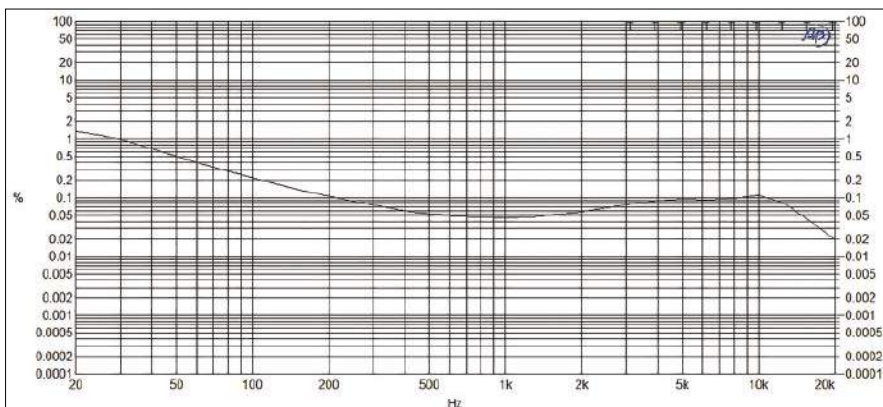
Wyślij e-mail na adres jrothman1962@gmail.com

Dlaczego warto to wszystko zrobić?

Dla historia audio i osoby z obsesją na punkcie komponentów jest to zdecydowanie warte tej pracy. Pokazuje, że komponenty z lat 70. XX wieku spełniły swoje zadanie, a pięć dekad później nadal są użyteczne. ■

Jake Rothman

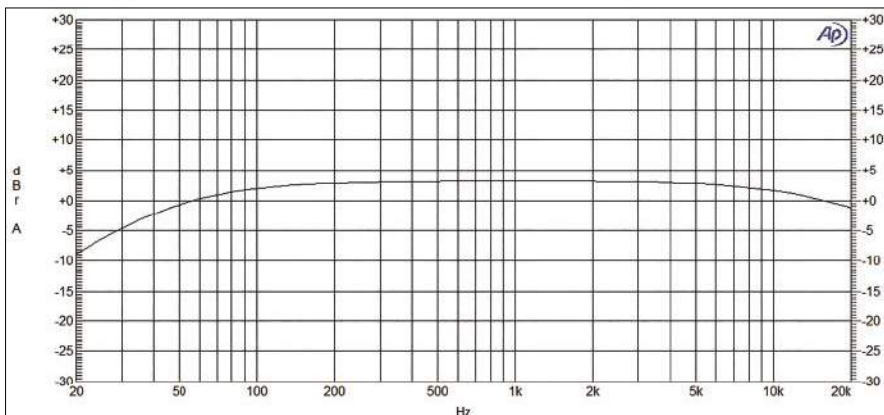
Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, luty 2021 (www.epemag3.com)



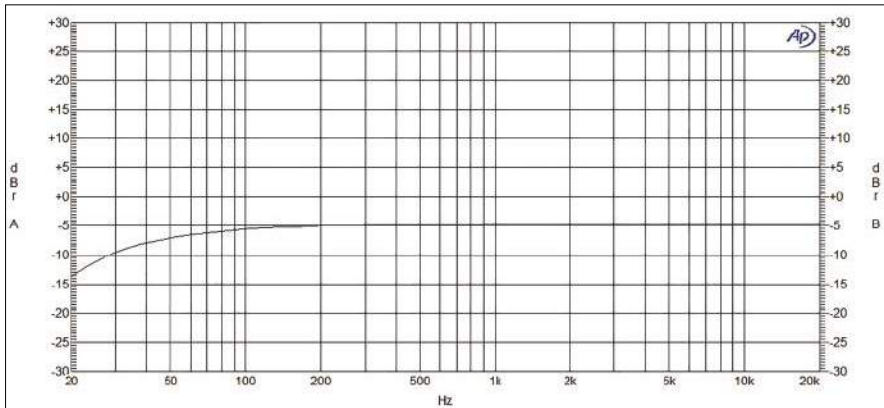
Rysunek 17. Kolejnym ulepszeniem projektu było zastosowanie tranzystora krzemowego BC557B w stopniu wejściowym TR1, przy zachowaniu reszty tranzystorów germanowych. Jako TR2 użyto 2SA12, a R2 ma 47 kΩ. Ta konfiguracja była najlepsza, wytwarzając zaledwie 0,045% zniekształceń dla środkowego pasma. Ponownie Leak zrobił to jako późniejszą modyfikację do wzmacniacza Stereo 30



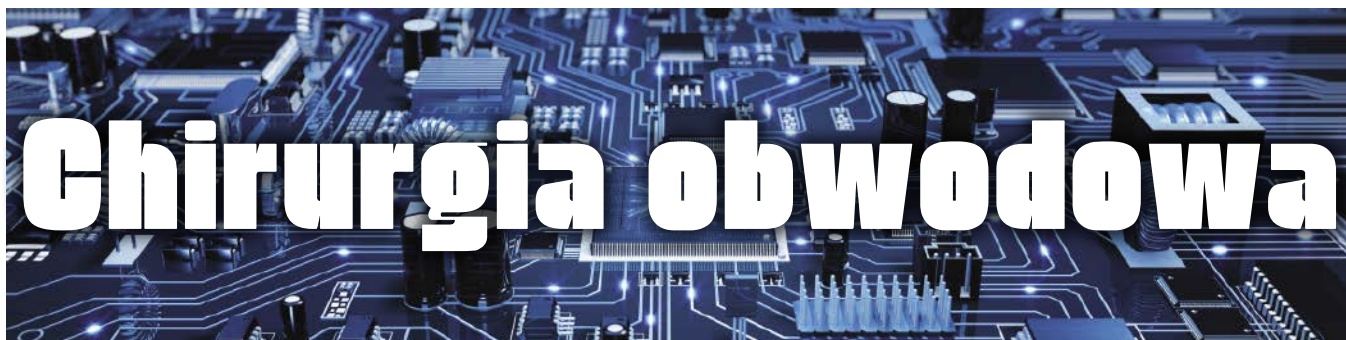
Rysunek 18. Odpowiednie głośniki (dostępne u autora) jako uzupełnienie wzmacniacza na tranzystorach germanowych. Nadają się szczególnie do gitary



Rysunek 19. Pasma przenoszenia wzmacniacza germanowego z wolnymi tranzystorami wejściowymi (NKT214) i podłączonym C10



Rysunek 20. Pasma przenoszenia wzmacniacza krzemowego. Wzmacniacze germanowe wykorzystujące szybko tranzystory dla TR1 i TR2 dawały tę samą charakterystykę przenoszenia



Chirurgia obwodowa

Czas i metastabilność w obwodach synchronicznych, część 2

Nasza dyskusja na temat taktowania cyfrowego i metastabilności rozpoczęła się kilka miesięcy temu, kiedy badaliśmy symulację cyfrowego dzielnika częstotliwości w programie Micro-Cap 12. Obwód oscylował podczas symulacji, ale działał poprawnie jako układ fizyczny. Symulowane zachowanie wynikało z faktu, że wszystkie bramki w symulacji miały dokładnie takie samo opóźnienie. Stworzyło to warunki dla oscylacji, co byłoby mało prawdopodobne w rzeczywistym obwodzie – jednak symulacja uwydatniła fakt, że w rzeczywistym wykonaniu układ może potencjalnie cierpieć z powodu problemów z synchronizacją.

Forum Micro-Cap 12

Jeśli chodzi o program Micro-Cap 12, to odkryliśmy, że niedawno zostało uruchomione forum online użytkowników: „Użytkownicy Micro-Cap EDA” pod adresem mc12.create-aforum.com.

Czytelnicy zainteresowani korzystaniem z tego oprogramowania mogą uznać to forum za przydatne źródło informacji. To niegdyś drogie oprogramowanie zostało udostępnione bezpłatnie w lipcu 2019 r. po zatrzymaniu prac rozwojowych, w związku z czym nie ma już oficjalnego wsparcia.

Podsumowanie synchronicznego taktowania obwodu

Obwód dzielnika był stosunkowo nietypowy, ponieważ został zaprojektowany przy użyciu minimalnej liczby bramek NAND i technik projektowania asynchronicznego, a nie tylko przy użyciu istniejącego gotowego przerzutnika. Dlatego w zeszłym miesiącu przyjrzelśmy się problemom synchronizacji, w znacznie bardziej powszechnym kontekście synchronicznych obwodów cyfrowych. Obwody synchroniczne są sterowane sygnałem zegarowym – regularnym ciągiem impulsów, który kontroluje ogólne taktowanie obwodu.

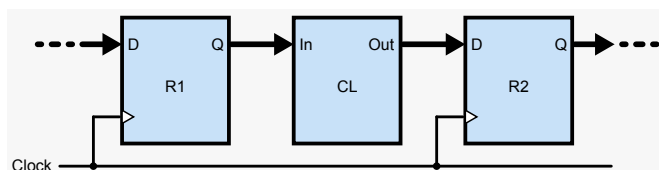
Nawet jeśli obwód synchroniczny ma złożoną ogólną strukturę, to zasadniczo obejmuje transfery między rejestrami, tak jak pokazano to na **rysunku 1** – dane przechowywane w rejestrze (R1) są przetwarzane przez logikę kombinacyjną (CL), a wynik jest zapisywany w (R2). W każdym cyklu zegara, (R1) ładuje nowe dane do przetwarzania, a (R2) przechowuje wynik przetwarzania danych, które były przechowywane w R1 w poprzednim cyklu zegara.

Obwód na **rysunku 1** nie jest nieskończenie szybki. Występuje w nim opóźnienie od chwili wystąpienia aktywnego zbocza zegara do zmiany wyjścia rejestru (T_{DR}) oraz opóźnienia od zmiany wejść logiki kombinacyjnej do chwili, w której możemy zagwarantować, że jej wyjścia są prawidłowe. Musimy również wziąć pod uwagę, że gdy dane się zmieniają, to wewnętrzne obwody przerzutnika potrzebują czasu, aby dostosować się do tej zmiany. Jeśli zegar zostanie aktywowany zbyt blisko zmiany danych, przerzutnik może nie działać poprawnie (mówimy wtedy, że nastąpiło naruszenie taktowania). Może on załadować niewłaściwą wartość lub stać się metastabilny, co może skutkować znacznie dłuższym niż zwykle opóźnieniem przed zmianą sygnału wyjściowego. Aby zapobiec naruszeniom

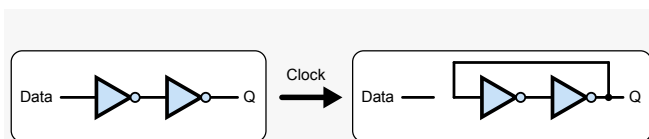
taktowania, przerzutniki mają określone czasy: ustalenia (T_{Setup}) i czas podtrzymania (T_{hold}) – tj. czas przed i po zboczu zegara, podczas którego dane nie mogą się zmieniać, aby zapewnić prawidłowe działanie.

Naruszenia taktowania

Jak to omówiono w zeszłym miesiącu, dla obwodu z **rysunku 1** minimalny okres zegara musi być większy niż $T_{DR} + T_{DC} + T_{Setup}$, aby mieć pewność, że dane załadowane do (R2) są prawidłowe. W obwodzie synchronicznym możemy to zapewnić konstrukcyjnie, co oznacza, że nie będą w nim występowały naruszenia taktowania. Niekoniecznie jest to łatwe, szczególnie w przypadku dużych projektów, gdzie wymagania dotyczące wydajności wymagają wysokich częstotliwości taktowania. Profesjonalne narzędzia do projektowania dużych obwodów cyfrowych (np. projektowania FPGA) zawierają analizatory taktowania, które pomagają zidentyfikować problemy z synchronizacją. Źródła problemów z synchronizacją są bardziej złożone niż tylko warunek odpowiedniego okresu zegara, o którym wspomnieliśmy powyżej. Na przykład w dużym projekcie zegar nie dotrze do każdego przerzutnika dokładnie w tym samym czasie (nazywa się to przesunięciem

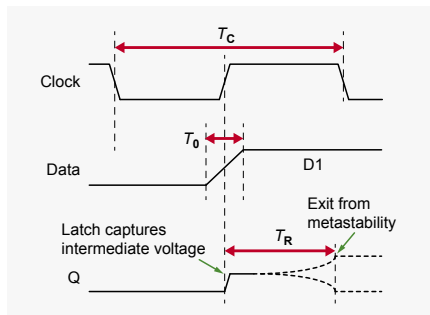


Rysunek 1. Transfer między rejestrami (R1, R2) poprzez blok logiki kombinacyjnej (CL) jest kluczową strukturą w obwodzie synchronicznym

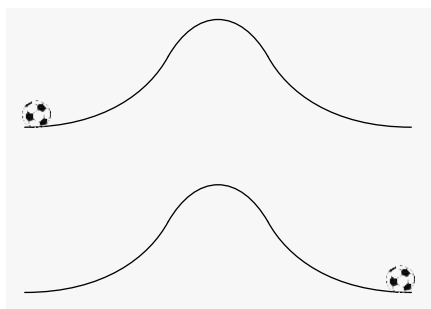


Rysunek 2. Obwód zatrasku przechwytuje 1 lub 0 w pętli pamięci. Metastabilność występuje, jeśli przechwytuje napięcie pośrednie

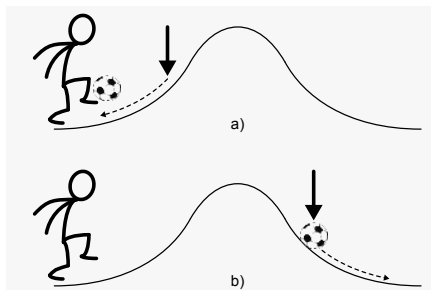
zegara), co również może powodować naruszenia synchronizacji. Niemniej jednak, przy pewnym wysiłku, można zapewnić, że w obwodzie synchronicznym z pojedynczym zegarem nie wystąpią naruszenia taktowania.



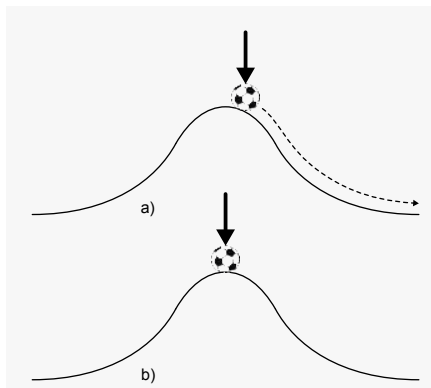
Rysunek 3. Przebiegi metastabilne w zatrasku



Rysunek 4. Analogia do piłki i wzgórza – na równinie po obu stronach znajdują się dwa stany stabilne



Rysunek 5. Analogia do normalnej pracy przerzutnika – piłka łąduje blisko stanu stabilnego i szybko go osiąga



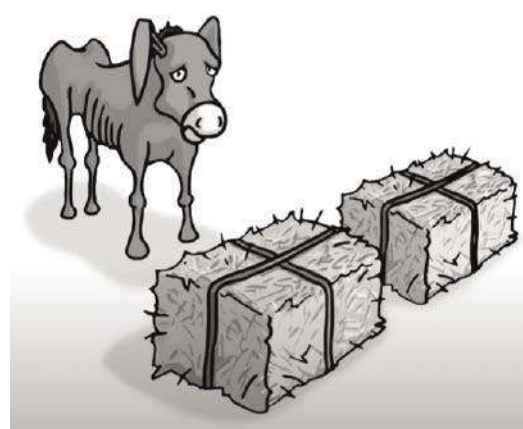
Rysunek 6. Analogia do metastabilnego działania przerzutnika – piłka łąduje blisko (a) lub dokładnie na (b) szczycie wzniesienia i powrót do stanu stabilnego zajmuje dużo czasu

„Gwarancja konstrukcyjna” nie ma zastosowania, gdy mamy zewnętrzne sygnały asynchroniczne – mogą one zmieniać się w dowolnym momencie cyklu zegara, co oznacza, że mogą zmieniać się także wystarczająco blisko aktywnego zbocza zegara, aby spowodować naruszenie taktowania. Podobnie w obwodach z wieloma zegarami (domenami zegara) istnieje możliwość naruszenia synchronizacji, gdy sygnały przekraczają domeny zegara. Istnieje pewien czas (okno metastabilności, T_0), w którym zmiany danych, podczas taktowania zatrasku spowodują metastabilność (patrz rysunki 2 i 3). Jeśli taka wystąpi to zatraskowi zajmie trochę czasu, zwanego „czasem rozdzielczości” (T_R), zanim powróci on do jednego ze stanów stabilnych. Teoretycznie wartość ta może być nieskończona, ale w praktyce jest bardziej prawdopodobne, że będzie się mieścić w zakresie do około dziesięciokrotności czasu opóźnienia propagacji (T_{DR}). W przypadku sygnałów asynchronicznych nie mamy kontroli nad względnym taktowaniem sygnału, dlatego nie możemy zagwarantować, że zapobiegniemy metastabilności. Musimy sobie z tym poradzić w kategoriach prawdopodobieństwa, które omówimy szerzej za chwilę.

Metastabilność: filozofia i analogie

Zanim przyjrzymy się prawdopodobieństwu metastabilności w obwodach, warto przyrzeć się jednej lub dwóm analogiom. Po pierwsze, w zeszłym miesiącu zauważyliśmy, że metastabilność jest jak niezdecydowanie przerzutnika – utknie on w połowie drogi między 0 a 1 i zajmie mu znacznie więcej czasu niż zwykle, zanim ostatecznie osiągnie jeden ze stanów stabilnych. Przypomina to paradoks znany w filozofii jako „osioł Buridana”. Idea jest taka, że zwierzę (osioł) znajduje się dokładnie w połowie drogi pomiędzy dwoma równie pożądanymi potrawami lub napojami i dlatego nie jest w stanie zdecydować, które z nich zjeść – podjęcie decyzji zajmuje mu tak dużo czasu, że umiera z głodu lub pragnienia. Oprócz tego, że idea ta pojawiła się w filozoficznych dysputach na temat rozumu i determinizmu od starożytności (sprzed Buridana), to była wielokrotnie używana w kulturze popularnej. Więcej informacji można znaleźć na stronie Wikipedii poświęconej osiołkowi Buridana (<http://bit.ly/pe-apr21-ass>).

Druga analogia – kula i wzgórze – jest powszechnie używana do omawiania obwodów metastabilnych. Pomaga nam to zrozumieć zmienność czasu rozdzielczości w zależności od taktowania sygnału wejściowego. Ideę ilustruje rysunek 4. Piłka może znajdować się w jednej z dwóch stabilnych pozycji po obu



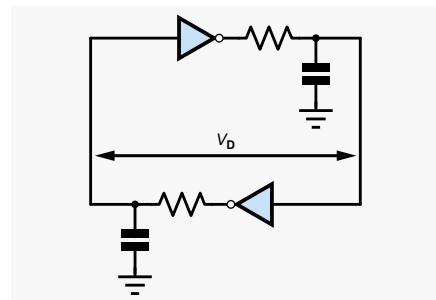
Nieszczęsny osiołek Buridana – dzięki uprzejmości Juliana Mayersa, YouTube

stronach wzgórza – odpowiada to obwodowi zatrasku trzymającemu 0 lub 1. Próba zapisania nowej wartości w zatrasku odpowiada kopnięciu piłki. Aby prawidłowo przeładować ten sam stan, piłka otrzymuje niewielkie kopnięcie i szybko cofa się z niestabilnej pozycji na zboczu do pierwotnego stanu stabilnego (rysunek 5a). Aby gładko zmienić stan, otrzymuje mocne kopnięcie, łąduje wtedy nisko po drugiej stronie i szybko toczy się do drugiej pozycji stabilnej (rysunek 5b).

Ta analogia nie opiera się na szczegółowej fizyce ruchu kopniętych piłek – zakładamy, że piłka spada pionowo na wzgórze. Jeśli piłka otrzyma kopnięcie o średniej sile, odpowiadaające zatraskowi przechowującemu napięcie pośrednie w zakresie od 0 do 1, to wyląduje ona w pobliżu szczytu wzgórza. Przetoczenie się do stanu stabilnego zajmie znacznie więcej czasu (rysunek 6a) lub, w najbardziej ekstremalnym przypadku, piłka będzie balansować dokładnie na szczycie wzgórza i osiągnięcie jednego ze stanów stabilnych zajmie jej potencjalnie nieskończony czas (rysunek 6b).

Analiza obwodu

Pętla inwerterów pokazana na rysunku 2 rejestruje dwa napięcia – na wyjściu każdego inwertera w momencie wystąpienia taktu zegara. Zwykle jeden inwerter ma stan

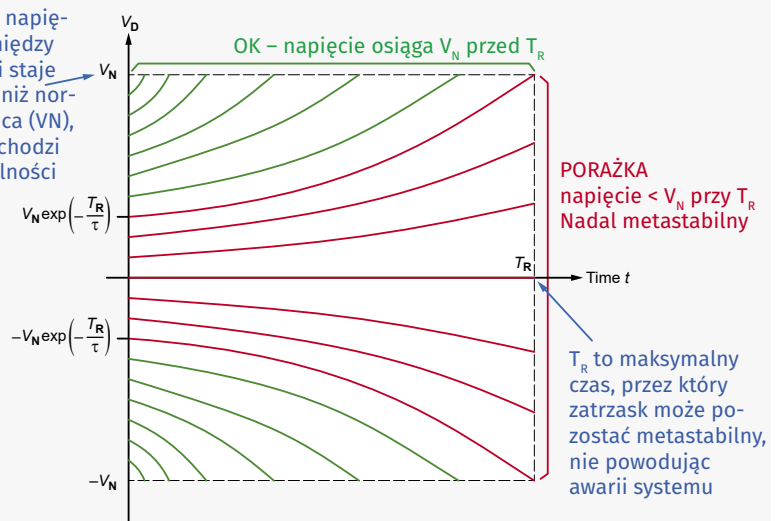


Rysunek 7. Pętla inwerterów (patrz rysunek 2) w zatrasku zachowuje się jak dwa wzmacniacze, każdy napędzający obwód RC

logiczny 0, a drugi 1, zatem różnica napięcia pomiędzy dwoma wyjściami (V_D) jest stosunkowo duża. Jeśli jednak dane zmieniają się w czasie zegara, jak pokazano na rysunku 3, pętla wychwyci niewielką różnicę napięcia. Możemy modelować to, co się dzieje, rozważając pętlę inwerterów jako dwa wzmacniacze podłączone do obwodów RC (okablowanie oraz pojemność oraz rezystancja wejściowa i wyjściowa inwertera) – inwertery działają jak wzmacniacze z pośrednimi napięciami wejściowymi (patrz **rysunek 7**). Wynikiem analizy jest równanie różniczkowe, ale nie będziemy się tutaj wdawać w szczegółowe obliczenia matematyczne. Jednakże czytelnicy zaznajomieni z ładowaniem kondensatora w obwodzie RC mogą nie być zaskoczeni, gdy dowiedzą się, że rozwiązaniem jest funkcja wykładnicza wiążąca V_D w czasie t po zboczu zegara, z początkową różnicą napięcia przechwyconą przez pętlę (V_{D0}). Im mniejsze jest V_{D0} , tym dłużej zatrzask potrzebuje na powrót do normalnego stanu – odpowiada to ładowaniu piłki bliżej szczytu wzniesienia w opisanej powyżej analogii.

Jeśli chodzi o projekt obwodów cyfrowych, chcielibyśmy, aby każdy przerzutnik, który stanie się metastabilny, działał wystarczająco szybko, aby nie powodować żadnych problemów. Zwykle oznacza to jeden cykl zegara, z uwzględnieniem odpowiednich parametrów, takich jak opóźnienia i czas ustalania, w sposób podobny do naszej wcześniejszej dyskusji na temat maksymalnej częstotliwości zegara. Określa to maksymalny czas rozdzielczości (T_R), który możemy tolerować. Rysunek 3 przedstawia dwa możliwe przebiegi napięcia na wyjściu zatrzaskowym (dla równych, ale przeciwnych napięć początkowych). Zostało to rozszerzone na **rysunku 8**, aby pokazać zakres przebiegów wynikających z różnych napięć początkowych. Zatrzask wychodzi z metastabilności, gdy różnica napięcia (V_D) przekracza minimum, które można uznać za „normalną” pracę zatrzasku – z cyfrowymi wartościami 0 i 1 na jego wyjściach – nazywamy to V_N . Z rozwiązania równania różniczkowego pętli można znaleźć związek pomiędzy początkową różnicą napięcia (V_{D0}), czasem rozdzielczości (T_R) a napięciem wyjściowym metastabilności (V_N). Granica między

Gdy różnica napięcia (V_D) pomiędzy inwerterami staje się większa niż normalna różnica (V_N), zatrzask wychodzi z metastabilności



Rysunek 8. Różnica napięć zmienia się w zatrzasku, który wchodzi w metastabilność w chwili $t = 0$

uszkodzonym i nieuszkodzonym obwodem występuje, gdy różnica napięcia osiąga właśnie V_N w T_R (patrz rysunek 8). Dzieje się tak przy określonej początkowej różnicy napięcia V_{D0N} . Z równania obwodu znajdujemy (jeśli rozwiążemy równanie różniczkowe):

$$V_{D0N} = V_N \exp\left(-\frac{T_R}{\tau}\right)$$

Tutaj τ jest stałą czasową pętli zatrzaskowej – zależy od wartości rezystorów i kondensatorów oraz wzmacnienia wzmacniacza (rysunek 7).

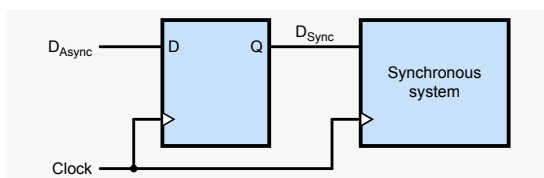
Prawdopodobieństwa i MTBF

Jeżeli zdarzy się, że zegar wystąpi w przedziale czasowym oznaczonym jako T_0 na rysunku 3, wówczas zatrzask przejdzie w stan metastabilny. W tym okresie założymy, że wszystkie napięcia początkowe (V_{D0}) występują z równym prawdopodobieństwem. Prawdopodobieństwo, że zatrzask ulegnie awarii (zakładając, że przede wszystkim stał się metastabilny) to prawdopodobieństwo, że zatrzask jest nadal metastabilny po akceptowalnym T_R – nazwijmy to P_S (prawdopodobieństwo „nadal metastabilne”). Jest to prawdopodobieństwo, że V_{D0} jest mniejsze niż V_{D0N} . Przy założeniu, że wszystkie wartości V_{D0} są jednakowo prawdopodobne, to prawdopodobieństwo awarii jest po prostu proporcją metastabilnych wartości V_{D0} mniejszych niż V_{D0N} . Maksymalna wartość V_{D0} dla metastabilności

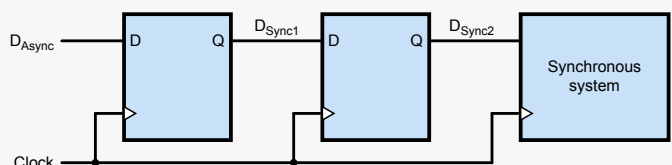
to V_N , ponieważ wyższe napięcia początkowe oznaczają normalną pracę. Zatem prawdopodobieństwo awarii po wejściu w metastabilność wynosi $P_S = V_{D0N}/V_N$. Z powyższego równania wykładniczego otrzymujemy:

$$P_S = \exp\left(-\frac{T_R}{\tau}\right)$$

Daje to prawdopodobieństwo, że zatrzask ulegnie awarii, jeśli stał się metastabilny, ale aby uzyskać ogólne prawdopodobieństwo uszkodzenia (P_F), musimy również znać prawdopodobieństwo, że zatrzask w pierwszej kolejności wejdzie w stan metastabilności (P_E). Jest to prostsze do obliczenia i zostało wspomniane w artykule z zeszłego miesiąca. Prawdopodobieństwo, że zatrzask stanie się metastabilny, to w zasadzie proporcja cyklu zegara zajmowana przez T_0 , czyli $P_E = T_0/T_C$ – zakładamy, że sygnał asynchroniczny może zmienić się w punkcie cyklu zegara z równym prawdopodobieństwem. Możemy to również zapisać jako $P_E = f_c T_0$, gdzie f_c to częstotliwość zegara. Prawdopodobieństwo uszkodzenia zatrzasku (P_F) to prawdopodobieństwo, że zarówno wejdzie on w stan metastabilny, jak i że nadal będzie metastabilny po upływie akceptowalnego czasu rozwiązania. Jeśli coś zależy od dwóch warunków występujących razem, mnożymy prawdopodobieństwa, aby znaleźć prawdopodobieństwo całkowite, więc $P_F = P_E P_S$.



Rysunek 9. Pojedynczy synchronizator przerzutnikowy chroniący system synchroniczny przed metastabilnością spowodowaną asynchronicznym wejściem



Rysunek 10. Synchronizator z dwoma przerzutnikami

Obwody cyfrowe przetwarzają informacje w sposób ciągły, więc pojedyncze prawdopodobieństwo awarii nie jest zbyt przydatne. Bardziej interesuje nas, jak często obwód będzie ulegał awariom. Biorąc pod uwagę asynchroniczne wejście do zatrzaśku synchronicznego, stopień awarii będzie określony przez szybkość zmiany danych (szybkość transmisji danych f_D) i prawdopodobieństwo awarii (P_F z góry), które występuje przy każdej zmianie danych. Otrzymujemy:

$$\text{czestotliwoscawarii} = f_D P_F = f_D P_E P_S = V_{D0N} = f_D f_C T_0 \exp\left(-\frac{T_R}{\tau}\right)$$

Często omawia się niezawodność systemu w kategoriach średniego czasu między awariami (MTBF), który jest po prostu odwrotnością wskaźnika awaryjności (1/wskaźnik awaryjności). MTBF dla przerzutnika wynosi:

$$MTBF = \frac{\exp\left(\frac{T_R}{\tau}\right)}{f_D f_C T_0}$$

Zwróć uwagę na zmianę znaku w funkcji wykładniczej w wyniku przyjęcia odwrotności.

Synchronizatory

Typową strategią pozwalającą uniknąć błędów wynikających z metastabilności powodowanej przez wejścia asynchroniczne do systemów synchronicznych jest dodanie przerzutnika, taktowanego przez zegar systemowy, pomiędzy sygnałem asynchronicznym a wejściem systemu (patrz **rysunek 9**). Nazywa się on „synchronizatorem”. Idea jest taka, że przerzutnik synchronizatora może stać się metastabilny, o ile powróci do stanu pierwotnego przed następnym cyklem zegara – dokładnie w tym scenariuszu, dla którego obliczyliśmy współczynnik MTBF powyżej. Co ważne, dodanie synchronizatora nie wyeliminuje możliwości awarii układu, ale będzie ona mniejsza niż w przypadku bezpośredniego wprowadzenia sygnału. Powyższe równanie MTBF wskazuje wydajność synchronizatora, ale musimy umieć poprawnie zinterpretować wyniki. Zazwyczaj 63% elementów ulegnie awarii w określonym czasie MTBF, więc generalnie musi on być znacznie dłuższy niż akceptowalny okres życia systemu bez błędów. Problem pogłębia się w przypadku dużych obwodów cyfrowych, które zawierają dużą liczbę synchronizatorów – system może ulec awarii, jeśli zawiedzie którykolwiek z nich.

Jeśli znamy wartości τ i T_0 , które zależą od konkretnej technologii i zastosowanych przerzutników, możemy obliczyć MTBF. Zegar i szybkości transmisji danych powinny być znane ze specyfikacji systemu, a T_R to zazwyczaj okres zegara minus czas konfiguracji wejścia

systemu synchronicznego i wszelkie opóźnienia propagacji od synchronizatora do systemu.

Przykład

Jako przykład obliczenia MTBF użyjemy $\tau=65$ ps i $T_0=400$ ps – nie dotyczą one konkretnej technologii, wybrano je tylko dla ilustracji. Rozważmy system z zegarem 500 MHz i wejściową szybkością transmisji danych 150 MHz. Cykl zegara wynosi 2 ns, więc T_R musi być mniejsze, powiedzmy 1,6 ns (ponownie, tylko w celach ilustracyjnych). Wkładając te liczby do równania MTBF otrzymamy około 3,5 godziny – zatem gdybyśmy zbudowali wiele kopii naszego układu, 63% uległoby awarii w ciągu pierwszych 3,5 godziny pracy. Jest to mało prawdopodobne, aby było to akceptowalne!

Jeżeli pojedynczy przerzutnik synchronizatora nie jest w stanie zapewnić wystarczającego współczynnika MTBF, to możliwym rozwiązaniem jest zastosowanie ich dwóch lub trzech w łańcuchu. Dla dwóch przerzutników (patrz **rysunek 10**) prawdopodobieństwo, że drugi przerzutnik jest metastabilny w momencie, gdy dane wchodzi do systemu, wynosi $P_F = P_{E1} P_{S2}$ – oznacza to, że pierwszy przerzutnik musi wejść w metastabilność i nadal być metastabilny po T_R , powodując, że drugi wejdzie w metastabilność i nadal musi być metastabilny po kolejnym czasie T_R . Jeśli oba mają te same parametry, otrzymamy nowe równanie MTBF z wykładniczym $2T_R$ zamiast T_R . Przeprowadzenie tych obliczeń na wartościach z poprzedniego przykładu daje współczynnik MTBF wynoszący około 150 milionów lat i znacznie mniejsze prawdopodobieństwo awarii w okresie eksploatacji obwodu. W praktyce znalezienie wartości τ i T_0 może być trudne, ale miejmy nadzieję, że te przykłady ilustrują fakt, że niekoniecznie jest oczywiste, ile stopni synchronizatora jest wymaganych.

Synchronizacja wielobitowa

Synchronizatory pokazane na rysunkach 9 i 10 mogą być używane tylko do danych jednobitowych. Jeśli mamy dane wielobitowe, może się wydawać, że moglibyśmy po prostu zastosować synchronizator na każdym bicie równolegle. Niestety charakter metastabilności i wpływ niewielkich różnic w taktowaniu wejściowym, przesunięciu zegara oraz zmienności poszczególnych przerzutników sprawiają, że jest to podejście bardzo ryzykowne. Rozważmy wprowadzenie pojedynczego bitu do przerzutnika synchronizatora; powiedzmy, że zmienia się on z 0 na 1, ale powoduje metastabilność, która ustępuje w czasie, więc nie powoduje awarii. Może zostać rozstrzygnięty na 0 lub 1, w zależności od tego, co dokładnie zostało przechwycone w zatrzaśku. W następnym cyklu zegara bit wejściowy

z pewnością ustawi się na 1, więc przerzutnik synchronizatora załaduje 1 bez metastabilności. W tym scenariuszu cyfra 1 wchodzi do systemu OK, ale nie ma pewności, w którym cyklu zegara to nastąpi. Jest to w porządku w przypadku pojedynczego bitu – jest asynchroniczne, więc nie jest oczekiwane w konkretnym cyklu zegara. W przypadku wartości wielobitowych, które zmieniają się blisko zegara, poszczególne bity w zestawie równoległych synchronizatorów mogą być rozdzielane w różnych cyklach zegara. Spowodowałoby to uszkodzenie danych w systemie przez cały cykl zegara, co mogłoby mieć katastrofalne skutki, w zależności od konsekwencji wprowadzenia błędnych wartości.

