

ELEKTRONIKA

dla wszystkich

nr 2/2023 (325) • luty • www.elportal.pl

Programowalny termostat

DIY PLUS
tylko dla prenumeratorów

PROJEKTY dla elektroników

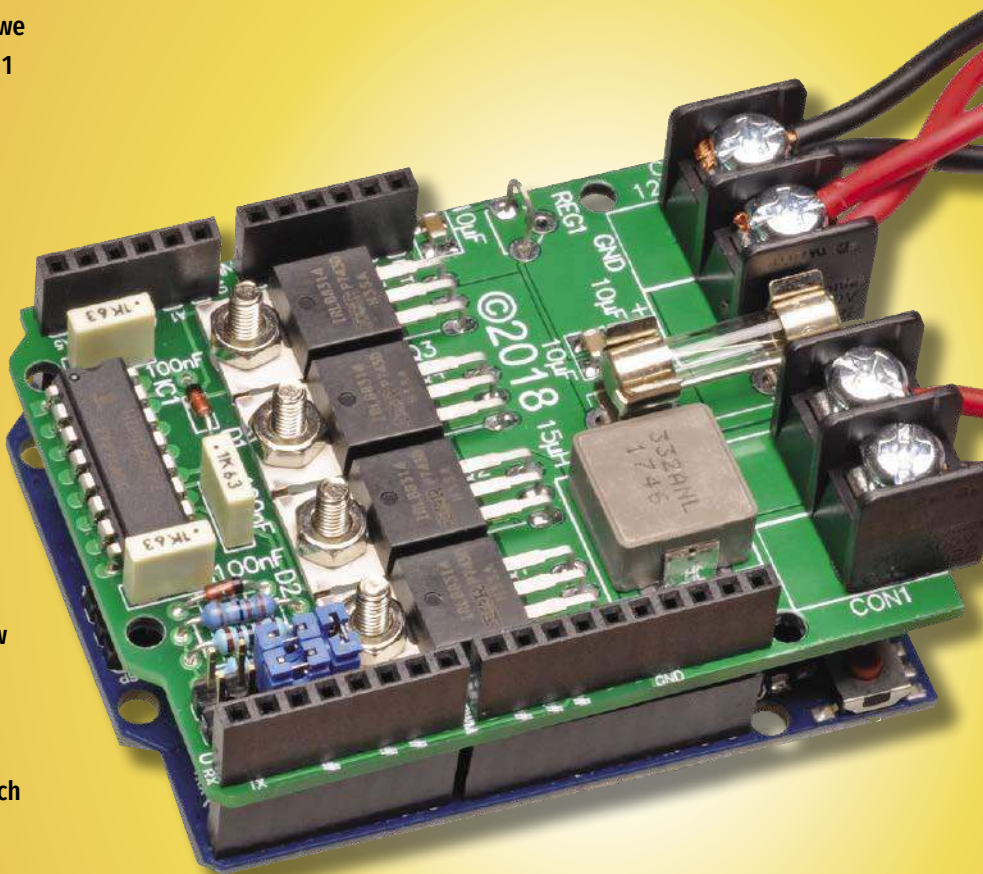
- ▶ Łatwe do budowy aktywne głośniki półkowe Hi-Fi z opcjonalnymi subwooferami, część 1
- ▶ Wysokościomierz samochodowy z ekranem dotykowym
- ▶ Analogowy automat perkusyjny

DIY dla wszystkich

- ▶ Projekty układów zwrotnic do zestawów głośnikowych
- ▶ Wykrywanie obiektów z użyciem modułu Lidar
- ▶ Sterowanie komputerem ruchem oczu
- ▶ Smart Glasses na wzór Google Glass z wykorzystaniem wyświetlacza OLED

TUTORIALE

- ▶ PE Mini-monitor – zwrotnica dla głośników Wavecor, część 1
- ▶ Poziomy logiczne, część 2
- ▶ Silniki krokowe w praktyce, część 3: Podstawowe sterowniki silników krokowych
- ▶ Pokój Nauczycielski



16,90 zł (w tym 8% VAT)



EP.com.pl

Największy
portal dla
elektroników
konstruktorów

FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

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

www.piekarz.pl



AVTEDU

Poznaj całą serię

Zupełnie nowa edukacyjna seria kitów AVTEDU. Wypróbuj je wszystkie i zostań mistrzem lutownicy, poznaj świat elektroniki i zgłębiaj go razem z nami

#AVTEDU #NaukaLutowania #KityAVT

Zestaw umożliwiający rozpoczęcie nauki techniki lutowania elementów elektronicznych. Wraz z serią kitów AVTEDU tworzy idealne uzupełnienie zagadnienia montażu prostych urządzeń elektronicznych.

Zestaw zawiera **lutownicę**, wysokiej jakości **podstawkę** z czyścikiem, **cynę** z topnikiem, **kalafonię**, **pęsety**, **odsysacz** do cyny oraz **szczypce** tnące boczne.

W komplecie na dobry początek znajduje się również **zestaw AVTEDU do zlutowania**.



AVTEDUSTART - zestaw narzędzi do nauki lutowania



sklep.avt.pl

AVT SPV Sp. z o.o., 03-197 Warszawa, ul. Leszczynowa 11
tel.: (22) 257 84 51 e-mail: handlowy@avt.pl



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

Zaprenumeruj
„Elektronikę dla Wszystkich”,
a dostaniesz bezpłatny
dostęp do archiwalnych
e-wydań EdW!

Nie dotyczy wydań z ostatnich 24 miesięcy.

na start
do 6* wydań gratis

po 5 latach
nieprzerwanej
prenumeraty
do 12* wydań gratis

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

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

Oferta dotyczy prenumeraty drukowanej.

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

Po opłaceniu prenumeraty przślemy Ci kod dostępu do projektów DIY plus na www.elportal.pl

prenumerata@avt.pl

AVT-Korporacja sp. z o.o., ul. Leszczynowa 11, 03-197 Warszawa,
konto 18 1050 1012 1000 0024 3173 1013

eprasa.pl 7c234590bd



30

Projekty dla elektroników:

Programowalny termostat	8
Łatwe do budowy aktywne głośniki półkowe Hi-Fi z opcjonalnymi subwooferami, część 1	19
Wysokościomierz samochodowy z ekranem dotykowym	30
Analogowy automat perkusyjny	37



37

Tutoriale:

PE Mini-monitor – zwrotnica dla głośników Wavecor, część 1	45
Poziomy logiczne, część 2	50
Silniki krokowe w praktyce, część 3:	
Podstawowe sterowniki silników krokowych	56
Edukacja w EdW dla szkół i uczelni: Wykład 3 „Tranzystory”	61
Pokój Nauczycielski	73



DIY dla wszystkich:

Projekty układów zwrotnic do zestawów głośnikowych	78
Wykrywanie obiektów z użyciem modułu Lidar	83
Smart Glasses na wzór Google Glass z wykorzystaniem wyświetlacza OLED	86
Sterowanie komputerem ruchem oczu	88

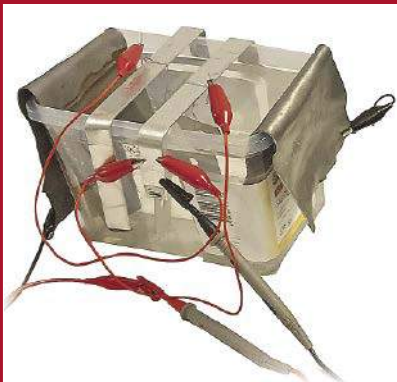
DIY PLUS

Uniwersalny wskaźnik włączonego biegu manualnej skrzyni biegów w motocyklu	91
Automatyczne Led-owe oświetlenie klatki schodowej z wykorzystaniem Arduino	91

Rubryki stałe:

Prenumerata	3
Od wydawcy	5
Pocztą	6

A za miesiąc w marcowym EdW



✳ **Myjka ultradźwiękowa dużej mocy**
To niesamowita gratka. Rzadko się zdarza projekt o tak wybitnych walorach użytkowych. Urządzenie do czyszczenia ultradźwiękami niezwykle skutecznie usuwa najbardziej nawet uciążliwe zabrudzenia z różnych przedmiotów – od zastaw domowych do częściowo przemysłowych. Usuwa smary, oleje, pozostałości metali, itp. Czyszczenie jest bezpieczne dla małych i delikatnych przedmiotów i penetruje w przestrzenie wewnętrzne.

✳ **Programowalny regulator temperatury, część 2**

Opublikowana przed miesiącem pierwsza część artykułu wywołała duże zainteresowanie. W tej części opisujemy dokładnie wszystkie szczegóły konstrukcyjne regulatora temperatury działającego na bazie elementów Peltiera ze sterownikiem na układzie Arduino.

✳ **Łatwe w budowie aktywne kolumny Hi-Fi z opcjonalnymi subwooferami, część 2**

Przed miesiącem opisaliśmy konstrukcję i działanie aktywnych kolumn, które się świetnie prezentują, wspaniale brzmią i są niedrogie. W drugiej części artykułu przedstawiamy wszystkie szczegóły dotyczące montażu i uruchomienia.

✳ **Plus zwykła porcja intrygujących projektów DIY.**

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

**W kioskach
od 1 marca**

Kto wynalazł tranzystor?

Pozornie proste pytanie. Oczywiście, że Amerykanie John Bardeen, Walter Brattain i William Shockley, którzy za ten wynalazek otrzymali Nagrodę Nobla z fizyki w 1956 roku.

Za datę wynalezienia tranzystora przyjmuje się 23.12.1947 r., gdy W. Brattain zademonstrował szefem firmy Bell Telephone Laboratories efekt wzmocnienia (a dzień później generacji) sygnału uzyskanego na eksperymentalnym układzie zbudowanym z kryształu germanu. Był to prototyp elementu, który nazwano **point-contact transistor** (polska nazwa **tranzystor ostrzowy**).

Gdy wnikiemy w szczegóły odpowiedzi na pytanie postawione w tytule nie jest tak oczywista. Po pierwsze, w eksperymentach nie uczestniczył W. Shockley, który był szefem W. Brattaina i J. Bardeena. Nie oznacza to, że nie należała mu się Nagroda Nobla. Jemu należała się najbardziej, bo miesiąc później wynalazł **tranzystor złączowy bipolarny** i opracował teorię półprzewodników, która obowiązuje do tej pory, a jego książka z 1950 r. *Electrons and Holes in Semiconductors* stała się biblią dla całej generacji elektroników. J. Bardeen (uhonorowany jeszcze jedną Nagrodą Nobla za prace dotyczące nadprzewodnictwa) był fizykiem teoretykiem, któremu eksperymenty wykonywane przez W. Brattaina służyły do potwierdzenia jego teorii stanów powierzchniowych. A W. Brattain wykonując eksperymenty dla J. Bardeena przypadkowo zauważył, że eksperymentalna konstrukcja daje wzmocnienie.

O co chodziło w tych eksperymentach? Otóż zespół W. Shockleya pracował nad **tranzystorem polowym**, wzorowanym na patencie J.E. Lilienfelda z 1925 r. Za chwilę jeszcze wrócę do tego patentu. Skonstatujmy tylko fakt, że tranzystor ostrzowy wynaleziono przypadkowo podczas badania stanów powierzchniowych, stanowiących zasadniczą przeszkodę dla opracowania działającego tranzystora polowego.

Gwoli sprawiedliwości zauważmy, że tranzystor ostrzowy wynaleziono niemal równocześnie we Francji, a Amerykę mogły wyprzedzić o kilka lat Niemcy. Już w 1944 r. niemiecki fizyk Herbert Mataré eksperymentował z tzw. **duodiodą**, tj. kryształem półprzewodnika z dwoma elektrodami ostrzowymi. W latach 1944–45 zaawansowane prace nad kryształami germanu prowadził inny fizyk niemiecki Heinrich Welker. Zawierucha wojenna przerwała te prace, a w roku 1946 agenci francuskich władz okupacyjnych odnaleźli obu fizyków i zaproponowali im pracę we francuskim laboratorium Westinghousa. Propozycję przyjęli, ale dopiero w 1947 roku wrócili do przerwanych prac i w czerwcu 1948 roku ogłosili o skonstruowaniu udoskonalonej duodiody, czyli tranzystora ostrzowego, który we Francji wszedł do produkcji pod nazwą **transistron**. Mataré i Welker na pewno nie wiedzieli o pracach w Bell Labs, bo do lata 1948 roku były one utajnione.

Jeszcze jeden kraj pretenduje do pierwszeństwa. W ZSRR kiedyś twierdzono, że w zespole kierowanym przez akademika A.F. Joffe wynaleziono tranzystor przed Amerykanami. Potem to twierdzenie osłabiono. W rosyjskiej Wikipedii można przeczytać „ЕСЛИ БЫ НЕ НЕОБХОДИМОСТЬ СОЗДАНИЯ АТОМНОГО ОРУЖИЯ, ОТКРЫТИЕ ТРАНЗИСТОРОВ МОГЛО ПРОИЗОЙТИ В СССР”. To znaczy: gdyby nie konieczność opracowania broni jądrowej, to odkrycie tranzystora nastąpiłoby w ZSRR.

Tu przypomnia mi się wykład z matematyki na Wydziale Łączności Politechniki Warszawskiej (lata sześćdziesiąte), gdy Docent Trajdos rysuje na tablicy funkcję skoku jednostkowego i pisze „funkcja Heaviside’a” zwracając się do studentów – nie pytajcie, jak się to nazwisko czyta – Hevisajda, czy Hivisajda. Kiedyś mnie student spytał, więc doradziłem mu, żeby zajrzał do książki rosyjskiej (były w bibliotece PW), to się dowie, bo Rosjanie piszą nazwiska fonetycznie. Student zajrzał i zameldował – to jest funkcja Zacharenko-Waszczenko.

Popatrzmy na pytanie tytułowe jeszcze z innej strony. Dlaczego zafiksowaliśmy się na wynalazku tranzystora ostrzowego? Przecież nie odegrał on żadnej roli, choć w niewielkich ilościach był produkowany przez kilkanaście lat. To była ślepa uliczka. Natomiast cała mikroelektronika dużej skali integracji działa na tranzystorach MOS, których protoplastą był patent Lilienfelda z 1925 r. (Kanada) i 1926 r. (USA). Dlaczego temu wynalazkowi nie przypisaliśmy pierwszeństwa światowego. Odpowiedź na poważnie jest prosta – wynalazek tranzystora ostrzowego zapoczątkował produkcję tranzystorów, której rozwój zmienił oblicze tego świata. Natomiast wynalazek Lilienfelda spoczął na półkach urzędów patentowych, choć po latach dał impuls do prac zespołowi W. Shockleya w Bell Labs. Dodam też, już nie całkiem serio, że z Lilienfeldem byłby problem, który kraj ma się nim szcycić. Urodził się w Lwowie w 1882 r. Po odzyskaniu przez Polskę niepodległości został obywatelem polskim, studiował i karierę naukową robił w Niemczech. Patent kanadyjski złożył jako obywatel Polski. W Wikipedii anglojęzycznej jest fizykiem austro-węgierskim i amerykańskim. W polskiej jest polskim fizykiem pochodzenia żydowskiego. W rosyjskiej jest fizykiem niemieckim.

Podobny problem był z Einsteinem, który przemawiając w 1930 r. w Sorbonie powiedział: „Jeśli moja teoria względności okaże się słuszna, to Niemcy ogłoszą, że jestem uczonym niemieckim, a we Francji będą twierdzić, że jestem obywatelem świata. Gdyby jednak moja teoria okazała się fałszywa, to Francja będzie twierdzić, że jestem uczonym niemieckim, a Niemcy ogłoszą, że jestem Żydem”. No comment.

Wiesław Marciniak

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

Szanowny Panie Profesorze

Jestem nauczycielem akademickim, czuję się wychowankiem niedawno zmarłego Profesora Andrzeja Jakubowskiego. Wiem z rozmów z Profesorem Jakubowskim, że przyjaźniliście się Panowie. Niedawno od studenta dowiedziałem się, że zaczął Pan publikować wykłady w miesięczniku Elektronika dla Wszystkich. Z zacięciem zapoznałem się z wykładem o układach scalonych. Zaskoczyło mnie, że w tak popularnym miesięczniku pojawiają się wykłady nie stroniące od trudnych problemów. Jak Pan zapewne pamięta Profesor Jakubowski przez wiele lat zajmował się problemami skalowania tranzystorów MOS i fundamentalnymi ograniczeniami zasilania układów CMOS poniżej 1 V, bo napięcie progowe trzeba by zmniejszyć poniżej 300 mV, aby zachować odpowiednią relację prądu drenu w stanie przewodzenia i nieprzewodzenia. Prąd przewodzenia powinien być co najmniej 5 dekad większy od prądu nieprzewodzenia, a trudno to zapewnić, gdyż przy zerowym napięciu bramki i niskim napięciu progowym (ok. 300 mV) nie dochodzi do pełnego zatkania tranzystora, bo tranzystor pracuje w tak zwanym zakresie podprogowym i płynnie liczący się prąd drenu. Jednak pomysłowość człowieka nie ma granic. W ostatnich latach wymyślono tranzystory z ujemną pojemnością w postaci warstwy ferroelektryka, co pozwala lepiej zatkać tranzystor przy zerowym napięciu bramki. Myślę jednak, że to niewiele pomoże i prawo Moore'a kończy swój żywot. A o ujemnej pojemności mogę napisać artykuł do Elektroniki dla Wszystkich, jeśli uzna Pan, że to nie jest za trudny temat dla tego pisma.

J.S.

Pytanie o tożsamość

...na początku lat siedemdziesiątych, jeszcze za wczesnego Gierka studiowałem w Wojskowej Akademii Technicznej. Wykładowcą półprzewodników był kapitan Wiesław Marciniak, młody, prawie rówieśnik słuchaczy, ale świeżo po doktoracie. Było to 50 lat temu. Pod wykładem o mikroprocesorach w grudniowym EdW zobaczyłem podpis autora – Wiesław Marciniak. Czy to możliwe, że Pan jest tą samą osobą, którą znałem z WAT-u? Jeśli to krępujące pytanie, to proszę nie odpowiadać, ale ciekaw jestem bardzo.

P.K.

Emerytowany oficer

Redakcja

Tak, to jestem ja. Ten sam, ale nie taki sam. Mój debiut autorski w EdW wywołał lawinę (no, może lawinkę) listów (lekką ponad 20) z różnymi pytaniami. Nie wypada mi publikować tych listów, bo kierowane są bardziej do mnie niż do Redakcji gazety. A propos Redakcji, często jestem pytany, dlaczego nie umieszczam siebie w stopce jako Redaktora Naczelnego. To proste, działam jako wydawca, wciąż pracuję nad profilem gazety i źródłami atrakcyjnych materiałów. Chyba idzie to nieźle, bo na koniec 2022 roku uzyskaliśmy wzrost sprzedaży EdW o 43% w stosunku do poprzedniego roku. Najbardziej cieszy wzrost prenumeraty. To stabilizuje naszą pozycję rynkową, więc każdego Czytelnika namawiamy, żeby został prenumeratorem. Wśród kilku korzyści, jakie daje prenumerata, chcę zwrócić uwagę na bezpłatny dostęp do wydań archiwalnych. Tradycyjnie otwieraliśmy dostęp do archiwaliów na ograniczone okresy, teraz otwieramy bezterminowo – ramka niżej.

Przypominamy, że

Prenumeratorki EdW

mają dostęp bezpłatny do e-wydań archiwalnych EdW (dostęp do e-wydań archiwalnych innych tytułów AVT z dużym rabatem). Tym razem otwieramy dostęp do e-wydań EdW bezterminowo. Nie dotyczy wydań z ostatnich 24 miesięcy.

Wydawnictwo AVT

Patronat AVT

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

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



Najbardziej popularne kity AVT

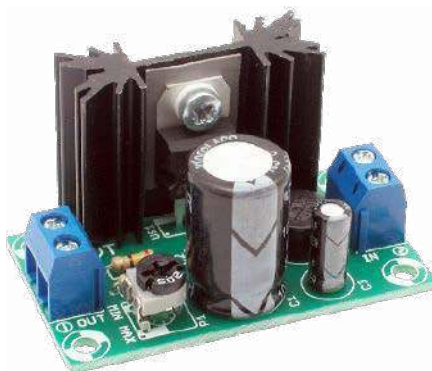
Poznaj listę **TOP 100** na www.elportal.pl/kityavt



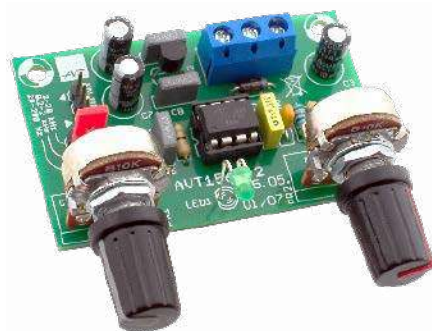
AVT1476 Automatyczny włącznik zmiernicowy
<https://sklep.avt.pl/avt1476.html>



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



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



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



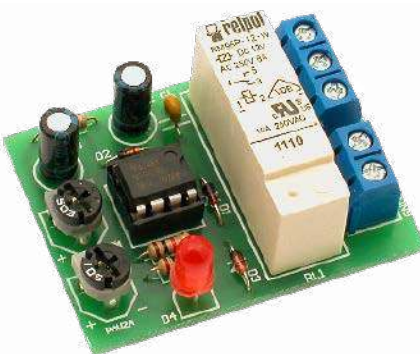
AVT1023 Przedwzmacniacz gramofonowy o charakterystyce RIAA
<https://sklep.avt.pl/avt1023.html>



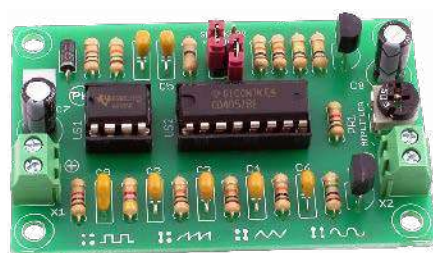
AVT5540 Radio FM z RDS
<https://sklep.avt.pl/avt5540.html>



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



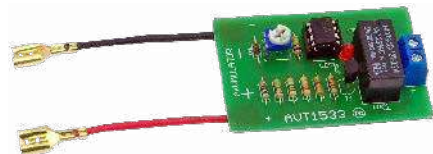
AVT1459 Uniwersalny układ czasowy
<https://sklep.avt.pl/avt1459.html>



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



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



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



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



Pełna oferta na: sklep.avt.pl

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



Materiały dodatkowe dostępne są na stronie
Silicon Chip: <http://bit.ly/3G6mccS>
Materiały dodatkowe są również dostępne
na stronie edw.elportal.pl: <https://bit.ly/3G0zrIn>

Programowalny termostat. Pomożemy Ci uwarzyć to piwo... lub cokolwiek innego, co wymaga stałej temperatury

Precyzyjna kontrola temperatury jest integralną częścią wielu procesów przemysłowych. Jeśli jesteś zainteresowany procesami warzenia czy fermentacji, aby jako hobbysta wytwarzać żywność na własny użytek, przekonasz się, że dla otrzymania jak najlepszych wyników niezwykle ważne jest utrzymanie odpowiedniej temperatury procesu. Nawet nam zdarzały się próby zrobienia własnego sera, piwa czy cydru (oczywiście nie w biurze!).

Cechy

- Aktywnie chłodzi i grzeje
- Steruje modułami Peltiera o łącznej mocy ponad 200 W
- Wykorzystuje wiele czujników temperatury
- Wykorzystuje moduł Arduino, co zapewnia elastyczność trybów pracy

W przypadku piwa, słód jęczmienny poddawany jest fermentacji drożdżowej w celu wytworzenia alkoholu i nadania smaku. Proces fermentacji nasycza również gotowy produkt dwutlenkiem węgla, dlatego piwo musi być i ma orzeźwiający smak.

W przypadku domowego piwa lub cydru fermentacja odbywa się w plastikowym pojemniku przystosowanym do kontaktu z żywnością. Dobre wyniki można osiągnąć po prostu trzymając naczynie w pomieszczeniu, w którym temperatura nie ulega większym zmianom, ewentualnie owijając je kocem w chłodniejszych miesiącach.

Jednak dla zachowania ciągłości procesu i pewności, że fermentacja zakończy się prawidłowo (w innym przypadku butelki mogą nawet eksplodować!), potrzebny jest sposób nadzorowania i sterowania temperaturą brzeczki. Właściwa regulacja temperatury jest jednym z czynników, dzięki którym komercyjne browary mogą gwarantować identyczny smak każdej partii piwa.

W przypadku domowej fermentacji nawet trzymanie naczynia do warzenia piwa w pomieszczeniu z termostatem czy klimatyzacją może nie wystarczyć.

W miarę postępu fermentacji aktywność drożdży raz rośnie, raz maleje. Ciepło wytwarzane w tym procesie może zmienić temperaturę piwa od wewnątrz, nawet jeśli warunki na zewnątrz są stałe.

Potrzebujemy więc narzędzia do pomiaru, jak i regulacji temperatury zacieru.

Wybraliśmy do tego celu ogniwa Peltiera, ponieważ mają one zdolność zarówno

ogrzewania, jak i chłodzenia; wymagają jedynie niskiego napięcia zasilania DC i łatwo się nimi steruje. Nie są to najbardziej wydajne urządzenia, ale do operacji na małą skalę nadają się idealnie.

Gotowanie metodą „Sous-vide”

Termostat znajduje także zastosowanie przy gotowaniu metodą „sous-vide”. Choć francuski

Możliwe zastosowania

- Produkcja serów
- Warzenie piwa/wina/cydru/kombucy (grzyba herbacianego)
- Temperowanie czekolady
- Gotowanie metodą „sous-vide”
- Chłodzenie komputera
- Laboratoryjna łaźnia wodna
- Akwaria (zwłaszcza duże tropikalne)
- Klimatyzator osobisty
- Ulepszone chłodzenie dla wycinarek laserowych

termin „sous-vide” tłumaczy się dosłownie jako „pod próżnią”, próżnia nie jest tu decydująca. Powodzenie gotowania „sous-vide” zależy głównie od precyzyjnej kontroli temperatury.

Później zajmiemy się tym nieco bardziej szczegółowo. Na razie chcemy podkreślić, że utrzymywanie temperatury na ściśle określonym poziomie prowadzi do oczekiwanych i powtarzalnych wyników.

Utrzymanie odpowiedniej temperatury przez wystarczająco długi czas daje gwarancję, że wszelkie bakterie zostaną zabite, a tym samym jedzenie będzie bezpieczne.

Do innych procesów kulinarnych, w których sprawdza się dobrze precyzyjna regulacja temperatury, należy temperowanie czekolady. Poddanie czekolady działaniu ściśle określonej temperatury zmienia jej strukturę i sprawia, że po stwardnieniu czekolada ma połysk i kruchość.

Jedną z fascynujących możliwości naszego urządzenia jest to, że można go użyć do utrzymywania żywności w bezpiecznej temperaturze przechowywania (około 4°C, jak we wnętrzu lodówki) przez wiele godzin, a następnie o zaprogramowanej porze podgrzać ją i ugotować, tak aby była gotowa do spożycia.

Jeśli chcesz zastosować termostat do tego celu, sugerujemy zmodernizowanie oprogramowania tak, aby alarmowało ono, jeśli temperatura jedzenia w trybie przechowywania znacznie przekroczy 4°C. Dzięki temu będziesz mieć pewność, że jest ono bezpieczne do spożycia.

Inne zastosowania

Osoby, które pracowały w laboratoriach, zetknęły się na pewno z laboratoryjną łaźnią wodną stosowaną do utrzymywania próbek w stałej temperaturze. Nasz termoregulator nadaje się w tym przypadku również do tego zastosowania.

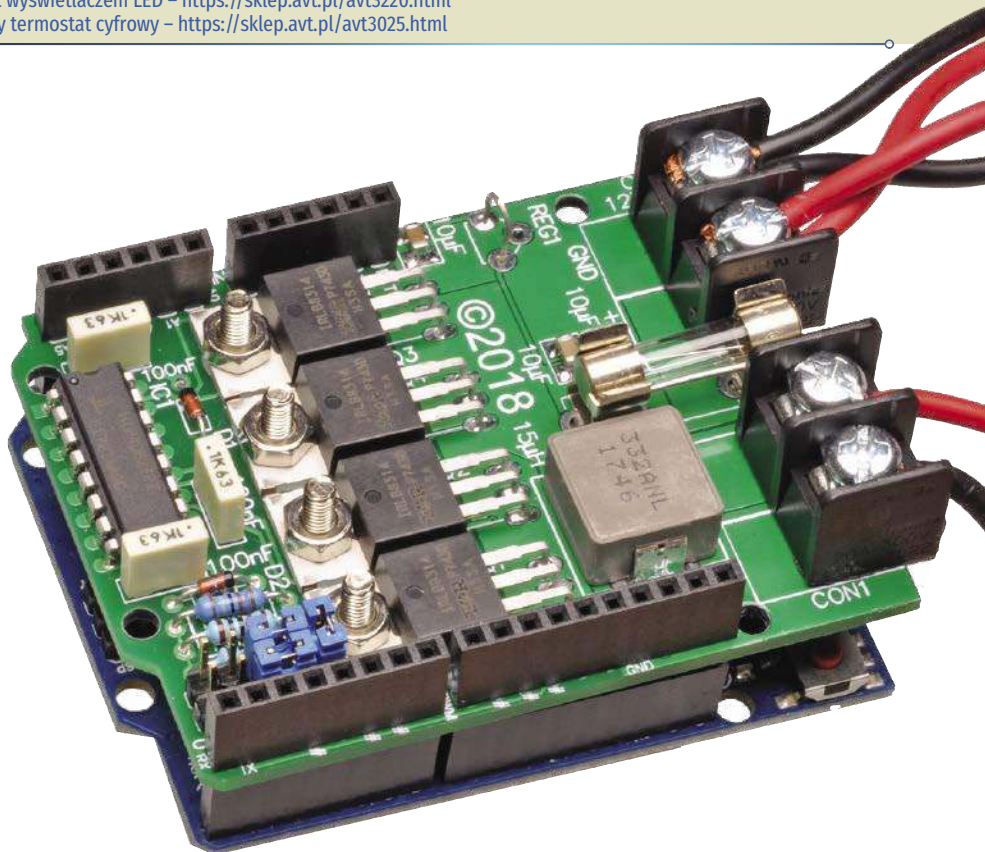
Żartowaliśmy nawet, że moglibyśmy użyć termoregulatora jako osobistego klimatyzatora. Rozgrzana z jednej strony płytka wytwarza po przeciwnej stronie orzeźwiający bryzę, uważamy więc, że takie zastosowanie byłoby całkiem sensowne.