Do przesyłania danych wielobitowych pomiędzy systemami synchronicznymi musimy zastosować różne podejścia. Jednym ze sposobów jest użycie zsynchronizowanych sygnałów potwierdzeń. Nadawca ustala nową wartość danych i dopiero po jej ustabilizowaniu wysyła jednobitowy sygnał uzgodnienia „Mam dane dla ciebie” za pośrednictwem synchronizatora do systemu odbierającego. Odbiornik wysyła z powrotem jednobitowy sygnał potwierdzenia, ponownie za pośrednictwem synchronizatora, po załadowaniu danych. Jest to skuteczne, ale stosunkowo powolne. Aby uzyskać większą szybkość transmisji danych, można zastosować specjalną pamięć FIFO (pierwsze weszło, pierwsze wyszło). Każde FIFO zawiera bank pamięci z dwoma portami (jest zapisywane i odczytywane przez różne porty, a nie przez pojedynczą magistralę) i działa jako bufor, w którym tempo zapisu i odczytu danych może się różnić (np. buforowanie filmów online). Jeśli obie strony używają tego samego zegara, wszystko jest proste, ale jeśli są asynchroniczne, problemem nie jest synchronizacja danych w pamięci, ale synchronizacja liczników, które śledzą, gdzie dane są zapisywane i odczytywane w pamięci. Należy je porównać, aby sprawdzić, czy FIFO jest pełne czy puste. Ponieważ z każdym systemem asynchronicznym jest powiązany jeden licznik, są to wartości wielobitowe, które należy zsynchronizować, aby możliwe było przeprowadzenie porównań. Sprytną sztuczką jest użycie dla liczników systemu liczbowego Graya. W kodzie Graya przyrost o jeden powoduje zmianę tylko jednego bitu, więc dla każdego bitu można zastosować synchronizatory równoległe, takie jak te na rysunkach 9 i 10. Ponieważ na raz zmienia się tylko jeden bit, nie ma możliwości zniekształcenia wartości licznika przez różne synchronizatory pracujące w różnych cyklach zegara. ■

Ian Bell

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, kwiecień 2021 (www.epemag3.com)



Komunikacja 433 MHz



Zgodnie z prawem można nawiązać połączenie bezprzewodowe na częstotliwości nośnej 433 MHz. Częstotliwość ta została dopuszczona do komunikacji bezprzewodowej na krótkich dystansach za pomocą tzw. telegramów.

Pasmo częstotliwości 433 MHz

Do czego można wykorzystać częstotliwość 433 MHz?

Pasmo częstotliwości 433 MHz jest dostępne dla każdego do bezprzewodowej cyfrowej komunikacji o niskim poborze mocy między urządzeniami, takimi jak czujniki, stacje pogodowe, otwieracze drzwi garażowych, bezprzewodowe dzwonki do drzwi i automatyka domowa. Dobrze znane tanie nadajniki i odbiorniki, między innymi firmy KlikAanKlikUit, które umożliwiają bezprzewodowe sterowanie urządzeniami i zdalne przyciemnianie oświetlenia, działają we wspomnianym paśmie 433 MHz.

Komunikacja 433 MHz działa w technologii Amplitude Shift Keying

Urządzenia te pracują ze 100% modulacją amplitudy. Fala nośna o częstotliwości 433 MHz może być włączona lub nie. Można na przykład uzgodnić, że obecność fali nośnej odpowiada transmisji cyfrowego „H”, a jej brak oznacza transmisję cyfrowego „L”. System ten nosi nazwę „Amplitude Shift Keying”, w skrócie ASK. Dzięki brakowi dwukierunkowej komunikacji i szyfrowania, protokoły komunikacyjne są z natury proste i łatwe do stworzenia, na przykład za pomocą Arduino.

Komunikacja 433 MHz służy do przesyłania krótkich komunikatów

Być może nie zdajesz sobie z tego sprawy, ale w Twojej okolicy jest bez wątpienia wiele urządzeń, które nadają i odbierają na częstotliwości 433 MHz. Pomyśl o bezprzewodowym dzwonku do drzwi sąsiadów, własnych ściemniaczach ściennych i pilocie do otwierania drzwi garażowych u innych sąsiadów. Urządzenia te mogą oczywiście zakłócać się nawzajem, dlatego zabronione jest ciągle przysyłanie przez nie sygnałów. Załóżmy, że masz stację pogodową, która komunikuje się bezprzewodowo z termometrem zewnętrznym. Na przykład termometr ten będzie szybko przysyłał chwilową temperaturę co dziesięć minut, mając nadzieję, że jego stacja bazowa odbierze tę transmisję i zanotuje

temperaturę. Taki ciąg impulsów na częstotliwości 433 MHz nazywany jest „pakietem”. Pakiet najczęściej składa się z dwóch słów. Pierwsze słowo zawiera unikalny identyfikator lub kod adresu, który informuje odbiornik, że „jego” nadajnik nadaje. Drugie słowo zawiera dane, które nadajnik musi wysłać, na przykład temperaturę. Nadawca może wysłać swój pakiet kilka razy z rzędu, dzięki czemu szansa na odebranie wiadomości przez odbiorcę wynosi prawie 100%.

Dzięki komunikacji za pomocą krótkich pakietów z identyfikatorami można używać w domu zarówno nadajników, jak i odbiorników automatyki domowej, zainstalować bezprzewodową stację pogodową i korzystać z elektronicznie otwieranej bramy garażowej bez wzajemnego przeszkadzania sobie tych trzech systemów.

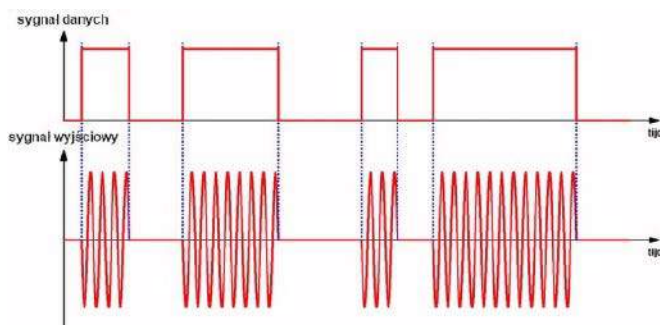
O czym należy pamiętać?

Komunikacja ASK na częstotliwości 433 MHz jest niezwykle prostym systemem. Jest idealny do samodzielnego eksperymentowania, ponieważ wymaga niewielkiej wiedzy i prostego sprzętu. System nie jest jednak w 100% niezawodny. Jeśli dwa nadajniki wyślą pakiety jednocześnie w tym samym czasie, istnieje duża szansa, że oba sygnały będą się wzajemnie zakłócać, a odbiorcy nie uznają pakietów za „swoje” i nic z nimi nie zrobią.

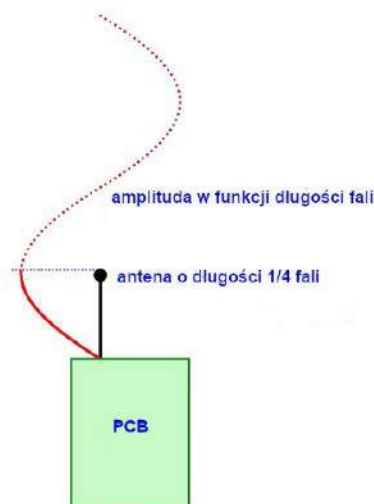
Korzystanie z anten

Jeśli używasz modułów 433 MHz bez dodatkowych anten, uzyskasz tylko bardzo ograniczony zasięg. Dlatego zaleca się podłączenie zarówno modułów nadajnika, jak i odbiornika do anteny zewnętrznej. Większość płytek drukowanych posiada odpowiednie złącze. Pozostaje pytanie, jak taka antena powinna wyglądać i jak duża powinna być.

Większość artykułów na ten temat zaleca użycie anteny prętowej o długości ćwierć fali, tak zwanej anteny „ćwierćfalowej”. Jak pokazano na poniższym rysunku, natężenie pola w górnej części takiej anteny jest maksymalne. Jednak w punkcie wejścia sygnału natężenie pola wynosi zero, a natężenie prądu jest maksymalne. W ten sposób maksymalna moc jest pompowana do anteny, a zasięg będzie największy.



Zasada kluczowania z przesunięciem amplitudy (ASK) (© 2018 Jos Verstraten)



Zasada działania anteny ćwierćfalowej (© 2018 Jos Verstraten)

Jak zapewne wiesz, długość fali jest równa prędkości światła podzielonej przez częstotliwość sygnału. Jeśli zastosujemy ten wzór do sygnału 433 MHz, okaże się, że długość fali takiego sygnału wynosi 69,28 cm. Czwierćfala dla tej częstotliwości musi zatem mieć długość 17,32 cm.

Duży wybór

Istnieją dziesiątki modułów, które generują sygnał 433 MHz i mają wyprowadzone wejście do załączania modulacji ASK. Można również wybierać spośród dziesiątek modułów do odbioru i demodulacji takiego sygnału. Po długich poszukiwaniach w Internecie znaleźliśmy prostą i niezwykle taną kombinację, która idealnie nadaje się do przeprowadzania pierwszych eksperymentów. Najniższa cena, jaką znaleźliśmy za ten zestaw, to tylko 1,42 €. Przy zamawianiu należy pamiętać o określeniu częstotliwości. Zestaw jest dostarczany z częstotliwością 433 MHz i 315 MHz. Częstotliwość 315 MHz jest dozwolona do swobodnej komunikacji w Japonii, ale nie w Europie.

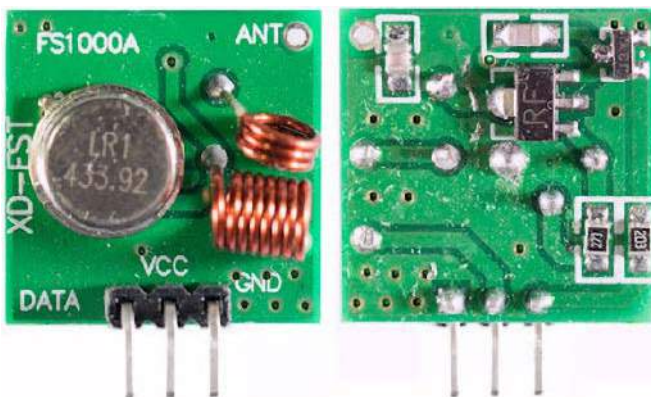
Moduł nadajnika FS1000A lub XD-FST

Prezentacja

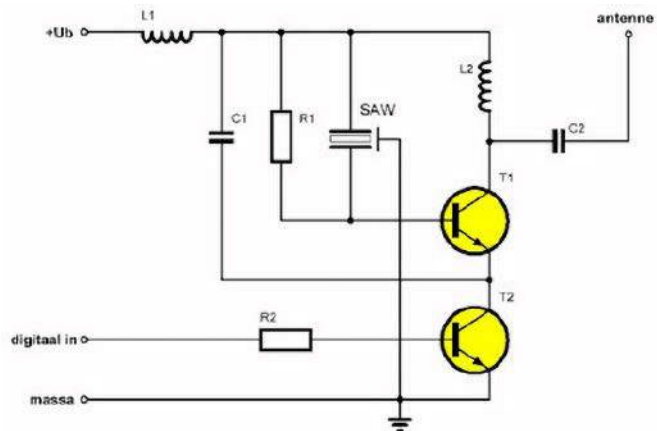
Moduł przedstawiony na poniższym rysunku jest dostarczany pod dwiema nazwami handlowymi, a mianowicie FS1000A lub XD-FST. Na płytce PCB, która ma wymiary zaledwie 19 mm na 19 mm, widać dwie cewki i okrągłą metalową część z kodem LR1 433.92 po jednej stronie. Okazuje się, że jest to bardzo egzotyczna część, a mianowicie rezonator SAW, który rezonuje z częstotliwością 433,92 MHz. SAW to skrót od **Surface Acoustic Wave**. Po drugiej stronie znajdują się dwa tranzystory SMD oraz kilka rezystorów i kondensatorów. Trzy piny połączeniowe nie pozostawiają wątpliwości co do ich funkcji. W prawym górnym rogu znajduje się otwór do podłączenia anteny.

Schemat modułu nadajnika

Schemat tego modułu pokazano na poniższym rysunku. Tranzystor T1 jest oscylatorem. Rezonator SAW znajduje się między kolektorem a bazą, dzięki czemu spełnia warunek oscylacji obwodu. Tylko szum o częstotliwości 433,92 MHz jest sprzężony i dlatego będzie coraz bardziej wzmacniany, aż do wytworzenia ładnego sygnału transmisji o tej częstotliwości. Sygnał ten jest sprzężony z anteną poprzez obwód LC. Sygnał nadawczy naturalnie przepływa również przez cewkę L2. Wokół tej cewki wytwarzane jest pole elektromagnetyczne, które odpowiada za to, że moduł ma stosunkowo niewielki zasięg transmisji, nawet bez anteny. Modulacja ASK jest zapewniana przez T2 w bardzo prymitywny sposób. Jeśli na bazę tego tranzystora podasz stan niski, tranzystor zostanie zatkany. Tranzystor T1 zostanie pozbawiony napięcia i naturalnie przestanie oscylować. Jeśli ustawisz stan wysoki na bazie T2, ten będzie przewodził, a oscylator wokół T1 zostanie uruchomiony.



Moduł nadajnika FS1000A lub XD-FST (© 2018 Jos Verstraten)



Schemat modułu nadajnika FS1000A lub XD-FST (© 2018 Jos Verstraten)

Specyfikacja techniczna

Specyfikacje techniczne tego modułu są w skrócie następujące:

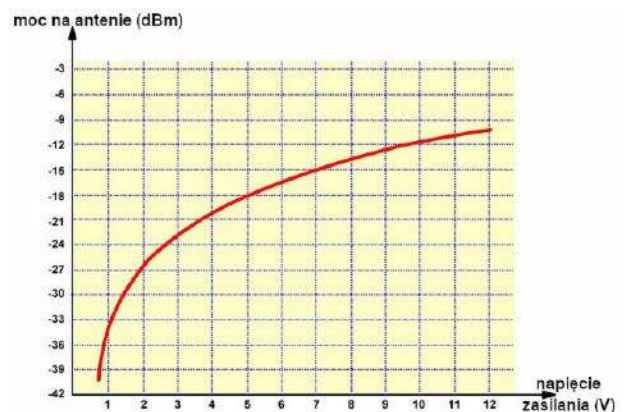
- Prędkość transmisji : 10 kB/s
- Zasięg: ok. 20 m (w zależności od zasilania i anteny)
- Napięcie zasilania: 3,0 V DC ~ 12,0 V DC
- Prąd zasilania: 28 mA maks. aktywny
- Wymiary: 19 mm × 19 mm
- Waga: 3 g

Moc nadawania

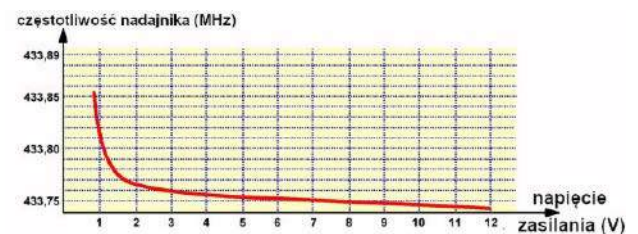
Moc nadawania modułu jest określona na 10 mW. W praktyce okazuje się jednak, że moc transmisji jest dość zależna od napięcia zasilania. Nasi koledzy z Gough's Tech Zone zmierzili rzeczywistą moc transmisji w funkcji napięcia zasilania, pomiary, które włączyliśmy do ładnego wykresu. Wynika z niego, że moc transmisji wzrasta o 13 dB przy zwiększeniu napięcia zasilania z 3,0 V do 12,0 V.

Stabilność częstotliwości

Częstotliwość transmisji w funkcji napięcia zasilania została również zmierzona przez kolegów z Gough's Tech Zone. Poniższy wykres pokazuje,



Moc nadawania w funkcji napięcia zasilania (© 2018 Jos Verstraten na podstawie danych z Gough's Tech Zone)



Częstotliwość transmisji w funkcji napięcia zasilania (© 2018 Jos Verstraten na podstawie danych z Gough's Tech Zone)

że zmienia się ona o około 100 kHz między napięciem zasilania 1,0 V a 12,0 V. To dziwne, ponieważ prawdopodobnie myślisz, tak jak my, że SAW ma bardzo stabilną częstotliwość rezonansową, która nie dba o napięcie zasilania obwodu, w którym jest używana.

Moduł odbiornika XD-RF-5V

Prezentacja

Jak pokazuje poniższe zdjęcie, odbiornik jest znacznie bardziej złożony niż nadajnik. Na płytce drukowanej o wymiarach 30 mm × 14 mm widać wiele komponentów. Ma to sens, ponieważ i tak już niski strumień mocy w antenie nadawczej zmniejsza się bardzo szybko wraz ze wzrostem odległości między nadajnikiem a odbiornikiem. Odbiornik prawdopodobnie odbierze tylko kilka μW mocy i będzie musiał znacznie wzmocnić ten niewielki sygnał. Ponadto odbiornik będzie musiał radzić sobie z szumami i oczywiście jest, że AVR, automatyczna kontrola wzmocnienia, musi być obecna, aby regulować wzmocnienie wzmacniacza 433 MHz. Widoczne są dwa tranzystory i podwójny wzmacniacz operacyjny. Dwa środkowe piny złącza są połączone ze sobą i tworzą wyjście obwodu. Na samym dole po lewej stronie, obok cewki, widać pole lutownicze, do którego można ewentualnie podłączyć antenę.

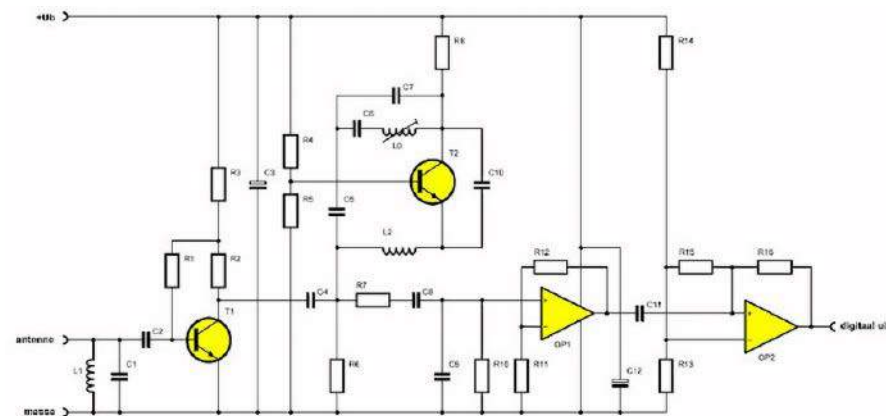
Schemat modułu odbiorczego

W obiegu znajdują się różne wersje tego modułu z niewielkimi różnicami w schemacie i płytce drukowanej. Jednak elektronika zasadniczo odpowiada temu, co narysowaliśmy na poniższym schemacie. Sygnał antenowy jest podawany do obwodu rezonansowego L1/C1, który jest dostrojony do częstotliwości nadajnika. Po tym następuje wzmacniacz tranzystorowy wokół T1. Wzmocniony sygnał jest wyprowadzany przez kondensator odcinający składową stałą C4. Tranzystor T2 wraz z otoczeniem stanowi prosty tłumik sterowany napięciem. Powoduje on zwarcie mniejszej lub większej części sygnału przez R6 do napięcia zasilania. W ten sposób układ będzie wzmacniał maksymalnie, gdy nie jest odbierany żaden sygnał i wzmacniał minimalnie, gdy odbierany jest ładny sygnał 433 MHz. Obwód można precyzyjnie dostroić do częstotliwości nadajnika za pomocą przestrajalnej cewki L0. Sygnał jest następnie ponownie wzmacniany przez obwód wokół OP1. Mała szerokość pasma tego wzmacniacza zapewnia, że tylko obwiednia impulsów 433 MHz jest sprzężona. Op-amp OP2 jest podłączony jako komparator z histerezą, który jest odpowiedzialny za konwersję sygnału modulowanego ASK na ładny impuls cyfrowy.

Specyfikacje techniczne

Specyfikacje techniczne modułu odbiornika są następujące:

- Kod produktu: XD-RF-5V
- Producent: nieznany
- Częstotliwość: 433,92 MHz
- Demodulacja: ASK
- Napięcie zasilania: 5 V DC



Schemat płytki odbiornika XD-RF-5V (© 2018 Jos Verstraten)

- Prąd zasilania: 4 mA
- Czułość: -105 dB
- Wymiary: 30 mm × 14 mm × 7 mm

Uwagi

Obwód jest bardzo wrażliwy na zmiany napięcia zasilania, dlatego należy bezwzględnie używać stabilizowanego zasilania 5 V. Odbiornik jest również bardzo wrażliwy na sygnały zakłócające. Dzieje się tak, ponieważ wzmocnienie jest zwiększane do maksimum, gdy nie jest odbierany żaden sygnał 433 MHz. Szum, który jest następnie wzmacniany, może spowodować przełączenie komparatora na wyjściu i wygenerowanie impulsu. Przed wysłaniem prawdziwego pakietu zaleca się najpierw wysłać serię fal sinusoidalnych 433 MHz, aby odbiornik mógł je odebrać i dostosować wzmocnienie do wielkości odbieranego sygnału. Biblioteka „VirtualWire dla Arduino” proponuje również metody programowe w celu rozwiązania problemu impulsów zakłócających.

Zastosowanie obu modułów

Nie jest to możliwe bez inteligentnego sprzętu

Jeśli myślisz, że możesz sterować diodą LED bezprzewodowo, włączając i wyłączając nadajnik za pomocą przełącznika i podłączając wyjście odbiornika do diody LED za pomocą tranzystora, bardzo się mylisz. Odbiornik odbiera każdy sygnał 433 MHz, a także reaguje na znaczące sygnały szumu. Dioda LED będzie migać w sposób ciągły bez konieczności wykonywania jakichkolwiek czynności.

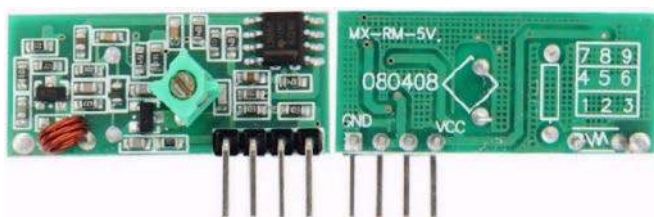
System działa bezbłędnie tylko wtedy, gdy rozbudujesz nadajnik w taki sposób, aby wysyłał omówione już pakiety. Odbiornik musi również posiadać inteligentną elektronikę na wyjściu, która rozpoznaje pakiety i wykonuje żądane działanie tylko wtedy, gdy taki pakiet zostanie odebrany.

Praca z Arduino lub Raspberry Pi

Jeśli potrafisz programować mikrokontrolery, możesz oczywiście samodzielnie napisać kod niezbędny do uczynienia nadajnika i odbiornika inteligentnymi. Jeśli nie potrafisz tego zrobić, ale masz w domu dwie dobrze znane płytki z mikrokontrolerami, możesz pobrać kod. W końcu internet jest pełen przykładów, które można załadować do Arduino lub Raspberry Pi i udoskonalić system.

Pracuj z układami kodującymi i dekodującymi HT12E i HT12D

Można jednak również pracować ze specjalizowanymi układami scalonymi, które zostały opracowane do kodowania i dekodowania telegramów. Jedną z najbardziej znanych i najtańszych par jest kombinacja HT12D/HT12E firmy Holtec. Para ta jest dostępna w sprzedaży w cenie od 2,20 €. Oba układy scalone mają odpowiednio osiem wejść adresowych i cztery wejścia i wyjścia danych. Obsługa jest niezwykle prosta. Ustawia się ten sam adres na obu układach scalonych, na przykład za pomocą przełączników DIP. Teraz można wysłać szesnaście różnych



Przód i tył płytki odbiornika XD-RF-5V (© 2018 Jos Verstraten)

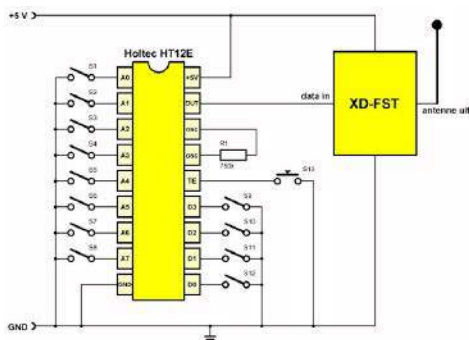
pakietów za pomocą HT12E, podłączając cztery wejścia danych do masy lub pozostawiając je otwarte. Ten pakiet jest nadawany co najmniej cztery razy z rzędu. HT12D odbiera te cztery pakiety i sprawdza, czy skład bitów adresu jest prawidłowy. Jeśli tak, pakiet jest zatwierdzany, a układ sprawdza, które bity danych to „L” lub „H”. Następnie wysyła swoje cztery wyjścia danych do tego samego stanu. Dane te są przechowywane w wewnętrznym zatrasku i pozostają tam, dopóki układ scalony nie otrzyma nowego zatwierdzonego pakietu. Dzięki tym dwóm układom scalonym można bardzo łatwo zdalnie włączać i wyłączać cztery obciążenia za pomocą czterech przełączników na nadajniku.

Przykładowy schemat nadajnika

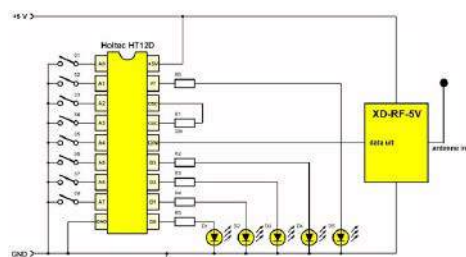
Schemat kompletnego nadajnika pokazano na poniższym schemacie. Piny adresowe od A0 do A7 są wewnętrznie podłączone do zasilania +5 V za pomocą rezystorów podciągających. Gdy te styki są otwarte, odpowiada to stanowi wysokiemu „H”. Można podłączyć te wejścia adresowe do masy za pomocą przełączników S1 do S8, tj. ustawić je w stan niski „L”. W ten sposób można ustawić adres nadajnika na 1 z 256 kodów. Cztery styki danych D0 do D3 są również wewnętrznie podłączone do zasilania +5 V za pomocą rezystorów. Wejścia te można podłączyć do masy za pomocą przełączników od S9 do S12. Zamknięty przełącznik oznacza zatem ustawienie odpowiedniej linii danych w stan niski „L”. HT12E zaczyna wysyłać telegramy, gdy na wyprowadzenie TE zostanie podany stan niski „L”. Odbywa się to za pomocą przycisku S13. Po jego naciśnięciu układ scalony wyśle telegramy i zatrzyma się po zwolnieniu przycisku. Telegramy są przekazywane na pin OUT. Oczywiście wyjście to podłącza się do wejścia danych nadajnika 433 MHz XD-FST. Całość zasilana jest napięciem 5 V. Jest to maksymalne napięcie dla HT12E. Jeśli chcesz zasilic nadajnik wyższym napięciem (zakres transmisji!), musisz włączyć stabilizator 5 V, aby zasilic układ. Częstotliwość wewnętrznego oscylatora zegara jest ustawiana na prawidłową wartość za pomocą rezystora R1.

Przykładowy schemat odbiornika

Schemat układu odbiornika jest równie prosty. Oczywiście należy ustawić ten sam adres za pomocą przełączników S1 do S8, co na nadajniku. Podłącz pin „wyjścia danych” odbiornika 433 MHz do wejścia D IN układu HT12D. Na tym schemacie cztery wyjścia danych są symbolicznie połączone z czterema diodami LED od D1 do D4. Będzie jasne,



Kompletny schemat nadajnika 433 MHz (© 2018 Jos Verstraten)



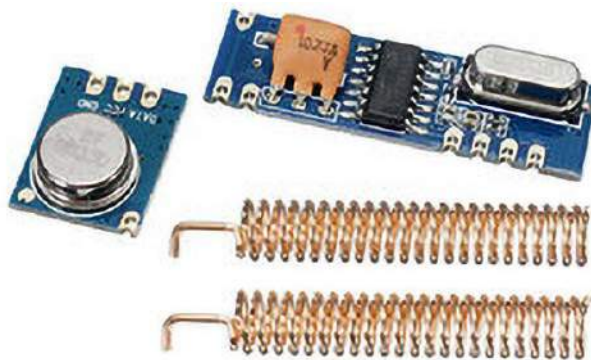
Kompletny schemat odbiornika 433 MHz (© 2018 Jos Verstraten)

że w praktyce można coś włączyć lub wyłączyć za pomocą tych czterech sygnałów danych za pomocą przełączników lub optotriaków. Wyjście VT dostarcza wysoki impuls, gdy HT12D odbiera prawidłowy pakiet. Wyjście to pozostaje w stanie „H” tak długo, jak długo naciśnięty jest przycisk nadajnika i kontynuowana jest transmisja poprawnych telegramów.

Zasięg stu metrów przy częstotliwości 433 MHz

Kombinacja STX882/SRX882

Z mocą 10 mW, kombinacja XD-FST/XD-RF-5V ma zasięg co najwyżej dwadzieścia do trzydziestu metrów. Jeśli potrzebujesz większego zasięgu, możesz użyć dwóch innych modułów: pary STX882/SRX882. Ta kombinacja NiceRF gwarantuje maksymalną odległość stu metrów między nadajnikiem a odbiornikiem. Moduły te są również bardzo tanie, znaleźliśmy cenę 2,49 € za zestaw. Dwie płytki PCB są dostarczane z antenami, które składają się z cewki z drutu miedzianego. Moduły te posiadają również wejście i wyjście danych, które można podłączyć do pary HT12E/HT12D.

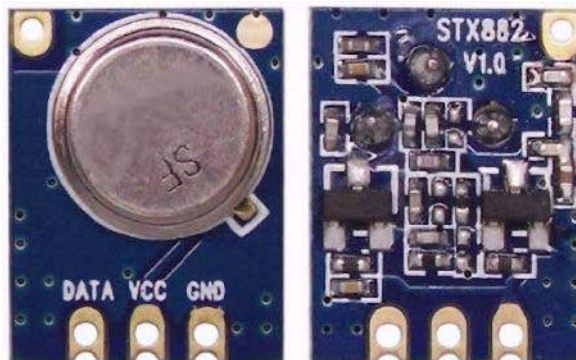


Połączenie STX882/SRX882 firmy NiceRF (© 2018 Jos Verstraten)

Moduł nadajnika STX882

Ten nadajnik jest również dostarczany zarówno dla częstotliwości 433 MHz, jak i 315 MHz. Należy więc zwrócić na to uwagę przy zamawianiu! Po jednej stronie tej płytki drukowanej o wymiarach 15,2 mm na 12,0 mm widać tylko cztery otwory połączeniowe i rezonator SAW. Po drugiej stronie znajdują się dwa tranzystory oraz kilka rezystorów i kondensatorów. Pomimo intensywnych poszukiwań, nie udało nam się znaleźć schematu tego modułu. Dziwne jest to, że najwyraźniej brakuje tu cewek. Połączenie dla anteny znajduje się w lewym górnym rogu, trzy pozostałe połączenia mówią same za siebie. Dane techniczne, dostarczone przez producenta NiceRF, są następujące:

- Napięcie zasilania: od 1,2 V do 6,0 V
- Pasywny prąd zasilania: mniej niż 0,01 μ A



Obie strony modułu nadajnika STX882 (© 2018 Jos Verstraten)

- Prąd zasilania nadajnika: 34 mA przy napięciu zasilania 3,3 V
- Częstotliwość: 433,82 MHz ~ 434,02 MHz
- Moc nadawania: 50 mW (zasilanie 3,6 V)
- Szybkość transmisji danych: 9,6 kb/s maks.
- Wymiary: 15,2 mm × 12,0 mm × 3,4 mm

Moduł odbiornika SRX882

Sercem tego układu jest chip PT4303-S firmy Princeton Technology Corp. Jest to wysoce zintegrowany odbiornik superheterodynowy o niskiej mocy dla pasm częstotliwości 433 MHz i 315 MHz. Częstotliwość jest określana przez zewnętrzny rezonator kwarcowy. Układ ten zawiera wzmacniacz o niskim poziomie szumów (LNA), lokalny oscylator sterowany napięciem (VCO), pętlę synchronizacji fazowej (PLL) i mikser do częstotliwości środkowej 10,7 MHz. Brązowa część widoczna na wydruku to filtr piezoceramiczny o częstotliwości 10,7 MHz. Ta część jest aż nazbyt znajoma, jeśli kiedyś lutowałeś razem radia FM. Następnie znajduje się demodulator ASK, filtr danych i komparator. Producent NiceRF udostępnia dane techniczne poniżej:

- Napięcie zasilania: od 2,4 V do 5,5 V
- Prąd zasilania w trybie działania: typowo 2,8 mA
- Prąd zasilania w trybie uśpienia: mniej niż 2 μ A
- Opóźnienie danych: 9 ms maks.
- Zakres częstotliwości: 433,82 MHz ~ 434,02 MHz
- Czułość: -110 dBm typowo
- Szerokość pasma: typowo 200 kHz
- Maksymalna szybkość transmisji danych: 9,6 kb/s
- Wymiary: 35 mm × 10,4 mm × 5 mm

Po lewej stronie płytki drukowanej znajduje się pole lutownicze na górze i na dole, do którego można podłączyć antenę. Na dole znajduje się



Obie strony modułu nadajnika SRX882 (© AliExpress)

również pole do uziemienia dowolnego kabla koncentrycznego między płytką drukowaną a anteną.

Po prawej stronie układu znajdują się cztery pola lutownicze, od lewej do prawej:

- Dodatnie źródło zasilania.
- CS.
- Wyjście danych.
- Masa.

CS (Chip Select) to połączenie, które umożliwia przełączenie modułu w tryb działania lub uśpienia. Jeśli podłączysz to połączenie do masy, moduł przejdzie w tryb uśpienia. Podłączenie do zasilania aktywuje moduł.

Anteny

Dwie dostarczone anteny są nawinięte z litego drutu miedzianego o średnicy 0,8 mm, zawierają 23,5 zwojów, mają długość 32 mm i średnicę 5,5 mm. Impedancja przy częstotliwości nadawania wynosi 50 Ω , a zysk 2,1 dB. Współczynnik fali stojącej (Voltage Standing Wave Ratio VSWR) wynosi mniej niż 1,5. ■

Jos Verstraten



Komunikacja 433 MHz

Rozwiązanie znajdziesz na www.elportal.pl/quizy

1. Z łączności na paśmie 433MHz może korzystać:

- ☐ każdy;
- ☐ każdy licencjonowany producent elektroniki;
- ☐ każdy radioamator z licencją krótkofalarską lub licencjonowany producent elektroniki.

2. Tanie moduły 433MHz używają modulacji:

- ☐ FSK;
- ☐ PSK;
- ☐ ASK.

3. Tanie moduły 433MHz przeznaczone są do:

- ☐ przesyłania niewielkich pakietów informacji sporadycznie i w jednym kierunku;
- ☐ przesyłania niewielkich pakietów informacji sporadycznie w obu kierunkach;
- ☐ bezprzewodowej komunikacji szeregowej SPI/I²C/UART.