Elektronika termoregulatora

Elektronika termoregulatora składa się z trzech głównych części. Płytki Arduino Uno (lub równoważna) zapewnia sterowanie całością, jak również zawiera niektóre obwody regulatora mocy.

Nakładka sterownika Peltiera (płytki dodatkowa do Arduino), nadzorowana przez Arduino, zawiera mostek H o dużej mocy. Jest on używany do zasilania ogniw Peltiera.

Druga nakładka (moduł interfejsu) posiada liczne wejścia i wyjścia; zajmuje się przede



Nie mogliśmy zmieścić wszystkich elementów na jednej nakładce, więc są dwie! Ta nakładka (dołączona do płytki Uno) jest zaprojektowana do zasilania ogniw Peltiera prądem do 20 A w trybie mostka H, co oznacza, że kierunek przepływu prądu może zostać odwrócony i moduł Peltiera może być użyty do grzania lub do chłodzenia. Na tej nakładce znajduje się wiele części lutowanych powierzchniowo, ale żadna z nich nie jest szczególnie mała, więc montaż nie jest trudny.

wszystkim wykrywaniem tego, co dzieje się z ogniwami Peltiera, i może również zasilac inne urządzenia, takie jak pompy i wentylatory.

O tych sprawach napiszemy później. Do wykonania projektu potrzebna jest znajomość oprogramowania Arduino IDE; można je pobrać za darmo z siliconchip.com.au/link/aaatq lub strony projektu Arduino.

Ponieważ moduł sterownika ma tak wiele potencjalnych zastosowań, zaprojektowaliśmy go tak, aby był jak najbardziej elastyczny. Przed zapoznaniem się z dalszą częścią artykułu warto przeczytać wyróżnioną ramkę w której opisujemy działanie ogniw Peltiera.

Inspiracja do tego artykułu

Pomysł na tę serię artykułów narodził się z rozważań o projektach takich jak chłodziarka Peltiera („Chłodziarka płytkowa”) z 2003 roku.

Tamten projekt obejmował dość duży radiator i wentylator, dołączone do pojedynczego modułu Peltiera, które miały sprawnie odprowadzić całe generowane ciepło i zapewnić ogniwu Peltiera wydajną pracę. Jeśli używasz kilku płytek Peltiera, aby wymoc przepompowanie większej ilości ciepła, będziesz potrzebować ogromnego radiatora.

Choć jest to najprostsze i stosunkowo tanie rozwiązanie, nie jest ono idealne.

Weźmy na przykład silnik samochodowy, który wytwarza ogromne ilości ciepła (w niektórych przypadkach setki kilowatów). Choć wczesne silniki były chłodzone powietrzem, większość producentów szybko przeszła na chłodzenie cieczą. O wiele łatwiej jest bowiem usunąć całe ciepło za pomocą odrobiny wody, która następnie trafia do dużej chłodnicy posiadającej wystarczającą powierzchnię, aby przekazać ciepło do powietrza.

Pomyśleliśmy więc: „Dlaczego nie zastosować tych samych zasad do ogniw Peltiera?” Małe chłodnice wodne używane w komputerach są obecnie łatwo dostępne w umiarkowanej cenie, a niezbędne wentylatory, pompy i rurki również nie kosztują dużo. Kupiliśmy więc kilka części i przeprowadziliśmy serię eksperymentów, które doprowadziły nas do wniosków i rozwiązań, które tutaj prezentujemy.

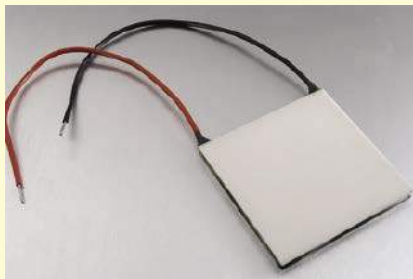
Przykład zastosowania

Dobrym przykładem do zademonstrowania tego, co może osiągnąć nasz sprzęt, jest gotowanie metodą „sous-vide”.

Jak już wspomnieliśmy, termin ten znaczy dosłownie „pod próżnią”. Ma to niewiele wspólnego z procesem obróbki, poza tym, że produkty, które mają być gotowane (zazwyczaj mięso, ryby lub jajka) przed zanurzeniem w kąpieli wodnej o kontrolowanej

Jak działają ogniwa Peltiera

Ogniwo Peltiera jest rodzajem elektrycznej pompy ciepła bez ruchomych części. Prąd elektryczny przepływający przez urządzenie powoduje przemieszczanie się ciepła z jednej strony na drugą. Ogniwo składa się z jednego lub wielu złączy różnych metali, do których przyłożone jest napięcie. Ogólna budowa takiego urządzenia przedstawiona jest na załączonym rysunku.



Prawa termodynamiki nie pozwalają na „twożenie” ciepła lub „chłodu” z niczego; procesy cieplne są jedynie konsekwencją przemieszczania energii cieplnej z jednego miejsca do drugiego lub transformacji pomiędzy energią cieplną a innymi rodzajami energii. Przykładowo, grzejniki elektryczne zamieniają energię elektryczną na energię cieplną w stosunku 1:1. Niestety, zasada wzrostu entropii oznacza, że musimy wydatkować dodatkową energię, aby przenieść tę energię cieplną. Dlatego proces ten nie może przebiegać ze 100% sprawnością.

Odwrotnością efektu Peltiera jest efekt Seebecka, w którym różnica temperatur zamieniana jest na napięcie. Energia dostarczana przez to napięcie pochodzi z energii cieplnej przepływającej od strony gorącej do zimnej. Jest to efekt wykorzystywany przez termopary i termostaty mierzące temperaturę.

Efekt Seebecka można zaobserwować również w urządzeniach Peltiera, choć nie są one projektowane z myślą o tym i dlatego nie są zbyt wydajne. Na przykład, jeśli do urządzenia Peltiera przyłożymy na kilka sekund prąd (wystarczający do wywołania różnicy temperatur obu stron), a następnie odłączymy go, na zaciskach płytki można zmierzyć napięcie. Wynika to z efektu Seebecka, czyli energii elektrycznej generowanej z resztkowej różnicy temperatur.

Urządzenie Peltiera składa się z układu naprzemiennych materiałów, w wyniku czego powstają naprzemienne złącza o przeciwstawnym zachowaniu. Są one ułożone tak, że ciepło jest przekazywane z jednej strony na drugą, poprzez utrzymywanie każdego typu złącza po właściwej stronie.

Ostatni projekt wykorzystujący ogniwo Peltiera opublikowaliśmy w 2003 roku (siliconchip.com.au/Article/3969). Polegał on na dodaniu aktywnego chłodzenia do małej chłodziarki, aby pomóc w schłodzeniu puszek z napojami. Urządzenie mogło być również użyte jako grzejnik; ponieważ jedną z zalet efektu Peltiera jest to, że jest on odwracalny. Jeśli kierunek prądu jest odwrócony, to ciepło płynie w przeciwnym kierunku.

Być może używałeś tego typu chłodziarki. Spisują się one całkiem nieźle, ale większość z nich nie stanowi konkurencji dla zwykłej domowej lodówki lub klimatyzatora, które wykorzystują sprężarkę i nie wykazują skutków ubocznych opisanych poniżej.

Zaletą ogniwa Peltiera jest ich odwracalność i brak ruchomych części. Mają one też jednak i swoje minusy. Największym jest to, że materiały, które zapewniają najsilniejszy efekt Peltiera, nie są dobrymi izolatorami termicznymi; w efekcie ciepło może uciekać z gorącej strony z powrotem do zimnej. Efekt ten jest tym silniejszy, im większa jest różnica temperatur w urządzeniu.

Praktyczne urządzenia Peltiera są zwykle wykonane z materiałów półprzewodnikowych o skończonej rezystancji. Jako takie, podlegają one również ogrzewaniu rezystancyjnemu z powodu przepływającego przez nie prądu. Jest to obliczane jako I^2R , więc podwojenie prądu spowoduje czterokrotne

zwiększenie strat rozpraszania. Jednak ilość ciepła, które jest pompowane, jest proporcjonalna do natężenia prądu, dlatego ogniwa Peltiera działają najlepiej, gdy wymagania wobec nich są niewielkie.

Ponadto ogniwa Peltiera są zazwyczaj wykonane z kruchej ceramiki. Jest ona niezbędna do zapewnienia izolacji elektrycznej, a jednocześnie pozwala na efektywne odprowadzanie ciepła od powierzchni roboczych. Przykładem takiej ceramiki jest

spiekany azotek boru BN, niestety bardzo drogi.

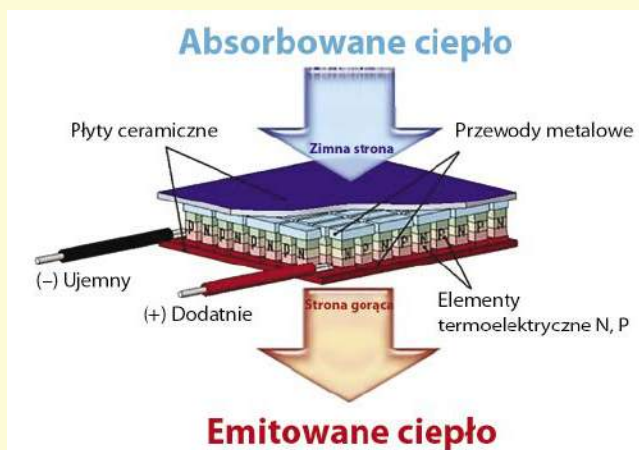
Bezpieczne używanie ogniw Peltiera

Szybkie zmiany natężenia prądu mogą powodować w ogniwie gradient temperatury; wynikające z tego zmiany temperatury mogą prowadzić do naprężeń termicznych, a nawet pęknięcia. Używając technik takich jak PWM (modulacja szerokości impulsu) do regulacji prądu należy postępować ostrożnie, aby uniknąć uszkodzeń. Niezbędnym minimum jest ustawienie wysokiej częstotliwości PWM, aby zapobiec takim efektom.

Wielu producentów urządzeń Peltiera podaje, że powinno być do nich dostarczane napięcie o niskiej zawartości tętnień (rzędu 5–10%). Aby uzyskać optymalne rezultaty, należy zastosować napięcie stałe bez składowej zmiennej.

Jest jeszcze jeden powód, aby unikać PWM. Rozważmy przypadek, gdy do ogniwa Peltiera podajemy napięcie 6 V DC nie zawierające tętnień, w porównaniu do 12 V DC z zasilacza PWM przy 50% wypełnieniu. Kiedy patrzymy na straty I^2R , widzimy, że są one podwojone w przypadku napięcia 12 V. Choć cykl pracy z 50% wypełnieniem oznacza, że moc jest stosowana przez połowę czasu, podwojenie napięcia oznacza, że efekt I^2R jest czterokrotnie większy.

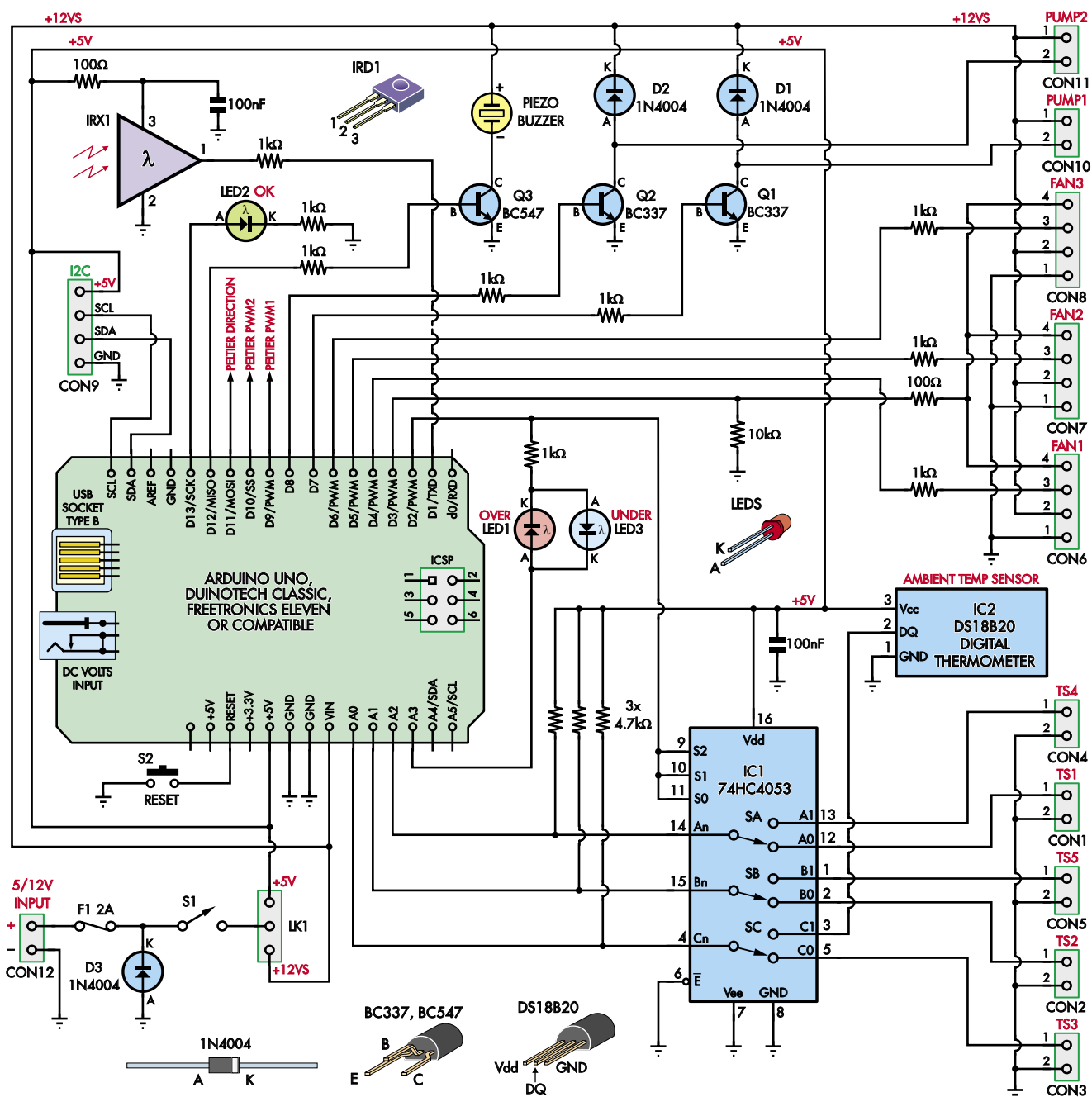
Nasza nakładka sterownika Peltiera została zaprojektowana z uwzględnieniem powyższych czynników. Dostarcza ona niemal jednostajny prąd stały nie zawierający tętnień w pełnym zakresie napięć dodatnich i ujemnych, umożliwiając zarówno grzanie jak i chłodzenie. Ułatwia to również pracę źródła zasilania!



Ogniwo Peltiera jest zwykle wykonane z matrycy półprzewodnikowej, której złącza są elektrycznie połączone szeregowo, ale termicznie równolegle, ze względu na sposób ułożenia. Dzięki temu, po przyłożeniu napięcia, ciepło przepływa z jednej strony na drugą, w zależności od polaryzacji napięcia. Źródło rysunku: <https://cpb-us-e1.wpmucdn.com/sites.suffolk.edu/dist/f/759/files/2014/02/2.jpg>



eprasa.pl 7c234590bd



NAKLADKA INTERFEJSU REGULATORA TERMICZNEGO

Rysunek 2. Nakładka interfejsu monitoruje do pięciu termistorów i może zasilać kilka pomocniczych urządzeń 12 V, które mogą być potrzebne, w tym wentylatory i pompy. Multiplexer IC1 pozwala poprzez wejścia analogowe odczytywać sygnały z sześciu czujników temperatury, ponieważ niektóre wejścia analogowe są zarezerwowane dla komunikacji szeregowej I²C.

scalonego sterującego MOSFET-ami w mostku H. Ma on również maksymalne VDD wynoszące 15 V, chociaż może sterować mostkiem, który przełącza napięcia do 80 V.

Wejścia sterujące IC1 ustawia się zworkami JP1, JP2, LK3, LK4. Pozwalają one na podłączenie styków wejściowych IC1 w różnych kombinacjach do różnych portów obsługujących PWM na płytce Uno. Dwa rezystory 10 kΩ polaryzujące wejścia ALI, BLI do masy zapewniają, że porty pozostają nieaktywne

(z wyłączonym mostkiem H), gdy moduł Uno jest zresetowany lub nie jest zaprogramowany, itp.

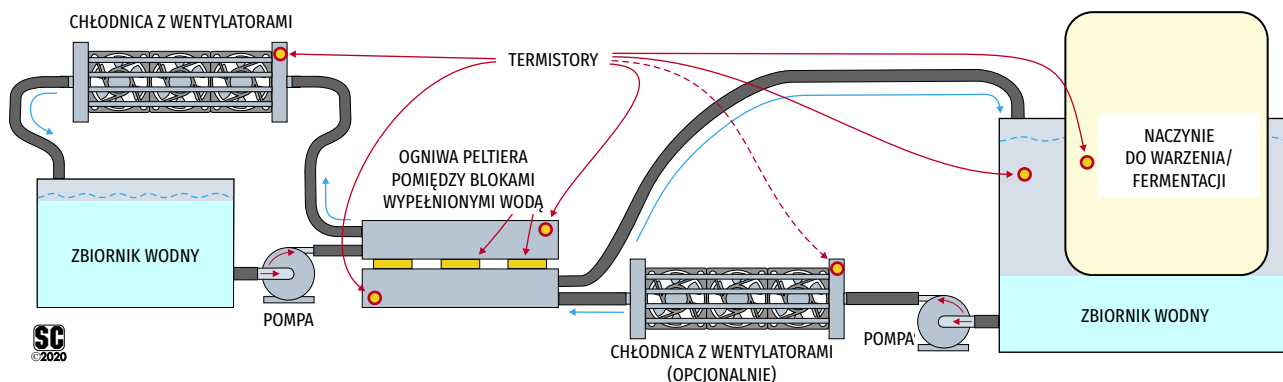
Rezystor 1,8 kΩ łączący port DEL układu IC1 (styk 5) z masą ustawia opóźnienie włączenia, a tym samym czas martwy MOSFET-ów na około 200 ns.

Diody D1 i D2 oraz związane z nimi kondensatory 100 nF tworzą obwody „startowe”, które zapewniają wystarczająco wysokie napięcia doysterowania bramek MOSFET-ów Q2

i Q4 usytuowanych po stronie zasilania, wykorzystując wyjściowy sygnał prostokątny do utworzenia pompy ładunkowej.

Od Red. EdW: zagrożenie przełączania MOSFET-ów po stronie zasilania i po stronie masy omówione zostało także w opisie regulatora silnika DC dużej mocy, EdW 10/2022, str. 9 i następne.

IC1 posiada również swój własny kondensator bocznikujący zasilanie o wartości 100 nF.



Rysunek 3. Ten schemat pokazuje jak termoregulator mógłby być użyty do sterowania kuchenką „sous-vide” albo do produkcji sera, fermentacji piwa lub wina. Chociaż dwa obiegi wodne sprawiają, że sprzęt jest nieco bardziej złożony, to jednak dzięki temu jest w stanie „przepompować” więcej ciepła, co jest niezbędne do osiągnięcia wyższych temperatur potrzebnych do gotowania.

Cztery N-kanalowe MOSFET-y typu IRLB8314 Q1-Q4 pracują w konfiguracji mostka H.

Mogą one przy odpowiednim chłodzeniu przełączać przy napięciu do 30 V i prądzie ponad 100 A, chociaż prąd jest ograniczony przez inne fragmenty modułu, takie jak złącza i obciążalność ścieżek PCB.

Zastosowanie mostka H oznacza, że kierunek przepływu prądu może być odwrócony, a cykl pracy może być również sterowany poprzez szybkie przełączanie mostka H pomiędzy oboma, przeciwnymi kierunkami przewodzenia.

Sterownik (IC1) jest potrzebny, ponieważ MOSFET-y przełączane po stronie zasilania są odmianami N-kanalowymi.

Dlatego też ich bramki muszą pracować przy potencjale wyższym od potencjału ich źródeł, tj. napięcia szyny zasilającej (w stanie przewodzenia). Realizację tego celu umożliwia obwód „podbicia” („bootstrap”). Sterownik zapewnia także, że pojemności bramek MOSFET-ów mogą być szybko ładowane i rozładowywane, tak aby zapewnić wysoką częstotliwość PWM. Dzięki temu możemy ją filtrować, aby uzyskać jednostajne, wygładzone napięcie na płytkach Peltiera, co jest warunkiem ich bezawaryjnej pracy – patrz objaśnienia.

Niska rezystancja MOSFET-ów w stanie przewodzenia, wynosząca około 2,4 mΩ, oznacza, że wymagane jest minimalne chłodzenie tranzystorów. Przy niewielkich prądach (do około 20 A) wystarczające chłodzenie zapewnia sama płytka drukowana.

Pomiędzy wyjściem mostka H a złączem wyjściowym CON1 znajduje się filtr dolnoprzepustowy LC składający się z cewki L1 o pojemności 3,3 μH i wielowarstwowego kondensatora ceramicznego (MLC) o pojemności 10 μF. Tworzy to rodzaj prostego wygładzania obniżonego napięcia z przetwornika DC/DC, jakim faktycznie jest układ PWM.

Gdy na wejścia sterujące układu IC1 zostanie podany sygnał PWM o odpowiednio

wysokiej częstotliwości (około 300 kHz), na wyjściu otrzymamy praktycznie prąd stały. Oznacza to również, że prąd pobierany z nominalnej szyny 12 V jest efektywnie na stałym poziomie, nie są więc wymagane żadne duże kondensatory filtrujące na płycie.

Jednym ze sposobów analizy tego układu jest założenie, że ogniwa Peltiera mają efektywną rezystancję około 1 Ω (12 A @ 12 V).

Możemy wtedy obliczyć, że sygnał PWM o częstotliwości 300 kHz jest tłumiony ponad 100 razy (około 40 dB), a więc tętnienia są utrzymywane znacznie poniżej zalecanej dla ogniwa wartości 5%.

Ta nakładka nadaje się w każdym przypadku, w którym wymagane jest ustawianie, stosunkowo dobrze wyglądzone wysokoprądowe niestabilizowane zasilanie DC.

W naszym prototypie cewka wybrana jako L1 ma obciążalność 19 A i nie ma sensu jej zwiększać, ścieżki PCB i złącza wytrzymują maksymalnie prąd około 20 A. Ze względu na typ zastosowanych MOSFET-ów maksymalne napięcie zasilania jest ograniczone do 30 V.

Nakładka interfejsu

Nakładka interfejsu (schemat ideowy pokazany na rysunku 2) łączy się maksymalnie z sześcioma czujnikami temperatury, może zasilать do trzech wentylatorów w trybie PWM i dwie małe pompy.

Jeden z czujników temperatury to DS18B20 zamontowany na płytce drukowanej, który wykrywa temperaturę otoczenia. Pozostałe pięć kanałów dopasowanych jest do czujników cyfrowych DS18B20 lub tanich termistorów NTC (poprzez złącza CON1-CON5).

Interfejs posiada również trzy diody LED stanu (czerwona, zielona i niebieska) a także brzęczyk i odbiornik podczerwieni do wprowadzania danych przez użytkownika.

Czterodrożne złącze CON9 wyprowadza szynę I²C Arduino. Pasuje do niej wiele czujników i modułów, my zdecydowaliśmy się tutaj wykorzystać alfanumeryczny moduł

LCD, podobny do tego, który opisaliśmy w SC w marcu 2017 roku (siliconchip.com.au/Article/10584).

Tego rodzaju wyświetlacz jest łatwy do wysterowania i dobrze nadaje się do pokazywania dużej liczby zmieniających się parametrów, takich jak temperatura czy prędkość wentylatora, w czasie zbliżonym do rzeczywistego.

Na płytce interfejsu nie ma żadnych buforów szyny I²C, ponieważ są one zamontowane w module interfejsu LCD.

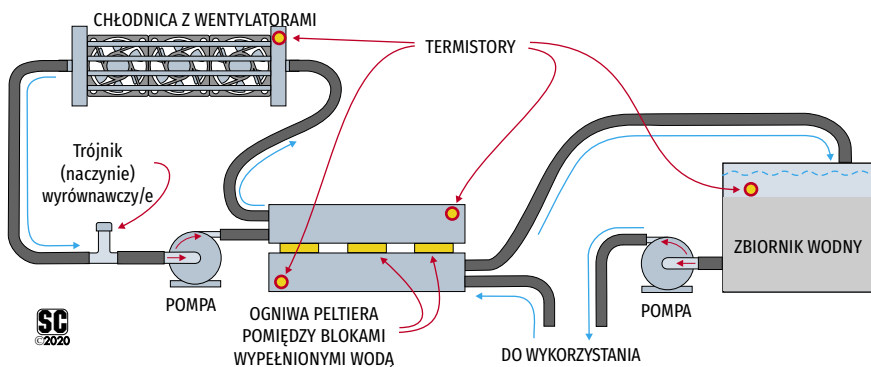
CON12 umożliwia doprowadzenie do nakładki zasilania 5 V lub 12 V (ustawiane przez LK1). D3 zapewnia ochronę przed odwrotną polaryzacją, przewodząc prąd wystarczający do przepalenia bezpiecznika F1 w przypadku odwrócenia zasilania. Przełącznik S1 może być użyty do włączenia lub wyłączenia zasilania.

Jeśli LK1 jest ustawiony w pozycji 12 V, zasilanie jest podawane do portu VIN Uno, który z kolei dostarcza regulowane napięcie 5 V z powrotem do nakładki poprzez regulator i wyjście 5 V Uno. Pozycja 5 V zworki LK1 podaje zasilanie bezpośrednio do końcówki 5 V.

Zwórka może być również pozostawiona otwarta, jeśli na przykład linie 12 V (VIN) i 5 V są dostępne z innego miejsca, takiego jak dołączona nakładka sterownika Peltiera.

Chociaż Uno ma sześć kanałów ADC (wejść analogowych), dwa z tych wejść są współdzielone z peryferiami I²C i dlatego nie mogą być użyte. Dlatego stosujemy IC1, potrójny dwukierunkowy multiplexer analogowy 74HC4053. Jest on używany do przełączania analogowych wejść A0, A1 i A2 pomiędzy CON2, CON3 i CON1 odpowiednio w jednym stanie, a IC2 (DS18B20), CON4 i CON5 odpowiednio w drugim stanie.

Wejścia sterujące dla wszystkich trzech kanałów multiplexera są połączone razem, do cyfrowego wyjścia D2 na module Uno. Końcówka „wyjście dostępne” (invE) jest podłączona do masy, więc trzy przełączniki w IC1 są zawsze aktywne.



Rysunek 4. Jest to wariant rysunku 3. Naczynie z lewej pętli zostało zastąpione trójnikiem wyrównawczym, którego otwarty koniec powinien znajdować się w najwyższej części obiegu, aby uniknąć rozlania wody. Wodę z prawej pętli możesz wykorzystać do chłodzenia lub ogrzewania czegokolwiek potrzebujesz (np. klimatyzatora osobistego wykonanego z innego grzejnika i kilku wentylatorów).

Styki A0, A1 i A2 (An, Bn i Cn układu IC1) mają oddzielne rezystory polaryzujące 4,7 kΩ do poziomu 5 V. Zapewnia to pośrednie zasilanie czujnika, jeśli zamontowany jest DS18B20 lub tworzy górną połowę obwodu dzielnika napięcia, jeśli zamontowany jest termistor NTC.

CON6, CON7 i CON8 to czterodrożne wtyczki do podłączenia wentylatorów obsługujących PWM. Ich zasilanie: 12 V i GND jest pobierane ze złączy VIN i GND nakładki.

Wyjścia tachometrów wentylatorów są doprowadzone przez rezystory 1 kΩ odpowiednio do wejść D4, D5 i D6 Arduino Uno. Można je ustawić jako wejścia cyfrowe do wykrywania prędkości obrotowej wentylatorów.