4. Niski zasięg tanich modułów radiowych można zwiększyć poprzez:

- ☐ umieszczenie nadajnika w metalowej puszcze;
- ☐ dodanie prostej anteny z przewodu od długości ćwierci fali 433MHz;
- ☐ dodanie scalonego wzmacniacza LNA między nadajnik, a antenę.

5. W module nadajnika FS1000A/XD-FST sterowanie nadawaniem polega na:

- ☐ włączaniu i wyłączaniu zasilania nadajnika;
- ☐ modulowanie sygnałem sterującym oscylatora;
- ☐ otwieranie i zatykanie tranzystora, który kontroluje połączenie kolektora tranzystora oscylatora do masy.

6. Moc nadawcza FS1000A/XD-FST zależy od:

- ☐ napięcia zasilania;
- ☐ napięcia na wejściu danych;
- ☐ temperatury.

7. Moduł XD-RF-5V jest bardzo wrażliwy na:

- ☐ temperaturę;
- ☐ napięcie zasilania;
- ☐ wstrząsy.

8. Układy HT12D i HT12E to:

- ☐ scalone nadajnik i odbiornik 433 MHz;
- ☐ mikrokontrolery z wbudowanymi nadajnikiem i odbiornikiem 433 MHz;
- ☐ koder i dekodek do budowy pilotów zdalnego sterowania.

9. STX882 to:

- ☐ koder do budowy pilotów zdalnego sterowania;
- ☐ mikrokontroler ośmiobitowy z modułem nadajnika 433MHz od STMicroelectronics;
- ☐ moduł nadajnika 433 MHz.

10. Wymienione w artykule moduły oferują maksymalny zasięg:

- ☐ 20–100 metrów;
- ☐ 50–200 metrów;
- ☐ 200–500 metrów.

Komunikacja radiowa 433MHz w pytaniach i odpowiedziach

Czy poza pasmem 433MHz są jakieś inne pasma radiowe dla elektroniki?

Tak. Ogólnie pasma przeznaczone do prostej komunikacji są nazywane pasmami ISM, jest to skrót od zastosowań tych pasm: Industrial, Scientific, Medical – przemysłowe, naukowe i medyczne. Poza pasmem 433 MHz popularne wśród hobbystów są pasma 868 MHz i 2,4 GHz. Pasma 433 MHz jest najpopularniejsze. Zakresy częstotliwości wszystkich pasm ISM dopuszczalnych do użytku w Polsce przedstawia poniższa tabela.

Tabela 1. Pasma ISM w Polsce i ich zastosowania.

Od	Do	Kategoria urządzeń
6765 kHz	6795 kHz	Urządzenia do zastosowań ISM
13553 kHz	13567 kHz	Urządzenia do zastosowań ISM i RFID
26.957 MHz	27.283 MHz	Urządzenia do zastosowań ISM
40.660 MHz	40.700 MHz	Urządzenia do zastosowań ISM i sterowania modelami
433.050 MHz	434.790 MHz	Urządzenia bliskiego zasięgu ogólnego stosowania
868 MHz	870 MHz	Urządzenia bliskiego zasięgu ogólnego stosowania i RFID
902 MHz	928 MHz	Urządzenia do zastosowań ISM
2400 MHz	2500 MHz	Urządzenia do zastosowań ISM i RFID
5725 MHz	5875 MHz	Urządzenia do zastosowań w RTTT oraz ISM
24.000 GHz	24.250 GHz	Urządzenia bliskiego zasięgu ogólnego stosowania
61.000 GHz	61.500 GHz	Urządzenia do zastosowań ISM
122 GHz	123 GHz	Urządzenia do zastosowań ISM
244 GHz	246 GHz	Urządzenia do zastosowań ISM

Przedstawiona w artykule modulacja ASK wygląda mi bardziej na modulację OOK.

Zgadza się. Modulacja ASK polega na zmianie amplitudy nadajnika między dwiema lub więcej wartościami. Modulacja OOK, czyli On-Off Keying to szczególny przypadek modulacji ASK, gdzie amplituda jest przełączana między wartością maksymalną, a zerem, inaczej pisząc nadajnik jest włączony lub wyłączony.

Dlaczego tanie moduły 433 MHz w praktyce mają tak niskie zasięgi?

Po pierwsze, pasmo 433 MHz jest zatłoczone, o czym wspominałem wyżej. Po drugie, najtańsze moduły są wręcz prymitywne w swojej budowie i działaniu, a modulacja ASK nie oferuje dużej odporności na zakłócenia. Ponadto większość z nich jest zestrojonych na jedną i tę samą częstotliwość, przez co mogą się wzajemnie zakłócać. Ponadto mimo wszystko parametry najtańszych modułów mogą być zawyżone.

Czy mogę dodać wzmacniacz mocy do nadajnika by zwiększyć zasięg?

Legalnie nie. Na paśmie 433 MHz można nadawać z maksymalną mocą 10 mW. W praktyce większa moc nadawania nie zawsze przekłada się na znaczący wzrost zasięgu.

Jak zatem mogę poprawić zasięg?

Jest pięć możliwości. Po pierwsze, można użyć zewnętrznej anteny. Jeśli nadajnik i odbiornik są stacjonarne, to można użyć anteny kierunkowej. Oczywiście samodzielne zaprojektowanie i wykonanie takiej anteny stanowi nie lada wyzwanie dla hobbysty. Jeśli moduły na to pozwalają, można przestroić je na mniej zajętej część pasma 433 MHz.

Drugą metodą jest redukcja prędkości transmisji, szczególnie jeśli zostanie połączona z inną formą kodowania komunikatu. Dwa popularne rozwiązania to użycie sygnałów długich i krótkich do oznaczania zer i jedynek oraz kodowanie Manchester. Radioamatorzy od lat osiągają ogromne zasięgi przy minimalnych mocach właśnie przez transmisję o ekstremalnie niskiej prędkości (czasem liczącej dziesiątki sekund na znak). Nazywa się to QRSS.

Trzecią metodą jest użycie bardziej zaawansowanego modułu łączności dwukierunkowej. Moduły te kontrolowane są zwykle przez magistralę SPI lub I²C, i oferują szereg zalet, jak łatwa selekcja kanału, adresy dla poszczególnych modułów pozwalające na prostą komunikację między wieloma układami, pakiety z sumą kontrolną. Ponieważ łączność jest dwukierunkowa i używa bardziej zaawansowanych metod modulacji, jak FSK, można uzyskać pewną łączność na większych dystansach – moduł nadający będzie ponawiał transmisję, póki moduł odbiorczy nie potwierdzi jej odebrania. Niestety, ceny

takich modułów bywają wyższe, a sposób ich użycia jest bardziej skomplikowany.

Czwartą drogą jest zmiana pasma na 868 MHz lub 2,4 GHz. Pasma 868MHz jest dużo mniej popularne, używane głównie przez systemy alarmowe. Pasma 2,4 GHz jest zatłoczone przez sieci Wi-Fi i komunikację Bluetooth, ale chińskie moduły transmisji dwukierunkowej oparte o klony układów nRF24L01+ nie są droższe od prostych modułów 433 MHz. Ponadto większość użytkowników sieci Wi-Fi w ogóle nie zmienia domyślnego kanału, przez co tylko wąski fragment całego pasma jest zatłoczony.

Po piąte, można rozważyć użycie modułów LoRa, które są zaprojektowane do zapewniania dużych zasięgów komunikacji. W testach praktycznych moduły te często osiągają zasięgi od kilometra wzwyż. Szczególnie wartymi polecenia są moduły z układem SX1278, kosztujące w Polsce 17–25 złotych.

Te zaawansowane moduły wydają się zbyt skomplikowane. Jak mogę zrealizować transmisję dwukierunkową „po tanioci”?

Możesz w każdym urządzeniu umieścić zarówno moduł nadawczy, jak i odbiorczy, przy czym mikrokontroler w każdym urządzeniu będzie kontrolować oba moduły. Dostępne są przy tym dwie podstawowe metody postępowania:

1. Wszystkie urządzenia nastuchują, gdy jedno ma potrzebę przekazać informację, zaczyna nadawać, po czym czeka na potwierdzenie odbioru. W razie jego braku nadaje ponownie.
2. Jedno urządzenie, stacja bazowa, cyklicznie nadaje zapytania do poszczególnych układów podległych. Te w odpowiedzi podają swój status, i ewentualnie przekazują potrzebne informacje.

W praktyce jest to realizacja „na piechotę” funkcji tych bardziej rozbudowanych modułów ISM.

Jakie są inne alternatywy do komunikacji radiowej na paśmie 433 MHz lub innym paśmie ISM?

Najprostsza metoda, to użyć łączności przewodowej. W przemyśle i w telekomunikacji stosuje się od dekad pętle prądowe, zarówno analogowe, jak i cyfrowe. Wysoki koszt początkowy jest rekompensowany dużymi odległościami, czasami wynoszącymi nawet dziesiątki kilometrów, oraz wysoką odpornością na zakłócenia. Dlatego też pętle prądowe 4–20 mA stały się standardem w przemysłowych systemach kontroli procesów.

Relatywnie prostą i tanią metodą na uzyskanie łączności bezprzewodowej małego zasięgu jest użycie podczerwieni. Jest to metoda stosowana od lat do kontrolowania różnych urządzeń RTV, ale nie tylko tam. Standard IrDA był dość popularną formą przesyłania niedużych plików między komputerami i urządzeniami mobilnymi. Niektóre lepsze multimetry, na przykład firmy Fluke, używają podczerwieni do komunikacji z komputerem zapewniającej wysoki poziom izolacji galwanicznej. Warto wspomnieć o tym, iż z pomocą podczerwonych diod LED dużej mocy oraz kolimatorów można w relatywnie prosty sposób znacząco zwiększyć zasięg takiej łączności. Ponad półtorej dekady temu testowałem użycie siedmiu diod LED IR o mocy 1 W każda oraz soczewek skupiających 5° w roli oświetlacza dla kamkordera. W teście praktycznym plama światła podczerwonego była widoczna na obiekcie oddalonym o ponad 150 metrów.

Pochodną powyższej metody może być użycie modułu lasera półprzewodnikowego zamiast zwykłej diody LED IR. To rozwiązanie nadaje się do zastosowania tam, gdzie oba urządzenia są stacjonarne. Dużym problemem może być wycelowanie i skupienie wiązki, zwłaszcza jeśli użyty zostanie laser podczerwony. Za to potrzebna moc może być niższa, a zasięg zaś większy.

Do łączności zasadniczo globalnego zasięgu można wykorzystać Internet. Urządzenia można podłączyć z siecią LAN przewodową. Dodatkowo używając technologii PoE (Power over Ethernet) możemy urządzenie zasilić. Można też użyć sieci Wi-Fi – ta metoda stosowana jest często w systemach automatyki domowej, często w połączeniu z protokołem MQTT. Trzecią, niegdyś popularną metodą jest użycie sieci telefonii komórkowej. Wiele telefonów sprzed ery smartfonowej pozwalało na komunikację szeregową z komputerem lub innym urządzeniem, łącznie z przesyłaniem informacji za pomocą wiadomości SMS lub przez GPRS. Dostępne są też moduły i modemy GSM, które też to potrafią.

Podzespoły: ogniwa Peltiera



Dzięki ogniwom Peltiera można całkowicie elektrycznie chłodzić lub ogrzewać obiekty lub pomieszczenia o dziesiątki stopni Celsjusza, bez użycia ruchomych części.

Wprowadzenie

Czym są ogniwa Peltiera?

Ogniwa Peltiera to komponenty termoelektryczne, które można uznać za pompy ciepła. Ogniwo Peltiera umożliwia przenoszenie ciepła z jednej strony na drugą. Kierunek przepływu ciepła przez ogniwo zależy od kierunku przepływającego przez nie prądu stałego. Różnica temperatur między jedną a drugą stroną może przekraczać 60°C.

Ogniwa Peltiera obniżają temperaturę

Ogniwa Peltiera można używać zarówno do chłodzenia, jak i ogrzewania. Ale do elektrycznego ogrzewania czegoś dostępne są oczywiście znacznie prostsze i tańsze części: zwykłe rezystory. Oczywiście jest zatem, że w praktyce ogniwa Peltiera są używane głównie do chłodzenia części, które rozpraszają tak dużo mocy na bardzo małej powierzchni, że normalne chłodzenie za pomocą radiatorów lub nawet wymuszone chłodzenie za pomocą wentylatorów nie jest już wystarczające.

Każdy, kto czyta „mała powierzchnia”, naturalnie od razu myśli o układach scalonych. Nie bez powodu, ponieważ ogniwa Peltiera odgrywają ważną rolę w chłodzeniu bardzo energochłonnych układów scalonych, takich jak nowoczesne procesory i pamięci. Problem z takimi układami polega na tym, że w chipie rozpraszana jest dość duża ilość energii. Zwykle powstające ciepło jest odprowadzane przez obudowę. Obudowa ta może być opcjonalnie wyposażona w radiator z wentylatorem. Jednak w nowoczesnych, wysoce zintegrowanych układach generowane jest tak dużo ciepła, że rozwiązania te nie są już wystarczające. Dlatego opracowano tak zwane „pokrywy lodowe”, bloki chłodzące, które są przymocowane do szybkich mikroprocesorów i skutecznie odprowadzają ciepło generowane w chipie na zewnątrz za pomocą efektu Peltiera.

Ogniwa Peltiera stabilizują temperaturę

Przez ogniwo Peltiera można skierować przepływ ciepła w dwóch kierunkach – po prostu odwracając polaryzację zasilania. W jednym przypadku można użyć ogniwa do chłodzenia pomieszczenia, w drugim do jego ogrzewania. Są zatem idealnymi komponentami do utrzymywania części, takich jak rezonatory kwarcowe, w stałej temperaturze.

Stare odkrycie

Chociaż efekt Peltiera jest znany od ponad 150 lat, to to zjawisko fizyczne nie było stosowane na dużą skalę aż do 2000 roku. Powodem tego była bardzo wysoka cena ogniwa Peltiera i stosunkowo niska ich wydajność. Jednak teraz, gdy ogniwa te są oferowane bardzo tanio, można coraz częściej za ich pomocą rozwiązywać problemy z temperaturą.

Zalety ogniwa Peltiera

Oczywiście istnieją inne systemy transportu dużych ilości ciepła z jednego miejsca do drugiego. Należą do nich techniki kompresji, systemy absorpcyjne i oraz odprowadzanie ciepła z mniejszych radiatorów do większych

bloków chłodzących za pomocą cieczy (chłodzenie wodne). Ogniwa Peltiera mają jednak sporo zalet w porównaniu z tymi tradycyjnymi systemami:

- ogniwa Peltiera chłodzą i ogrzewają.
- ogniwa Peltiera nie mają ruchomych części.
- ogniwa Peltiera działają całkowicie bezobsługowo.
- ogniwa Peltiera działają bezgłośnie.
- ogniwa Peltiera działają we wszystkich pozycjach i kierunkach ruchu. Wypompowane ciepło można bardzo łatwo uwolnić do powietrza zewnętrznego za pomocą tradycyjnych bloków chłodzących, ewentualnie uzupełnionych wentylatorami zaś moc cieplną pompy ciepła można regulować za pomocą dość prostych układów elektronicznych.

Fizyczne podstawy działania ogniwa Peltiera

Nazwa

Nazwa „ogniwo Peltiera” pochodzi oczywiście od nazwiska wynalazcy efektu Peltiera. To zjawisko fizyczne zostało odkryte w 1834 roku przez francuskiego fizyka Jeana Charlesa Athanase’a Peltiera. W codziennej praktyce często można spotkać się z alternatywnymi nazwami:

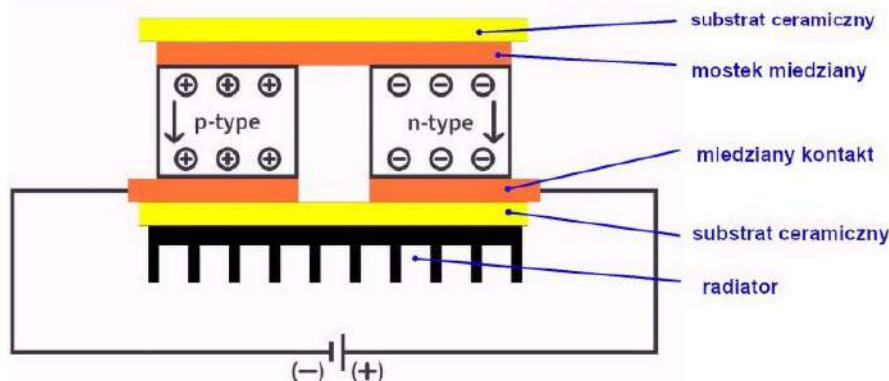
- Frigistor, pochodzi od angielskiego słowa oznaczającego chłodzenie.
- TED, skrót od „Thermo Electric Device” (urządzenie termoelektryczne).

Zasada działania efektu Peltiera

Jean Peltier ustalił podczas eksperymentów, że punkt styku dwóch różnych metali, przez które przepływał prąd stały, stawał się zimniejszy lub cieplejszy niż otoczenie. Kierunek przepływu prądu przez punkt styku determinował wzrost lub spadek temperatury punktu styku. W tamtym czasie odkrycie to było uważane za czysto akademickie i zostało dodane do bardzo obszernego zbioru dziwnych zjawisk fizycznych, które zostały odkryte przypadkowo. Nie widziano żadnych praktycznych zastosowań. Jednak dzięki technologii półprzewodnikowej efekt Peltiera stał się bardzo aktualny i praktycznie użyteczny.

Nowoczesne ogniwo Peltiera

Nowoczesne ogniwo Peltiera nie składa się z dwóch metali, ale z materiałów półprzewodnikowych. W większości stosuje się tellurek bizmutu (Bi_2Te_3), który jest silnie domieszkowany dodatnio i ujemnie. Dwa takie przeciwnie domieszkowane bloki materiału półprzewodnikowego n^- i p^+ są połączone ze sobą z jednej strony za pomocą miedzianego mostka. Zostało to przedstawione w górnej części poniższego rysunku. Dwie wolne strony są połączone z dwoma biegunami źródła napięcia stałego za pomocą



Skład nowoczesnego ogniwa Peltiera (© 2023 Jos Verstraten)

miedzianych powierzchni kontaktowych. Warstwa graniczna między półprzewodnikami a miedzianym mostkiem jest tak cienka ze względu na bardzo silne domieszkowanie, że nie może wystąpić rektyfikacja. Oznacza to, że prąd stały może przepływać przez strukturę w obu kierunkach.

Przy dość dużym prądzie (rzędu amperów), generacja swobodnego nośnika ładunku ma miejsce w obszarze granicznym górnej miedzi i materiału półprzewodnikowego. Elektrony walencyjne zrywają swoje wiązania atomowe. Wymaga to jednak energii, która jest pobierana z miedzi. W rezultacie temperatura mostka miedzianego spada. W pokazanym kierunku prądu mostek miedziany tworzy zimną stronę ogniwa Peltiera. W dolnej części ogniwa dochodzi oczywiście do rekombinacji elektronów i jonów dodatnich. Tworzy to ciepło, które jest uwalniane do dwóch miedzianych powierzchni kontaktowych. Konstrukcja jest umieszczona pomiędzy dwoma ceramicznymi substratami, bardzo cienkimi płytkami, które zapewniają izolację elektryczną ogniwa. Ciepła strona musi być zawsze wyposażona w radiator, który zapewnia odprowadzanie generowanego ciepła do powietrza. W ten sposób konstrukcja może działać jak pompa ciepła.

Efekt odwracalny

Z omówienia działania ogniwa Peltiera natychmiast wynika, że efekt jest odwracalny. Jeśli kierunek prądu zostanie odwrócony, generacja swobodnych nośników ładunku będzie miała miejsce na dwóch miedzianych powierzchniach kontaktowych, a rekombinacja na miedzianym mostku. W takim przypadku mostek tworzy gorącą stronę ogniwa, a powierzchnie kontaktowe zimną stronę. Jest to „efekt Seebecka”. Z konstrukcji ogniwa wynika jednak, że miedziany mostek ma pozostać zimną stroną ogniwa! Strona ta ma większą powierzchnię styku i dlatego może być najbardziej efektywnie wykorzystywana do pochłaniania ciepła z chłodzonej części.

Jednak w praktyce można znaleźć wiele systemów, w których odwracalność ogniwa Peltiera jest wykorzystywana do chłodzenia obiektu, gdy jest zbyt gorący lub do ogrzewania tego obiektu, gdy jest zbyt zimny.

Druga odwracalność

W zasadzie można traktować ogniwo Peltiera jako dużą termoparę. Nie będzie więc zaskoczeniem, że efekt Peltiera jest odwracalny również w tym obszarze. Można zmierzyć bardzo małe napięcie stałe na ogniwie, którego wielkość jest proporcjonalna do różnicy temperatur między górną i dolną częścią ogniwa.

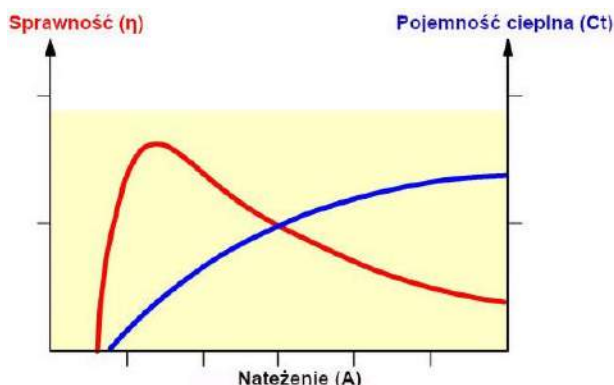
Energia pompy ciepłej

Ilość energii cieplnej przepływającej przez ogniwo w ciągu sekundy można wyrazić wzorem:

$$P_k = \alpha \cdot T_k \cdot I$$

gdzie:

- α jest równe współczynnikowi Seebecka materiału półprzewodnikowego



Sprawność i pojemność cieplna w funkcji natężenia prądu
(© 2023 Jos Verstraten)

- T_k jest temperaturą kryształu
- I jest wielkością prądu stałego.

Jednak w praktyce sprawność jest nieco niższa, ponieważ trzeba oczywiście wziąć pod uwagę powstawanie ciepła w wewnętrznej rezystancji materiału. Można to wyrazić za pomocą dobrze znanego wzoru:

$$P = I^2 \cdot R$$

Ponadto, w ogniwie występuje również przepływ ciepła z cieplej strony na zimną. Ponieważ jednak opór cieplny materiału półprzewodnikowego jest dość wysoki, jest on praktycznie pomijalny w porównaniu z energią pobieraną ze strony zimnej.

Na poniższym wykresie sprawność η jest pokazana po lewej stronie, a pojemność cieplna C_t jest pokazana po prawej stronie zaproponowanego ogniwa Peltiera. Zwrot najpierw wzrasta, a następnie ponownie maleje. Jest to wynikiem rosnącej mocy cieplnej generowanej przez rosnący prąd w ogniwie. Aby uzyskać najwyższą sprawność ogniwa Peltiera, należy przesłać przez ogniwo prąd o natężeniu niższym niż maksymalne dopuszczalne.

Ogniwa Peltiera w praktyce

Budowa

Ogniwo Peltiera ma rozmiar zaledwie kilku mm² i w praktyce nie można z nim wiele zrobić. Dlatego też ogniwa są zawsze łączone w bloki lub ogniwa chłodzące, składające się z dużej liczby identycznych ogniw. Podstawowy skład takiego ogniwa praktycznego zastosowania przedstawiono na poniższym rysunku. Jak wynika z rysunku konstrukcyjnego, poszczególne ogniwa są elektrycznie połączone szeregowo. Jednak termicznie ogniwa są równoległe, ponieważ wszystkie zimne mostki kontaktowe znajdują się po tej samej stronie konstrukcji. To samo dotyczy oczywiście ciepłych mostków stykowych. Zwiększa to powierzchnię chłodzenia i wydajność pompy ciepła.

Podłączenie poszczególnych ogniw szeregowo ma tę zaletę, że napięcia ogniw są również szeregowo i można zasilać całą konstrukcję bezpośrednim napięciem od 12 V do 15 V.

Całkowicie odwracalny

Omówienie konstrukcji praktycznie użytecznego ogniwa Peltiera pokazuje, że działa ono całkowicie odwracalnie. Powierzchnia miedzianych płytek kontaktowych po jednej stronie jest prawie taka sama jak po drugiej stronie. Przepływ ciepła w jednym kierunku jest zatem prawie równy przepływowi ciepła w drugim kierunku.

Maksymalizacja wydajności

Przepływ ciepła przez ogniwo Peltiera zależy między innymi od różnicy temperatur między stroną zimną i ciepłą. Aby zmaksymalizować tę różnicę, strona ciepła musi mieć jak największą powierzchnię. Można to osiągnąć poprzez przymocowanie tej strony do tradycyjnego radiatora, który może być dodatkowo chłodzony za pomocą wentylatora. Na poniższym zdjęciu widać „zestaw”, który jest oferowany na różnych chińskich stronach w cenie około 5,80 euro. Zestaw zawiera ogniwo Peltiera, radiator i wentylator.

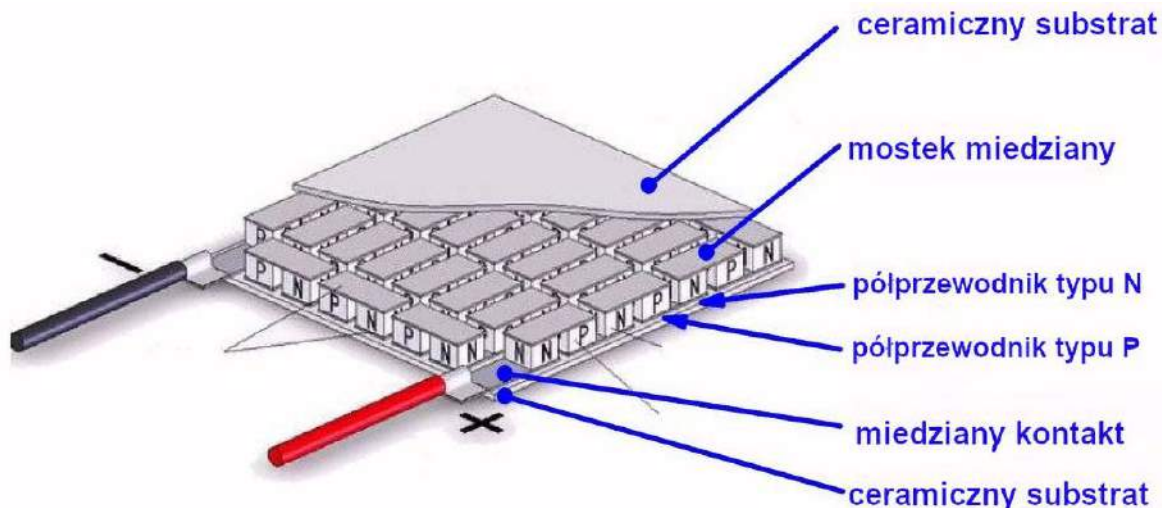
Praktyczne wersje ogniw Peltiera

Wprowadzenie

Chińscy producenci zdobyli również duży udział w rynku ogniw Peltiera. Na stronach takich jak Banggood i AliExpress można zamówić ogniwa Peltiera za kilka euro. Wyglądają one tak, jak na poniższym zdjęciu. Jest to TEC1-12706, który można kupić na AliExpress za nieco ponad dwa euro.

Oznaczenia ogniw Peltiera

Poniższy rysunek pokazuje, jak interpretować oznaczenia. Nie są one dowolne, lecz wynikają z pewnych ustaleń między producentami.



Oto jak zbudowane są praktyczne ogniwa Peltiera (© Microwaves101.com, pod redakcją Josa Verstratena)



Ogniwo Peltiera chłodzone aktywnie (© AliExpress)

Należy pamiętać, że określony jest tylko prąd, a nie napięcie. Należy więc zasilac ogniwo Peltiera stałym prądem zamiast stałym napięciem.

Czasami za ostatnimi dwiema cyframi występuje dodatkowe oznaczenie w formie T120, T160 czy T200. Określa ono maksymalną temperaturę pracy w stopniach Celsjusza. (Przyp. Tłum.)

Wydajność chłodzenia

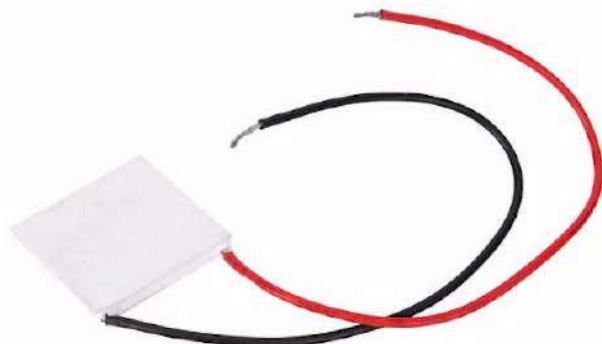
Jest to najważniejszy parametr i jest on wyrażany w watach przy równych temperaturach po obu stronach ogniwa, tj. natychmiast po włączeniu.

Strona gorąca i zimna

Jeśli podłączysz czerwony przewód do dodatniego bieguna zasilania, strona z nadrukiem będzie zimną stroną ogniwa.

Model TEC1-12703

- Wymiary: 40 mm × 40 mm × 4,2 mm

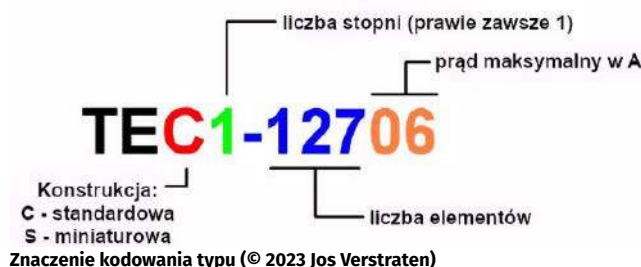


Typowy przykład taniej chińskiej kopii (© AliExpress)

- Długość przewodu: 310 mm
- Opór: 4,0 Ω ~ 4,3 Ω
- Prąd pracy: 3 A
- Wydajność chłodzenia: 18 W
- Temperatura pracy: -40°C ~ +60°C
- Maksymalna temperatura różnicowa: 65 °C
- Cena: 2,39 € (AliExpress)

Model TES1-12704

- Wymiary: 30 mm × 30 mm × 3,2 mm
- Długość przewodu: 200 mm
- Opór: 3, 37 Ω ~ 3,80 Ω



Znaczenie kodowania typu (© 2023 Jos Verstraten)

- Prąd pracy: 4 A
- Wydajność chłodzenia: 36 W
- Temperatura pracy: -40°C ~ +60°C
- Maksymalna temperatura różnicowa: 60°C
- Cena: 2,02 € (AliExpress)

Model TEC1-12705

- Wymiary: 40 mm × 40 mm × 3,4 mm
- Długość przewodu: 100 mm
- Rezystancja: 2,4 Ω ~ 2,7 Ω
- Prąd pracy: 5 A
- Wydajność chłodzenia: 42,5 W
- Temperatura pracy: -40°C ~ +60°C
- Maksymalna temperatura różnicowa: 61°C
- Cena: 2,11 € (AliExpress)

Model TEC1-12706

- Wymiary: 40 mm × 40 mm × 3,75 mm
- Długość przewodu: 100 mm
- Rezystancja: 2,1 Ω ~ 2,4 Ω
- Prąd pracy: 4,6 A

- Wydajność chłodzenia: 50 W
- Temperatura pracy: $-40^{\circ}\text{C} \sim +60^{\circ}\text{C}$
- Maksymalna temperatura różnicowa: 61°C
- Cena: 2,11 € (AliExpress)

Praktyczne obwody z ogniwami Peltiera

Sterowanie stałym prądem lub modulacją szerokości impulsu

Jak wynika ze wzoru na moc, moc cieplna ogniwa Peltiera jest wprost proporcjonalna do prądu przesyłanego przez ogniwo. Ta idealna zależność oznacza, że bardzo łatwo jest elektronicznie sterować mocą pompy ciepła. Wystarczy opracować obwód elektroniczny, który przesyła pewien stały prąd przez ogniwo Peltiera. Rozszerzając to źródło zasilania o system sprzężenia zwrotnego, z czujnikiem temperatury jako wejściem i ogniwem Peltiera jako wyjściem, można ustabilizować temperaturę w pomieszczeniu na określonej wartości. Na przykład, za pomocą prostego obwodu elektronicznego i wysokowydajnego ogniwa Peltiera można zaprojektować komorę chłodniczą, w której temperatura pozostaje stała z tolerancją $\pm 1/10^{\circ}\text{C}$. Za pomocą takiej komory można badać zachowanie termiczne komponentów elektronicznych.

Zamiast sterować stałym prądem, można również sterować stałym napięciem dostarczającym do ogniwa Peltiera za pomocą sterowania PWM za pośrednictwem tranzystora MOSFET. Taka modulacja szerokości impulsu zapewnia, że średni prąd przepływający przez ogniwo Peltiera można regulować w zakresie od 0% do 100% wartości maksymalnej.

Schemat elektronicznej lodówki turystycznej

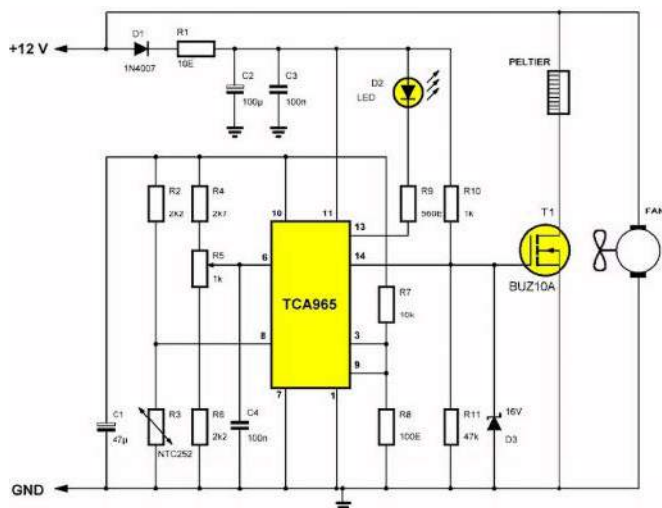
Poniższy rysunek przedstawia schemat elektronicznej lodówki turystycznej, opublikowany w Funkschau. Temperaturę w chłodzarnie można ustawić w zakresie od -2°C do $+25^{\circ}\text{C}$ za pomocą potencjometru.

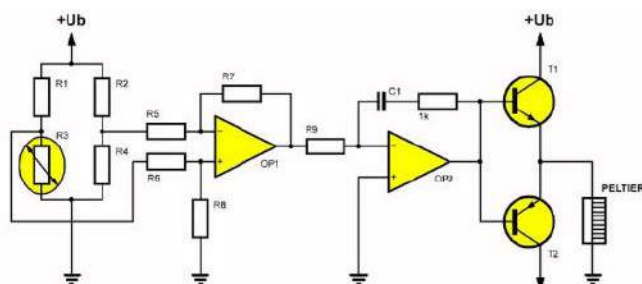
Schemat ten nie stosuje proporcjonalnego sterowania ogniwem Peltiera. Prąd przepływający przez ogniwo jest włączany lub wyłączany, dzięki czemu sterowanie to można porównać do sterowania termostatem kotła centralnego ogrzewania. Jako element sterujący zastosowano stary, ale wciąż dostępny układ TCA965 firmy Siemens. Ten układ scalony komparatora ma wbudowany generator odniesienia, który generuje bardzo stabilne napięcie około 7,9 V na pinie 10. To napięcie odniesienia jest wykorzystywane do zasilania czujnika temperatury i potencjometru. Czujnik temperatury, rezystor o ujemnym współczynniku temperaturowym NTC252, jest oczywiście umieszczony w chłodzonej przestrzeni. NTC znajduje się w dzielniku napięcia z rezystorem R2 z folii metalowej o wartości 2,2 k Ω . Napięcie na styku obu rezystorów zależy od temperatury w lodówce. Napięcie to jest podawane na wejście 8 układu TCA965. Napięcie z potencjometru trafia na wejście 6. Układ scalony porównuje te dwa napięcia i steruje ogniwem Peltiera za pośrednictwem tranzystora MOSFET BUZ10A. Jeśli napięcie na pinie 6 jest większe niż napięcie na pinie 8, ogniwo jest włączone. W przeciwnym razie prąd płynący przez ogniwo spadnie do 0.

Abym zapobiec ciągłemu przełączaniu komparatora w obszarze U6=U8, zapewniona jest histereza. Próg ten jest ustawiany przez stosunek rezystorów R7 i R8 na pinach 3 i 9. Przy stosunku 10 k Ω i 100 Ω histereza wynosi około 1°C . Jeśli obniżysz dolną rezystancję do 56 Ω , histereza spadnie do około $0,5^{\circ}\text{C}$.

Pin 13 przechodzi w stan niski, gdy tranzystor MOSFET jest wysterowany. Można zatem użyć tego wyjścia jako wskaźnika działania za pomocą diody LED D2.

Tranzystor MOSFET BUZ10A może obsługiwać maksymalny prąd 12 A przy maksymalnym napięciu 50 V. W tym zastosowaniu maksymalny prąd wynosi 5 A, a maksymalne napięcie 14 V, więc tranzystor MOSFET może być używany w tym obwodzie bez żadnych problemów.





Ogrzewanie i chłodzenie za pomocą ogniw Peltiera (© 2023 Jos Verstraten)

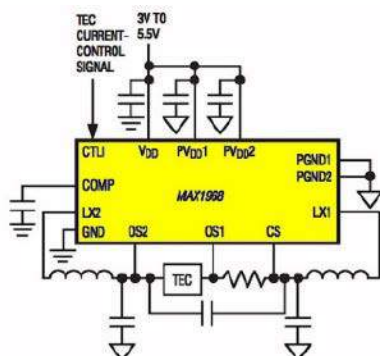
i ilość prądu dostarczanego do ogniw Peltiera. Ponieważ tranzystory T1 i T2 działają liniowo, mogą zużywać znaczną ilość energii podczas przewodzenia. Chłodzenie jest zatem konieczne!

Układ MAX1968 firmy Maxim zasila ogniwo Peltiera prądem ± 3 A

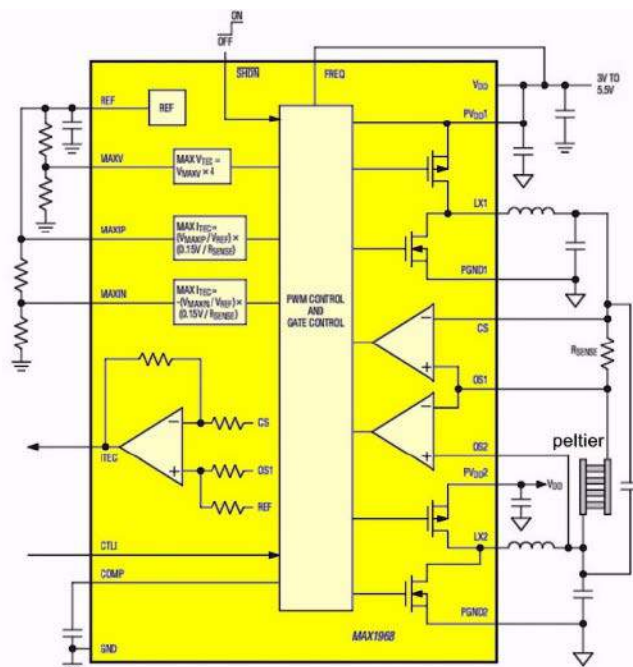
Firma Maxim opracowała bardzo drogi układ scalony, który jest całkowicie skoncentrowany na odwracalnym zasilaniu ogniwa Peltiera z pojedynczego napięcia zasilania. Oznacza to, że ten układ scalony zawiera obwody, które tworzą dwa konwertery buck, które generują zarówno dodatnie jak i ujemne napięcia do przesyłania prądu w dwóch kierunkach przez ogniwo Peltiera. W RS-Components układ ten jest oferowany w cenie około 25,00 € za sztukę. Poniższy rysunek przedstawia podstawowy schemat układu MAX1968. Ogniwo Peltiera jest tutaj określane jako „TEC”. Układ scalony wykorzystuje sterowanie prądem stałym w celu wyeliminowania szczytów prądu w elemencie Peltiera. Obwód jest zasilany napięciem zasilania o maksymalnej wartości +5 V. Dwie synchroniczne przetwornice buck działają z częstotliwością od 500 kHz do 1 MHz, dzięki czemu jako elementy wyglądające można zastosować bardzo niskie pojemności kondensatorów. Dwie cewki połączone szeregowo z ogniwnem Peltiera stanowią najważniejsze części przetwornic buck. Zapewniają one podnoszenie napięć zasilających dostarczanych do ogniwa Peltiera.

Układ MAX1968 jest w stanie zasilac ogniwo Peltiera prądem o natężeniu od -3 A do $+3$ A bez użycia zewnętrznych komponentów. Pozwala to uniknąć „martwych stref” wokół punktu zerowego. Prąd przepływający przez ogniwo jest po prostu regulowany przez przyłożenie analogowego napięcia sterującego do wejścia CTLI układu scalonego.

Typowa aplikacja dla układu MAX1968 zalecana przez producenta pokazano na schemacie poniżej. Z tego schematu można również wywnioskować wewnętrzną strukturę tego układu scalonego. Szczegółowe omówienie tego schematu wykraczałoby jednak poza zakres artykułu wprowadzającego do tematu ogniw Peltiera. Kartę katalogową układu MAX1968 można pobrać za pośrednictwem poniższego łącza: <https://tiny.pl/cspcz>.



Podstawowy schemat układu MAX1968 (© 2015 Maxim Integrated Products, Inc.)



Przykładowy obwód wokół MAX1968 (© 2015 Maxim Integrated Products, Inc.)

Instalacja ogniwa Peltiera

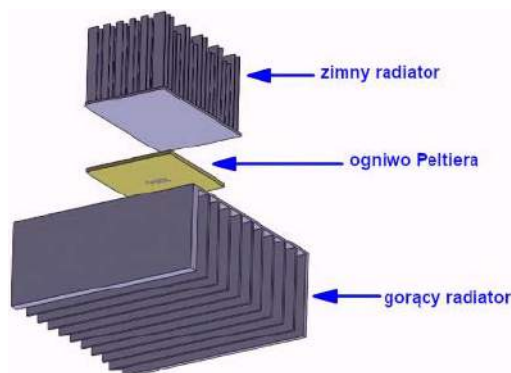
Montaż w przypadku użycia ogniwa jako radiatora

Jeśli chcesz użyć ogniw Peltiera do chłodzenia czegoś, montaż nie stanowi problemu. Należy upewnić się, że zimna strona ogniwa ma dobry kontakt termiczny z chłodzonym obiektem. Po ciepłej stronie należy zapewnić dobre odprowadzanie ciepła. Po tej stronie montujesz radiator, z ewentualnym wentylatorem. Takie konstrukcje są w sprzedaży wszędzie, zdjęcie takiego zestawu pokazano już w innym miejscu tego artykułu.

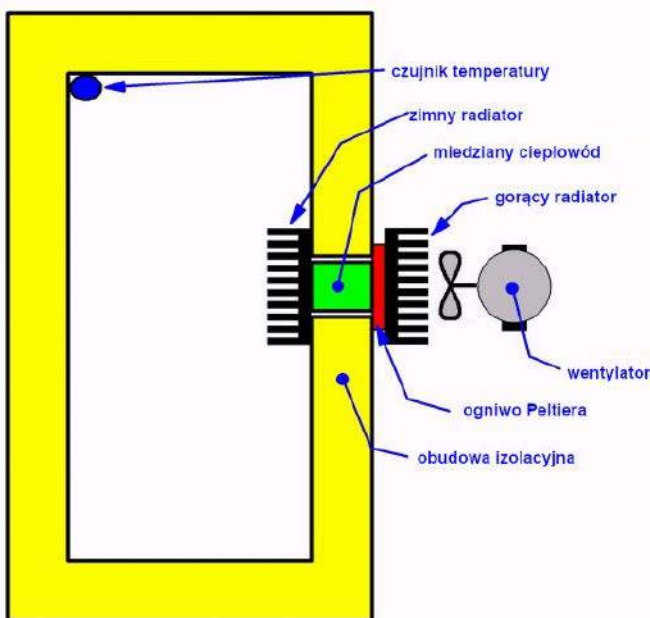
Instalacja przy zastosowaniu jako stabilizator temperatury

Jeśli używasz ogniwa Peltiera do stabilizacji temperatury w pomieszczeniu, sprawy stają się nieco bardziej skomplikowane. Radiatory są niezbędne zarówno po stronie zimnej, jak i ciepłej, aby zmaksymalizować transport ciepła w obu kierunkach.

Ale... jest większy problem! Podczas usuwania ciepła z jednej strony na drugą, należy upewnić się, że usunięte ciepło nie może przepływać z powrotem do zimnej strony za pomocą wszelkiego rodzaju skrótów. Należy więc zapewnić najlepszą możliwą barierę termiczną pomiędzy stroną zimną i gorącą. Poniższy rysunek pokazuje, jak to zrobić. Pomieszczenie, w którym chcemy utrzymać stałą temperaturę, musi być zaizolowane grubą warstwą styropianu (kolor żółty). W warstwie tej znajduje się otwór, w którym instaluje się pompę



Schemat zasady stabilizacji temperatury (© 2023 Jos Verstraten)



Izolacja termiczna między gorącą i zimną stroną (© 2023 Jos Verstraten)

ciepła. Ogniwo Peltiera (czerwony) jest montowane na zewnątrz, razem z ciepłym profilem z wentylatorem. Zimny profil znajduje się po wewnętrznej stronie styropianu. Przestrzeń między tym profilem a ogniwem Peltiera musi być wypełniona blokiem z litej miedzi (zielony), który służy jako przewodnik ciepła między zimnym profilem a ogniwem Peltiera. Zamontuj czujnik temperatury (niebieski) jak najdalej od pompy ciepła, w tym przypadku w lewym górnym rogu izolowanego pomieszczenia.



Zależność między napięciem na elemencie Peltiera a różnicą temperatur po obu stronach ogniwa (© www.ydom.ru, edytuj Jos Verstraten)

Ogniwo Peltiera jako źródło napięcia

Podczas omawiania działania ogniwa Peltiera wskazaliśmy już, że takie ogniwo jest odwracalne. Na ogniwie mierzy się niewielkie napięcie stałe, którego wartość zależy od różnicy temperatur Δt między obiema stronami ogniwa. Jest to „efekt Seebecka”. Jednak napięcie na jednym ogniwie jest tak małe, że niewiele można z nim zrobić. Zmienia się to w przypadku pomiaru napięcia na elemencie Peltiera. TEC1-12706 zawiera 127 ogniw, a napięcia na tych ogniwach są połączone szeregowo. Na takim elemencie Peltiera powstaje zatem doskonale mierzalne napięcie, które jest w przybliżeniu równe 20 mV na każdy stopień różnicy temperatur. Jak pokazuje poniższy wykres, zależność między napięciem a różnicą temperatur jest nawet liniowa. Prosty system, który pozwala zmierzyć temperaturę obiektu! Należy jednak zawsze pamiętać, że pomiar ten nie mierzy temperatury bezwzględnej, ale różnicę między dwiema temperaturami! ■

Jos Verstraten



Ogniwa Peltiera

Rozwiązanie znajdziesz na www.elportal.pl/quizy

1. Ogniwa Peltiera służą do:

- ☐ wytwarzania energii elektrycznej z różnicy temperatur między stronami;
- ☐ wytwarzania różnicy temperatur między stronami za pomocą energii elektrycznej;
- ☐ wymuszania przepływu energii cieplnej między stronami za pomocą energii elektrycznej.

2. Wydajność pracy ogniwa Peltiera zależy od:

- ☐ napięcia;
- ☐ prądu;
- ☐ różnicy temperatur.

3. Współczesne półprzewodnikowe ogniwa Peltiera są wykonane z:

- ☐ selenku bizmutu;
- ☐ tellurku bizmutu;
- ☐ arsenku bizmutu.

4. Gdy ogniwo Peltiera podłączy się odwrotnie:

- ☐ strona zimna stanie się stroną gorącą, a strona gorąca stanie się stroną zimną;
- ☐ nic się nie stanie, bo ogniwo nie przewodzi w kierunku zaporowym;
- ☐ ogniwo ulegnie uszkodzeniu.

5. Sprawność ogniwa jest ograniczona przez:

- ☐ niską rezystancję cieplną elementów;
- ☐ wysoką rezystancję cieplną elementów;
- ☐ rezystancję elektryczną elementów.

6. Maksymalna różnica temperatur między stronami ogniwa, jaką da się uzyskać w praktyce wynosi:

- ☐ 50–55 stopni;
- ☐ 60–65 stopni;
- ☐ 70–75 stopni.

7. Czy ogniwo Peltiera może pracować jako źródło energii?

- ☐ tak, ale wydajność takiego rozwiązania jest zbyt niska, by miało to praktyczny sens;
- ☐ tak, ale różnica temperatur między stronami musi być bardzo wysoka, co czyni to rozwiązanie trudnym w realizacji;
- ☐ nie, zjawiska dzięki którym ogniwo Peltiera pracuje jako pompa ciepła, nie są odwracalne.

8. Jednym z praktycznych zastosowań ogniw Peltiera w elektronice nie jest:

- ☐ chłodzenie półprzewodników;
- ☐ stabilizacja temperatury wrażliwych na jej zmiany elementów;
- ☐ ogrzewanie katod w specjalizowanych lampach elektronowych, jak klistrony.

9. Jedną z nazw dla ogniw Peltiera spotykanych w literaturze zagranicznej jest:

- ☐ frigistor;
- ☐ chilliac;
- ☐ thermactor.

10. Czy zimna strona ogniwa Peltiera może osiągnąć temperatury poniżej zera stopni Celsjusza?

- ☐ nie;
- ☐ tak;
- ☐ tak, ale tylko do -5°C , bo w niższych temperaturach użyty półprzewodnik pęka.

Patronat EdW nad szkołami i uczelnianymi Kołami Naukowymi rozkwita i daje redakcji EdW impulsy zachęcające do wspierania edukacji szkolnej i uczelnianej. Działa sprzężenie zwrotne. Dostajemy mnóstwo wiadomości od uczniów, nauczycieli i studentów. Dla nich jest ta rubryka.



Co to jest multimetr cyfrowy?

Multimetr, zwany również miernikiem uniwersalnym lub miernikiem AVO, to przyrząd pomiarowy służący do pomiaru różnych wielkości elektrycznych. Wynalezienie multimetru przypisuje się brytyjskiemu technikowi PTT Donaldowi Macadie. Zmęczony noszeniem przy sobie wielu przyrządów pomiarowych, w 1927 roku zaprojektował miernik, za pomocą którego mógł mierzyć napięcia, prądy i rezystancję. Pierwszym z komercyjnych multimetrów jest bez wątpienia Model 8 angielskiej firmy AVO. Model ten został wprowadzony w 1951 roku i był produkowany do 2008 roku. Ostatnią wersją był Mark 7. Prawie każdy inżynier elektroniki mający więcej niż 40 lat pracował z tym multimetrem. Miernik ten był powszechnie używany na zajęciach technicznych o różnym poziomie zaawansowania. Za jego pomocą demonstrowano podstawowe prawa elektryczności. Model 8 mógł mierzyć napięcie stałe, napięcie przemiennie, prąd stały, prąd przemienny i rezystancję. Nazwa AVO składała się z trzech jednostek: amper, volt i om. Starsi eksperci od elektroniki nadal mówią o „mierniku AVO” zamiast o multimetrze.

Pierwszy multimetr cyfrowy został wprowadzony na rynek w 1955 roku przez American Non Linear Systems. Był to oczywiście ciężki i nieporęczny przyrząd stołowy. Pierwszy multimetr ręczny został opracowany przez Intron Electronics i pojawił się około 1977 roku.

Rodzaje multimetrów

Nowoczesne multimetry cyfrowe mierzą znacznie więcej niż A, V i O. Nowoczesna technologia zapewnia również możliwość wyboru spośród wielu wzorów i setek typów urządzeń. Wszystkie przyrządy pomiarowe można podzielić na dziewięć kategorii. Na poniższym zdjęciu zestawiliśmy te dziewięć typów, od lewej do prawej:

- multimetry regulowane ręcznie,
- multimetry półautomatyczne (z automatycznym zakresem),
- multimetry w pełni automatyczne (inteligentne),
- multimetry długopisowe (piórowe),
- multimetry stołowe,
- multimetry szczypcowe (cęgowe),
- multimetry pęsetowe (pincetowe),



Multimetr jest tak samo niezbędny dla elektronika hobbysty jak lutownica. W tym wykładzie omawiamy rodzaje, specyfikacje i budowę tego bardzo przydatnego przyrządu.



Śłynny Model 8 od AVO (© Peter Vis)



Dziewięć typów multimetrów (© 2023 Jos Verstraten)

- multimetry smartfonowe,
- multimetry graficzne.

Multimetr z ręczną zmianą zakresów

Ten typ charakteryzuje się dużym przełącznikiem obrotowym z wieloma pozycjami, czasami ponad trzydziestoma. Każda pozycja odpowiada jednemu zakresowi pomiarowemu. Podczas obsługi takiego miernika należy zawsze zwracać szczególną uwagę na to, czy przełącznik obrotowy jest ustawiony w żądanej pozycji.

Ponieważ takie przełączniki nie są produkowane standardowo, a zaprojektowanie i wyprodukowanie takiej części dla jednego typu multimetru byłoby zbyt kosztowne, są one zawsze projektowane jako przełączniki drukowane. Segmenty miedzianych okręgów są wytrawiane na płytce drukowanej, a pod przyciskiem przełącznika obrotowego umieszcza się szereg przesuwanych styków, które łączą odpowiednie segmenty okręgów dla każdej pozycji przełącznika. Poniższe zdjęcie pokazuje, jak taki przełącznik wygląda w praktyce. Pytanie brzmi, czy kontakt pomiędzy ścieżkami na płytce drukowanej a stykami ślizgowymi sprawdza się w dłuższej perspektywie? W droższych modelach ścieżki i styki ślizgowe są połączane, aby ten słaby punkt był nieco bardziej niezawodny.



Ręcznie regulowany multimetr (© AliExpress)



Konstrukcja przełącznika (© 2023 Jos Verstraten)



Półautomatyczny multimetr (automatyczny zakres pomiarowy)

„Automatyczny zakres pomiarowy” to funkcja, dzięki której urządzenie automatycznie wybiera właściwy zakres pomiarowy dla wybranej wielkości. Oznacza to, że wystarczy wybrać mierzoną wielkość za pomocą przełącznika obrotowego, a miernik sam wykona wszystkie czynności regulacyjne. Jeśli mierzona wartość nagle się zmienia, multimetr automatycznie dostosuje zakres, aby zachować najwyższą rozdzielczość pomiaru.

Na poniższym zdjęciu widać typowy przykład takiego półautomatycznego lub automatycznego miernika zakresu. Przełącznik obrotowy ma tu tylko dziewięć pozycji, co znacznie ułatwia obsługę miernika i eliminuje błędy.

Niektóre multimetry z automatycznym zakresem pomiarowym mają dodatkowy przycisk „RANGE”. Wyłącza on funkcję automatyczną i umożliwia ręczne ustawienie zakresu pomiarowego poprzez kilkukrotne naciśnięcie przycisku.

W pełni automatyczny multimetr (smart)

W kolejnej kategorii mierników automatyzacja poszła o krok dalej. Nie tylko zakres, ale także funkcja pomiarowa jest ustawiana przez urządzenie. Tak przynajmniej twierdzą reklamodawcy takich „inteligentnych” multimetrów. W praktyce funkcja ta działa tylko częściowo i czasami bywa zawodna. Jeśli podłączysz do miernika napięcie stałe, napięcie przemienne lub rezystancję, urządzenie bez wątpienia to wykryje i włączy odpowiednią funkcję pomiarową. W przypadku wszystkich innych wielkości, które możemy zmierzyć takim multimetrem, takich jak pojemności lub częstotliwości, nadal będziesz musiał ustawić wielkość ręcznie. Ale ponieważ multimetr jest najczęściej używany do pomiaru napięcia i rezystancji, funkcja „smart” często działa zadowalająco. Oczywiście większość inteligentnych multimetrów oferuje opcję wyłączenia automatyki i ustawienia wszystkiego ręcznie.

Multimetr piórowy

Trudno jest zamocować piny miernika we właściwych miejscach na płytce drukowanej i odczytać miernik, zwłaszcza teraz, gdy części na płytce drukowanej stały się tak małe. Aby rozwiązać problem oczu biegających w tę i z powrotem między płytką drukowaną a wyświetlaczem multimetru, wynaleziono multimetr piórowy. Taki miernik ma tylko jeden przewód pomiarowy, a drugim biegunem jest ostra igła zamontowana na końcu miernika. Takie mierniki są zawsze półautomatyczne (z automatyczną zmianą zakresów). Wielkość



Multimetr z automatycznym zakresem pomiarowym (© Fluke)



W pełni automatyczny multimetr (© Mustek)

mierzoną ustawia się za pomocą małego pokrętki. Aby dokonać pomiaru należy podłączyć jeden przewód pomiarowy do masy obwodu, w którym chcemy dokonać pomiaru. Należy użyć ostrej sondy na czubku multimetru, aby zlokalizować punkt, w którym chcemy dokonać pomiaru. Ponieważ można skupić wzrok na tym punkcie, w końcu widać także wyświetlacz, szansa na ześlizgnięcie się igły z punktu pomiarowego jest znacznie mniejsza. Większość mierników piórowych ma ograniczony wybór mierzonych wielkości. Prawie nigdy nie można nimi mierzyć prądu.



Typowy multimetr piórowy (© Zangzhou)

Multimetr stołowy

Przyrządy te są większymi wersjami multimetrów ręcznych i mają tę wielką zaletę, że nawet urządzenia dla elektroników – amatorów są solidne i stabilne. Znacznie większa obudowa zapewnia miejsce na zamontowanie większego wyświetlacza na płycie czołowej. Przyciski sterujące są również nieco większe. Wymiary obudów są (mniej więcej) dostosowane do wymiarów obudów generatorów funkcyjnych, mierników RLC i niektórych płaskich zasilaczy, co umożliwia układanie sprzętu pomiarowego w stosy. Można też zamocować w urządzeniu poręczny uchwyt do przenoszenia w takiej pozycji, aby można było umieścić obudowę pod kątem na stole, tak aby wyświetlacz był łatwy do odczytania.

Dzięki ekstremalnej miniaturyzacji nowoczesnej elektroniki, obudowy takich multimetrów są w dużej mierze puste. Niektórzy producenci wykorzystują tę pustą przestrzeń do stworzenia schowka na przewody pomiarowe. Udogodnieniem, które jest obecne w większości multimetrów stołowych, ale nie w ręcznych, jest to, że bezpiecznik chroniący najwyższy zakres prądu jest wbudowany w gniazdo 4 mm na panelu przednim. W razie przypadkowego przepalenia tego bezpiecznika nie trzeba otwierać obudowy, ale można go łatwo wymienić.

Urządzenia te mogą być obsługiwane ręcznie lub półautomatycznie. Główną różnicą w stosunku do wersji ręcznych jest to, że istnieją urządzenia stołowe, które nie są obsługiwane za pomocą przełączników obrotowych, ale za pomocą przycisków (podświetlanych lub nie). Niektórzy uważają, że jest to znacznie wygodniejsze i bardziej przejrzyste niż przełączniki obrotowe.

Kolejną istotną różnicą jest to, że wszystkie urządzenia stacjonarne są zasilane z sieci, choć niektóre modele mają również wbudowane baterie lub akumulatory.



Typowy multimetr piórowy (© Zangzhou)



Multimetr cęgowy

Pomiar prądu za pomocą omówionych dotychczas wersji wymaga przerywania obwodu, w którym ma być mierzony prąd. Następnie należy podłączyć miernik szeregowo między zasilaniem, a obciążeniem. Jest to zawsze trudne, a czasami nawet niemożliwe. Aby rozwiązać ten problem, opracowano multimetry cęgowe, zwane po angielsku clamp meters. Na poniższym zdjęciu widać typowego przedstawiciela tego modelu. Po prawej stronie miernika znajdują się dwa złącza 4 mm, za pomocą których można mierzyć rezystancję, napięcie i ewentualnie kondensatory. Aby pomiar prądu był możliwy, należy użyć zacisku po lewej stronie miernika. Zacisk ten można otworzyć, aby możliwe było zamocowanie go wokół przewodu przewodzącego prąd. Prąd elektryczny przepływający przez przewód wytwarza wokół niego pole magnetyczne. Pole to jest mierzone przez zacisk i przekształcane w niewielkie napięcie, które jest mierzone przez miernik. Zazwyczaj używany jest czujnik Halla wbudowany w cęg. Napięcie wyjściowe tego czujnika może zostać przekształcone przez procesor multimetru na wartość prądu, która następnie pojawia się na wyświetlaczu.

Należy pamiętać, że pomiar prądu płynącego przez przewód zasilający w ten sposób nie ma sensu! W końcu w kablu znajdują się dwa przewody, w których prądy płyną w przeciwnych kierunkach. Wynikowe pole magnetyczne wokół kabla wynosi wtedy zero. Mierniki cęgowe działają tylko wokół jednego przewodu przewodzącego prąd.

Multimetr pęsetowy (pęseta)

To dość nowe osiągnięcie w dziedzinie multimetrów pokazano na poniższym zdjęciu. Za pomocą takiego urządzenia można mierzyć wszystkie komponenty, które można zacisnąć między



Typowy wygląd miernika cęgowego (© UNI-T)

dwiema szczękami pęsety: baterie, akumulatory, rezystory, diody i kondensatory. Większość mierników pęsetowych nie może mierzyć napięć i prądów przemiennych. Oczywiście mierniki te działają również z automatycznym zakresem.

Multimetr w kształcie smartfona!

Producenci sprzętu pomiarowego (szczególnie chińscy) projektują obecnie multimetry, które do złudzenia przypominają smartfony. Poniższe zdjęcie przedstawia model 683 firmy Aneng. Multimetr ma zaledwie 2,5 cm grubości, a ekran wypełnia cały przód urządzenia. Po raz pierwszy ekran dotykowy służy również do ustawiania multimetru.

Multimetr graficzny

Droższe wersje multimetrów oferują opcję wyświetlania na wyświetlaczu danych innych niż alfanumeryczne. Poniższy model firmy Fluke umożliwia rejestrowanie danych i tworzenie wykresów zarejestrowanych danych na ekranie. W ten sposób można natychmiast zobaczyć, jak mierzona wielkość zmienia się w czasie. Inne, głównie chińskie, multimetry graficzne zawierają prosty oscyloskop, dzięki czemu można również wyświetlać kształt mierzonego sygnału na ekranie.

Inny wygląd multimetrów

Wyświetlacz

Jeśli chodzi o sposób wyświetlania wartości mierzonej na wyświetlaczu, można wybierać pomiędzy różnymi technologiami i konstrukcjami. Najprostsze wyświetlacze składają się z ekranu LCD pokazującego tylko cztery siedmiosegmentowe wskaźniki, które pokazują wartość mierzonej wielkości. Te zaliczane do najbardziej zaawansowanych mają ekran graficzny o wysokiej rozdzielczości, który nie tylko wyświetla liczby bardzo starannie określonej czcionką, ale także pokazuje wiele innych informacji alfanumerycznych. Wyświetlacz wskazuje na przykład, do czego służą przyciski poniżej lub obok wyświetlacza. Poniższy rysunek, na którym zebrano wiele wyświetlaczy multimetrów, dobrze ilustruje ten ogromny wybór!

Obecnie standardem w lepszych miernikach jest pokazywanie na wyświetlaczu dwóch wskaźników numerycznych.

Pierwszy duży wskaźnik wyświetla wartość mierzonej wielkości, drugi mniejszy wskaźnik wyświetla drugą wielkość mierzonego sygnału. W ten sposób można wizualizować nie tylko wielkość napięcia przemiennego, ale także częstotliwość. Na przykład, można zobaczyć aktualną wartość napięcia stałego, ale także minimalną, maksymalną lub średnią wartość podczas cyklu pomiarowego. Drugim udogodnieniem, które można znaleźć w prawie wszystkich multimetrach, jest opcja wyświetlania zmierzonej wartości nie tylko liczbowo, ale także analogowo na skali przypominającej analogowy termometr.

Czasami na wyświetlaczu pojawia się tak dużo danych, że odczyt miernika przestaje być czytelny. Niestety nigdy nie znaleźliśmy opcji pozwalającej wybrać, które dane chcemy widzieć, a które nie.

Połączenia

W tym względzie wyłonił się swego rodzaju standard, co przedstawiono na poniższym rysunku. Należy pamiętać, że nie wszyscy producenci przestrzegają tego standardu! Multimetry, które spełniają ten standard, mają cztery gniazda 4 mm rozmieszczone w odległości 19 mm od siebie. Najbardziej wysunięte na prawo czerwone gniazdo jest przeznaczone do pomiaru napięcia, rezystancji, kondensatorów, diod i temperatury. Obok znajduje się czarne gniazdo, „COM”, które jest wspólnym połączeniem dla



Pęseta lub miernik pęsety od UNI-T (© UNI-T)



Multimetr do smartfona Model 683 od Aneng (© Banggood)



Przykład multimetru graficznego (© Fluke)



Zbrano różne wyświetlacze multimetrów (© 2023 Jos Verstraten)



Faktyczny standard połączeń (© 2023 Jos Verstraten)

wszystkich pomiarów. Dwa lewe złącza są również czerwone i przeznaczone do pomiaru prądów w zakresie μA /mA lub A. Jak już opisano, niektóre multimetry stacjonarne mają bezpieczniki wbudowane w te gniazda w celu ochrony pomiarów prądu.

Złącza „SENSE”

Multimetry lepszej jakości mają również dwa tak zwane złącza „SENSE”. Niektóre mierniki nazywają to „ $\Omega 4\text{W}$ ”. Służą one do pomiaru małych rezystancji zgodnie z metodą czteroprzewodową, do czego wrócimy w dalszej części tego artykułu. Te dwa złącza „SENSE” znajdują się zawsze obok złącz „V/Ohm” i „COM”.

Funkcje multimetrów cyfrowych

Pomiar wielkości

Głównym zadaniem multimetru jest oczywiście pomiar dokładnych wartości różnych wielkości. Z biegiem lat liczba wielkości, które można mierzyć, znacznie się rozszerzyła. W przypadku pierwszych multimetrów można było mierzyć tylko napięcia, prądy i rezystancje. Nowoczesne mierniki pozwalają mierzyć:

- Napięcia stałe w mV lub V
- Napięcia przemienne w mV lub V
- Prądy stałe w μA , mA lub A
- Prądy przemienne w μA , mA lub A
- Rezystory w Ω , k Ω lub M Ω
- Kondensatory w pF, nF, μF lub mF
- Temperatury w $^{\circ}\text{C}$ lub $^{\circ}\text{F}$
- Indukcje w μH lub mH
- Wzmocnienia/tłumienia w dB
- Częstotliwości w Hz, kHz lub MHz
- Wypełnienie dla sygnałów PWM w %

Indukcje, czyli cewki? Tak, znaleźliśmy co najmniej jeden multimetr, który to potrafi – HP-770C od HoldPeak. Oprócz tych podstawowych funkcji, prawie wszystkie multimetry mają szereg dodatkowych funkcji:

- test diody
- test ciągłości
- test NCV
- test LIVE
- test HFE
- funkcja HOLD
- funkcja REL
- funkcja True-RMS
- funkcja MIN/MAX
- funkcja LOG
- funkcja automatycznego wyłączania
- funkcja podświetlenia
- funkcja Low-Z
- funkcja Signal-out.

Należy pamiętać, że nie wszystkie multimetry mają wszystkie te funkcje!

Test diody

Funkcja ta pozwala określić, czy diody podłączone w jednym kierunku przewodzą prąd i czy nie przewodzą w kierunku zaporowym. Może być również używana do testowania diod LED.

Podczas tego „testu diody” bardzo ważne jest napięcie między otwartymi zaciskami multimetru. Jeśli jest to tylko 2,5 V, nie można przetestować wszystkich typów diod LED. Lepszy rodzaj multimetru zapewnia 3,5 V dla tego testu, dzięki czemu można również testować niebieskie i ultrafioletowe diody LED.

Test ciągłości

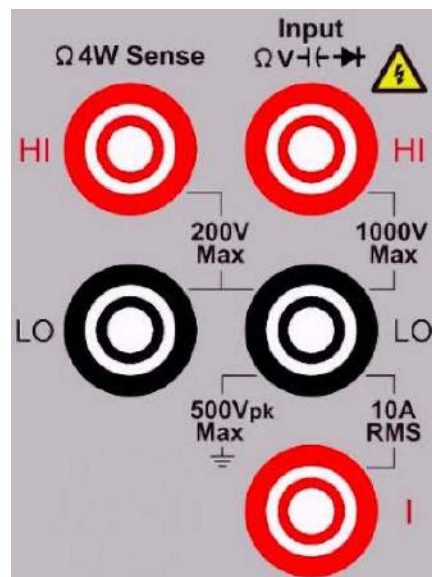
Ten test sprawdza, czy między dwoma przewodnikami istnieje połączenie o niskiej impedancji. Jeśli rezystancja między dwoma przewodami jest mniejsza niż, na przykład, 50 Ω , w mierniku zabrzmi brzęczyk. Wartość graniczna różni się w zależności od miernika, czasami stosuje się 100 Ω .

Test NCV

NCV jest akronimem *Non Contact Voltage*. Funkcja ta pozwala zlokalizować przewody AC w ścianie. W obudowie miernika znajduje się mała cewka. Podczas skanowania ściany za pomocą miernika, odbiera ona pole elektromagnetyczne wokół przewodnika, a miernik emituje sygnał dźwiękowy. Wskazanie siły pola często pojawia się na wyświetlaczu. Jednak jak dotąd w naszych testach multimetrów nigdy nie byliśmy pod wrażeniem tej funkcji.