Wspólny sygnał PWM dla wentylatorów doprowadzony jest z wyjścia D3 Arduino Uno poprzez rezystor 100 Ω. Ta linia ma również rezystor zerujący 10 kΩ, więc wentylatory są wyłączone podczas resetowania.

Złącza CON10 i CON11 służą do zasilania małych pompek wodnych 12 V. Każda z nich jest przełączana po stronie masy przez tranzystory NPN (Q1 i Q2), sterowane z wyjść Uno D7 i D8 przez rezystory 1 kΩ.

Diody tłumiące D1 i D2 są podłączone bezpośrednio do wyjść CON10 i CON11, aby gasić wszelkie skoki samoindukcji, gdy pompy się wyłączają.

Brzęczyk piezoelektryczny PB1 jest zasilany analogicznie przez tranzystor NPN Q3. Jego baza jest sterowana z wyjścia D12 Arduino poprzez rezystor 1 kΩ ograniczający prąd.

Z trzech wbudowanych diod LED, dioda LED2 jest zasilana przez rezystor 1 kΩ z wyjścia D13, gdy jest ono w wysokim stanie. Diody LED1 i LED3 są połączone antyrównolegle pomiędzy końcówkami D2 i A3 z rezystorem szeregowym 1 kΩ. LED1 świeci, gdy A3 jest w stanie wysokim, a D2 w niskim; LED3 świeci, gdy D2 jest w stanie wysokim, a A3 w niskim.

Oczywiście obie nie mogą świecić w tym samym czasie. Taki układ oznacza, że, gdy

D2 jest przełączane w celu skanowania czujników temperatury, diody mogą migotać, ale trwa to krótko.

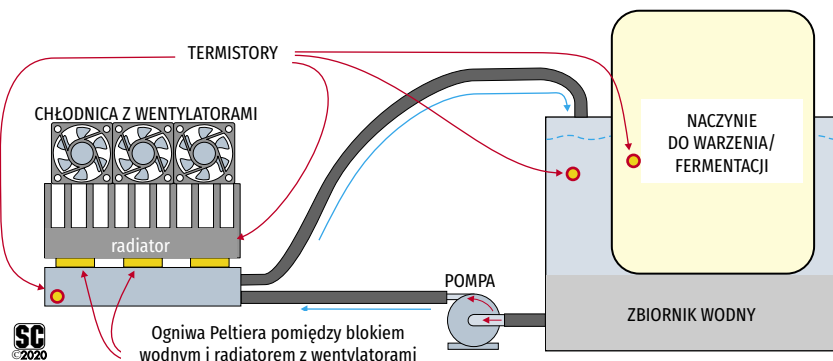
Odbiornik podczerwieni IRL1 zasilany jest przez rezystor 100 Ω i bocznikowany przez kondensator 100 nF. Jego wyjście jest podawane na wejście D1 Arduino przez rezystor 1 kΩ. Peryferia UART również korzystają z D1, więc nie mogą być używane jednocześnie z odbiornikiem.

Końcówki D9, D10 i D11 są pozostawione wolne i służą do sterowania nakładką sterownika Peltiera.

Napisaliśmy kilka procedur i funkcji do sterowania nakładką interfejsu, między innymi krzywe kalibracji termistora i przetwarzanie sygnału tachometru wentylatora, z wykorzystaniem przerywania.

Pewnym ograniczeniem napisanego kodu jest to, że obsługuje on tylko pojedynczy układ DS18B20 zamontowany na płytce. Możliwe jest, poprzez zmianę kodu, odczytanie temperatury z innych czujników DS18B20 pracujących na pośrednim zasilaniu z CON1-CON5, ale znacznie spowolni to pomiar temperatury.

Wybraliśmy takie rozwiązanie, ponieważ uznaliśmy, że dokładność tanich termistorów NTC jest wystarczająca.



Rysunek 5. Praktycznie minimalny układ obiegu wodnego. Dla uproszczenia, zamiast drugiej pętli wodnej używamy kombinacji wentylatora i radiatora. Choć nie jest to tak efektywne jak duża chłodnica, to taka konfiguracja może przenosić kilkaset watów ciepła.

Moc

Wszystko co wiąże się z przemieszczaniem znacznych ilości ciepła, wymaga sporej ilości energii. W naszym prototypie użyliśmy czterech ogniw Peltiera o prądzie 5 A. Wentylatory, pompy i nakładki dodają pobór nie więcej niż jednego ampera.

Większość płytek Peltiera jest przystosowana do pracy przy napięciu do 15 V. Tak więc potrzebujemy zasilania prądem około 21 A przy napięciu około 15 V. Zmniejszenie strat I²R jest dobrym pretekstem do użycia nieco niższego napięcia, jak 12 V, które jest również łatwiej dostępne.

Do naszego prototypu użyliśmy zasilacza komputerowego ATX zdolnego dostarczyć 22 A ze swojej szyny 12 V.

Choć takie obciążenie wydaje się być blisko maksymalnej mocy, pozostałe linie wyjściowe zasilacza (5 V, 3,3 V itd.) nie mają praktycznie żadnego obciążenia. W związku z tym zasilacz mieści się w granicach specyfikacji i nie wykazuje żadnych oznak przeciążenia przy ciągłej pracy.

Alternatywnie, można użyć modułu zasilacza 15 V lub 13,8 V do zabudowy lub wysoko-prądowego zasilacza laboratoryjnego.

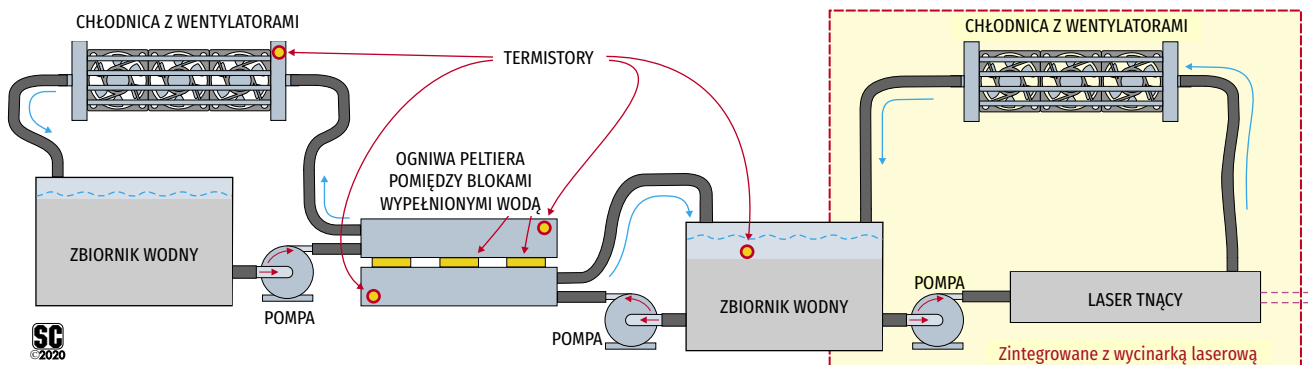
Przeprowadziliśmy nawet kilka wstępnych testów z wykorzystaniem naszego zasilacza 45 V/8 A z okresu październik-grudzień 2019 (siliconchip.com.au/Series/339), choć to niepełne wykorzystanie jego możliwości!

Pokażemy, jak podłączyliśmy zasilacz ATX; w porównaniu z nim inne opcje będą zapewne dość proste.

Pozostały sprzęt

Jak się zapewne domyślasz, ten projekt zawiera trochę więcej elementów niż tylko elektronika. Na szczęście większość części jest łatwo dostępna na stronach internetowych, takich jak AliExpress i eBay.

Radzimy, abyś przed przystąpieniem do budowy zapoznał się dokładnie z naszymi schematami i sprawdził, co będzie Ci potrzebne,



Rysunek 6. Jest to układ, który zainstalowaliśmy przy naszej wycinarkę laserowej, aby pomóc „zwiększyć” chłodzenie lasera w gorące dni. Zmniejsza on temperaturę lasera o około 6°C w porównaniu z chłodzeniem czysto pasywnym (co jest całkiem dobrym wynikiem, jeśli weźmiemy pod uwagę, że przy chłodzeniu pasywnym laser pracuje w temperaturze 10°C powyżej temperatury otoczenia).

ponieważ margines swobody realizacji projektu jest dość duży.

Jak wspomnieliśmy powyżej, naszym głównym nośnikiem ciepła jest woda. Ma ona wysoką pojemność cieplną (4,2 J na stopień i gram) i dobrze przenosi ciepło na drodze konwekcji i przepływu, dobrze również pobiera i oddaje energię cieplną w kontakcie z innymi materiałami, zwłaszcza metalami. Ponadto, istnieje wiele gotowych urządzeń nadających się do pracy z wodą.

Na przykład pompy, których używamy, są podobne do tych, które stosuje się przy cyrkulacji wody w akwarium.

Dużym minusem wody jest natomiast to, że nie jest ona izolatorem elektrycznym. Dlatego należy zadbać o to, aby woda nie dostała się do elektroniki (lub odwrotnie).

Obieg wody

Temperaturą kąpeli wodnej zarządzamy poprzez cyrkulację wody przez jedną lub więcej pętli. Przepływ ten ma na celu ciągłe mieszanie wody, tak aby w instalacji nie było gorących czy zimnych punktów. Rysunki 3–6 pokazują kilka wariantów „obiegów” wody, które można zbudować przy użyciu naszego sprzętu.

Rysunek 3 pokazuje układ, którego można użyć do fermentacji, rysunek 4 – ogólne zastosowanie do ogrzewania/chłodzenia, a rysunek 5 – uproszczone zastosowanie do fermentacji (które byłoby tańsze w budowie, ale prawdopodobnie mniej efektywne).

Rysunek 6 przedstawia jak zastosowaliśmy termoregulator do wstępnego schłodzenia wody do naszej wycinarki laserowej, zmniejszając temperaturę pracy lasera w gorące dni (więcej o tym później).

Patrząc na te rysunki można zauważyć, że dzięki pętli wodnej (obiegom wodnym) możemy utrzymywać grzejniki, radiatory i wentylatory, które wyrzucają „ciepło odpadowe” do powietrza, z dala od zbiornika, którego temperaturę próbujemy regulować.

Jest to kluczowa korzyść z używania wody do przekazywania ciepła.

Użycie większej objętości wody oznacza, że konfiguracja będzie bardziej odporna na zmiany zewnętrzne, ale osiągnięcie docelowej wartości zadanej temperatury zajmie więcej czasu. Celem jest jak najbardziej efektywne dostarczenie lub odprowadzenie ciepła z miejsca, które nas interesuje. Pętle umożliwiają łatwe transportowanie ciepła.

Potrzebne części

Wiele z użytych przez nas części zostało zakupionych jako elementy zestawu. Zestawy te są zazwyczaj przeznaczone do chłodzenia wodnego komputerów (np. do podkręcania – overclockingu). Musieliśmy również zdobyć kilka innych, różnych elementów.

Woda jest przetłaczana przez małe 12 V pompy zanurzeniowe. Są one tanie i pobierają około 300 mA każda. Woda nie jest pompowana na dużą wysokość, ponieważ pozostaje w obiegu zamkniętym, tak więc wysokie ciśnienie nie jest potrzebne. Ogólnie rzecz biorąc, tak długo jak woda jest przetłaczana, możemy utrzymać przenoszenie ciepła na poziomie, którego potrzebujemy.

Aby połączyć wszystko razem, użyliśmy elastycznych rurek silikonowych. Otrzymaliśmy je jako część naszego zestawu, choć można je również dostać w sklepach



Te pompy są małe i pobierają prąd o natężeniu tylko około 300 mA. Są uszczelnione, a więc w pełni zanurzalne (rotor silnika jest sprężony z wirnikiem pompy za pomocą magnesów). Ponieważ nie podnoszą wody na jakąś dużą wysokość, nie potrzebują dużej mocy. Najważniejsze, aby wlot był zawsze całkowicie zanurzony, gdyż nie są to urządzenia samozasysające.



Mosiężne złączki są dobrze dopasowane do przezroczystego węża, którego użyliśmy, i nie wykazywały żadnych oznak wysuwania się. Jednak dla pewności użyliśmy jeszcze opasek zaciskowych. Rurka, która została dostarczona z naszym zestawem była dość miękka, więc zastąpiliśmy ją grubszą rurką kupioną na miejscu.

z narzędziami, takich jak Bunnings, sklepach kempingowych, akwarystycznych czy ogrodniczych. Z naszego doświadczenia wynika, że najbardziej przydatne są rurki o średnicy wewnętrznej około 8 mm, które są dobrze dopasowane do karbowanych złączy na innych częściach.

Chociaż po naciągnięciu rurka jest ciasno dopasowana, nie ufaliśmy temu całkowicie. Aby zabezpieczyć rurkę, użyliśmy małych (6 mm–16 mm) opasek zaciskowych do węża.

Tam, gdzie musielibyśmy zgiąć rurkę pod ostrym kątem, zamiast tego użyliśmy małych kolanek i trójników. One również powinny zostać zabezpieczone zaciskami.

Ostatnią częścią naszego obwodu pierwotnego jest blok wodny. Jest to pusty w środku blok aluminiowy z dwiema karbowanymi złączkami na jednym końcu. Zapewnia on dobry kontakt termiczny pomiędzy wodą a ogniwami

Peltiera, umożliwiając łatwe przekazywanie ciepła.

Woda wpływa przez jeden koniec, przepływa w górę i z powrotem wzdłuż bloku, a potem wraca i wypływa przez drugą złączkę. Chociaż aluminium nie jest najlepszym przewodnikiem ciepła, jest tanie i łatwe w obróbce.

W typowej konstrukcji, ogniwa Peltiera są przymocowane do płaskich powierzchni bloku wodnego pokrytych pastą termiczną. Zapewnia ona dopasowanie termiczne na dużej powierzchni i dobrze przewodzi ciepło.

Oczywiście ogniwa Peltiera mają dwie strony i kiedy ciepło jest pobierane z jednej strony, musi być odbierane, i to z dużym, ok. 100% nadmiarem, po drugiej stronie. Najprostszym rozwiązaniem jest zastosowanie bloku radiatora, który jest aktywnie chłodzony przez wentylatory.

W naszym projekcie zasilacza 45 V/8 A (patrz wcześniejszy link), użyliśmy na radiatorze pary wentylatorów o dużej mocy i okazało się,

że są one w stanie rozproszyć kilkaset watów ciepła.

Przeprowadziliśmy kilka prób z użyciem tej techniki z ogniwami Peltiera i wypadła ona dobrze, ale nie tak dobrze jak odpowiednia chłodnica.

Lepszą metodą usuwania ciepła z drugiej strony płytki Peltiera jest wykorzystanie drugiej pętli wodnej.