Test LIVE

Dzięki tej opcji można wykryć, który przewód 230 V jest fazą, a który przewodem neutralnym. Funkcja ta wykrywa niezwykle mały pojemnościowy prąd upływowy, który przepływa między fazą a uziemieniem przez miernik i ciało użytkownika. Nie jesteśmy zbyt



Połączenia „SENSE” (© 2023 Jos Verstraten)

entuzjastycznie nastawieni do tej funkcji w testowanych przez nas multimetrach. Nie polegaj na niej! Używaj starego typu testera napięcia z neonówką, zawsze działa i jest bardzo niezawodny.

Test HFE

Niektóre mierniki mają gniazdo, do którego można podłączyć trzy nóżki tranzystora. Miernik wyświetla wtedy wzmacnienie prądowe półprzewodnika. Jednak ta wielkość tranzystora zależy od wielu parametrów. Dokładność wyświetlanej wartości jest zatem wysoce wątpliwa. Funkcja ta jest w rzeczywistości przydatna tylko do porównywania wzmacnienia identycznych tranzystorów.

Funkcja HOLD

Po naciśnięciu tego przycisku na wyświetlaczu pozostanie aktualnie zmierzona wartość.

Funkcja REL

Interesująca funkcja, w której aktualnie zmierzona wartość jest zapisywana w pamięci i odejmowana od wszystkich kolejnych pomiarów. Po naciśnięciu tego przycisku odczyt zostaje wyzerowany. Funkcji tej można użyć na przykład do kompensacji rezystancji przewodów pomiarowych podczas pomiaru bardzo niskich rezystancji.

Funkcja True-RMS

Mierzy rzeczywistą wartość skuteczną napięcia przemiennego, a nie wartość średnią. Rzeczywista wartość skuteczna to wartość, która generuje taką samą ilość mocy cieplnej w rezystorze, jak napięcie stałe o tej samej wartości.

Funkcja MIN/MAX

Dzięki tej funkcji minimalne i maksymalne zmierzone wartości są zapisywane w pamięci i wyświetlane na wyświetlaczu po naciśnięciu przycisku. Funkcja ta pozwala na przykład znaleźć minimalne i maksymalne wartości napięcia sieciowego dla jednego dnia. Następnie należy pozwolić multimetrowi mierzyć napięcie sieciowe przez cały dzień. Funkcja ta jest aktywna tylko podczas bieżącego cyklu pomiarowego.

Funkcja LOG

Multimetr zapisuje zmierzoną wartość w pamięci co kilka sekund lub minut. Zapisane wartości pomiarowe można następnie wyświetlić w formie wykresu na ekranie lub w formie tabeli zawierającej zapisane w i czasie, w których zostały zmierzone.

Funkcja automatycznego wyłączania

Po włączeniu tej funkcji multimetr wyłączy się po określonym czasie, jeśli układ elektroniczny miernika nie wykryje żadnej aktywności.

Funkcja podświetlenia

Wyświetlacz LCD jest podświetlany z tyłu, dzięki czemu można go łatwo odczytać nawet w ciemności.

Funkcja Low-Z

Funkcja ta jest ważna podczas pomiaru napięć przemiennych. W większości przypadków multimetr musi mieć najwyższą możliwą rezystancję wejściową podczas pomiaru napięcia (więcej informacji na ten temat można znaleźć w tym artykule). Jeśli zmierzysz napięcie przemiennie na przewodzie, który nie jest podłączony do niczego o tak wysokiej rezystancji wejściowej, nadal możesz zmierzyć na nim napięcie. Napięcie to jest wytwarzane przez indukcyjne i/lub pojemnościowe połączenia przewodów pod napięciem, które biegają równoległe do przewodu mierzonego. Może to być mylące: „Dlaczego mierzę napięcie na tym przewodzie? Myślałem, że i tak go odłączyłem!”. Niektóre multimetry mają przycisk „Low-Z”. Jego naciśnięcie powoduje zmniejszenie rezystancji wejściowej do znacznie niższej wartości.

Napięcia w punkcie pomiarowym, które wynikają wyłącznie ze sprzężeń indukcyjnych i pojemnościowych, znikają jak śnieg na słońcu. W końcu rezystancja prądu przemiennego tych sprzężeń jest znacznie, znacznie wyższa niż niska rezystancja wejściowa multimetru w funkcji „low-Z”. Jeśli napięcie pozostaje obecne, w punkcie pomiarowym występuje „prawdziwe” napięcie przemiennie. Oczywiście przycisk ten można nacisnąć tylko na krótki czas. Obniżenie rezystancji wejściowej miernika może spowodować rozproszenie dużej ilości energii cieplnej w mierniku.

Dla informacji, te indukowane napięcia nazywane są „napięciami fantomowymi”.

Funkcja wyjścia sygnału

Niektóre multimetry mają wyjście z symetrycznym napięciem kwadratowym. Jego użyteczność jest minimalna, ponieważ ani częstotliwość, ani rozmiar nie są regulowane. Naszym zdaniem jest to całkiem niepotrzebna funkcja multimetru!

Specyfikacje multimetrów cyfrowych

Specyfikacje określają jakość multimetru.

Do porównania multimetrów potrzebna jest pewna liczba dobrze zdefiniowanych parametrów, co do których nie ma wątpliwości, co dokładnie oznaczają. Parametry te są podsumowane w specyfikacjach multimetrów. Jest ich wiele, a najważniejsze z nich to:

- Liczba cyfr,
- Liczba zliczeń,



NO	MODE	VALUE
1	DCV	-00.362mVDC
2	DCV	-00.362mVDC
3	DCV	-00.362mVDC
4	DCV	-00.362mVDC
5	DCV	-00.362mVDC
6	DCV	-00.362mVDC
7	DCV	-00.362mVDC
8	DCV	-00.362mVDC
9	DCV	-00.362mVDC

Funkcja LOG w XDM1041 (© 2023 Jos Verstraten)

- Precyzja,
- Rozdzielczość,
- Dokładność,
- Częstotliwość próbkowania,
- Rezystancja wejściowa R_i ,
- Pojemność wejściowa C_i ,
- Zakres częstotliwości dla pomiarów napięcia AC,
- Współczynnik szczytu dla pomiarów true-RMS,
- Napięcie obciążenia dla pomiarów prądu,
- Współczynnik odrzucenia sygnału wspólnego (*Common Mode Rejection Ratio*).

Liczba cyfr

Liczba cyfr oznacza liczbę cyfr składających się na wyświetlacz. Jeśli skrajna lewa cyfra może reprezentować tylko jedną cyfrę, jest ona nazywana 1/2 cyfry. Licznik z wyświetlaczem do 9999 jest zatem licznikiem czterocyfrowym. Podobne urządzenie z wyświetlaczem do 1999 jest licznikiem 3 1/2-cyfrowym. Istnieją jednak również liczniki, w których lewa cyfra może wskazywać 0, 1, 2 i 3. Taki licznik mierzy do 3999 i jest nazywany licznikiem 3 3/4.

Liczba zliczeń

Określa największą liczbę, jaką może pokazać wyświetlacz. Oczywiście jest, że licznik z 9999 zliczeniami ma więcej opcji niż licznik z 1999 zliczeniami. Za pomocą pierwszego miernika można zmierzyć napięcie do 9,999 V z rozdzielczością 1 mV, za pomocą drugiego można zmierzyć tylko napięcie do 1,999 V. W przypadku niektórych mierników liczba zliczeń zależy od mierzonej wielkości. Na przykład licznik 3 3/4-cyfrowy może mierzyć napięcie do 3999 zliczeń, ale częstotliwość do 9999 zliczeń.



Trzy liczniki ze zliczaniem do 1999, 3999 i 9999 (© 2023 Jos Verstraten)

Precyzja

Precyzja odnosi się do zdolności multimetru cyfrowego do wielokrotnego pomiaru identycznej wielkości, co zawsze skutkuje taką samą wartością pomiarową na wyświetlaczu. Wydaje się to oczywiste, ale tak nie jest. Precyzja multimetru jest związana między innymi z jego długoterminową stabilnością, wilgotnością, temperaturą miernika i wieloma innymi właściwościami jego projektu i konstrukcji. Włącz miernik i zmierz napięcie baterii. Upewnij się, że temperatura w pomieszczeniu pozostaje stała i zmierz napięcie ponownie po godzinie. Załóżysz się, że zmierzysz nieco inną wartość?

Rozdzielczość

Wspomnieliśmy już o słowie rozdzielczość w poprzednim akapicie. Jest to najmniejsza zmiana mierzonej wartości, która może być nadal widoczna. Jeśli mierzysz napięcie za pomocą miernika z 4999 zliczeniami, możesz mierzyć w zakresie pomiarowym 5 V do napięcia 4,999 V z rozdzielczością 1 mV.

Dokładność

Dokładność odnosi się do największego możliwego błędu pomiarowego, który może wystąpić podczas pomiaru. Dokładność jest wyrażana w procentach i wartość ta wskazuje, jak blisko zmierzona wartość jest rzeczywistej wartości wielkości. Ponieważ zmierzona wartość może być niższa lub wyższa od wartości rzeczywistej, dokładność jest zawsze poprzedzona znakiem $\pm$. Miernik o dokładności $\pm 1\%$ zmierzy napięcie o rzeczywistej wartości 1,000 V jako leżące między 0,990 V a 1,010 V. W końcu jeden procent z 1 V to 10 mV. Czasami do wartości procentowej dodawana jest pewna liczba zliczeń. Nazywa się to „błędem statycznym”. Dokładność 4-cyfrowego/2999 zliczeń licznika jest na przykład określona jako $\pm(2\% + 2)$. Jeśli zmierzysz dokładne napięcie 100,00 V za pomocą takiego miernika, wyświetlana wartość powinna wynosić od 97,8 V do 102,2 V.

Dokładność multimetru nie jest identyczna dla wszystkich mierzonych zmiennych. Jako przykład podajemy określone dokładności multimetru FNIRSI S1:

- Napięcie stałe: $\pm(0,8\% + 3)$,
- Napięcie przemienne: $\pm(0,8\% + 3)$,
- Rezystancja: $\pm(1,2\% + 3)$,
- Pojemność: $\pm(4,5\% + 5)$,
- Częstotliwość: $\pm(0,1\% + 3)$,
- Temperatura: $\pm(5\% + 4)$.

Związek między rozdzielczością a dokładnością

Te dwa pojęcia często są przez laików używane zamiennie. To błąd, ale oczywiście istnieje pewien związek między tymi dwoma pojęciami. Załóżmy, że masz miernik o dokładności $\pm 1\%$. Jeśli mierzysz napięcie 100 V, zmierzona wartość może wynosić od 99 V do 101 V. W tym przykładzie pomiar z rozdzielczością 1 mV nie ma sensu. W końcu margines błędu wyniku pomiaru wynosi 2 V! Jedyną rzeczą, dla której wysoka rozdzielczość może być przydatna w tym przypadku, jest porównanie dwóch napięć lub pomiar zmiany wartości napięcia w funkcji czasu lub temperatury. Jednak ta wysoka rozdzielczość nie odgrywa żadnej roli w określaniu bezwzględnej wartości przyłożonego napięcia.

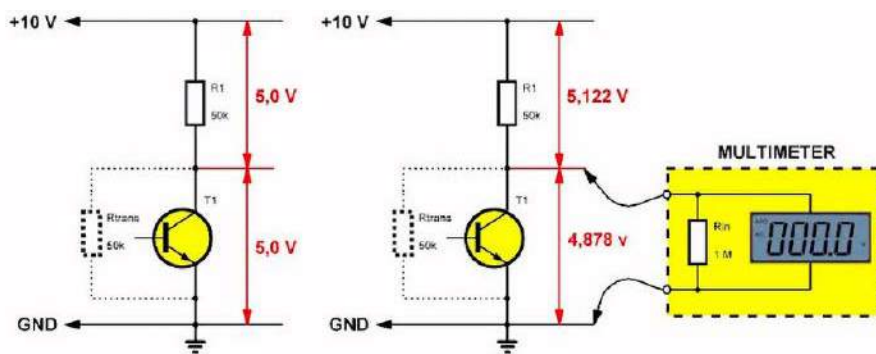
Częstotliwość próbkowania

Miernik cyfrowy musi konwertować zmienną wielkość analogową na wartość liczbową. Robi to poprzez pobranie próbki mierzonej wielkości i przekształcenie jej w wartość liczbową. Ten proces konwersji zajmuje pewien czas. Dlatego multimetr może pobierać tylko ograniczoną liczbę próbek na sekundę. Ta właściwość multimetru nazywana jest „częstotliwością próbkowania”. Większość multimetrów ma częstotliwość próbkowania od 2 do 5 próbek/s.

Rezystancja wejściowa R i dla pomiaru napięcia stałego

Jest to jedna z najważniejszych specyfikacji multimetru! Rezystancja wejściowa w dużej mierze decyduje o tym, czy zmierzone napięcie w danym punkcie jest równe napięciu w tym punkcie bez podłączonego miernika. Poniższy rysunek przedstawia praktyczną sytuację pomiarową po lewej stronie. Kolektor tranzystora T1 ma napięcie 5,0 V. Obwód jest zasilany napięciem 10 V, a rezystancja kolektora wynosi 50 k Ω . Z tych danych wynika, że rezystancja przewodzącego tranzystora również wynosi 50 k Ω . Jest to zaznaczone kropką. To samo napięcie spada na oba identyczne rezystory, tj. połowa napięcia zasilania. Teraz zmierzysz napięcie kolektora za pomocą multimetru, który ma rezystancję wewnętrzną 1 M Ω . Sytuacja ta została przedstawiona po prawej stronie. Rezystancja 1 M Ω miernika jest teraz umieszczona równolegle do rezystancji 50 k Ω tranzystora. Rezystancja zastępcza tych dwóch równolegle połączonych rezystorów wynosi 47,62 k Ω . Dzielnik napięcia między zasilaniem 10 V a masą wygląda teraz inaczej, górna rezystancja nadal wynosi 50 k Ω , a dolna 47,62 k Ω . Oznacza to, że rozkład napięcia również ulega zmianie. 5,122 V spada na górny rezystor, a 4,878 V spada na dolny rezystor. W wyniku podłączenia multimetru do kolektora napięcie na nim spada z 5,0 V do 4,878 V. Nie ma więc sensu projektować multimetru o rezystancji wejściowej 1 M Ω z dokładnością $\pm 0,1\%$. Dokładność pomiaru jest całkowicie stracona z powodu niskiej rezystancji wejściowej. Wniosek z tej historii jest taki, że multimetr o wysokiej rozdzielczości i wysokiej dokładności musi mieć jak największą rezystancję wejściową. Większość tanih multimetrów ma rezystancję wejściową 10 M Ω . W przedstawionym przykładzie zmierzone napięcie kolektora wyniosłoby 4,987 V. Duża poprawa!

Profesjonalne multimetry, takie jak nasz Fluke 8842A z 199999 zliczeniami i dokładnością $\pm 0,01\%$, mają jeszcze wyższą rezystancję wejściową, a mianowicie 10000 M Ω ! Przy niższej wartości tak wysoka rozdzielczość i dokładność stałyby się bez znaczenia.



Wpływ rezystancji wejściowej na pomiar (© 2023 Jos Verstraten)

Pojemność wejściowa Ci podczas pomiaru napięcia przemiennego

Podczas pomiaru napięcia przemiennego należy wziąć pod uwagę, że między wejściem multimetru a „COM” występuje pewna pojemność pasywna Ci. Ma ona impedancję, rezystancję prądu przemiennego, która zależy od częstotliwości. Jest ona równoległa do R i miernika i dlatego zapewnia, że rezystancja wejściowa podczas pomiaru napięcia przemiennego jest znacznie niższa niż podczas pomiaru napięcia stałego. Aby zminimalizować zależny od częstotliwości wpływ tej pojemności, zwykle wybiera się zmniejszenie rezystancji wejściowej dla napięcia przemiennego do 1 M Ω . Spadek napięcia przemiennego, który występuje w danym punkcie podczas pomiaru, jest zatem znacznie większy niż w przypadku pomiaru napięcia stałego.

Zakres częstotliwości pomiaru napięcia przemiennego

Większość multimetrów radzi sobie bardzo słabo w tym obszarze. Dokładny pomiar napięcia przemiennego zatrzymuje się średnio przy 10 kHz. Niektóre tanie mierniki radzą sobie jeszcze gorzej i zawodzą już przy kilku kHz!

Współczynnik szczytu dla pomiarów true-RMS

Większość nowoczesnych multimetrów mierzy rzeczywistą wartość skuteczną napięcia przemiennego. Dla napięcia sinusoidalnego wartość skuteczna jest łatwa do obliczenia. Stosunek wartości maksymalnej do wartości skutecznej wynosi 1,41. Stosunek ten jest bardzo różny dla innych form sygnału. Jest on nazywany współczynnikiem szczytu napięcia przemiennego. Multimetry nie są w stanie obliczyć napięcia skutecznego dowolnej formy sygnału. Większość mierników jest w stanie dobrze wykonać to obliczenie do współczynnika szczytu wynoszącego około 3. W przypadku napięć przemiennych o większym współczynniku szczytu wyświetlanie wartości skutecznej staje się mało wiarygodne!

Napięcie obciążenia w pomiarach prądu

Multimetry mierzą prąd, przesyłając go przez dokładną rezystancję wewnętrzną i mierząc spadek napięcia na tej rezystancji. Korzystając z prawa Ohma, procesor w mierniku może obliczyć wartość prądu i pokazać ją na wyświetlaczu. Napięcie, które odkłada się na tym rezystorze podczas pomiaru prądu w pełnej skali, nazywane jest „napięciem obciążenia” multimetru. Oczywiście jest, że pojawienie się tego napięcia w obwodzie, w którym dokonywany jest pomiar, spowoduje zakłócenie. Dlatego napięcie obciążenia musi być jak najmniejsze.

Współczynnik CMRR (Common Mode Rejection Ratio) multimetru

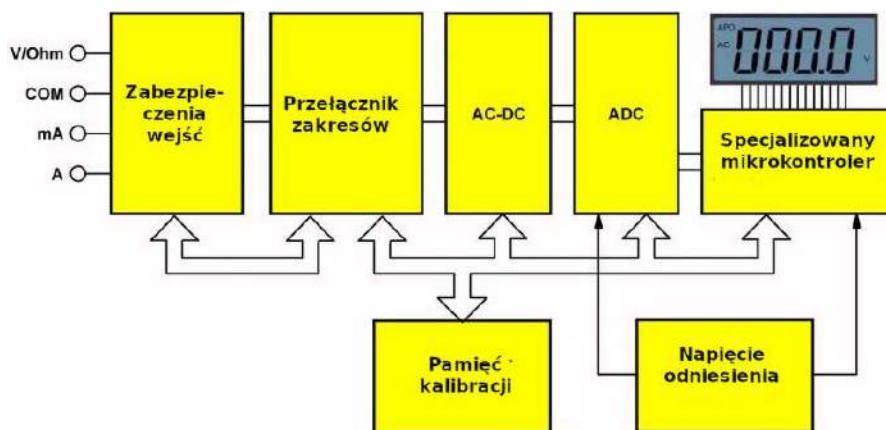
Nawet podczas pomiaru napięcia stałego miernik odbiera zmienne sygnały zakłócające w kształcie napięcia. Mogą to być zakłócenia, które docierają do zacisków wejściowych miernika na różne sposoby, na przykład przez przewody pomiarowe. Pola elektromagnetyczne pochodzące od napięcia sieciowego i jego harmonicznych są również obecne wszędzie i trafiają na zaciski wejściowe podczas pomiaru. Te sygnały zakłócające mogą wpływać na wartość próbek pobieranych przez miernik. Z tego powodu zakłócającego próbka może być nieco większa lub mniejsza niż poprzednia próbka. Elektronika w mierniku mierzy nieco większe lub mniejsze napięcie, w wyniku czego prawa cyfra odczytu nie jest stabilna, ale będzie przeskakować między kilkoma cyframi. Nazywa się to „jitterem” odczytu. Te sygnały zakłócające muszą być zatem jak najlepiej tłumione. Aby to zdefiniować, wprowadzono pojęcie „Współczynnika Odrzucania Trybu Wspólnego”, w skrócie CMRR. Parametr ten wskazuje, o ile dB tłumione są sygnały zakłócające. Wartości 120 dB dla dobrego multimetru nie są wyjątkiem!

Jak działają multimetry cyfrowe

Schemat blokowy

Poniższy rysunek przedstawia ogólny schemat blokowy każdego multimetru. Bardzo ważnym blokiem jest „Zabezpieczenia wejść”. To on gwarantuje, że ani miernik, ani operator nie zostaną uszkodzeni, jeśli coś zostanie ustawione nieprawidłowo. Przykładem może być wybranie opcji „pomiar prądu”, a następnie omyłkowe podłączenie napięcia sieciowego do miernika. Blok ten zapewnia również, że nic w elektronice nie zostanie uszkodzone podczas automatycznego przełączania zakresów. Blok „Przełącznik zakresów” może być ręczny, półautomatyczny lub automatyczny. Blok „AC/DC” przekształca napięcia i prądy przemienne w napięcia stałe, które mogą być próbkowane przez kolejny blok: „Przetwornik ADC”. Elektronika jest kontrolowana przez specjalny mikrokontroler opracowany do tego celu. Niektórzy chińscy producenci chipów zaprojektowali do tego celu układy, które można znaleźć w dziesiątkach tanich multimetrów i są oferowane bardzo tanio ze względu na ich masowe zastosowanie.

Dwa z bloków są obecne tylko w droższych wersjach. „Napięcie odniesienia” zapewnia niezwykle stabilne i dokładne napięcie stałe, które stanowi podstawę wszystkich pomiarów. W tańszych wersjach występuje specjalna dioda Zenera w mikrokontrolerze, która to robi, ale z mniejszą stabilnością i dokładnością. Blok „Pamięć kalibracji” przechowuje dane używane do fabrycznej kalibracji poszczególnych multimetrów. W końcu kalibracja poprzez obracanie potencjometrów regulacyjnych to już przeszłość! Dla każdego zakresu pomiarowego w tej pamięci przechowywany jest współczynnik korekcji, który kompensuje niedokładności rezystancji częściowych w „Przełączniku zakresów”.



Schemat blokowy multimetru cyfrowego (© 2023 Jos Verstraten)

Pomiar napięcia stałego

W przypadku multimetrów z obsługą ręczną poniższy schemat służy do pomiaru napięcia stałego. Rezystancyjny dzielnik napięcia zapewnia współczynniki podziału 1/1, 9/1, 99/1 i 999/1. Tworzy to cztery zakresy pomiarowe 1,999 V, 19,99 V, 199,9 V i 1,999 V. Każde napięcie wejściowe jest redukowane przez dzielnik napięcia do napięcia gdzieś pomiędzy 0,000 V a $\pm 1,999$ V. To właśnie napięcie jest używane przez ADC. jest próbkowane i konwertowane na kody cyfrowe, które sterują wyświetlaczem. Ponieważ wszelkiego rodzaju przepisy zabraniają przedstawiania takiego miernika jako odpowiedniego do pomiaru 2 kV, ostatnia pozycja jest zwykle nazywana „600 V”. Działanie jest łatwe do wyjaśnienia. Załóżmy, że chcesz zmierzyć napięcie 100 V. Następnie ustaw przełącznik S1 w pozycji „199,9 V”. Tworzy to dzielnik napięcia z R1+R2 na górze i R3+R4 na dole. Innymi słowy, 9900 k Ω i 100 k Ω . Górna rezystancja jest zatem 99 razy większa niż dolna rezystancja i występuje na niej 99 razy większe napięcie. Ze 100 V na wejściu, 99 V spada na R1+2 i tylko 1 V spada na R3+R4. To 1 V pojawia się na wyświetlaczu jako 100,0 V. Jest to kwestia ustawienia prawidłowej kropki dziesiętnej.

Bardzo dokładne rezystory, o tolerancji $\pm 0,1\%$ lub nawet $\pm 0,01\%$, są standardem i często znajdują się w specjalnych tablicach rezystorów. Na poniższym zdjęciu widać taką część z wyprowadzeniami przewlekаныmi przez płytkę PCB (THT), ale istnieją one również jako komponenty SMD.

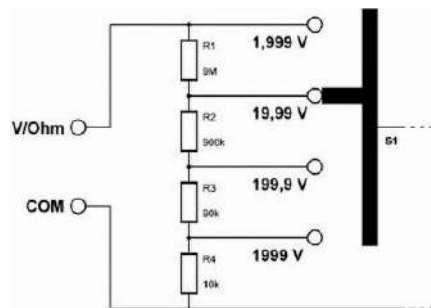
Pomiar prądu stałego

Większość multimetrów ma sześć zakresów pomiaru prądu:

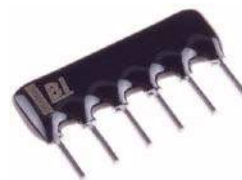
- 2 μ A
- 3 mA
- 1 A

Podstawowy schemat przedstawiono na poniższym rysunku. Aby zminimalizować użycie drogich precyzyjnych rezystorów o niskiej impedancji, często stosuje się stopień wzmacniający $\times 10$. Przełączniki S1 i S2 służą do przełączania między zakresami niskimi i wysokimi w zakresach pomiarowych μ A i mA. W ten sprytny sposób potrzebne są tylko trzy rezystory. Mierzony prąd przepływa przez rezystancję 0,01 Ω , 1 Ω lub 100 Ω i, zgodnie z prawem Ohma, generuje napięcia, które są numerycznie identyczne z prądem i dlatego mogą być mierzone za pomocą przetwornika ADC.

Rezystory te nazywane są „rezystorami bocznikowymi”. Aby mierzyć duże prądy, należy stosować bardzo niskie boczniki. Wartości 0,1 Ω , a nawet 0,01 Ω nie są wyjątkiem. Nie jest łatwo dokładnie zmierzyć spadek napięcia na takim rezystorze. Wymaga to dużego doświadczenia



Zasada pomiaru napięcia stałego (© 2023 Jos Verstraten)



Tablica z bardzo precyzyjnymi rezystorami (© 2023 Jos Verstraten)

w projektowaniu układu PCB! Takie rezystory wyglądają również nieco nietypowo. Poniższe zdjęcie pokazuje przykłady takich komponentów od chińskiej firmy, która specjalizuje się w produkcji tak dokładnych rezystorów o niskiej rezystancji.

Pomiar napięć przemiennych

Podczas pomiaru tej wielkości ważną rolę odgrywają pojemności pasożytnicze obecne między ścieżkami płytki drukowanej, gniazdami połączeniowymi, stykami przełączników itp.

Pojemności te nie są duże, ale należy pamiętać, że pojemność 100 pF przy częstotliwości 10 kHz ma już rezystancję prądu przemiennego około 160 kΩ. Dzielnik napięcia z górną rezystancją 9 MΩ będzie zatem dość mylony przez rezystancję prądu przemiennego pojemności pasożytniczych. Przy częstotliwościach w zakresie kHz niewiele pozostaje z pięknego podziału napięcia 999/1, 99/1 itd. Rozwiązaniem tego problemu jest drugi dzielnik napięcia, ale pojemnościowy. Na każdym rezystorze opisanego już dzielnika napięcia stałego znajduje się kondensator. Wartości tych kondensatorów zostały dobrane w taki sposób, aby miały identyczny współczynnik podziału dla napięcia przemiennego, jak rezystory dla napięcia stałego. Najmniejsze kondensatory są umieszczane w poprzek najwyższych rezystorów. Są one często projektowane w formie trymerów. Zmodyfikowany schemat został przedstawiony na poniższym rysunku.

Pomiar prądów przemiennych

Boczniki do pomiaru prądów stałych mogą być również używane do pomiaru prądów przemiennych bez kompensacji pojemnościowej. Rezystancje te są tak małe, że pasożytnicze pojemności mają znikomy wpływ.

AC czy DC+AC?

Droższy typ multimetru zawiera również kondensator szeregowy C5, patrz schemat powyżej. Odcina on napięcie stałe obecne w napięciu przemiennym. To musi być duży kondensator wysokiego napięcia. Są one drogie i muszą być włączane lub wyłączane z obwodu za pomocą równie drogiego przełącznika. Na próżno szukać tego w niskobudżetowych miernikach. Nie mierzą one czystego napięcia przemiennego, lecz „DC+AC”. W rezultacie za pomocą takiego miernika bez kondensatora separującego nie można na przykład zmierzyć tętnienia napięcia przemiennego na kondensatorze wyglądającym.

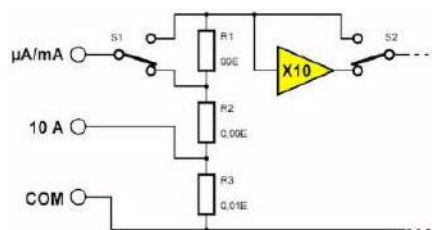
Prostownik jest zawsze niezbędny do pomiaru napięć i prądów przemiennych. W tanich miernikach stosowany jest prosty prostownik diodowy, który przenosi średnią wartość DC napięcia AC na kondensator i dostarcza ją do przetwornika ADC. Taki prosty prostownik można skalibrować w wartościach skutecznych dla czystego sinusoidalnego napięcia przemiennego. Jednak w przypadku pomiaru napięć przemiennych o bardzo różnych kształtach, takich jak przebiegi prostokątne, spłaszczone sinusoidy lub przebiegi sinusoidalne ucięte przed przejściem przez zero (przez fazowy regulator mocy), nie można już polegać na dokładności pomiaru. Dlatego obecnie wiele multimetrów jest reklamowanych hasłem „true-RMS”. Oznacza to, że w układzie elektronicznym znajduje się obwód, który oblicza rzeczywistą wartość skuteczną każdego napięcia przemiennego przyłożonego do miernika. W lepszych multimetrach stosowane są do tego układy scalone firmy Analog Devices, takie jak AD736 lub AD5361. Istnieje jednak ograniczenie dokładności takich układów. Jest ona określana przez współczynnik szczytu napięcia przemiennego. Każdy dobry miernik mierzący prawdziwą wartość skuteczną ma w specyfikacji tabelę, z której można wywnioskować, jaki dodatkowy błąd pomiaru powoduje współczynnik szczytu odbiegający od 1,414. 1,414 to współczynnik szczytu czysto sinusoidalnego napięcia przemiennego. Nowoczesne multimetry sterowane mikrokontrolerem wykorzystują rozwiązania cyfrowe. Sygnał napięcia przemiennego jest próbkowany z najwyższą możliwą częstotliwością, aby jak najdokładniej uchwycić kształt fali. Wartość RMS jest następnie obliczana przy użyciu pierwiastka kwadratowego ze średniej wartości kwadratów poszczególnych pomiarów. Należy jednak wziąć pod uwagę, że istnieją również ograniczenia współczynnika szczytu, przy którym taki multimetr nadal daje wiarygodne wskazania.

Pomiar rezystancji za pomocą dwóch przewodów pomiarowych

Ta metoda jest stosowana w prawie wszystkich multimetrach ręcznych. Rezystancję mierzy się, podłączając ją między gniazdami „COM” i „V/Ohm”. Stosowana metoda pomiaru nazywana jest „pomiaru stosunku”. Mierzona rezystancja jest podłączona szeregowo z dokładnie znaną rezystancją do wewnętrznego napięcia odniesienia U_{ref} . Przetwornik ADC mierzy spadek napięcia na nieznanej rezystancji. Ponieważ napięcie odniesienia i rezystancja szeregową są znane, wartość rezystancji można wyprowadzić z pomiaru spadku napięcia. Dla każdego zakresu pomiarowego, inny precyzyjny rezystor jest włączany szeregowo z mierzoną rezystancją. Schemat działania przedstawiono na poniższym rysunku.

Pomiar rezystancji za pomocą czterech przewodów pomiarowych

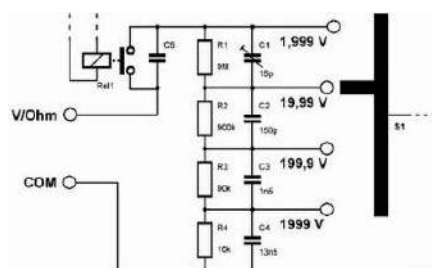
Wadą metody dwuprzewodowej jest to, że mierzona jest również rezystancja dwóch przewodów łączących. W przypadku multimetrów z funkcją „REL” można to oczywiście



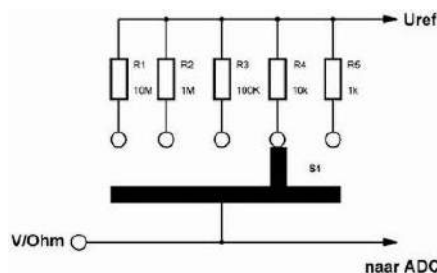
Podstawowy obwód do pomiaru prądów stałych
(© 2023 Jos Verstraten)



Specjalne rezystory bocznikowe o bardzo niskiej impedancji
(© 2023 Jos Verstraten)

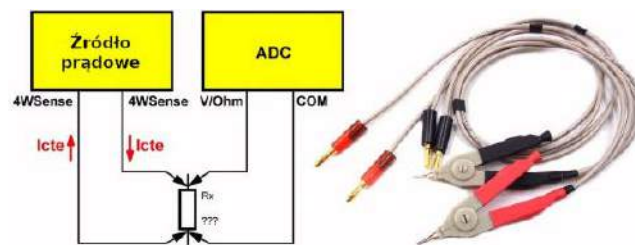


Pojemnościowy dzielnik napięcia do pomiaru napięć przemiennych
(© 2023 Jos Verstraten)



Pomiar rezystancji za pomocą „pomiaru stosunku”
(© 2023 Jos Verstraten)

skompensować, ale nie jest to idealne rozwiązanie. Prawie wszystkie multimetry stołowe wykorzystują czteroprzewodową metodę pomiaru rezystancji. Jest ona również nazywana „metodą Kelvina”. Zasada jest dość prosta. Bardzo dokładny stały prąd jest przesyłany do mierzonej rezystancji za pośrednictwem dwóch przewodów pomiarowych podłączonych do dwóch gniazd „4WSense”. Spadek napięcia, który generuje prąd na rezystorze, jest mierzony za pomocą dwóch innych przewodów pomiarowych podłączonych do gniazd „COM” i „V/Ohm”. Jeśli cztery zaciski czterech przewodów zostaną umieszczone jak najbliżej rezystancji, nawet najmniejsza rezystancja zostanie dokładnie zmierzona. Na rynku dostępne są specjalne przewody pomiarowe do tego pomiaru, tak zwane zaciski pomiarowe Kelvina. Składają się one z dwóch szczypek z połączanymi szczykami. Każda szczyka łączy jeden przewód źródła prądu z rezystorem i odprowadza jeden biegun generowanego napięcia.



Pomiar rezystancji „metodą Kelvina” (© 2023 Jos Verstraten)

Pomiar kondensatorów

Wszystkie multimetry mierzą kondensatory w identyczny sposób. Zasadę przedstawiono na poniższym rysunku. Mierzony kondensator C x jest podłączony do elektronicznego przełącznika S1. W pokazanej pozycji 1 kondensator jest rozładowywany przez rezystor R1 prądem rozładowania I aż do całkowitego rozładowania. Nieco później miernik przełącza przełącznik S1 do pozycji 2, a mierzony kondensator jest teraz ładowany stałym prądem ładowania I, dostarczającym przez bardzo stabilne źródło prądu. Ładowanie trwa do momentu, gdy napięcie na kondensatorze wzrośnie do określonej wartości U ref. W tym momencie obwód pomiarowy przywraca przełącznik elektroniczny do pozycji 1. Czas wymagany do naładowania kondensatora od 0 V do wartości odniesienia U ref jest wprost proporcjonalny do pojemności elementu. W końcu z teorii elektryczności wiadomo, że: $Q=U \cdot C$.