Ten obieg wykorzystuje drugą pompę i związaną z nią sieć rurek podobnie do kąpieli wodnej w pierwszej pętli. Woda w drugiej pętli przechodzi przez chłodnicę chłodzoną wentylatorem.

Chłodnica ta jest jakby mniejszą wersją chłodnicy samochodowej. Woda przepływa przez chłodnicę, a powietrze jest przemieszczane nad nią przez wentylatory.

Jeśli woda jest cieplejsza od powietrza, to jest schładzana (a powietrze ogrzewane). Jeśli woda jest chłodniejsza od powietrza, zostaje ona ogrzana.

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

Programowalny regulator termiczny (Arduino/Peltier)

- 1 zestaw sprzętu do obiegów wodnych (patrz tekst i poniżej)
- 1 płytka Arduino Uno R3 lub kompatybilna (ATmega328)
- 1 nakładka sterownika Peltiera (patrz poniżej)
- 1 nakładka interfejsu (patrz poniżej)
- 1 wysokoprądowy zasilacz DC (patrz tekst)
- 1 alfanumeryczny ekran LCD 20×4 z interfejsem I²C [SILICON CHIP ONLINE SHOP nr kat. SC4203]
- 1 odcinek płaskiego kabla (do ekranu LCD)
- 1 czterostykowe gniazdo kierunkowe plus styki, np. typu 402-4 (do ekranu LCD)
- 1 uniwersalny pilot na podczerwień [Jaycar XC3718, Altronics A1012]

Zestaw sprzętu do obiegu wodnego (pojedyncza pętla)

- 4 ogniwa Peltiera 5 A
- 1 naczynie na wodę (wg własnego wyboru)
- 1 mata pompa wodna 12 V DC [np. www.aliexpress.com/item/32810010753.html]
- 1 aluminiowy blok wodny 40×200 mm [np. www.aliexpress.com/item/4000599538658.html]
- 2 zestawy montażowe bloku wodnego [np. www.aliexpress.com/item/32323128854.html]
- 1 radiator o długości 200 mm (pasujący do bloku wodnego) [Jaycar HH8530, Altronics H0536]
- 2 wentylatory 80 mm 12 V albo dopasowane do radiatora [Jaycar YX2512, Altronics F1050]
- osprzęt montażowy pasujący do wentylatorów
- kilka metrów elastycznej rurki silikonowej o średnicy wewnętrznej 8 mm
- kilka kolanek i trójników dopasowanych do rurki
- co najmniej 4 zaciski do węża 6-16 mm
- 1 tubka pasty termicznej
- różne opaski kablowe

Zestaw sprzętu do obiegu wodnego (podwójna pętla)

- 4 ogniwa Peltiera 5 A
- 2 zbiorniki na wodę dostosowane do planowanego użycia
- 2 maty pompy wodne 12V DC [np. www.aliexpress.com/item/32810010753.html]
- 2 aluminiowe bloki wodne 40×200mm [np. www.aliexpress.com/item/4000599538658.html]
- 2 zestawy do montażu bloków wodnych [np. www.aliexpress.com/item/32323128854.html]
- kilka kolanek i trójników dopasowanych do rurki
- co najmniej 8 szt. 6-16 mm opasek zaciskowych
- 1 tubka pasty termicznej
- różne opaski kablowe
- 1 radiator do wentylatorów 120 mm, zalecany typ 360 mm [np. <https://www.aliexpress.com/item/32992836396.html>]
- 1-3 wentylatory 12 V pasujące do radiatora (np. 120 mm) [Jaycar YX2574, Altronics F1165]
- osprzęt montażowy pasujący do wentylatorów
- Elementy nakładki sterownika Peltiera
- 1 dwustronna płytka drukowana o kodzie 21109182, 53,5×68,5 mm
- 1 szt. 10-stykowe gniazdo jednorzędowe żeńsko/męskie (do łączenia w stos), (11 mm wysokości styków)
- 1 szt. 8-stykowe gniazdo jednorzędowe żeńsko/męskie (do łączenia w stos), (11 mm wysokości styków)
- 2 szt. 6-stykowe gniazda jednorzędowe żeńsko/męskie (do łączenia w stos), (11 mm wysokości styków)
- 2 szt. 2-drożne wysokoprądowe listwy zaciskowe, raster 8,3 mm (CON1, CON2)
- 1 listwa kołkowa jednorzędowa raster 2,54 (JP1, JP2, LK3, LK4)
- 4 zworki (do złączy powyżej)
- 2 oprawki bezpieczników M20×5 do montażu na PCB (F1)
- 1 bezpiecznik 25 A M20×5 (F1)

- 1 diodzik 3,3 μH 19 A SMD, 14,0×12,8 mm [np. Pulse PA4343.332ANLT; Digi-Key 553-4025-1-ND]
- 4 śruby M3×9
- 4 nakrętki sześciokątne M3

Półprzewodniki:

- 2 szt. diod 1N4148 (D1, D2)
- 1 sterownik mostka H układ HIP4082, DIP-16 (IC1) [Digi-Key HIP4082IPZ-ND]
- 1 stabilizator napięcia 78L12, TO-92 (REG1; opcjonalnie – patrz tekst)
- 4 szt. N-kanalowe MOSFET-y IRLB8314, TO-220 (Q1-Q4) [Digi-Key IRLB8314PBF-ND]

Kondensatory:

- 3 szt. kondensatory 100 nF MKT lub wielowarstwowe ceramiczne
- 4 szt. kondensatory ceramiczne 10 μF/16 V*, MLC X7R SMD 3216/1206 [Digi-key 1276-6641-1-ND]

* wymagane są wersje o wyższym napięciu, jeśli zasilanie DC >15 V

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

- 2 szt. 10 kΩ 1 szt. 1,8 kΩ

Elementy nakładki interfejsu Peltiera

- 1 dwustronna płytka drukowana o kodzie 21109181, 53,5×68,5 mm
- 1 szt. 10-kołkowa męska listwa prosta jednorzędowa
- 1 szt. 8-kołkowa męska listwa prosta jednorzędowa
- 2 szt. 6-kołkowe męskie listwy proste jednorzędowe
- 1 uchwyt bezpiecznika ostrzowego do montażu na płytce drukowanej (F1; opcjonalnie, typu samochodowego)
- 1 bezpiecznik samochodowy 2 A (F1)
- 5 szt. wtyków 403-2 kołkowych prostych do druku męskich (CON1-CON5)
- 4 szt. wtyki 403-4 kołkowe proste do druku męskie (CON6-CON9)
- 3 złącza śrubowe 2 pola, raster 5,08 mm do montażu na PCB (CON10-CON12)
- 1 przełącznik SPDT R/A montowany na PCB (S1; opcjonalnie) [Altronics S1320]
- 1 odcinek 3-kołkowy listwy męskiej do druku + zworka (LK1)
- 1 przycisk zwrotny chwilowy 6 mm (S2)
- 1 brzęczyk piezoelektryczny (PB1) [Jaycar AB3459, Altronics S6104]
- 5 termistorów NTC 10 kΩ/100 kΩ z przewodami [np. www.aliexpress.com/item/32916207487.html lub <https://www.aliexpress.com/item/33014111002.html>]
- 5 szt. obudów gniazda 402-2 stykowych na przewód + styki (jeśli termistory nie są dostarczane z odpowiednią wtyczką)
- kabel płaski (jeśli przewody czujników nie są wystarczająco długie)

Półprzewodniki:

- 1 potrójny 2-kanalowy multiplexer analogowy 74HC4053, DIP-16 (IC1)
- 1 cyfrowy czujnik temperatury DS18B20, TO-92 (IC2)
- 2 tranzystory NPN BC337, TO-92 (Q1, Q2)
- 1 tranzystor NPN BC547, TO-92 (Q3)
- 1 czerwona dioda LED 5 mm (LED1)
- 1 zielona dioda LED 5 mm (LED2)
- 1 niebieska dioda LED 5 mm (LED3)
- 3 diody 1N4004 400 V 1 A (D1-D3)
- 2 kondensatory 100 nF MKT lub ceramiczne MLC

Rezystory: (wszystkie 1/4 W axial 1% metal film)

- 3 szt. 4,7 kΩ 1 szt. 10 kΩ 9 szt. 1 kΩ 2 szt. 10 Ω

Od Red. EdW: oferty na portalu AliExpress często się zmieniają, podane zostały linki przykładowe

Taki system działa lepiej, bo chłodnica wodna ma większą powierzchnię transportu ciepła, a wentylatory poruszają większą masę powietrza. Jest to jednak znacznie bardziej skomplikowany układ.

Zamontowaliśmy go w naszej wycinarce laserowej.

Na zdjęciu widać termistory użyte w bloku wodnym ogniw Peltiera. Na podstawie odczytów temperatury możemy wiele powiedzieć o tym, jak działa system. W szczególności, temperatury po gorącej i zimnej stronie ogniw Peltiera wskazują jak bardzo są one obciążone i sugerują najlepszą strategię dla uzyskania optymalnej sprawności cieplnej.

W dalszej części wyjaśnimy dokładniej, dlaczego czasami korzystne jest wyłączenie zasilania ogniw Peltiera. I oczywiście pamiętajmy, aby do pomiarów temperatury w naszej łaźni wodnej użyć dedykowanych czujników, które wskażą nam, czy – i kiedy osiągnęliśmy temperaturę docelową.

Naczynie na wodę do warzenia piwa...

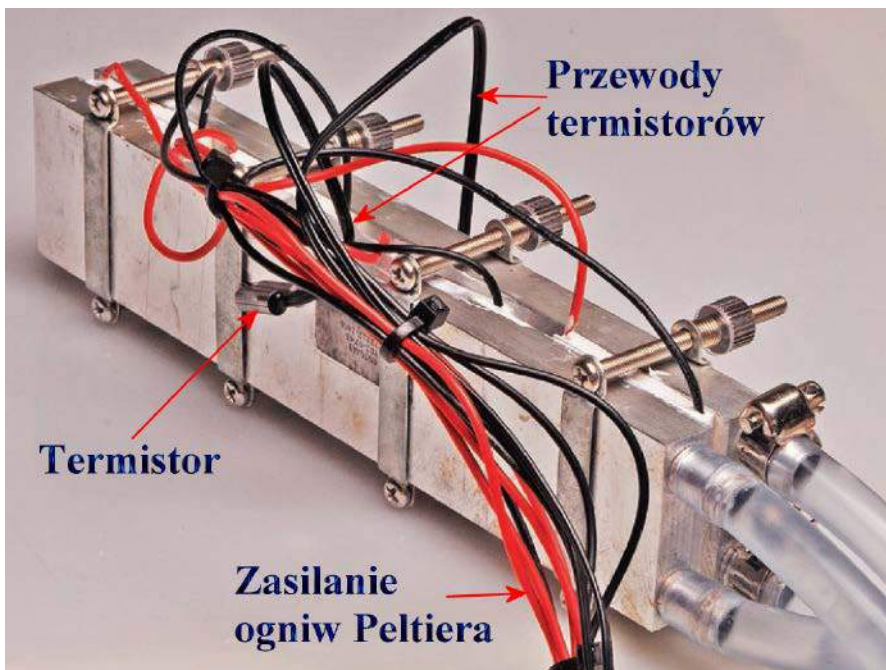
Kolejną częścią, której nie ma w typowych zestawach chłodzenia wodnego komputera, jest zbiornik na wodę. Dobieramy go w zależności od zastosowania.

W przypadku ostatecznej wersji chłodnicy dużej wydajności („boost”) dla naszej wycinarki laserowej, po prostu wykorzystaliśmy istniejący zbiornik wody, którym był plastikowy pojemnik obiadowy. Oryginalny system pasywnego chłodzenia, który zbudowaliśmy dla naszej wycinarki laserowej, można zobaczyć w naszym artykule w SC z czerwca 2016 roku (siliconchip.com.au/Article/9960).

Chociaż kuszący może być pomysł, by w przypadku warzenia piwa cyrkulować brzeczkę fermentacyjną w obiegu bezpośrednio w kontakcie z ogniwami Peltiera, zdecydowanie to odradzamy. Nie widzieliśmy nigdzie zapewnień, że części, których użyliśmy, są bezpieczne dla żywności, a w każdym razie resztki piwa przepompowywanego jako ciecz obiegową byłoby bardzo trudne do późniejszego usunięcia z instalacji. No i nie zapominajmy o musowaniu i obecności słoju.

Piwo jest lekko kwaśne, a wiele roztworów czyszczących jest silnie alkalicznych. Armatura może nie wytrzymać tego rodzaju czyszczących środków chemicznych.

Dlatego w przypadku zastosowań związanych z warzeniem i fermentacją sugerujemy umieszczenie naczynia do warzenia w dużej łaźni wodnej. Zakładając, że używasz plastikowej jednostki o pojemności 25 litrów, najprostszą opcją będzie plastikowy baniak, taki jak te dostępne w dyskontach i sklepach z narzędziami.



Ten panel jest trzymany razem przez zaciski, z czterema ogniwami Peltiera umieszczonymi pomiędzy blokami wodnymi. Pod drugim zaciskiem od lewej widoczny jeden z termistorów, pozostałe również są utrzymywane na miejscu przez zaciski. Nie widać niewielkiej ilości pasty termicznej pomiędzy powierzchniami przewodzącymi ciepło. Chociaż bloki wodne są potączone elektrycznie ze sobą poprzez zaciski, to ceramiczne płytki na roboczych powierzchniach ogniw Peltiera zapobiegają zwarcu.

Większy pojemnik działa jak płaszcz wodny. Nie musi on całkowicie otaczać mniejszego naczynia, ale powinien umożliwić w miarę głębokie zanurzenie zbiornika fermentacyjnego, aby zwiększyć powierzchnię, przez którą przekazywane jest ciepło. Otwór wycięty w pokrywie większego pojemnika (tworząc swego rodzaju uszczelnienie wokół naczynia) zmniejsza potencjalne parowanie wody, a tym samym zmniejsza moc potrzebną do utrzymania temperatury.

Tak duże naczynie, ze względu na swoją dużą powierzchnię, może tracić (lub zyskiwać) ciepło z otoczenia, dlatego przyda się nawet niewielka ilość izolacji; powinien wystarczyć zwykły ręcznik.

... i naczynie do gotowania

Jak już wspomnieliśmy, wyższe temperatury używane do gotowania „sous-vide” będą bardziej obciążać moduły Peltiera. Do tego zastosowania zalecamy użycie małej chłodziarki piankowej. Podczas testów używaliśmy takiej, która została zaprojektowana do przechowywania sześciu puszek z napojami. Jej mały rozmiar minimalizuje powierzchnię, przez którą tracone jest ciepło, jak również objętość cieczy, która ma być podgrzewana. Jest jednak wystarczająco duża, aby zmieścić większość potraw, które można by ugotować.

Chłodziarki takie można znaleźć w Internecie, w sklepach z artykułami AGD

REKLAMA

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

OBWODY DRUKOWANE

Produkcja, Projektowanie, Montaż

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



Ta chłodnica, trochę podobna do samochodowej, jest bardziej efektywna w odprowadzaniu ciepła niż prosty radiator i wentylatory. Wynika to z jej większej powierzchni aktywnej.

lub na wyprzedających. Sprawdź, czy jest dostarczana z pokrywą, ponieważ przy stosowanych przez nas temperaturach mogą wystąpić dość duże straty parowania. Należy również zachować ostrożność podczas użytkowania, ponieważ łatwo o oparzenie.

Ponieważ podczas procesu „sous-vide” żywność jest zamykana w wodoodpornych torebkach, ryzyko skażenia spowodowane stykiem z częściami bez atestu do kontaktu z żywnością jest minimalne. Dla pewności warto użyć dwóch torebek.

Do realizacji naszego projektu w wariancie dwuobiegowym potrzebne będzie drugie naczynie. Izolacja w tym drugim obiegu nie jest tak krytyczna, ponieważ chłodnica i wentylatory w drugiej pętli starają się utrzymać temperaturę wody zbliżoną do temperatury otoczenia. Przydatna może być jednak pokrywa, aby zapobiec długotrwałej utracie wody przez parowanie. W naszych testach jako drugiego naczynia na wodę użyliśmy plastikowej wanienki do łodów.

Kolejną sprawą, na którą należy uważać, jest możliwość skażenia bakteriami/glonami, szczególnie jeśli używasz naszego termoregulatora do gotowania. W ciepłej wodzie mogą rozwijać się bakterie i glony. Przykładowo, przypadki „choroby legionistów” występujące w przemysłowych wieżach chłodniczych, zostały powiązane z cyrkulacją ciepłej wody wystawionej na działanie powietrza.

Oczywiście należy zadbać o to, aby woda z pętli chłodzących nie znalazła się w pobliżu niczego, co może zostać spożyte. Należy również regularnie wylewać i wymieniać wodę w pętli, co pomoże ograniczyć gromadzenie się patogenów.

Jeśli jesteś zaznajomiony z procesem warzenia piwa, będziesz wiedział, jak kluczowe dla dobrych wyników jest zachowanie odpowiedniej czystości.

Zmierzona wydajność

Stwierdziliśmy, że w dobrze izolowanych warunkach, przy temperaturze otoczenia 18°C, nasza łaźnia wodna osiąga temperaturę

około 70°C. W trakcie tych testów, nasze główne naczynie zawierało około dwóch litrów wody w izolowanej chłodziarce piankowej. Druga pętla zawierała około litra wody w (czystym) pojemniku na lody.

W naszych testach wytworzona została duża ilość pary wodnej, co prowadziło do chłodzenia przez pobieranie ciepła parowania, które zmusza moduły Peltiera do cięższej pracy.

Podczas chłodzenia zeszliśmy do temperatury około 2°C. Czynnikiem ograniczającym jest zbliżenie się do punktu zamarzania wody. Widzieliśmy szron na płytkach Peltiera, jasne więc było, że niektóre części schłodziły się do ujemnej temperatury.

Typowy czas osiągnięcia tych ekstremalnych temperatur to około pół godziny przy użyciu czterech standardowych ogniw 5 A pracujących przy napięciu około 11 V. Widać więc, że zmiana temperatury nie jest szybka. Do osiągnięcia ekstremów temperaturowych niezbędna jest dobra izolacja termiczna.

Obliczyliśmy, że wtórna pętla wodna z chłodnicą i wentylatorami ma około dwukrotnie większą wydajność usuwania ciepła niż proste rozwiązanie z radiatorem.

Weźmy pod uwagę, że we wszystkich przypadkach próbujemy skutecznie przemieścić ciepło pomiędzy powietrzem otoczenia (powietrze przedmuchiwane przez wentylatory i przez radiator) a kąpielą wodną. Im bardziej zbliżone są ich temperatury, tym łatwiejsze będzie nasze zadanie. Nasze doświadczenie potwierdza, że przy dużych różnicach temperatur pomiędzy stronami ognia Peltiera mają one największe problemy podczas pracy.

Przykładowo, podczas niektórych naszych początkowych testów, gdy próbowaliśmy schłodzić gorącą wodę, zauważyliśmy, że bardziej efektywne było wyłączenie modułów Peltiera i pozwolenie na przewodzenie ciepła przez moduł. Zasilanie urządzeń Peltiera dodało do systemu jeszcze więcej ciepła (straty I²R), które również musiało zostać usunięte.



Mały pojemnik/chłodnica piankowa, taki jak pokazany tutaj, jest dobrym wyborem do gotowania „sous-vide” z termoobiegami. Potrzebny jest wysoki stopień izolacji termicznej, a dopasowana pokrywa minimalizuje ilość wody i ciepła traconego na skutek parowania.

Gotowanie metodą „sous-vide” to zastosowanie, która według naszych przewidywań wymaga najbardziej ekstremalnych temperatur, dlatego dla uzyskania dobrych wyników niezbędna jest doskonała izolacja termiczna. W niektórych przypadkach można wstępnie podgrzać wodę w czajniku, a następnie pozwolić termoregulatorowi osiągnąć temperaturę docelową i utrzymywać ją na tym poziomie; będzie to szybsze niż rozpoczynanie pracy z zimną wodą.

W przyszłym miesiącu

W tym miejscu musimy zrobić przerwę.

W przyszłym miesiącu opiszemy, jak zbudować obie nakładki, zaprogramować Arduino i złożyć cały system w całość. Jeśli przymierzasz się do budowy termoregulatora, w międzyczasie możesz ustalić potrzebną konfigurację układu i zamówić części. Jeśli te dotrą szybko, będziesz mógł rozpocząć budowę obiegów wodnych i zespołów wymiany ciepła.

Oprogramowanie, które zaprezentujemy w przyszłym miesiącu, posiada kilka różnych trybów pracy, takich jak ustawienie temperatury docelowej, którą urządzenie utrzymuje, zapewniając maksymalne ogrzewanie lub chłodzenie, a także tryb, w którym po uruchomieniu urządzenie podąża za wstępnie ustawionym „profilem” temperatury. ■

Tim Blythman
Nicholas Vinen

Łatwe do budowy

Aktywne głośniki półkowe Hi-Fi z opcjonalnymi subwooferami, część 1

Te monitory – głośniki o wysokiej jakości – są przeznaczone do użytku z telewizorami, komputerami lub sprzętem audio. Są nie- drogie i łatwe w budowie, a jednocześnie mają doskonałą jakość dźwięku, z niskimi zniekształceniami i dość płaską charakterystyką częstotliwościową. Jeśli więc szukasz głośników półkowych DIY i wiernej reprodukcji dźwięku, bez wydawania majątku, ten projekt jest dla Ciebie. Opcjonalnie, dopasowane subwoofery znacznie zwiększają poziom najniższych tonów i zapewniają znacznie wyższą głośność.

Nowoczesne telewizory cały czas stają się cieńsze i zgrabniejsze. O ile ta tendencja świadczy o wielkim skoku w technologii wyświetlania obrazu, o tyle istnieje kilka praw fizyki, które ograniczają możliwości wewnętrznych głośników, dopasowanych wzorniczo do nowoczesnych telewizorów i umieszczonych w niewielkiej dostępnej dla nich przestrzeni.

Nie ma się co oszukiwać – głośniki w prawie wszystkich nowoczesnych telewizorach, jeśli w ogóle są, brzmią fatalnie, a niektóre nie zapewniają nawet zrozumiałości mowy.

Idealnym rozwiązaniem jest zewnętrzny zestaw głośników i wzmacniacz podłączony do telewizora. Dla zwiększenia wygody wzmacniacz może znajdować się w obudowie samych głośników.

Głośniki można podłączyć do wyjścia liniowego (lub słuchawkowego) telewizora, dzięki czemu nadal można korzystać ze zdalnej regulacji głośności pilotem.

Głośniki, które dobrze współpracują z telewizorem, nadają się również bardzo dobrze do zapewnienia wysokiej jakości dźwięku z komputera, do oglądania filmów i słuchania muzyki, grania w gry lub do edycji dźwięku i filmów.

Dzięki wbudowanym wzmacniaczom i wysokiej jakości głośniki idealnie pasują do tych zadań.

Zaprojektowałem je tak, aby miały niewielkie wymiary i nie zajmowały zbyt wiele miejsca. Jednak w niektórych przypadkach, szczególnie przy oglądaniu telewizji i filmów, możesz potrzebować więcej niskich tonów niż może to zapewnić mała obudowa.

Dlatego też w razie potrzeby, opcjonalne, dopasowane wzorniczo, zestawy bass-reflex poszerzają pasmo przenoszenia, a ponieważ zawierają własny wzmacniacz, zwiększają głośność maksymalną.



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

Materiały dodatkowe są również dostępne na stronie edw.elportal.pl: <https://bit.ly/3XgqC1j>

Przedstawione tutaj dwudrożne głośniki półkowe, z opcjonalnymi subwooferami, które pełnią również funkcję podstawek pod głośniki, są tanie i łatwe w budowie.

Cele projektowe

Przy projektowaniu tych głośników chciałem osiągnąć:

1. niewielkie rozmiary głośników półkowych (monitorów), około 200 mm szerokości, 300 mm głębokości i 400 mm wysokości.
2. płaską i akceptowalną krzywą impedancji.
3. przyzwoitą głośność maksymalną, co najmniej 100 dB SPL (Sound Pressure Level = poziom głośności) w odległości 1 metra bez nadmiernych zniekształceń.
4. pasmo przenoszenia ze spadkiem nie większym jak -3 dB w paśmie od 40 Hz do 20 kHz dla samych głośników półkowych.
5. płaską charakterystykę dźwięku, nominalnie ± 3 dB w zakresie od 40 Hz do 20 kHz.
6. dopasowane do monitorów subwoofery, rozszerzające pasmo przenoszenia basów i pracujące w zakresie poniżej jakichś 90 Hz.
7. konstrukcję drewnianą, pozwalającą na wykorzystanie łatwo dostępnych materiałów.
8. prostotę konstrukcji, mającą ułatwić pracę majsterkowiczom.
9. niski koszt; poniżej 300 dolarów za podstawowy system półkowy stereo, i dodatkowo nie więcej niż 150 dolarów za dwa dodatkowe subwoofery.
10. dla uproszczenia zintegrowane wzmacniacze mocy.

W przypadku opcjonalnych subwooferek, moimi dodatkowymi celami były:

1. pasmo przenoszenia w dół do około 35 Hz, wymagające objętości około 35 litrów i przetwornika 8-calowego (20 cm).
2. możliwość wykorzystania subwooferek jako podstawek pod monitory.
3. (alternatywnie) możliwość zbudowania subwooferek w kształcie prostopadłościanów, tak aby można było je schować pod biurkiem.
4. aktywna zwrotnica, która rozdziela sygnał pomiędzy głośniki półkowe i subwoofery.
5. zintegrowane wzmacniacze mocy dla subwooferek.
6. maksymalne wymiary: około 200 mm szerokości, 300 mm głębokości i 800 mm wysokości.

W rezultacie wymiary końcowe to: 210×296×280 mm dla głośników i 210×296×800 mm dla subwooferek.

Kompromisy

Podczas projektowania musieliśmy iść na kompromis pomiędzy kosztami a parametrami. Istnieje kilka bardzo kosztownych opcji przetworników, które kuszą wyjątkową wydajnością. Poważni audiofile mogą być zadowoleni z wydania wielu setek dolarów na pojedynczy przetwornik dźwięku, natomiast my uważamy, że taki wydatek nie jest konieczny dla uzyskania doskonałych parametrów.

Wyniki, jakie uzyskaliśmy, potwierdzają tę opinię. Dzięki zastosowaniu łatwo dostępnych, niedrogich przetworników i prostej zwrotnicy, pomiary i testy odsłuchowe pokazują, że oczarowują one w małym dwudrożnym systemie typu monitor.

Jakość dźwięku samych kolumn półkowych jest bardzo dobra, ale brakuje im na dole nieco basów, więc można oczekiwać bardziej „wiernego” dźwięku, jeśli zbudujemy również opcjonalne subwoofery. Zarówno głośniki półkowe, jak i subwoofery są „aktywne”, tzn. w każdej parze znajduje się dedykowany wzmacniacz. Dzięki temu można je podłączyć bezpośrednio do telewizora lub komputera bez konieczności budowania osobnego wzmacniacza.

Niektóre z kompromisów, na które musiałem przystać podczas pracy nad tym projektem, to:

- Rozmiar: chciałem, aby głośniki były stosunkowo małe, co ograniczyło rozmiar przetworników i objętość obudowy, a to oznacza, że nie emitują one naprawdę głębokiego basu.

- Materiał obudowy: wybrałem sklejkę (lub płytę MDF) o grubości 15 mm, która jest tania i łatwa do zdobycia, chociaż wolałbym użyć grubszego materiału.
- Wykończenie: zdecydowałem się na wykończenie z bejcowanego lub lakierowanego drewna, aby utrzymać koszty na niskim poziomie i uprościć konstrukcję. W razie potrzeby ze względów estetycznych można zastosować lakier lub okleinę.
- Przetworniki: przetworniki, które wybrałem, choć są tanie i zapewniają doskonałą jakość dźwięku, mają pewne cechy charakterystyczne, które sprawiają, że projektowanie zwrotnic jest nieco kłopotliwe i wymaga sporo uwagi. To sprawia, że zwrotnice są nieco droższe, ale koszt przetworników jest na tyle niski, że to się w sumie opłaca.

Elektronika

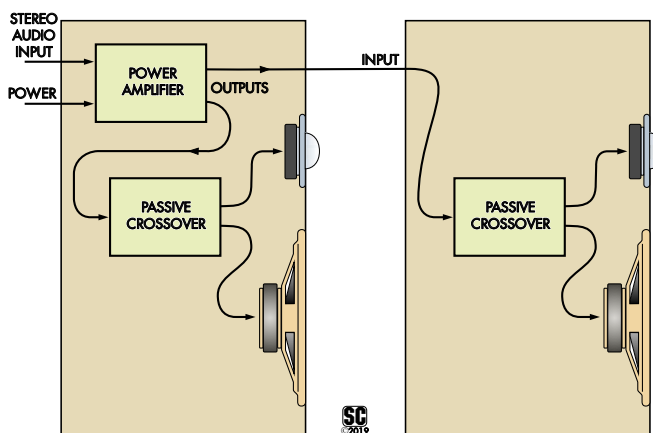
Dla uproszczenia, jedna kolumna półkowa zawiera wzmacniacz stereo zasilający oba głośniki, z pasywnymi zwrotnicami w każdej obudowie. Dzięki temu para głośników jest w pełni samodzielna, z wyjątkiem zasilacza (patrz **rysunek 1**). Używamy pudełkowego („brick”) zasilacza sieciowego AC-DC jak w laptopach, więc wystarczy tylko włożyć wtyczkę do gniazdka sieciowego. Są one dość tanie i wydajne w stosunku do oferowanej mocy.