Ładunek w kondensatorze jest równy pojemności kondensatora pomnożonej przez napięcie na nim.

Ale także: $Q=I \cdot t$

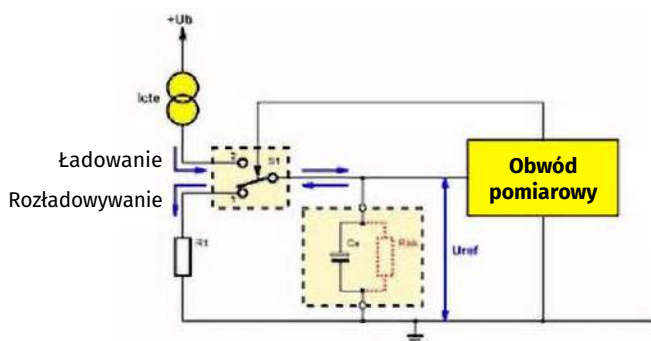
Ładunek w kondensatorze jest równy prądowi, który płynie w kondensatorze, pomnożonemu przez czas, w którym ten prąd płynie.

Zatem: $U \cdot C = I \cdot t$

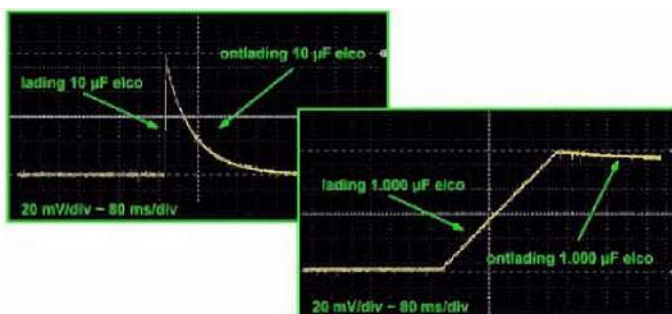
lub: $C = [I \cdot t] / U = [I / U] \cdot t = cte \cdot t$

W omawianym obwodzie prąd ładowania I i napięcie odniesienia U ref są stałe. Istnieje zatem wprost proporcjonalna zależność między wartością C kondensatora a czasem ładowania t wymaganym do naładowania całkowicie rozładowanego kondensatora stałym prądem do napięcia odniesienia. Mierzając ten czas t, procesor multimetru może obliczyć wartość kondensatora i pokazać ją na wyświetlaczu.

System ten działa doskonale i bardzo dokładnie podczas pomiaru kondensatorów nieelektrolitycznych. Mają one całkowicie pomijalny prąd upływu. W przypadku kondensatorów elektrolitycznych należy przedstawić ten duży prąd upływu za pomocą pasywnego rezystora upływu R, który jest równoległy do kondensatora C x (patrz schemat powyżej). Oczywiście przez ten rezystor również będzie płynął prąd – prąd upływu kondensatora. Ponieważ źródło zasilania dostarcza stały prąd, ten prąd upływu jest odejmowany od prądu płynącego przez idealny kondensator C x. W praktyce czas ładowania kondensatora zależy więc od wielkości prądu upływu. Jeśli zmierzmy dwa kondensatory elektrolityczne o identycznych pojemnościach, ale o różnych prądach upływu, multimetr cyfrowy wyświetli na ekranie dwie bardzo różne wartości pojemności.



Pomiar pojemności kondensatorów (© 2023 Jos Verstraten)



Pomiar temperatury

Wszystkie multimetry mierzące temperaturę działają zgodnie z zasadą pomiaru termoparą. Termopara składa się z dwóch przewodów wykonanych z różnych stopów metali, które stykają się ze sobą w jednym punkcie. Zwykle stosuje się termoparę typu K, która składa się z drutu nichromowego (stop niklu (Ni) i chromu (Cr)) i drutu alumelowego (stop niklu (Ni) z aluminium (Al), manganem (Mn) oraz krzemem (Si)). Na poniższym zdjęciu widać taką termoparę dostarczoną z multimetrem.

Taka sonda wytwarza bardzo niskie napięcie termoelektryczne, a mianowicie 40,4 µV/°C, potrzebne jest zatem duże wzmocnienie! To nie jest jedyny problem. Ponieważ większość multimetrów ma tylko gniazda 4 mm, termopara musi być podłączona do miernika za pomocą wtyków bananowych 4 mm. Tworzy to szereg tak zwanych „zimnych złączy”. Chromowany przewód jest wkręcany w metal wtyczki bananowej, a wtyczka bananowa znajduje się w gnieździe multimetru. Są to również punkty, w których dwa różne metale

stykają się ze sobą i powstają napięcia termoelektryczne. Są one dodawane lub odejmowane od napięcia termoelektrycznego dostarczanego przez „właściwą” termoparę i powodują duży błąd pomiaru.

Aby dokładnie mierzyć temperaturę za pomocą termopar, należy wykonać „kompensację zimnego końca”. Mierzona jest temperatura otoczenia, a napięcia termiczne wytwarzane przez zimne złącza w tej temperaturze są kompensowane w mierniku.

Dobre multimetry posiadają czujnik temperatury pomiędzy gniazdami 4 mm „COM” i „V/Ohm”, który mierzy wewnętrzną temperaturę miernika. Napięcie wyjściowe tego czujnika jest wykorzystywane przez procesor do kompensacji napięć zimnych złączy. Czasami czujnik ten znajduje się w specjalnym układzie procesora zaprojektowanym dla danego modelu multimetru. Na poniższym zdjęciu, pobranym z lygte-info, widać taki czujnik temperatury w pobliżu gniazd wejściowych na płycie drukowanej drogiego multimetru Fluke.

Jak działają mierniki z automatycznym zakresem pomiarowym

Zasadniczo multimetry z automatycznym zakresem pomiarowym działają podobnie do multimetrów ręcznych. Zakresy są wybierane poprzez wybranie odpowiedniego węzła dzielnika rezystancyjnego i odłączenie go do przetwornika ADC. Obecnie jednak wybór węzła odbywa się za pomocą przełączników kontakttronowych i/lub przełączników elektronicznych. Na przykład poniższy schemat pokazuje, w jaki sposób multimetr z automatycznym zakresem wybiera właściwy zakres napięcia. Układy scalone z rodziny CMOS 4000 są zwykle używane jako przełączniki elektroniczne. Rezystancja ON tych przełączników jest znikoma, jeśli porównać ją z wartością rezystorów na schemacie. Przełącznik kontakttronowy jest używany tylko do przełączania na najniższy zakres pomiarowy.

Bezpieczeństwo multimetru

Wprowadzenie

Dobry multimetr jest zaprojektowany tak, aby był „odporny na błędy”.

Oznacza to, że miernik przetrwa każde nieprawidłowe ustawienie i co najwyżej wewnętrzny bezpiecznik ulegnie awarii. Jednak w przypadku multimetrów cyfrowych, które kosztują kilkanaście złotych, należy założyć, że tak nie jest i że należy dokładnie przemyśleć sposób ustawienia miernika przed wykonaniem pomiaru.

Ochrona zakresów prądowych za pomocą bezpieczników

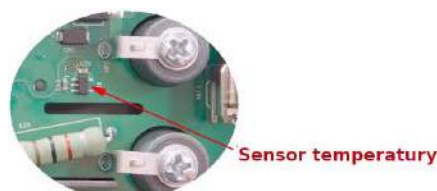
Wszystkie multimetry posiadają jeden lub dwa bezpieczniki, które chronią zakresy pomiarowe prądu przed przeciążeniem. W tanich miernikach stosowane są dobrze znane bezpieczniki szklane o wymiarach 5,0 mm × 20,0 mm. Wysokiej klasy multimetry cyfrowe są wyposażone w specjalne bezpieczniki o dużej pojemności energetycznej, które są zaprojektowane do pochłaniania bardzo dużej mocy, która może być generowana, jeśli na przykład napięcie sieciowe zostanie podłączone w pozycji „10,00 A”. Są to tak zwane „bezpieczniki HRC”, skrót od „High Rupturing Capacity”. Bezpieczniki te kosztują kilkadziesiąt złotych, a pokusą jest zastąpienie ich tańszymi bezpiecznikami o identycznych wymiarach. Nigdy tego nie rób! Zwiększasz ryzyko eksplozji łuku elektrycznego, jeśli miernik, przełączony na zakres prądu, zostanie przypadkowo podłączony do napięcia sieciowego. Na zdjęciu poniżej widać taki bezpiecznik. To, co od razu rzuca się w oczy, to fakt, że takie bezpieczniki są znacznie większe niż znane nam małe bezpieczniki o wymiarach 5,0 mm × 20,0 mm. Standardowe wymiary to 10 mm × 35 mm lub nawet 10 mm na 38 mm. Wynika to nie tylko z ich większej pojemności energetycznej, ale także z faktu, że bezpieczniki te są przeznaczone do wyższych napięć niż 250 V standardowych bezpieczników szklanych.

Ochrona zakresów napięcia

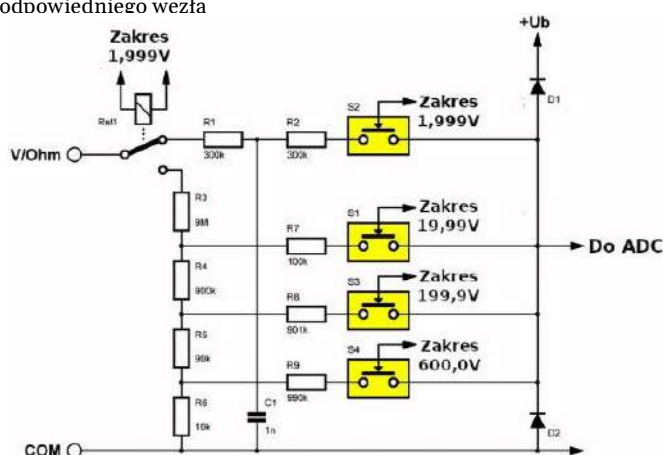
Jeśli ustawisz multimetr na 199,9 mV i podłączysz go do napięcia sieciowego, teoretycznie nic złego nie powinno się wydarzyć. Aby być tego pewnym, opracowano szereg systemów zabezpieczeń, które czasami są używane osobno, a czasami razem. Wszystkie te zabezpieczenia



Typowy przykład termopary dostarczanej z multimetrem (© 2023 Jos Verstraten)



Czujnik kompensacji zimnego złącza (© lygte-info, pod redakcją Josa Verstratena)



Obwód automatycznego przełączania dla napięcia stałego (© 2023 Jos Verstraten)



Bezpiecznik o wysokiej zdolności rozerwania (© 2023 Jos Verstraten)

mają za zadanie ograniczyć napięcie docierające do wewnętrznej elektroniki do bezpiecznej wartości. Warunkiem jest, aby te dodatkowe obwody nie wpływały na dokładność pomiarów. Zabezpieczenia można znaleźć poniżej:

- Tranzystory używane jako diody Zenera,
- MOV,
- PTC.

Ochrona za pomocą tranzystorów używanych jako diody Zenera

Jeśli kolektor krzemowego tranzystora NPN będzie ujemny w stosunku do emitera, część ta, przy określonym napięciu, będzie działać jako dioda Zenera o bardzo niskim prądzie upływu. Napięcie Zenera wynosi około 10 V. Zabezpieczenie wygląda wtedy tak, jak naszkicowano na poniższym rysunku.

Bezpieczeństwo dzięki MOV i PTC

Inną metodę ochrony elektroniki miernika przed nadmiernym napięciem przedstawiono na poniższym rysunku. Dzielnik napięcia jest utworzony z PTC i kilku MOV. PTC to termistor o dodatnim współczynniku temperaturowym. W temperaturze pokojowej taki element ma bardzo niską rezystancję. Powyżej pewnej temperatury progowej rezystancja nagle gwałtownie wzrasta, czasami od 10 Ω do 100 k Ω . MOV to „warystor z tlenku metalu” który ma bardzo wysoką rezystancję przy niewielkim napięciu na komponencie. Wraz ze wzrostem napięcia rezystancja spada do niskiej wartości. Obie części są podłączone między wejściami multimetru, jak pokazano na poniższym rysunku. Jeśli do wejścia miernika zostanie przyłożone zbyt wysokie napięcie, duży prąd, który popłynie, spowoduje nagrzanie PTC, powodując ogromny wzrost jego wartości. Rezystancja MOV zmniejsza się. Oba te działania skutecznie chronią elektroniczne podzespoły multimetru przed niszczącymi przepięciami.

Ochrona czułych zakresów prądowych za pomocą TVS

Wreszcie, schemat często spotykany w tanich multimetrach do ochrony zakresów prądowych. Po bezpieczniku, TVS jest podłączony między wejściem zasilania a „COM”. TVS to akronim od „Transient Voltage Suppressor”. Taki element zachowuje się jak dwie diody Zenera połączone przeciwobnie. Jeśli przypadkowo przyłożysz wysokie napięcie między wejściem „mA/ μ A” a wejściem „COM”, TVS ulegnie uszkodzeniu i utworzy zwarcie. Bezpiecznik F1 natychmiast się przepali.

Połączenie z komputerem

Rejestrowanie danych za pomocą komputera

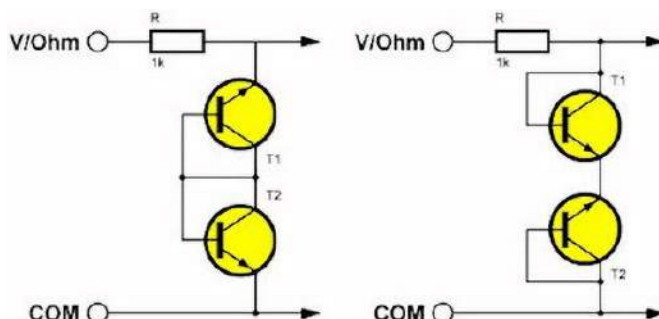
Wiele multimetrów jest oferowanych ze złączem USB, które umożliwia podłączenie miernika do portu USB w komputerze. Następnie można obsługiwać miernik za pomocą myszy, korzystając z dostarczonego oprogramowania. Nie wydaje nam się to zbyt interesujące. Bardziej przydatna jest opcja zapisywania wyników pomiarów miernika w pliku na dysku twardym komputera, dzięki czemu można również używać multimetru jako rejestratora danych.



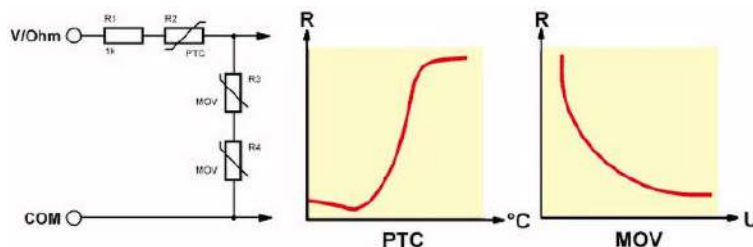
Brak standardowego interfejsu

Każdy producent multimetrów oferuje własny

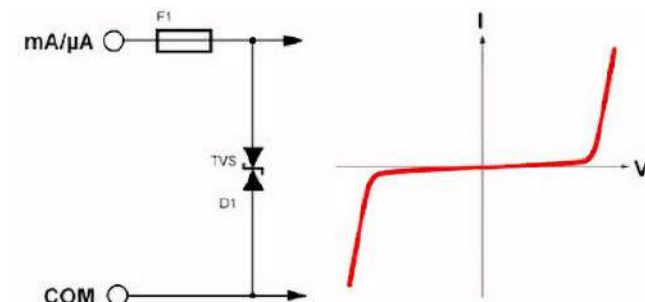
interfejs do ekranu komputera. Niektóre z nich są niezwykle prymitywne i ledwo użyteczne, inne są bardzo dobrze przemyślane i stanowią przydatne rozszerzenie możliwości multimetru. ■



Zabezpieczenie z użyciem tranzystorów (© 2023 Jos Verstraten)



Zabezpieczenie z użyciem MOV i PTC (© 2023 Jos Verstraten)



Zabezpieczenie z użyciem TVS (© 2023 Jos Verstraten)



Dwa z wielu interfejsów PC dostarczanych z multimetrem (© 2023 Jos Verstraten)

Konsultacje do wykładu „Multimetry”

P. Który multimetr jest najlepszy: ręczny, półautomatyczny czy automatyczny?

O. W teorii wszystkie trzy typy oferują podobne możliwości w obrębie jednej klasy pomiarowej. Częściowa czy pełna automatyka nie ma znaczącego wpływu na dokładność pomiarów. W praktyce jednak im bardziej multimetr jest zautomatyzowany, tym mniej wymaga uwagi ze strony użytkownika. Z drugiej strony użycie przełącznika trybów (i opcjonalnie zakresów) upraszcza projekt i produkcję multimetru, a przy tym zapewnia nieco lepszą ochronę instrumentu w razie awarii lub przeciążenia. Dodatkowo multimetry ręczne i półautomatyczne mają mechaniczny wyłącznik zasilania, dzięki czemu nawet po długim „leżakowaniu” w szufladzie multimetr powinien być gotowy do pracy.

P. W Internecie niektórzy mocno zachwalają multimetry analogowe, jako lepsze instrumenty od cyfrowych. Czy to prawda?

O. Absolutnie nie. Zwłaszcza w czasach, gdy multimetr cyfrowy 4,5 cyfry kosztuje mniej niż 200 złotych. Multimetry analogowe wymagają ciągłych regulacji, impedancja wejściowa na większości zakresów napięciowych jest dużo niższa, pomiar rezystancji jest w skali logarytmicznej, a sam odczyt wymaga dobrego wzroku. Przy tym nawet tani multimetr cyfrowy wypada świetnie. Nie mogę jednak zaprzeczyć, iż wskaźniki wychyłowe mają swój niepowtarzalny urok.

P. Czy warto mieć multimetr cęgowy?

O. Tak, jeśli w planach są częste pomiary prądów powyżej 1–2 A, praca z instalacjami samochodowymi i instalacjami napięcia sieciowego. Hobbysta rzadko kiedy potrzebuje mierzyć duże prądy, więc multimetr cęgowy raczej mu się nie przyda.

P. Dlaczego w wykładzie twierdzi się, że zwykłym multimetrem należy mierzyć prądy tylko do 1–2 A? Przecież większość ma zakres do 10 A, a niektóre aż do 20 A.

O. Zrób prosty eksperyment: ustaw na zasilaczu warsztatowym lub laboratoryjnym napięcie 10 V i prąd 2 A. Podłącz swój multimetr w trybie pomiaru prądu w amperach. Trzymając multimetr w ręku obserwuj, jak wartość na nim powoli się zmienia, a sam miernik staje się cieplejszy. Czy możesz zgadnąć, co tu się dzieje?

Multimetry mierzą prąd za pomocą rezystora bocznikowego o małej wartości. Wartość tego rezystora musi być na tyle duża, by miernik bez problemu mógł zmierzyć spadek napięcia na nim występujący zgodnie z prawem Ohma. Jednocześnie na tym rezystorze wydzieli się pewna moc w postaci ciepła. Ciepło to nie tylko rozgrzeje sam rezystor zmieniając nieznacznie jego parametry, ale też rozgrzeje całe wnętrze multimetru, łącznie ze źródłem napięcia odniesienia dla przetwornika, wzmacniaczami napięcia i układem ADC. Na parametry pracy tych elementów wpływ ma właśnie temperatura. Multimetr zacznie zakłamywać wynik. Jak dużo mocy wydzieli się na boczniku multimetru możemy obliczyć ze wzoru:

$$P=RI^2$$

Zmierzyłem opór bocznika w tanim multimetrze z owadziego dyskontu. Wyniósł on ponad 0,22 Ω . A to oznacza, że przy pomiarze prądu 5 A wydzieli się na nim 5,5 W (sic!). Ciekawe, po jakim czasie bocznik się odlutuje od płytki.

P. Nie mam multimetru cęgowego, ale muszę zmierzyć duży prąd. Czy mogę to jakoś zrobić?

O. Ależ oczywiście, metodą tradycyjną. Potrzebować będziesz 10 rezystorów 0,1 Ω /5 W. Łączysz je równolegle i takiego bocznika używasz do pomiarów, używając zakresu napięcia DC multimetru. Na każde 10 A napięcie na rezystorze wyniesie 100 mV. Tą metodą można mierzyć prądy do około 70 A. Oczywiście dokładność takiego pomiaru będzie w okolicy 5–10%, no ale coś za coś.

P. Ile multimetrów potrzebuje hobbysta?

O. Dwa to rozsądna liczba na początek. Dlaczego dwa? Albowiem czasem zachodzi potrzeba pomiaru dwóch punktów układu naraz. Albo jednoczesnego pomiaru napięcia i prądu celem oszacowania pobieranej mocy. Co do wyboru, to jeden multimetr powinien być nieco lepszej klasy. Drugi może być tańszym modelem, jak wspomniany wyżej dyskontowy wyrób. Tańszy multimetr przeznaczony jest do pomiarów, które nie wymagają wysokiej dokładności i pomiarów w trudnych warunkach. Instrument lepszej klasy, jak choćby popularny od jakiegoś czasu Aneng AN870, przyda się do dokładniejszych pomiarów, szczególnie gdy musimy dobrać wartości rezystorów lub dokładnie zmierzyć czy ustawić jakieś napięcie w obwodzie. Możemy go nawet użyć, by sprawdzić zbieżność wyników tańszego multimetru.

Z czasem można rozważyć zakup dodatkowego multimetru, sugeruję wybrać miernik innego typu, niż posiadane. Moim zdaniem dobrym wyborem jest multimetr piórowy, gdyż przyda się bardziej w przypadku układów z dużą ilością komponentów SMD. W następnej kolejności warto rozważyć dopięcie drugiego multimetru lepszej klasy.

P. Czy hobbysta powinien mieć multimetr stołowy?

O. Generalnie nie ma takiej potrzeby. Multimetry stołowe, zwłaszcza instrumenty lepszej klasy za kilka tysięcy złotych nie są przeznaczone dla hobbystów, tylko do prac laboratoryjnych. Amator nie wykorzysta ich możliwości, a tylko niepotrzebnie rozstanie się z pieniędzmi. Jedyny wyjątek stanowią multimetry stołowe chińskiej produkcji, które oferują typowo 20000 zliczeń (4,5 cyfry) oraz funkcję czytania wyniku pomiaru na głos. Cena wynosi między 450, a 750 złotych, zależnie od tego, czy zamawiamy z Chin, czy kupujemy w Polsce, i od wybranego modelu.

P. Jak dokładne są najtańsze multimetry?

O. Zwykle między 3, a 5% dla najtańszych modeli. Generalnie jeśli miernik zasilany jest baterią 6F22 (9 V), to nie należy oczekiwać od niego wysokiej dokładności. Producenci jednak nawet w tych modelach starają się doregulować je na tyle, by oferowały znośną dokładność.

P. Jak bezpieczne są najtańsze multimetry?

O. Wcale nie są bezpieczne. Pojęcie norm bezpieczeństwa często bywa obce producentom najtańszych chińskich bubli. Przykład z życia: mój pierwszy, tani multimetr został zamordowany po tym, jak mój ojciec podłączył go bez mojej wiedzy w obwód alternatora celem pomiaru prądu ładowania. Nie dość, że bocznik stał się wydajnym grzejnikiem, to jeszcze przewody pomiarowe się stopiły, potwierdzając przy okazji iż alternator jednak działa. Dlaczego tak się stało? Ponieważ na zakresie 10 A nie było żadnego bezpiecznika. Zakresy niższych prądów zaś były zabezpieczone malutkim bezpiecznikiem szklanym, w dodatku wlutowanym w płytkę.

Tu warto poruszyć jeszcze jeden aspekt bezpieczeństwa: czasem multimetry są obciążone napięciem przekraczającym możliwości ich zabezpieczeń. Niektóre komponenty mogą nawet eksplodować. W multimetrach lepszej klasy obudowy są zaprojektowane tak, by w razie rozsadzenia jakiegoś elementu nic nie wydostało się przez obudowę i nie stanowiło zagrożenia dla użytkownika, jeśli ten właśnie trzyma miernik w ręku. Tanie multimetry są zaprojektowane tak, by się nie rozlecieć w transporcie. Obudowa nie zapewni żadnej ochrony, a napisy w rodzaju CAT III 1000 V to pobożne życzenia producenta, a nie fakty.

P. Dlaczego funkcja hFE (pomiar wzmocnienia tranzystora bipolarnego) występuje w tanich multimetrach, a w drogiej już nie?

O. Głównie ze względu na jej wątpliwą użyteczność. Po pierwsze, większość układów tranzystorowych jest projektowana tak, by wzmocnienie prądowe tranzystora nie miało znaczenia. Są sytuacje, gdzie trzeba dobrać tranzystory w pary o takim samym wzmocnieniu. Przy takim pomiarze jednak trzeba wziąć pod uwagę szereg czynników. Wzmocnienie tranzystora zależy, nieliniowo od prądu kolektora, a do tego jeszcze zmienia się wraz z temperaturą. W tanich miernikach funkcję pomiaru realizuje sprytnie poprowadzony układ ścieżek i trzy rezystory: dwa polaryzują bazy dla tranzystorów NPN i PNP. Trzeci funkcjonuje jako bocznik, przy czym dla tranzystorów PNP mierzy prąd kolektora, ale dla tranzystorów NPN już prąd emitera. Przy doborze pary komplementarnej ma to znaczenie, gdyż:

$$I_e = I_c + I_b$$

Na domiar złego użycie rezystorów zamiast poprawnych źródeł prądowych sprawia, iż wynik pomiaru będzie się zmieniał wraz ze spadkiem napięcia baterii. Dlatego funkcja pomiaru wzmocnienia nadaje się raczej tylko do sprawdzania, czy tranzystor w ogóle działa.

P. Czy są jakieś alternatywy do zwykłych multimetrów dające więcej funkcji?

O. Tak! Na przykład skopometry, czyli połączenie multimetru z oscyloskopem. Zwykle część multimetrowa jest półautomatyczna, ale tryby pracy przełączane są przyciskami zamiast pokrętką. Część oscyloskopowa w tańszych skopometrach miewa zawyżone parametry, zwłaszcza pasmo przenoszenia, same dostępne funkcje są ograniczone względem pełnoprawnych oscyloskopów cyfrowych, ale w zamian takie urządzenie pozwala nam wykonywać pomiary oscyloskopowe w miejscach obwodu, gdzie podłączenie zwykłego oscyloskopu bez specjalnej, drogiej sondy różnicowej lub izolowanej skończyłoby się uszkodzeniem instrumentu. Ponadto skopometry są zasilane przez wbudowany akumulator, więc możemy zabrać je w teren.

Innym typem użytecznego instrumentu są testery komponentów. Pozwalają na automatyczne rozpoznawanie większości powszechnie stosowanych elementów. Bez problemu odróżniają typowy tranzystor bipolarny od polowego, identyfikując przy okazji układ wyprowadzeń, w miarę dokładnie mierzą rezystory, kondensatory i cewki oraz pozwalają sprawdzić różne typy diod. Nie są przy tym drogie – nasi przyjaciele z Chin wręcz zalewają rynek tymi użytecznymi maleństwami.

Ostatnim typem instrumentu wartym uwagi są mierniki/mostki LCR/LRC. Te instrumenty mierzą rezystory, kondensatory i cewki z większą dokładnością, niż oferują multimetry czy zwykłe testery komponentów. Nie są to raczej instrumenty dla początkujących hobbystów, ale ci bardziej zaawansowani, jak choćby radioamatorzy budujący własne urządzenia, mogą znaleźć mostki LCR niezwykle użytecznymi.

P. Co sądzisz o starych, używanych multimetrach i mostkach pomiarowych LC polskiej produkcji? Jest ich trochę na portalach aukcyjnych i ogłoszeniowych. Czy warto coś takiego kupić?

O. To dość ciekawe pytanie. Z jednej strony instrumenty te były często produkowane na rynek profesjonalny, przez co oferowały bardzo przyzwoite parametry, i spokojnie mogą konkurować z współczesnymi miernikami lepszej klasy. Niestety, spora część mierników padła już ofiarą cwaniaczków, którzy wyciągają z nich lampy Nixie i sprzedają je oddzielnie. W takiej sytuacji zakup zamienników może być kosztowny, a wsadzenie wyświetlaczy siedmiosegmentowych wymaga kilku przeróbek. Druga sprawa to kalibracja takiego instrumentu – po 20–30 latach leżenia w magazynie wiele komponentów się zestarzało i zmieniło swoje parametry przez co sam instrument się zwyczajnie rozkalibrował. Do tego przed „emeryturą” mógł pracować 10–20 lat, i nie zawsze był kalibrowany. A i nie do każdego modelu jest łatwo zdobyć instrukcję ze schematami. Tak że „okazja życia” dość szybko może zmienić się w „wielki projekt serwisowy”. Pomijam już to, że te instrumenty ze względu na przestarzałą technologię są zwyczajnie duże i ciężkie.

P. Wspominałeś o oscyloskopach. Czy oscyloskop może zastąpić multimetr?

O. Teoretycznie tak. Oscyloskopem z odpowiednimi przystawkami możemy mierzyć napięcie, prąd, opór, pojemność kondensatorów i indukcyjność cewek, a nawet wyznaczać charakterystyki różnych półprzewodników. W praktyce zajmuje to dużo czasu, i trzeba jeszcze mieć te przystawki, wiedzieć jak ich używać, a często też dokonywać obliczeń. W skrócie można mierzyć wiele rzeczy oscyloskopem zamiast multimetru, tylko po co się męczyć?

P. W filmach na YouTube, niektórych artykułach i niekiedy w opisach multimetrów liczba zliczeń jest podawana jako 2000, 20000, 6000, itp. W tekście jednak podawane jest 1999, 39999, itp. O co chodzi?

O. Obie formy zapisu są poprawne i znaczą to samo. Po prostu łatwiej jest powiedzieć „dwa tysiące zliczeń”, niż „tysiąc dziewięćset dziewięćdziesiąt dziewięć zliczeń”. Dla przykładu Aneng AN870 ma na obudowie napisane „19999 counts”, ale spotyka się w opisach na stronach sklepów i w wielu wideorecenzjach zaokrąglenie do „20000 counts”. Tak się wyraża m.in. Dave Jones na kanale EEVBlog na YouTube.

P. W wykładzie twierdzi się, że termopara wytwarza napięcie, ale coś mi tu nie pasuje. Wyjaśnisz?

O. Chodzi tu zapewne o zjawisko Seebecka, o które opiera się działanie termopary. W opisach podręcznikowych zjawisko Seebecka definiowane jest jako zjawisko termoelektryczne polegające na wytwarzaniu siły elektromotorycznej na styku dwóch różnych metali lub półprzewodników. W opisie pomiaru temperatury termoparą zostało to uproszczone do terminu „napięcie termoelektryczne” dla poprawienia czytelności.

Warto tutaj nadmienić, iż zjawisko Seebecka występuje też przy połączeniach lutowanych na płytkach drukowanych, i przy używaniu bardzo precyzyjnych wzmacniaczy operacyjnych do pomiaru bardzo niskich napięć orientacja wzmacniacza względem komponentów wytwarzających ciepło ma znaczenie i różnica ułamków stopnia między pinami układu ma wpływ na pracę układu. ■

Paweł Kowalczyk



Migające diody LED i śliniący się inżynierowie (4)

Może pamiętacie z części pierwszej tej serii, jednym z moich projektów hobbystycznych jest Inamora Prognostication Engine (zobacz to zdjęcie na Instagramie autorstwa @PracticalElectronics: <https://bit.ly/2UUV2Y0>).

NeoPixel to specjalna forma trójkolorowych diod LED, którym przyjrzymy się w następnym odcinku tego cyklu. Powodem, dla którego o tym wspominam, jest to, że oprócz dwóch przełączników nożowych, ośmiu przełączników dwustabilnych, dziesięciu przełączników przyciskowych, pięciu potencjometrów z napędem umożliwiającym dokonywanie zdalnych nastaw, sześciu mierników analogowych i różnych czujników (temperatury, ciśnienia atmosferycznego, wilgotności, odległości), to maleństwo ma 83 NeoPiksele w górnej obudowie i 116 NeoPikseli w dolnej obudowie. Ojej! Prawie zapomniałem (patrzac na zdjęcie, przypomniało mi się), że jest jeszcze 155 NeoPikseli zasilających pięć potwornych lamp próżniowych zamontowanych na górze.

W trakcie moich pierwszych testów obsługiwałem różne podsystemy (piec, panele przednie, lampy próżniowe) za pomocą kolekcji płytek Arduino Uno i Mega, jednocześnie zdawałem sobie sprawę, że szybko zbliża się czas, kiedy będę potrzebował mocniejszego rozwiązania do przetwarzania.

Rozwiązanie te właśnie się pojawiło w postaci ShieldBuddy (<https://bit.ly/2xLZaBq>), który został stworzony przez specjalistów z firmy Hitex (hitex.com). Pierwszą rzeczą, którą można zauważyć w ShieldBuddy, jest to, że ma taką samą powierzchnię jak Arduino Mega (**rysunek 1**).

Nadszedł czas, abyś usiadł i czytał uważnie, ponieważ ten fragment jest ważny. Standardowe Arduino Mega bazuje na 8-bitowym mikrokontrolerze Microchip ATmega pracującym z częstotliwością 16 MHz z 256 KB pamięci Flash i 8KB pamięci RAM. Dla porównania, ShieldBuddy bazuje na mikrokontrolerze Infineon Aurix TC275. Te potężne „piękności” zwykle można znaleźć tylko w najnowocześniejszych systemach wbudowanych; rzadko wychodzą na światło dzienne w świecie hobbystów.

32-bitowy rdzeń mikrokontrolera TC275 działa z częstotliwością 200 MHz i ma 4 MB pamięci Flash oraz ~500 KB pamięci RAM. Zatrzymaj się na chwilę, aby porównać te liczby z tymi z Arduino Mega. Innymi słowy, rdzeń Arduino Mega wykonuje tylko około szesnastu 8-bitowych instrukcji na mikrosekundę (μs). Dla porównania, rdzeń TC275 ma cykl 5 ns, co oznacza, że zazwyczaj może wykonać od 150 do 200 32-bitowych instrukcji/ μs ($1 \mu s = 1000 ns$).

Och, czekaj! Czy powiedziałem „rdzeń” (liczba pojedyncza)? Głuptas ze mnie. Chciałem powiedzieć, że TC275 może pochwalić się trzema 32-bitowymi rdzeniami, z których każdy pracuje z częstotliwością 200 MHz. Co więcej, w przeciwieństwie do Arduino Mega, każdy z tych rdzeni ma własną jednostkę zmiennoprzecinkową (FPU), co oznacza, że korzystanie ze zmiennoprzecinkowych nie spowalnia znacząco działania.

Sprawy mają się coraz lepiej, po pierwsze, można programować ShieldBuddy za pomocą zintegrowanego środowiska programistycznego Arduino IDE i po drugie, biblioteka NeoPixel jest dostępna dla ShieldBuddy. Moja filiżanka jest pełna.

Dialog z Davidem

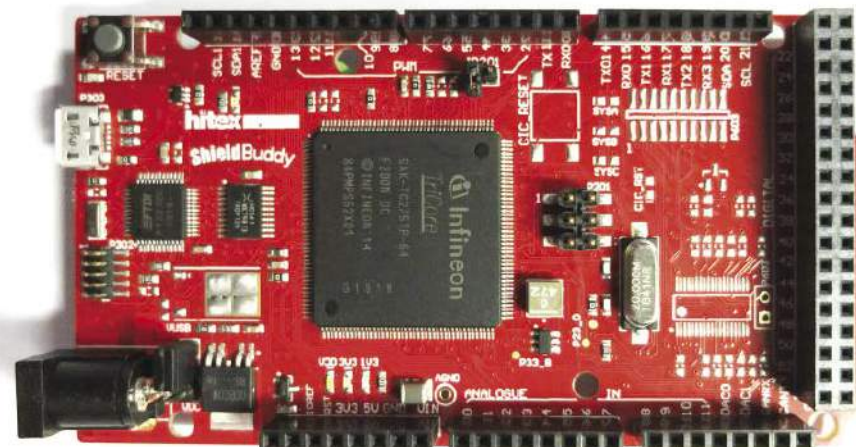
Kilka tygodni temu otrzymałem e-mail od czytelnika PE, Davida R. Humricha, który mieszka w Perth w Australii. David napisał mi, że właśnie skończył czytać pierwszą część tej miniserii i że skłoniło go to, do zagłębienia się w ogromną kolekcję Arduino, które zbierał przez lata.

Kilka dni później David ponownie wysłał e-mail z pytaniem, czy znam Duinotech 8x5 RGB LED Shield do Arduino (<https://bit.ly/2Jz5mQ0>). Nie znałem, ale wcześniej bawiłem się nieco podobnym modulem opartym o NeoPixel 8x8 i odpisałem Davidowi, że interesującym małym programem, którym mógłby chcieć poeksperymentować, byłby „wąż” pełzający po wyświetlaczu.

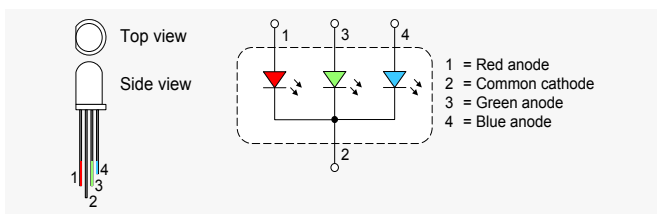
Pomysł polega na tym, że masz jeden piksel na „głowę” węża i kilka dodatkowych pikseli na jego „ciało”. Podświetlasz głowę i ciało różnymi kolorami i ustawiasz węża tak, aby losowo wędrował po wyświetlaczu. Oprócz tego, że jest to atrakcyjne wizualnie, dodatkowo to świetne małe zadanie, które może ułatwić naukę wielu sztuczek programistycznych języka C. Wysłałem Davidowi link do filmu przedstawiającego właśnie taki program uruchomiony na moim wyświetlaczu 8x8 (<https://bit.ly/2R5TfOn>).

Weźmy czerwoną pigułkę

Pamiętasz pierwszą część filmu „Matrix”, w którym Neo musi wybrać między zażyciem niebieskiej lub czerwonej pigułki (<https://bit.ly/39AW2Wo>)? Jak mówił Morfeusz: „Bierzesz niebieską pigułkę – historia się kończy, budzisz się w swoim łóżku i wierzysz w to,



Rysunek 1. ShieldBuddy ma wystarczającą moc obliczeniową, by wywołać zjawienie oczu



Rysunek 2. Zwykła trójkolorowa dioda LED

w co chcesz wierzyć. Bierzesz czerwoną pigułkę – zostajesz w Krainie Czarów, a ja pokazuję ci, jak głęboko sięga królicza nora”.

Cóż, obawiam się, że wybrałem czerwoną pigułkę, ponieważ mój dialog z Davidem skłonił mnie do zanurzenia się w mojej własnej króliczej norze, aby zbudować wspianą matrycę opartą na piłeczkach pingpongowych podświetlanych przez NeoPixela – coś w rodzaju „ściany wideo”, którą można zobaczyć na YouTube (<https://bit.ly/3aG1itl>).

Moim pierwszym projektem będzie mały prototyp ping-ponga $12 \times 12 = 144$. W przyszłości zamierzam zbudować znacznie większą wersję. Właśnie odebrałem dostawę pięciu metrów taśmy NeoPixel 30 pikseli na metr od Adafruit (<https://bit.ly/3dOa5v4>). Otrzymałem również 288 piłeczek pingpongowych (zawsze wierzę w posiadanie części zamiennych) z Amazona, gdzie kosztują tylko 11 USD za paczkę 144 sztuk w USA.

Nie są to piłki o jakości do gier, ale są więcej niż wystarczająco dobre do tego, co zamierzam z nimi zrobić. Będę oczywiście informował o nich w kolejnych artykułach.

Ponad tęczą

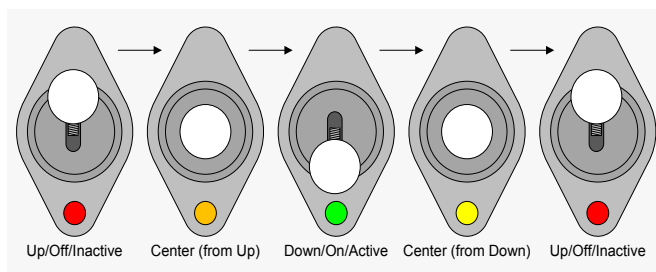
Czy widziałeś kiedyś wideo Izraela „IZ” Kamakawiwo'ole śpiewającego „Somewhere Over the Rainbow” grającego na ukulele (<https://bit.ly/34cx0f5>)? W rzeczywistości to właśnie obejrzenie tego wideo skłoniło mnie do zbudowania własnego ukulele, ale to historia na inny dzień.

Po pierwsze, nie zapomnieliśmy, że przyjrzelśmy się dwukolorowym diodom LED w poprzednim odcinku tego cyklu i że nadal jestem ci winien szkic (plik to CB-Jun20-01.txt – teraz dostępny do pobrania ze strony PE czerwiec 2020) i wideo (<https://bit.ly/2XiVli0>) dotyczące elementu z 2 zaciskami.

Kolejnym krokiem jest użycie trójkolorowej diody LED, która zawiera czerwoną, zieloną i niebieską diodę LED (rysunek 2). Używam diod LED RGB Chanzon 5 mm, które można kupić w opakowaniach po 100 sztuk w Amazon UK za jedyne 5,18 GBP (<https://amzn.to/2x1mc7l>).

Zgodnie z notą katalogową, czerwona dioda ma spadek napięcia przewodzenia od 2,0 do 2,2 V (przyjmijmy 2,0 V), podczas gdy zielona i niebieska dioda mają spadki napięcia przewodzenia od 3,0 do 3,2 V (przyjmijmy 3,0 V). Co więcej, nota katalogowa informuje, że wszystkie trzy diody mają maksymalne wartości prądu przewodzenia 20 mA. Z poprzednich odcinków wiemy, że oznacza to, że będziemy musieli użyć rezystora 150Ω ograniczającego prąd włączonego szeregowo z czerwoną diodą i rezystorów 100Ω dla zielonej i niebieskiej diody.

REKLAMA



Rysunek 3. Użycie trójkolorowej diody LED z przełącznikiem SPCO i dwoma kolorami dla pozycji środkowej

Jeśli po prostu włączymy i wyłączymy nasze trzy diody, możemy uzyskać $2^3 = 8$ różnych kolorów: czerwony, zielony, niebieski, żółty (czerwony + zielony), cyjan (zielony + niebieski), magenta (czerwony + niebieski), biały (czerwony + zielony + niebieski) i czarny (wszystkie wyłączone).

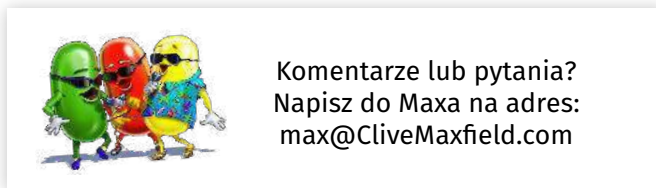
Alternatywnie, jeśli użyjemy 8-bitowej modulacji szerokości impulsu (PWM) do sterowania jasnością każdej diody, oznacza to, że każda dioda może mieć 256 różnych poziomów. W ten sposób mieszanie wszystkich trzech diod pozwala nam uzyskać $256 \times 256 \times 256 = 16\,777\,216$ różnych kolorów.

Na potrzeby tego odcinka, zakładając użycie jednobiegowego przełącznika SPCO, użyjemy naszej trójkolorowej diody LED do wygenerowania tylko czterech kolorów: czerwonego, zielonego, żółtego i pomarańczowego. Użyjemy koloru czerwonego, aby wskazać, kiedy przełącznik jest wyłączony/nieaktywny, zielonego, aby wskazać, kiedy przełącznik jest włączony/aktywny, oraz pomarańczowego lub żółtego, gdy przełącznik znajduje się w pozycji środkowej, aby wskazać jego poprzedni stan (rysunek 3). Możesz pobrać szkic (plik CB-Jun20-01.txt – dostępny na stronie PE z czerwca 2020 r.) i obejrzeć film (<https://bit.ly/3b6hDYx>), aby zobaczyć to wszystko w akcji.

W kolejnym odcinku

Standardowe trójkolorowe diody LED mogą być świetną zabawą, ale mają też kilka wad, między innymi to, że każda z nich wymaga trzech pinów wyjściowych z naszego mikrokontrolera, aby jeysterować.

Dla porównania, NeoPixel, którym przyjrzymy się w następnym odcinku, mają tylko cztery piny: 0 V, 5 V, Data-In i Data-Out. Każdy NeoPixel zawiera mały kontroler wraz z trzema 8-bitowymi generatorami (po jednym dla czerwonej, zielonej i niebieskiej diody LED). Jak zobaczymy, możemy połączyć te małe piękności w łańcuch, pozwalając kontrolować setki pikseli za pomocą jednego pinu z naszego mikrokontrolera.

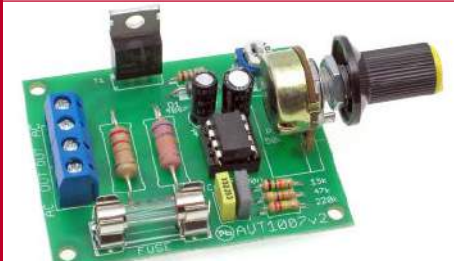


Komentarze lub pytania?
Napisz do Maxa na adres:
max@CliveMaxfield.com

AVT1007

Regulator obrotów silnika elektrycznego

sklep.avt.pl



Sprytne porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania



W poprzednim odcinku (EdW 12/2023) obiecałem, że w tym miesiącu zastanowimy się, w jaki sposób operatory << i >> wykonują swoją magię. Niestety, będziemy musieli odłożyć to na później, ponieważ czytelnik David R. Humrich, mieszkający w Perth (Australia), wysłał mi wiadomość z dość interesującym pytaniem dotyczącym użycia nawiasów klamrowych { }.

Zanim zajmiemy się pytaniem Davida, przypomnijmy, że { } może być używane do tworzenia tak zwanych „instrukcji złożonych”. Jest to mechanizm używany przez język programowania C do grupowania wielu instrukcji w coś, co można uznać za pojedynczą instrukcję.

Rozważmy na przykład, co się stanie, gdy zdefiniujemy funkcję o nazwie MyFunction ():

```
void MyFunction ()
{
    // Założmy, że ten komentarz jest stwierdzeniem
    // Założmy, że ten komentarz jest stwierdzeniem
    // Założmy, że ten komentarz jest stwierdzeniem
}
```

Kiedy wywołujemy tę funkcję z innego miejsca w programie, komputer „widzi” wszystkie instrukcje w funkcji jako tworzące jedną logiczną całość.

To samo dzieje się, gdy używamy { } wraz z instrukcją sterującą, taką jak if (). Po pierwsze, założmy, że jeśli warunek jest prawdziwy, chcemy wykonać tylko jedną akcję, w którym to przypadku moglibyśmy napisać to w następujący sposób:

```
if (done == true) fred = fred + 1;
```

W języku C nie ma znaczenia, ile odstępów między znakami użyjemy:

```
if (done == true)
    fred = fred + 1;
```

Zakładamy oczywiście, że zmienne done i fred zostały zadeklarowane w innym miejscu programu. Założmy teraz, że chcemy wykonać kilka akcji, jeśli nasz warunek jest prawdziwy. Można to zrobić w następujący sposób:

```
if (done == true) fred = fred + 1;
if (done == true) jane = jane - 1;
if (done == true) bert = fred + jane;
```

Oprócz tego, że wygląda to głupio, jest to nieefektywne, ponieważ wykonujemy ten sam test trzy razy. Ten przykład wzywa nas do użycia instrukcji złożonej w następujący sposób:

```
if (done == true) // Tożsame z powyższym
{
    fred = fred + 1;
    jane = jane - 1;
    bert = fred + jane;
}
```

W tym przypadku instrukcja złożona jest bardziej efektywna i sprawia, że staje się czytelniejsze, co chcemy osiągnąć.

Powrót do Davida

Wracając do pytania Davida. Zapytał on, co by się stało, gdyby ktoś użył { }, a tym samym utworzył samodzielną instrukcję złożoną w środku funkcji – przez „samodzielną” rozumiemy, że nie jest ona powiązana z instrukcją sterującą, taką jak if () lub for ().

W rzeczywistości jest to całkowicie legalne. Jeśli chcesz, możesz po prostu użyć { }, aby zebrać grupę instrukcji razem, aby wyjaśnić

sobie i innym, że uważasz te instrukcje za powiązane. Są one często określane jako „blok”, a korzystanie z tej techniki można nazwać „programowaniem blokowym”.

Można również zagnieźdzać { } do dowolnego poziomu. Ale naprawdę interesujące jest to, że oprócz instrukcji, bloki mogą również zawierać deklaracje zmiennych.

Dlaczego jest to interesujące? Cieszę się, że o to zapytałeś. W poprzednim odcinku rozmawialiśmy o zmiennych globalnych i lokalnych. Zauważyliśmy, że zmienne globalne są deklarowane poza dowolną funkcją i mogą być widoczne i modyfikowane przez dowolną funkcję. Dla porównania, zmienne lokalne są deklarowane wewnątrz funkcji i mogą być widziane i modyfikowane tylko przez funkcję, w której zostały zadeklarowane.

Rozmawialiśmy również o „zakresie” zmiennej, który odnosi się do zakresu, w jakim zmienna może być widoczna. Na przykład, jeśli zadeklarujemy zmienną jako część pętli for (), zakres zmiennej jest ograniczony do tej pętli (EdW 12/2023). Cóż, jeśli zmienna jest zadeklarowana wewnątrz bloku, jej zakresem jest ten blok; to znaczy, że nie może być „widziana” poza blokiem, nawet przez inne części funkcji, w której znajduje się blok. Rozważmy przykład kodu poniżej:

```
int John = 6;
```

```
void MyFunction ()
{
    int jane = 9;

    { // Pierwszy blok
        int bert = jane + John
        // Więcej rzeczy
    }

    { // Drugi blok
        int jack = jane - John;
        // Więcej rzeczy
    }
}
```

Pamiętaj, że używam początkowych wielkich i małych liter odpowiednio dla moich zmiennych globalnych i lokalnych. Tak więc John jest zmienną globalną, której zakresem jest każda funkcja w programie, podczas gdy jane jest zmienną lokalną, której zakres jest ograniczony do MyFunction (). Dla porównania, zakres zmiennej bert jest ograniczony do pierwszego bloku, a zakres zmiennej jack jest ograniczony do drugiego bloku.

Podzielenie dużego programu na kilka mniejszych, dobrze zdefiniowanych funkcji ułatwia testowanie każdej funkcji oddzielnie i ponowne wykorzystanie funkcji w różnych programach. Podobnie, jedną z zalet korzystania z techniki blokowej w celu ograniczenia zakresu zmiennych jest to, że ułatwia ona ponowne wykorzystanie tych bloków w innych funkcjach i innych programach. ■

Clive „Max” Maxfield

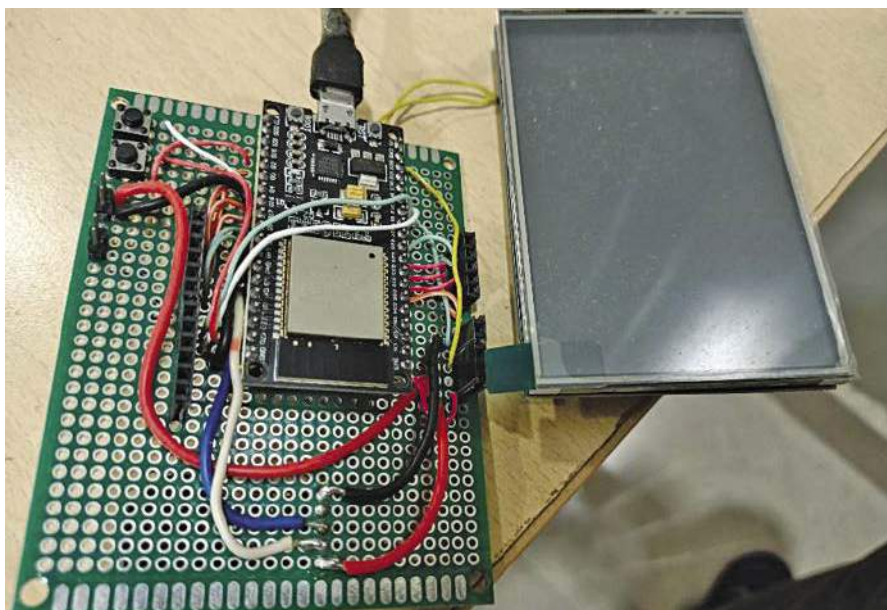
Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, czerwiec 2020 (www.epemag3.com)

Dokładny zegar ze wskazaniem milisekund z użyciem ESP32 RTC

Autor tego projektu jest entuzjastą budowy przeróżnych zegarów. Wiele z nich zbudował, ale jeszcze nigdy nie wykonał czasomierza z dokładnością milisekundową. Wbrew pozorom, taki zegar wymaga sporej mocy obliczeniowej.

Od Red. EdW: Warto zauważyć, iż dokładność i wskazanie z dokładnością do milisekundy, to dwie różne rzeczy. Osobną sprawą jest sens i potrzeba milisekundowej dokładności. W przypadku zegara czasu rzeczywistego RTC (Real Time Clock) wbudowanego w różnego rodzaju systemy, taka, a nawet większa dokładność i rozdzielczość może być wymagana. Inaczej sprawa wygląda, jeśli to zegar, na którym po prostu odczytujemy bieżący czas. Również nieco inaczej to wygląda, jeśli ów czasomierz ma pracować w trybie stopera.

Możliwość budowy zegara pokazującego czas z rozdzielczością milisekundową zrodziła się w chwili, gdy w zasięgu ręki pokazały się takie moduły jak trzy i pół calowy (8,9 cm) wyświetlacz TFT ILI 9488 i mikrokontroler ESP32. Wyświetlacz ten może pracować w kilku trybach. Tu wykorzystano tryb 8-bitowy. Wówczas interfejs między mikrokontrolerem i wyświetlaczem angażuje 12 linii GPIO. Istotne jest natomiast to, że zarówno mikrokontroler jak i transmisja danych jest na tyle szybka, że budowa zegara pokazującego godziny, minuty, sekundy i milisekundy stała się możliwa. Wyświetlacz ILI9488 jest kompatybilny z Arduino. W handlu dostępne są wersje, które jest bardzo łatwo połączyć z płytą Arduino Uno. Niestety, moc obliczeniowa mikrokontrolera na Arduino jest niewystarczająca. Aby sprostać wymaganiom narzuconym przez „milisekundowy clock” sięgnięto po moduł mikrokontrolera ESP32 oraz zegar RTC w postaci układu scalonego DS3231. Zegar prezentowany w tym projekcie przewyższa możliwości każdego innego zegara, który pewnie masz w swoim mieszkaniu. Wyświetla nie tylko datę, dokładną godzinę i dzień tygodnia, także wskazanie milisekundowe nie jest jedynym atutem tego zegara. Może on również pracować w trybie stopera. A wskazanie czasu pokazane jest zarówno w postaci zegara analogowego jak i cyfrowego. Ponadto, zegar ten pokaże także temperaturę w pomieszczeniu, w którym się znajduje. Na **rysunku 1** pokazano prototyp zmontowany przez autora. Na **rysunku 2** jest fotografia użytego w projekcie wyświetlacza. Dzięki wykorzystaniu



Rysunek 1. Prototyp zegara wykonany przez autora projektu

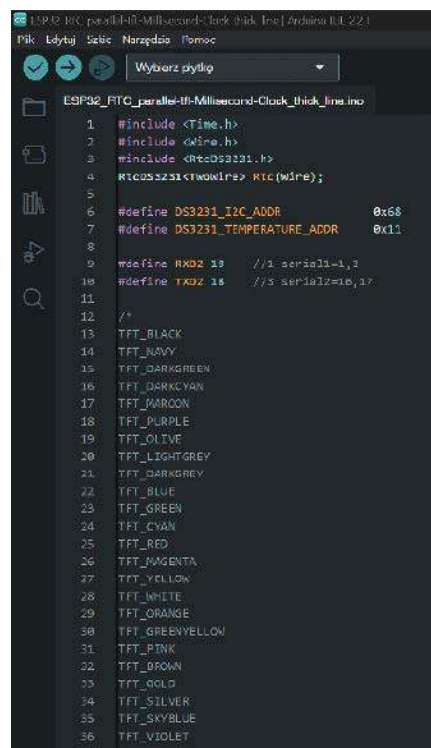
specjalizowanych modułów, użytych podzespołów jest niewiele. Zebrano je w **tabeli 1**.

Oprogramowanie

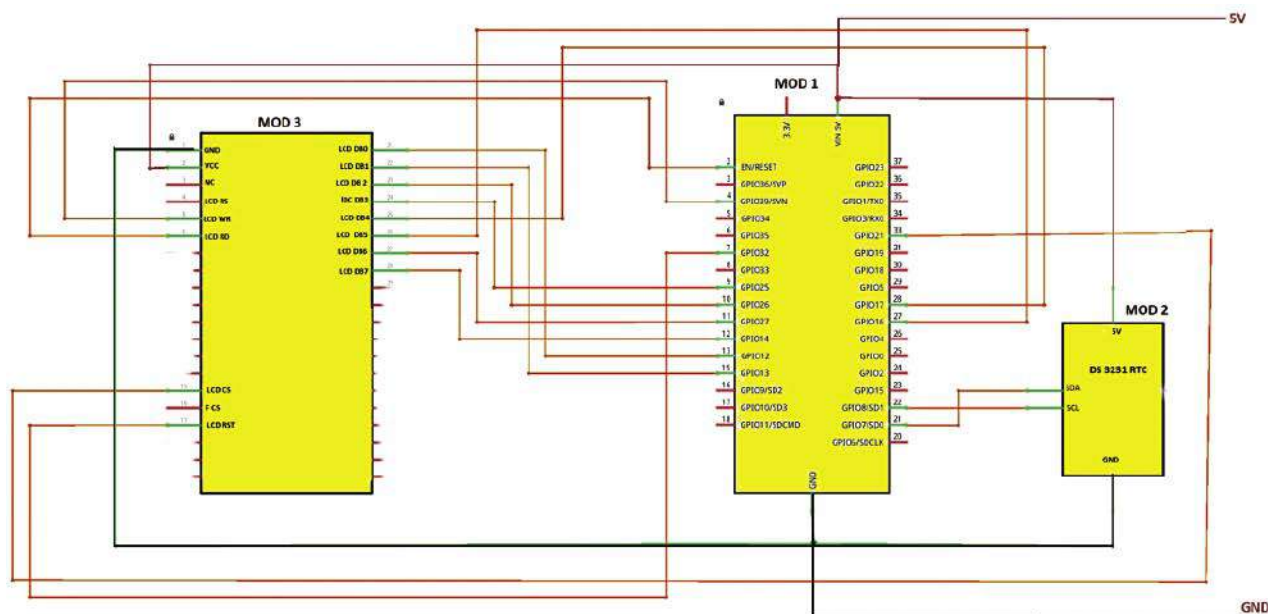
Pierwszą częścią kodu źródłowego jest adaptacja komunikacji z zegarem RTC. Mikrokontroler ESP odczytuje bieżący czas z modułu DS3231. Kolejnym zadaniem programu użytego w tym projekcie



Rysunek 2. Wygląd wyświetlacza użytego w projekcie



Rysunek 3. Zrzut ekranowy fragmentu kodu źródłowego



Rysunek 4. Schemat ideowy zegara

Tabela 1. Spis materiałów

Wyświetlacz TFT ILI9488 (MOD3) – 1 szt.
 Płytki ESP32S (MOD1) – 1 szt.
 Moduł zegara RTC DS3231 (MOD2) – 1 szt.
 Kabel mikro-USB – 1 szt.
 Bateria „coin cell” – 1 szt.

jest komunikacja z wyświetlaczem TFT. Mikrokontroler pracuje tu w prostej pętli programowej, odczytując dane zegara po magistrali I²C i przeliczając je do wyświetlacza. Wymóg milisekundowej dokładności i rozdzielczości narzuca dużą moc obliczeniową, co nie jest typowe w normalnej aplikacji zegara RTC. Jak wyżej stwierdzono, ten

właśnie wymóg skłonił autora do wykorzystania ESP32 zamiast prostszej płytki Arduino. Aczkolwiek, w celu wgrania kodu źródłowego do ESP można wykorzystać Arduino. Zrealizowano to pod kontrolą oprogramowania Arduino IDE. W procesie wgrzywania kodu źródłowego, należy jedynie poprawnie wybrać typ płytki i numer portu pod którym jest ona widziana w programie IDE. Zrzut ekranowy fragmentu kodu widzimy na **rysunku 3**.

Po wgraniu programu do mikrokontrolera ESP, należy podłączyć zegar RTC do linii magistrali I²C. Podobnie, należy połączyć wyświetlacz z modulem mikrokontrolera. Tutaj linii interfejsu jest więcej i zebrano je w **tabeli 2**.

Tabela 2. Linie interfejsu między ESP i wyświetlaczem TFT

Pin wyświetlacza TFT ILI9488	Pin na płytce ESP32
3,3 V	3,3 V
GND	GND
D0	7
D1	8
D2	1
D3	2
D4	3
D5	4
D6	5
D7	6
CS	15
DC	0 (pozostaw niepodłączony)
RST	14
WR	17
RD	18

Schemat układu i jego testowanie

Schemat ideowy pokazano na **rysunku 4**. Składa się on praktycznie tylko z trzech modułów. MOD1 to mikrokontroler ESP32S, MOD2 to właściwy zegar RTC DS3231, oraz MOD3 to wyświetlacz TFT ILI9488.

Zasilanie zegara jest 5-cio voltowe. Można wykorzystać baterijkę lub adapter 5 V ze złączem mikroUSB. Po poprawnym montażu, nasz układ zegara nie wymaga uruchamiania. Zegar DS3231 wyposażony jest także w funkcję pomiaru temperatury. Te dane także są odczytywane po dwuprzewodowej magistrali I²C. Oprogramowanie tego projektu uwzględnia także tę możliwość. Temperatura wyświetlana jest w rogu ekranu TFT. Każda sekunda dzielona jest przez 1000 i milisekundy wyświetlane są tylko w postaci zegara cyfrowego w dolnej linii. Dla poprawy efektu widocznego na ekranie TFT, linie zostały pogrubione. Parametr ten ujęty jest



Rysunek 5. Efekt projektu zegara wyświetlanego na ekranie TFT

w zmiennych programu kodu źródłowego, co zwiększa elastyczność projektu na ostateczny efekt widoczny na ekranie. Efekt końcowy zegara wg projektu autora widzimy na **rysunku 5**. ■

Somnath Bera

Ten artykuł był wcześniej opublikowany na łamach „EFY”, czerwiec 2023 (efymag.com)

Automatyczny dozownik odmierzonej ilości cieczy

Bieżący projekt dozownika cieczy wykorzystuje mikrokontroler Arduino oraz czytnik pastylek RFID. Układ można wykorzystać do dozowania dowolnej cieczy, w tym wody lub np. benzyny. Chcąc odmierzyć założoną ilość cieczy, należy do czytnika RFID przyłożyć odpowiednio opisaną pastylkę.

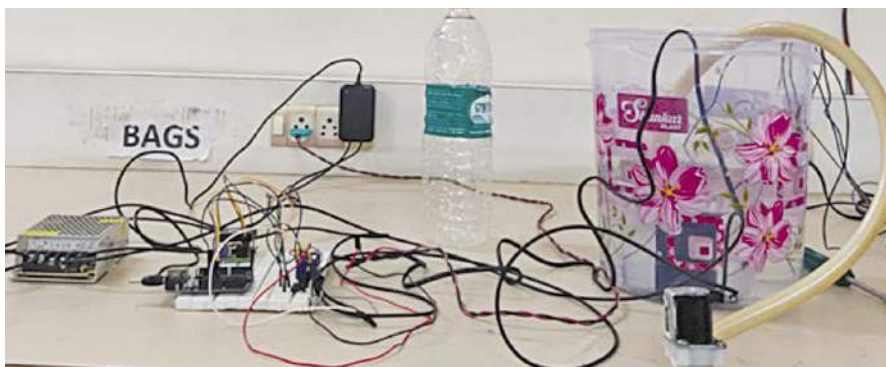
Zarówno czytnik jak i wykonawczy przełącznik obsługiwany jest przez Arduino po wgraniu do mikrokontrolera odpowiedniego szkicu programu. Jako driver przełączników wykorzystano układ scalony ULN2003. Proponowany system można wykorzystać w wielu aplikacjach. Jak np. dla oszczędności wody lub zabezpieczenie na wypadek nieuczciwego korzystania z dystrybutorów paliwa. Odmierzanie zadanej ilości cieczy bazuje w istocie na odmierzaniu czasu otwarcia zaworu i/lub włączenia pompy. W projekcie autora przewidziano 5 tagów RFID. Każdemu przypisano inną ilość cieczy w odstępach 1/2 litra. Czyli: 0,5 l, 1 l, 1,5 l, 2 l i 2,5 l. Prototyp wykonany przez autora widzimy na **rysunku 1**. Natomiast w **tabeli 1** jest spis wykorzystanych podzespołów.

Oprogramowanie

Pierwszą czynnością powinno być sczytanie kodów z wykorzystanych tagów RFID. W tym celu należy połączyć czytnik EM18 RFID (moduł MOD2) z Arduino zgodnie ze schematem na **rysunku 3** oraz do pamięci mikrokontrolera wgrać krótki program wg **rysunku 2**, który pozwoli na odczytanie unikalnego dwunastocyfrowego kodu przypisanego do każdej pastylki tag RFID. Odczyt tych danych polega na prostym zbliżeniu tagu do modułu czytnika RFID. W **tabeli 2** pokazano wynik tej operacji przeprowadzonej przez autora. Odczytany kod to 12 cyfr heksadecymalnych – 16¹² możliwych liczb, więc nie ma obawy, iż kod ten może się powtórzyć w dwóch różnych tagach.

Tabela 1. Spis materiałów wykorzystanych w projekcie

Arduino Uno (MOD1)	1 szt.
Czytnik EM18 RFID (MOD2)	1 szt.
Karty lub Tagi RFID (125 kHz)	5 szt.
ULN2003 (IC1)	1 szt.
Pompa wody 230 VAC (Pump1)	1 szt.
Elektrozawór 24VDC	1 szt.
Zasilacz	Adapter 230 VAC/12 VDC
Przełącznik SPST 5 VDC (R1, R2)	2 szt.



Rysunek 1. Prototyp wykonany przez autora

Właściwym kodem programu dla obsługi dozownika cieczy jest szkic w języku Arduino pokazany na **rysunku 4**. W programie tym musisz uzupełnić kody identyfikacyjne twoich tagów. Należy zaznaczyć, której pastylce tag chcesz przypisać jaką ilość dozowanej cieczy. Działanie programu polega na identyfikacji przypisanych tagów oraz na włączeniu wykonawczego przełącznika na określony czas.

Działanie układu i jego testowanie

Przygotowane podzespoły wg tabeli 1, należy połączyć zgodnie ze schematem na **rysunku 5**.

Zbliżenie jednego z tagów do czytnika RFID skutkuje jednoczesnym otwarciem dwóch przełączników widocznych na schemacie z **rysunku 5**. Przełącznik R1 włącza pompę wody (lub innej cieczy), podczas gdy relay R2 steruje cewką elektrozaworu. Jako driver cewek obu

Tabela 2. Zdekodowane wartości numerów przypisanych tagom użytym w projekcie autora

RFID Tag	Image	Decoded value from serial monitor (12-digit)
0001948765 029,48221		5A001DBC5DA6
0013946179 212,52547		5900D4CD4303
0013942651 212,49019		5900D4BF7B49
0001956746 029,56202		5A001DDB8A16
0001949492 029,48948		5A001DBF34CC

```
//RAKESH JAIN PROGRAM TO READ RFID
TAG
void setup() {
  Serial.begin(9600);
}

void loop() {
  if(Serial.available())
  {
    Serial.print((char)Serial.read());
  }
}
```

Rysunek 2. Kod źródłowy programu pozwalającego odczytać numery ID tagów RFID

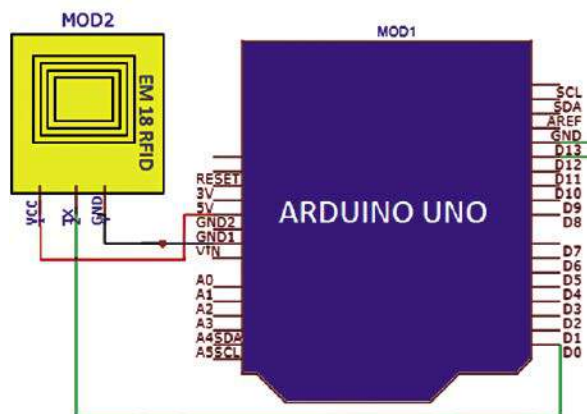
przełączników wykorzystano układ scalony ULN2003. Aktywnym stanem jest stan niski na wyjściach n.16 i n.15 IC1. W prototypie autora zmierzono, iż wydajność pompy wynosi pół litra cieczy w przeciągu 16 sekund. Dlatego, chcąc odmierzyć 1 litr, należy pompę i zawór otworzyć na 32 sekundy. Tym samym, pastylce RFID której chcemy przypisać dozowanie 2,5 litra, powinno towarzyszyć otwarcie zaworu i włączenie pompy na czas 80 sekund. ■

Rakesh Jain

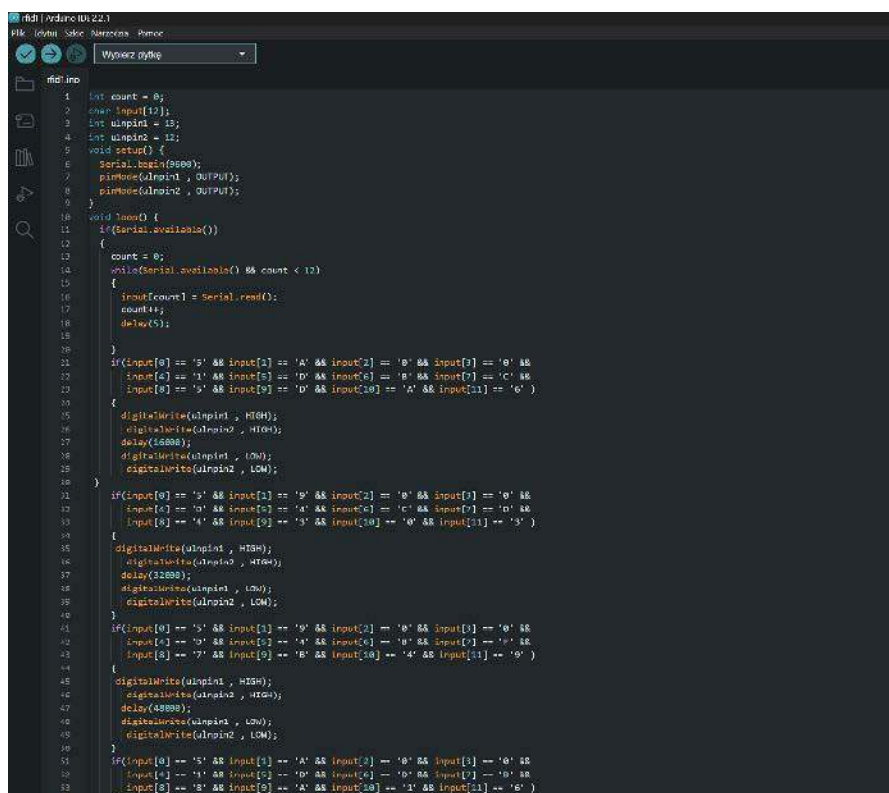
Ten artykuł był wcześniej opublikowany na łamach „EFY”, czerwiec 2023 (efymag.com)

Od Red. EdW: Oczywiście jest, że wiele parametrów tego projektu jest uzależnionych od konkretnej aplikacji. Zapewne Czytelnik będzie miał inne numery identyfikacyjne swoich tagów. Nie musi ich być 5. Może to być mniejsza bądź większa liczba. Prawdopodobnie będzie chciał przypisać swoim tagom inne wartości odmierzonej objętości cieczy. Jeśli nawet będą to wartości zbliżone z potrzebami w prezentowanym prototypie, to prawdopodobnie inna będzie wydajność posiadanej pompy. W prototypie pompa i elektrozawór otwierają i zamykają się równocześnie. W takiej sytuacji obu przełącznikom można by przydzielić tylko jeden driver. Jednakże potrzeby mogą być nieco inne. Także, niekoniecznie

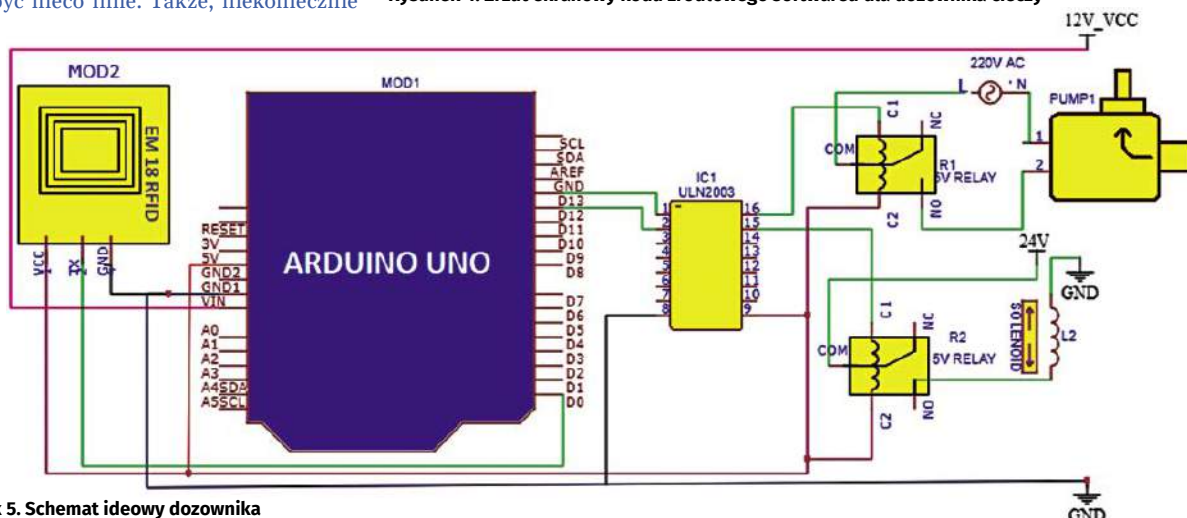
w każdej aplikacji będzie potrzebna pompa i zawór. Wiele parametrów będzie zależało od tego, jaką ciecz będziemy dozować. Duża elastyczność projektu w zależności od konkretnego zastosowania wynika z wykorzystania w projekcie mikrokontrolera na Arduino. Zmienne wyniki, jakące potrzeb Czytelnika, wystarczy odpowiednio uzupełnić w szkicu programu którego zrzut ekranowy pokazano na rysunku 4.



Rysunek 3. Konfiguracja interfejsu czytnika EM-18 RFID z Arduino



Rysunek 4. Zrzut ekranowy kodu źródłowego oprogramowania dla dozownika cieczy



Rysunek 5. Schemat ideowy dozownika

Prosty wskaźnik temperatury wody oraz sterownik z użyciem Arduino

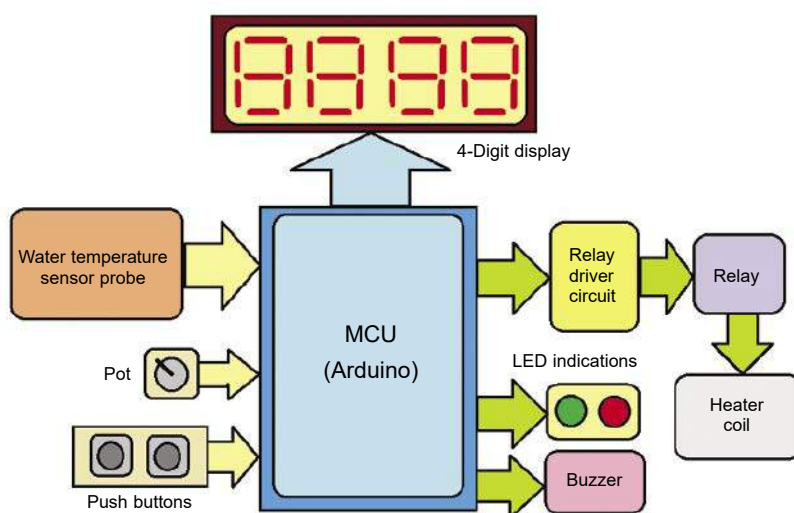
Bieżący projekt powstał z myślą o aplikacjach gdy trzeba kontrolować temperaturę wody (lub innej cieczy) i ew. podgrzewać ją do z góry założonej temperatury. Konstrukcja układu stwarza następujące możliwości: ustawienie maksymalnej temperatury do której woda ma się podgrzewać, wyświetlenie temperatury aktualnej, wyłączenie grzałki gdy zostanie osiągnięta temperatura zadana oraz indykacja tego faktu diodą LED oraz za pośrednictwem dźwięku buzzera. Co prawda, wiele podgrzewaczy wody dla użytku domowego takie możliwości ma, jednak chcąc wykorzystać typową grzałkę zanurzeniową takiej kontroli mieć nie będziemy.

Woda podgrzewa się tak długo, aż użytkownik ją wyłączy. Takie „ręczne sterowanie” skutkuje niepotrzebnie nadmiernym zużyciem energii elektrycznej oraz przegrzaniem wody ponad temperaturę której oczekujemy.

Na **rysunku 1** pokazano schemat blokowy urządzenia, a na **rysunku 2** schemat ideowy.

Jako czujnik temperatury wykorzystano układ scalony DS18B20. Do wskazywania temperatury wykorzystano cztero-cyfrowy siedmio-segmentowy multipleksowany wyświetlacz znakowy. Mózgiem zarządzającym jest mikrokontroler ATmega328 na płycie Arduino Nano. Pozostałymi podzespołami są: transoptor MCT2E, potencjometr, dwie diody LED, miniaturowy buzzer oraz przełącznik. Mikroprocesor na bieżąco odczytuje temperaturę z czujnika DS18B20 i wartość tą wyświetla na podłączonym czterocyfrowym wyświetlaczu znakowym. Docelową temperaturę można dostroić za pomocą dołączonego do Arduino potencjometru. Następnie można wystartować proces grzania wody. Od tego momentu automat przejmie kontrolę nad tym procesem. Wyłączy grzałkę, gdy żądana temperatura zostanie osiągnięta, równocześnie sygnalizując o tym fakcie.

W projekcie wykorzystano termometr Dallas-a DS18B20. Ten trzy-nóżkowy element oferuje bardzo cenne cechy. Komunikacja odbywa się po jednym przewodzie z wykorzystaniem protokołu One-Wire. Nie stwarza to problemu jeśli w projekcie i tak wykorzystano mikrokontroler. Odczyt jest w postaci cyfrowej. Mikrokontroler bez problemu przetworzy odczytaną z czujnika wartość temperatury i wyświetli ją użytkownikowi za pomocą wspomnianego wyświetlacza. To zadanie leży po stronie softwareowej projektu. Ustawienia żądanej temperatury także dokonuje program. Po stronie sprzętowej znajduje się prosty potencjometr. Mikrokontroler czyta analogową wartość napięcia i tutaj duża precyzja nie jest wymagana. Regulacja potencjometrem „przeziara” programowo założony zakres temperatury, którą natychmiast widać na wyświetlaczu. Wyświetlacz DIS1 jest typowym LEDowym wyświetlaczem

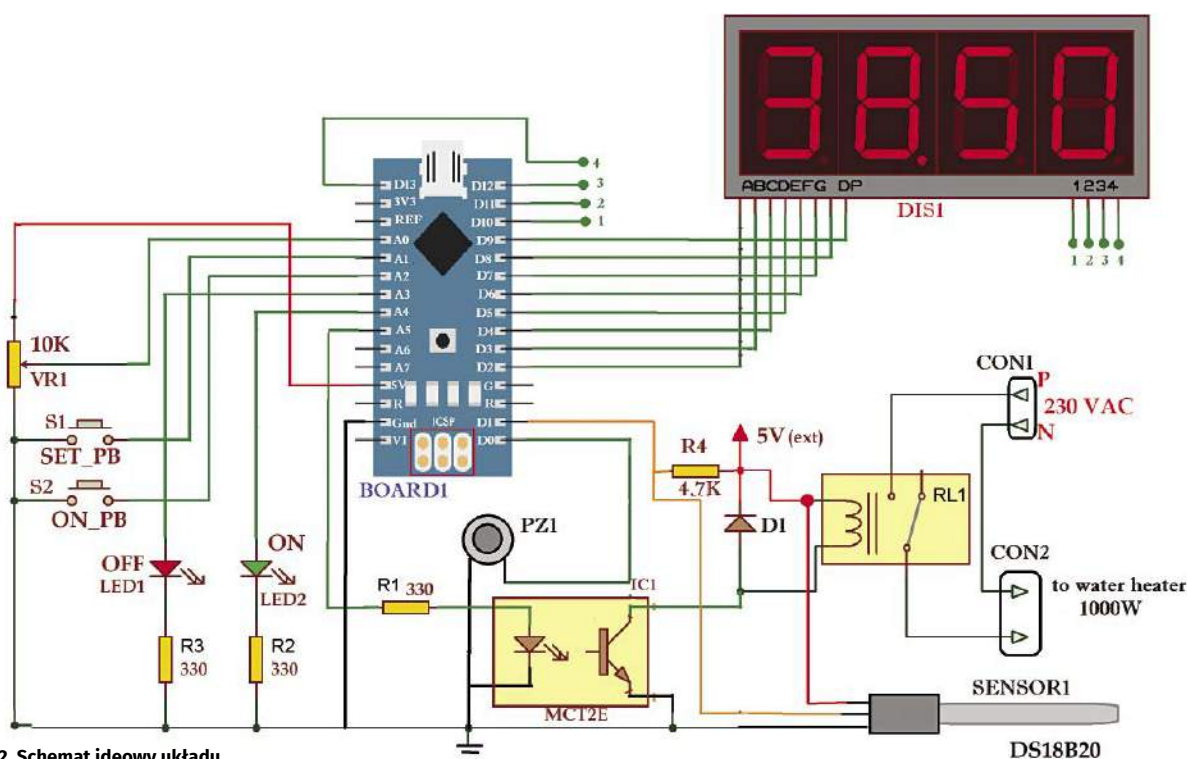


Rysunek 1. Schemat blokowy układu nadzorującego temperaturę wody

składającym się z cyfr, a każda z cyfr składa się jedynie z siedmiu segmentów (diod LED). Dla wyświetlacza nie przewidziano żadnego rejestru który miałby zapamiętywać wartość każdej cyfry. W zamian za to, mikroprocesor przeziara kolejne cyfry w trybie multipleksowania. Proces ten odbywa się z częstotliwością przekraczającą percepcję ludzkiego wzroku i jest to rozwiązanie stosowane od początku ery techniki cyfrowej. Informacja z wyświetlacza wzbogacona jest dwoma diodami LED które informują czy włączona jest grzałka, czy też została już osiągnięta temperatura zadana. Ta informacja jest także potwierdzona dźwiękiem buzzera. Obsługa wszystkich komponentów widocznych na rysunku 1 leży po stronie programowej, co stwarza łatwą możliwość adaptacji do konkretnego zastosowania. Dźwięk buzzera może być ciągły, lecz tutaj przewidziano krótki ‘beep’ tuż po wyłączeniu grzałki. Mikrokontroler na bieżąco wyświetla temperaturę odczytaną z termometru, a jedynie w trybie ustawiania wartości zadanej odczytuje wartość analogową z potencjometru z wykorzystaniem przetwornika ADC. Elementem wykonawczym systemu jest przełącznik włączający grzałkę. Jako driver wykorzystano transoptor MCT2E.

Izolacja na tej ścieżce ma charakter zabezpieczenia mikrokontrolera przed ew. uszkodzeniem przepięciami ze strony wysokoprądowej projektu (dodatkowe uwagi na temat drivera przełącznika na końcu artykułu).

W Arduino wykorzystano zarówno wejścia/ wyjścia cyfrowe jak i analogowe. Wejście analogowe wykorzystano do odczytywania napięcia ustawionego na potencjometrze. VR1 (10 kΩ) podłączono między masę i zasilanie +5 V płytki Arduino. Suwak potencjometru podłączono na wejście A0, które tym samym otrzymuje wartość analogową z pełnego zakresu zasilania płytki. O rozdzielczości decyduje już przetwornik analogowo-cyfrowy obecny w mikrokontrolerze ATmega328, lub program który go obsługuje. Switche S1 i S2 także są czytane na wejściach analogowych, choć generują sygnał binarny. Wciśnięcie przycisku generuje stan niski, co oznacza (lecz niekoniecznie) zero logiczne. S1 i S2 podłączone są do wejść, odpowiednio A1 i A2 Arduino. Dwie diody LED (czerwona i zielona) także podłączono do portu A, odpowiednio A3 i A4. Skonfigurowano je jako wyjście i generują sygnał binarny. Każda z diod LED zaświecana jest stanem wysokim, gdyż katody podłączono w stronę masy.



Rysunek 2. Schemat ideowy układu

Szeregowe rezystory ($330\ \Omega$) ograniczają prąd do ok. 10 mA. Głównym sygnałem wyjściowym jest linia A5. Tędy sterowany jest driver grzałki. Driverem jest transpotor MCT2E, a jego dioda sterowana jest w ten sam sposób jak LED1 i LED2 (stan aktywny wysoki z ograniczeniem prądu przez rezystor $330\ \Omega$). Cewka przekaźnika RL1 włączona jest bezpośrednio w obwód kolektora fototranzystora w transporcie. Dioda D1 podłączona równolegle z cewką przekaźnika jest konieczna ze względu na obciążenie indukcyjne (zweira ona prąd indukujący się w cewce przekaźnika podczas zwolnienia kotwicy). Miniaturowy buzzer PZ1 podłączono natomiast do wyjścia cyfrowego D0. Tu stanem aktywnym jest także stan wysoki, ponieważ PZ1 podłączono względem masy.

Najważniejsza informacja biegnie po linii D1. To cyfrowe wejście skonfigurowane do pracy w protokole One-Wire. Tędy mikroprocesor odczytuje aktualną temperaturę. DS18B20 to element 3-nóżkowy. Dwie nóżki to zasilanie i masa (choć potrafi on pracować „bez zasilania”, „kradnąc” je z linii danych). Linia wejścia/wyjścia oznaczona jest OP i tu obowiązuje pełny protokół magistrali 1-Wire. Na linii danych podłączono rezystor podciągający do zasilania $R4=4,7\ k\Omega$, zgodnie z zaleceniami karty katalogowej, aczkolwiek nie jest to konieczne, jeśli podłączono linię zasilania termometru DS18B20.

Cztero-cyfrowy wyświetlacz jest typu „ze wspólną katodą”. Jest on obsługiwany przez port cyfrowy Arduino. Linie segmentów A, B, C, D, E, F, G i DP (Decimal Point) połączono

z D2 do D9 Arduino i tu też obowiązuje logika, gdzie aktywny jest stan wysoki (w szeregu dla każdego z 8-miu segmentów również powinien znajdować się rezystor o wartości $330\ \Omega$, podobnie jak ma to miejsce dla wszystkich pozostałych diod LED w projekcie – przypis red.). Linie wyboru cyfry (wspólne katody) obsługują Arduino z wyjść cyfrowych D10 do D13. I tutaj aktywny jest stan niski. Wszystkie wymienione wyżej zadania realizuje tu Arduino w wersji Nano. A więc: odczytuje bieżącą temperaturę z DS18B20 (skonfigurowano tu programowo protokół 1-Wire); odczytuje interfejs użytkownika w postaci switchy S1 i S2 oraz analogową wartość z potencjometru VR1; obsługuje wyświetlacz posyłając na niego odczyt z termometru lub „przetworzoną” wartość z potencjometru; obsługuje indykację na diodach LED1 i LED2 oraz buzzer; i przede wszystkim włącza bądź wyłącza grzałkę za pośrednictwem przekaźnika i jego drivera.

Mimo sporych zadań zleconych mikrokontrolerowi na Arduino, struktura pracy układu jest względnie prosta. Początkowo (po włączeniu zasilania) grzałka jest wyłączona i zaświeca się dioda czerwona która o tym fakcie właśnie informuje. Można teraz ustawić temperaturę żadaną. Naciskamy przycisk S1 (SET_PB) informując mikroprocesor o tym, że ma odczytać wartość z potencjometru. W projekcie autora przyjęto wartość nastawy w przedziale od 40°C do 70°C . Wartość ta jest natychmiast widoczna na wyświetlaczu (jak np. 45, 48, 55, 65, itd.). Praca w trybie termostatu wystartuje po naciśnięciu przycisku

S2 (ON_PB). Grzałka zostanie włączona o ile temperatura wody jest niższa od ustawionej. Wskazaniem faktu włączenia grzałki jest zaświecenie diody zielonej (lepiej by się kojarzyło jeśliby grzanie było wskazane diodą czerwoną, a stan osiągnięcia temperatury zadanej zieloną – przypis red.). Wyświetlacz pokazuje cały czas temperaturę aktualną. A więc podczas grzania wskazanie na wyświetlaczu powinno rosnąć. Po osiągnięciu temperatury, którą wcześniej ustawiliśmy, grzałka się wyłącza i ponownie zaświeca się dioda czerwona. Równocześnie ‘beep’ z buzzera sygnalizuje, że woda jest już nagrzana.

Oprogramowanie

Program napisano w języku C/C++ z wykorzystaniem Arduino IDE. Środowisko IDE (dostępne bezpłatnie) pozwala także skompilować szkic do postaci „języka maszynowego”

Wykaz elementów:

Półprzewodniki:

IC1: transpotor MCT2E
LED1, LED2: diody LED 5 mm
D1: 1N4007 – dioda prostownicza

Rezystory: (wszystkie 0,25W/±5%)

R1...R3: $330\ \Omega$
R4: $4,7\ k\Omega$
VR1: $10\ k\Omega$ potencjometr

Inne:

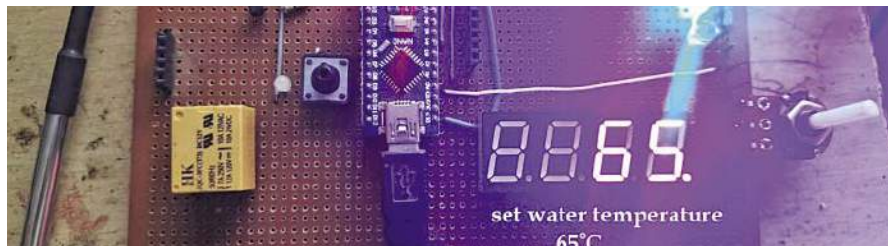
BOARD1: Arduino Nano
DIS1: czterocyfrowy wyświetlacz siedmiosegmentowy
S1, S2: przyciski „push-to-on”
RL1: przekaźnik SPDT 5 V
PZ1: buzzer piezoelektryczny
CON1, CON2: złącze 2-pinowe
SENSOR1: termometr DS18B20
grzałka zanurzeniowa (1000 W)

i wgrać go do pamięci flash mikrokontrolera ATmega328 na płytce Arduino. W szkicu dla tego projektu wykorzystano 3 biblioteki: OneWire.h, DallasTemperature.h i SevSeg.h. Pierwsza z tych bibliotek pozwala obsłużyć magistralę zgodnie z protokołem One-Wire. DallasTemperature jest biblioteką dedykowaną konkretnie dla termometru DS18B20, która zwraca odczytaną wartość w stopniach Celsjusza. Siedmiosegmentowy wyświetlacz znakowy obsługiwany jest przez bibliotekę SevSeg.h. Ta sekcja programu zamienia cyfrową wartość w pamięci mikroprocesora na czytelną wartość zaświecenia odpowiednich segmentów kolejnych cyfr wyświetlacza.

Konstrukcja i testowanie pracy układu

W pierwszej kolejności należy wgrać kod źródłowy programu do pamięci flash mikrokontrolera na płytce Arduino. Następnie część sprzętową można zmontować na płytce uniwersalnej zgodnie ze schematem ideowym na rysunku 2. Należy pamiętać o połączeniu linii oznaczonych na schemacie 1, 2, 3 i 4. Pomocny może być **rysunek 3** pokazujący układ zmontowany przez autora tego projektu.

Po zmontowaniu płytki PCB należy przygotować dla niej odpowiednią obudowę. Na czołowej stronie należy umieścić wyświetlacz, diody LED i dwa switchy S1 i S2. Sensor-termometr należy ułożyć w obudowie wodoszczelnej na wystarczająco długich przewodach (tyle na ile to jest konieczne).



Rysunek 3. Prototyp wykonany przez autora pokazuje temperaturę nastawioną 65°C

Po wykonaniu czynności montażowych można przystąpić do testu działania układu. Istotnym fragmentem jest ustawienie temperatury zadanej. Tutaj, cały zakres regulacji potencjometrem powinien pokryć przedział temperatury od 40°C do 70°C. Regulacja potencjometrem daje w istocie napięcie 0 V do 5 V na wejściu analogowym A0. Mikrokontroler ATmega328 wyposażony jest w dziesięcio-bitowy przetwornik analogowo-cyfrowy. Zatem oddaje on wartość cyfrową z przedziału 0 do 1023. Tą wartość należy przełożyć na przedział 40 do 70 interpretowany już jako temperatura w stu-stopniowej skali Celsjusza (program prawdopodobnie wykorzysta tablicę LUT Look Up Table, lub inny sposób przeskalowania tych wartości). Należy sprawdzić wskazanie wyświetlacza, jak np. na rysunku 3 temperatura zadana = 65°C. Następnie należy sprawdzić reakcję switchy SET i ON. Naciśnięcie każdego z nich skutkuje stanem zera logicznego (który jest tu stanem aktywnym) na pinach 14 i 15 Arduino (wejścia analogowe A1 i A2). Naciśnięcie przycisku SET oznacza – ustaw temperaturę, i powinno mu towarzyszyć ustawienie

odpowiedniej wartości analogowej potencjometrem VR1. Ustawiona wartość powinna być zapamiętana jako temperatura z którą będzie porównywana temperatura odczytana z czujnika-termometru. Przekroczenie tej wartości powinno zakończyć proces grzania. Rozpoczęcie grzania jest zaś inicjowane naciśnięciem przycisku ON. Gdy kontroler odczyta logiczne zero na linii A2, powinien wystawić stan wysoki na A5. To uruchamia driver i przekaźnik wykonawczy. Przewidziano grzejnik zasilany jednofazowym napięciem sieci 230 VAC. Jest on włączany wykonawczymi stykami przekaźnika RL1. Wysokiemu stanowi na linii A5 powinien towarzyszyć także wysoki stan linii A4 (pin 18 płytki Arduino). Tutaj jest dioda zielona która ma informować, że proces grzania jest w toku. Mimo analogowej regulacji potencjometrem, cały proces realizowany jest po stronie cyfrowej. Tym samym dokładność obciążona jest jedynie błędem termometru DS18B20 (który gwarantuje błąd w przedziale temperatur –10°C do +85°C nie przekraczający $\pm 0,5^{\circ}\text{C}$).



Multimetry

Rozwiązanie znajdziesz na www.elportal.pl/quizy

1. Pomiar metodą Kelwina służy do mierzenia:

- ☐ bardzo dużych rezystancji, rzędu dziesiątek megaomów;
- ☐ bardzo małych rezystancji, rzędu dziesiątek miliomów;
- ☐ temperatury.

2. Cęgi w multimetrze cęgowym służą do:

- ☐ pomiaru prądu;
- ☐ pomiaru napięcia;
- ☐ wieszania multimetru.

3. Czy za pomocą multimetru smartfonowego można wykonywać połączenia telefoniczne?

- ☐ tak;
- ☐ tak, ale tylko na numery alarmowe;
- ☐ nie.

4. Multimetr półautomatyczny:

- ☐ automatycznie wybiera zakresy pomiarów, ale ich rodzaj trzeba wybrać ręcznie;
- ☐ automatycznie wybiera rodzaj pomiarów, ale ich zakresy trzeba wybrać ręcznie;
- ☐ automatycznie się wyłącza, ale włączyć trzeba go ręcznie.

5. Multimetr piórowy:

- ☐ drukuje wykresy piórkami na taśmie papierowej;
- ☐ ułatwia pomiary w trudnodostępnych miejscach na płytce drukowanej;
- ☐ jest bardzo lekki.

6. Elementy MOV i PTC zabezpieczają multimetr przed:

- ☐ zbyt wysoką składową stałą przy pomiarach sygnałów zmiennych;
- ☐ zbyt wysokim prądem;
- ☐ zbyt wysokim napięciem.

7. Funkcja hFE:

- ☐ mierzy transkonduktancję tranzystorów polowych;
- ☐ mierzy wzmocnienie tranzystorów bipolarnych;
- ☐ mierzy napięcie przebicia tranzystorów.

8. Test diody:

- ☐ sprawdza napięcie przewodzenia diod różnych typów;
- ☐ sprawdza napięcie Zenera;
- ☐ sprawdza napięcie przewodzenia diod LED.

9. Multimetr ręczny:

- ☐ jest przeznaczony do trzymania w ręku, jak sama nazwa wskazuje;
- ☐ pozwala mierzyć drobne komponenty na płytce drukowanej jedną ręką dzięki końcówkom pęsetowym;
- ☐ pozwala na ręczny wybór typu pomiaru i zakresu za pomocą dużego pokrętła.

10. Kompensacja zimnego końca termopary:

- ☐ podgrzewa złącza multimetru do stabilnej temperatury, by kompensować napięcia wytwarzane na połączeniach różnych metali;
- ☐ mierzy temperaturę przy złączach, by skompensować efekt termoelektryczny na złączach.
- ☐ odejmuje temperaturę otoczenia od temperatury „gorącego końca” by podać wynik w kelwinach.

Gdy Arduino odczyta z DS18B20 temperaturę równą lub wyższą od zadanej, wysyła zero logiczne na linię A5, co wyłącza grzałkę. Równocześnie wysyła stan wysoki na linię A3 (pin 17 Arduino) co zaświeci diodę czerwoną. Stan wysoki zostanie też wystawiony na linię D0. Tu przewidziano, iż jedynie na czas 2...3 s. Linia D0 stanowi wprost zasilanie piezoelektrycznego buzzera. Zatem, ten krótki beep sygnalizuje, że woda jest już podgrzana do założonej temperatury. ■

Ashutosh M. Bhatt

Ten artykuł był wcześniej opublikowany na łamach „EFY”, czerwiec 2023 (efymag.com)

Od Redakcji EdW: To ciekawy projekt z wykorzystaniem Arduino. Przede wszystkim dlatego, że mikrokontroler jest nieco bardziej wykorzystany aniżeli w prostszych projektach. Należy powiedzieć, nie jest całkiem bezrobotny. Aczkolwiek, oczywiście potrafi wiele więcej. Całość zaprezentowanego układu

obsługiwana jest programowo, i tu nie można mieć zastrzeżeń. Dobrze jest także, że autor postanowił zastosować izolację galwaniczną w torze włączającym grzałkę. Widzimy tu transoptor, który nie ma wielkiego sensu, gdyż izolację zapewnia sam przekaźnik. Jednak stwierdzenie nie ma wielkiego sensu to mało powiedziane. Cewka przekaźnika wysterowana jest bezpośrednio z kolektora tranzystora w transoptorze. Jaki przekaźnik potrafi on wysterować? Może taki mały jak widać na zdjęciu – rysunek 3, to zda egzamin. Dopuszczalny prąd kolektora w transoptorze MCT2E to 50 mA. Czy to wystarczy? Jeśli zastosujemy przekaźnik który potrafi włączyć grzałkę 230 V/1000 W to wątpliwe. Ale te 50 mA to też wartość mocno zawyżona. W obwodzie pierwotnym, diody transoptora jest rezystor $R1=330\ \Omega$. W stanie wysokim wyjścia A5 Arduino, przez diodę tę popłynie prąd ok. 10 mA. I jest to wartość bezpieczna. Ale CTR (Current Transfet Ratio) transoptora MCT2E jest na poziomie 50%. To daje prąd w obwodzie C-E ok. 5 mA. A i to jest wartość zawyżona, gdyż

katalog transoptora podaje, że CTR (minimum) może wynosić 20%. Prąd wyjściowy transoptora należałoby konieczności wzmocnić dodatkowym tranzystorem. Może Darlington, lub lepiej dołożyć tranzystor pnp i cewkę przekaźnika sterować od góry (to znaczy od 5 V). Praktycznie transoptor można wyrzucić. On i tak izolacji galwanicznej tutaj nie realizuje. Katoda diody i emiter tranzystora są na wspólnej masie. Można zastosować pojedynczy tranzystor npn (o odpowiednio dużym wzmocnieniu β), choć w przypadku dużego przekaźnika, także jego driver należałoby rozbudować. Nie oznacza to jednoznacznie, że w niektórych aplikacjach rozwiązanie tak asekuracyjne (z transoptorem) można usprawnić, jednak driver nie może się ograniczać do samego optocouplera.

Autor nie podaje też wielu informacji zaszytych w programie. Z jaką histerezą układ pracuje w trybie termostatu. Rozsądna jest też opcja, gdy układ nie utrzymuje zadanej temperatury. W zamian za to, za każdym razem czeka na ręczne wystartowanie przyciskiem ON_PB.

Prezentujemy 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**.

DIY PLUS

tylko dla prenumeratorów



3-fazowy sterownik silnika bezszczotkowego wykorzystujący L6235

Prezentowany projekt to sterownik 3-fazowego silnika bezszczotkowego. Projekt składa się z układu L6235, który jest w pełni scalonym 3-fazowym sterownikiem silnika DMOS z zabezpieczeniem nadprądowym. Wyprodukowany w technologii BCD układ łączy tranzystory mocy z izolowaną bramką DMOS z układami CMOS i bipolarnymi na tym samym chipie. Układ zawiera wszystkie obwody potrzebne do napędzania 3-fazowego silnika BLDC, w tym 3-fazowy mostek DMOS, kontroler prądu PWM o stałym czasie wyłączenia oraz logikę dekodowania dla pojedynczych czujników Halla, która generuje wymaganą sekwencję dla stopnia mocy. Przy włączaniu układu w trybie autonomicznym, trzy piny sterujące włączają układ, ustawiając kierunek obrotów, przerywając silnik, a zworka wybiera pracę w trybie momentu obrotowego lub prędkości. Płytką jest odpowiednia dla małych silników stosowanych w wielu urządzeniach domowych i może obsługiwać kontrolę momentu obrotowego i prędkości z wysoką wydajnością do 95%.



Sterownik solenoidu, przekaźnika, zaworu z regulacją prądu

Prezentowany projekt to sterownik prądu PWM dla solenoidów, przekaźników i zaworów. Płytką reguluje prąd za pomocą dobrze kontrolowanego kształtu fali, aby aktywować i jednocześnie zmniejszyć rozpraszanie mocy. Prąd solenoidu jest szybko zwiększany, aby zapewnić otwarcie zaworu lub przekaźnika. Po początkowej rampie, prąd solenoidu jest utrzymywany na wartości szczytowej, aby zapewnić prawidłowe działanie, po czym jest redukowany do niższego poziomu podtrzymania w celu uniknięcia problemów termicznych i zmniejszenia rozpraszania mocy.

Miesięcznik „Elektronika dla Wszystkich” (12 numerów w roku) jest wydawany we współpracy z kilkoma redakcjami zagranicznymi



Wydawnictwo:
AVT-Korporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Adres redakcji:
03-197 Warszawa, ul. Leszczyńska 11
e-mail: edw@elportal.pl, www.elportal.pl

Redaktor merytoryczny:
Mariusz Ciszewski, Paweł Sujko

Dział Reklamy:
Katarzyna Gugala
katarzyna.gugala@elportal.pl, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobanski
jakub.sobanski@elportal.pl

Copyright AVT-Korporacja Sp. z o.o., Warszawa, ul. Leszczyńska 11. Projekty publikowane w „Elektronice dla Wszystkich” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki dla Wszystkich”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice dla Wszystkich” jest dozwolone wyłącznie po uzyskaniu pisemnej zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice dla Wszystkich”.

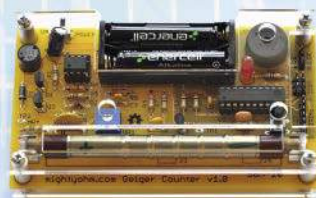
DTP, okładka, redakcja strony internetowej www.elportal.pl:
MAD Sp. z o.o.

Prenumerata:
W Wydawnictwie AVT, e-mail: prenumerata@avt.pl
tel. 22 257 84 22, (godz. 10:00–14:00)

W RUCH S.A., e-mail: prenumerata@ruch.com.pl
tel. 801 800 803, 22 717 59 59, www.prenumerata.ruch.com.pl

Elektor Bestsellers

SAVE UP TO
26% NOW!



www.elektor.com/sale/deals