Podobnie, jeśli budujesz opcjonalne subwoofery, jeden subwoofer zawiera stereofoniczny wzmacniacz mocy do zasilania siebie i drugiego (pasywnego) subwoofera, plus aktywną zwrotnicę, która rozdziela odpowiednio sygnały do obu subwooferek, oraz do pary głośników półkowych. Układ ten pokazany jest na **rysunku 2**. Do zasilania wzmacniacza subwoofera służy osobny zasilacz tego samego typu, co oznacza, że do całego systemu potrzebne są dwa.

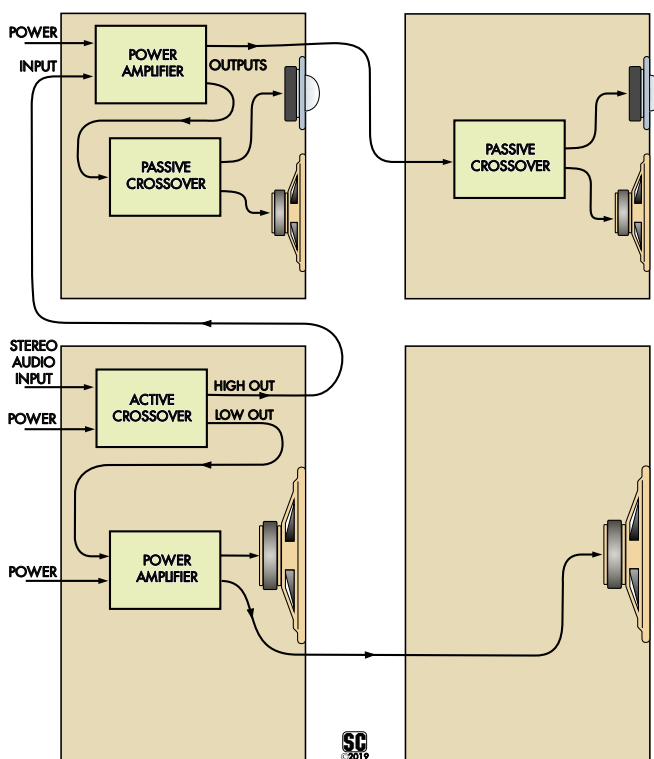
Używane przez nas moduły wzmacniaczy to wzmacniacze klasy D, korzystające z układu scalonego TDA7498. Wytwarzają one dużą moc bez nadmiernego drenażu Twojej kieszeni. Rozważaliśmy zastosowanie do tych głośników wzmacniacza ze znakomitymi układami LM3886 lub wzmacniacza dyskretnego, ale nie byliśmy w stanie uzasadnić związanego z tym wzrostu komplikacji i kosztów.

Ten typ wzmacniacza, którego używamy (zmontowany na jednej płycie drukowanej i przeznaczony do montażu wewnątrz kolumny), jest często opisywany jako „wzmacniacz płytowy” („plate amplifier”).

Zdecydowaliśmy się na zasilacz pudełkowy do głośników, ponieważ znacznie upraszcza to konstrukcję i eliminuje potrzebę



Rysunek 1. Konfiguracja podstawowego systemu głośników półkowych. Sygnały audio lewego i prawego kanału oraz zasilanie 24 V DC są doprowadzane do jednego z głośników (może być lewy lub prawy, w zależności od tego, jak go okablujemy wewnątrz). Jeden z kanałów wewnętrznej wzmacniacza mocy zasila poprzez pasywną zwrotnicę głośnik wysoko- i niskotonowy, podczas gdy drugi kanał wzmacniacza zasila przewodami połączeniowymi drugi głośnik. Ten również posiada wewnętrzną pasywną zwrotnicę, która rozdziela te sygnały przed przekazaniem ich do głośnika wysoko- i niskotonowego.



Rysunek 2. Jeśli budujemy pełną wersję z subwooferami, wnętrze głośnika półkowego jest identyczne, jak na poprzednim rysunku. Jednakże sygnał wejściowy trafia teraz do pierwszego subwoofera, gdzie jest dzielony na składową sygnału o niskiej i wysokiej częstotliwości przy pomocy zwrotnicy aktywnej. Oba sygnały z wyjścia dolnozaporowego trafiają do wejścia stereo w pierwszym głośniku półkowym, a następnie do drugiego głośnika półkowego, jak poprzednio. Sygnały z wyjścia dolnoprzepustowego po zmiksowaniu trafiają do drugiego wzmacniacza mocy w pierwszym subwooferze, a wyjścia tego wzmacniacza zasilają bezpośrednio dwa większe głośniki niskotonowe w każdej z kolumn niskotonowych.

jakiegokolwiek okablowania sieciowego w projekcie. Tak więc, jeśli jesteś pewny swoich umiejętności stolarskich i potrafisz zbudować i podłączyć wzmacniacze, może to być dobry projekt do wypróbowania.

Należy zauważyć, że wyjście liniowe z subwoofera jest filtrowane dolnozaporowo, więc kiedy używane są subwoofery, monitory nie muszą przetwarzać sygnałów o niskiej częstotliwości. W tej konfiguracji, wychylenie membrany w monitorach jest o wiele niższe niż w konfiguracji samodzielnej. W rezultacie, średnica pasma jest czystsza, a system jest w stanie uzyskać znacznie wyższy poziom natężenia dźwięku.

Uwagi dotyczące konstrukcji monitorów

Wybrany przetwornik niskotonowy to Altronics C3038 o średnicy 130 mm (5 cali) z aluminiową membraną. Po wielu testach i analizach zdecydowaliśmy się na niego, ponieważ dobrze spisywał się sam w obudowie o skromnej objętości.

Przetwornik ten może być również stosowany w systemie dwudrożnym z podziałem częstotliwości około 3 kHz, czyli powyżej normalnego zakresu częstotliwości wokalu, co skutkuje mniej słyszalnymi zniekształceniami. Przetwornik ma również doskonały stosunek jakości do ceny.

Zdecydowaliśmy się na niego po zbadaniu kilku mniejszych 100 mm (4-calowych) głośników niskotonowych. Wszystkie one wypadły kiepsko w zakresie basów. Rozważaliśmy również większe przetworniki, o średnicy 150–180 mm (6–7 cali).

Wiele z nich jest w stanie emitować dobry bas, ale wszystkie wymagają rozmiar obudowy znacznie przekraczający 16 litrów, na który się zdecydowaliśmy. Już tylko to jest kompromisem, ponieważ naszym pierwotnym celem projektowym było zmniejszenie obudowy do 10 litrów.

Przetwornik Altronics C3038 ma deklarowaną moc 20–40 W, pasmo przenoszenia od 46 Hz do 10 kHz, średnicę cewki 25 mm, średnicę całkowitą 130 mm i efektywność 87 dB @ 1 W/1 m.

Te dane są typowe dla przetwornika tej wielkości. Natomiast jego waleorem jest bardzo szerokie pasmo przenoszenia, aż do 10 kHz. Dzięki temu możemy go łatwo zintegrować w systemie dwudrożnym z głośnikiem wysokotonowym.

Mimo to, najlepiej jest unikać podawania na głośnik niskotonowy sygnałów aż do 10 kHz, ponieważ stwierdziliśmy, że zachowuje się on tam dość nieprzewidywalnie, co wymagało odpowiedniej elektroniki zwrotnicy.

Przy dostarczonej mocy 30 W można osiągnąć ponad 100 dB SPL w odległości jednego metra. W warunkach domowych jest to naprawdę głośno. Z bliska odpowiada to głośności młota pneumatycznego. Mimo niewielkich rozmiarów, te przetworniki mogą zapewnić solidną moc wyjściową.

Modelowanie działania tego przetwornika w proponowanej obudowie pokazało, że możemy uzyskać wyrównanie „rozszerzonego zakresu basu” (rysunek 3), gdzie wymuszamy niewielkie poszerzenie zakresu niskich tonów kosztem płaskiej charakterystyki przy niższych częstotliwościach. To dobry kompromis dla mniejszych głośników. Należy pamiętać, że po dodaniu do systemu opcjonalnych subwooferów, przejmują one odtwarzanie częstotliwości poniżej 90 Hz, dzięki czemu uzyskujemy bardziej płaską ogólną charakterystykę.

Wybraliśmy dość wąską obudowę, o szerokości zewnętrznej 21 cm. Dzięki temu głośnik może stać na biurku lub półce, nie zajmując dużo miejsca. Wysokość i głębokość głośnika zostały dobrane tak, aby zapewnić wymaganą objętość wewnętrzną 16 litrów. Pozostałe wymiary to 297 mm głębokości i 390 mm wysokości.

Głębokość 297 mm pozwala na przecięcie na pół standardowego arkusza skleiki o wymiarach 1200×600 mm w celu wykonania paneli: bocznych i górnego, co minimalizuje ilość odpadów i koszty.

Drugim aspektem obudowy jest rozmieszczenie głośnika niskotonowego i wysokotonowego. Zauważcie, że głośnik wysokotonowy został umieszczony tuż przy głośniku niskotonowym. Powodem tego jest zminimalizowanie odstępów pomiędzy środkami przetwornika nisko- i wysokotonowego.

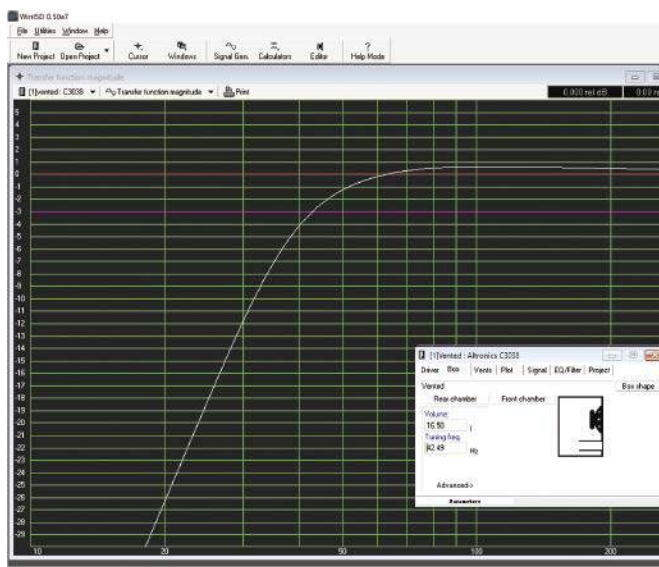
W miarę jak słuchacz porusza głową, usytuowanie tych przetworników blisko siebie minimalizuje różnice w odległości każdego z nich od ucha słuchacza. W rezultacie dźwięk z głośników pozostaje spójny w całym obszarze odsłuchu. Innymi słowy, głośniki te zapewniają dobre nagłośnienie także poza ich osią.

Zwrotnica

Przetwornik niskotonowy C3038 całkiem dobrze radzi sobie z wyższymi częstotliwościami. Zdecydowaliśmy się na zastosowanie częstotliwości podziału z głośnikiem wysokotonowym przy około 3,2 kHz, co pozwala na pokrycie krytycznego zakresu ludzkiego głosu obejmującego 300–3000 Hz.

Niestety, przetwornik ten ma poważne obszary poszarpanej charakterystyki w zakresie częstotliwości 9–11 kHz, co jest wynikiem zastosowania bardzo sztywnej membrany. Tworzy to grupę maksymalnych i minimalnych obszarów efektywności w zakresie górnych częstotliwości. Początkowo wypróbowaliśmy zwrotnicę, która nie uwzględniała tego faktu i szybko zrozumieliśmy nasz błąd!

Druga wersja zwrotnicy zawierała specjalne filtry do „wycinania” tych maksimów. Działało to, ale sprawiło, że zwrotnica była bardzo duża i droga. Postanowiliśmy więc wypróbować zwrotnicę drugiego rzędu i połączyliśmy filtr ze zwrotnicą. Zwrotnica, a właściwie impedancja sekcji niskotonowej, została zaprojektowana tak, aby tłumić szczyty 9–11 kHz bardziej niż zwykle.



Rysunek 3. Zaimplementowaliśmy parametry głośnika niskotonowego Altronics C3038 do programu WinISD i eksperymentowaliśmy z wymiarami małej wentylowanej obudowy, uzyskując odpowiedź pokazaną na wykresie. Zapewnia to nieco rozszerzone pasmo niskich tonów kosztem trochę mniejszej liniowości w tym zakresie. Biorąc pod uwagę, że odchylenie jest mniejsze niż 1 dB, raczej nie da się go zauważyć na słuch. A odpowiedź niskotonowa jest poszerzona w dół o około 10 Hz, co jest bardzo korzystne.

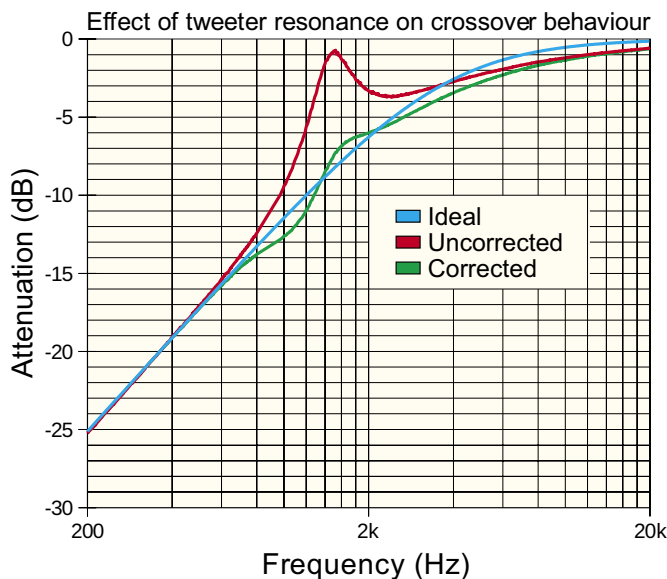
Jedną z konsekwencji tego zabiegu jest spadek impedancji kolumny do około 4 Ω w zakresie 2,5–5 kHz. Nie będzie to przeszkadzać w przypadku większości wzmacniaczy. Końcowy dźwięk przetwornika niskotonowego jest bardzo czysty i nie ma w sobie nic z szorstkości „surowego” brzmienia głośnika.

Głośnik wysokotonowy

Bardzo chcieliśmy wybrać dobry głośnik wysokotonowy, ponieważ gdy głośnik wysokotonowy ma zbyt poszarpaną charakterystykę lub duże zniekształcenia, rezultatem jest kolumna, która powoduje zmęczenie przy dłuższym słuchaniu. Wybrany głośnik wysokotonowy musi również przetwarzać niższe, jak to tylko możliwe, częstotliwości wynikające z podziału pasma w zwrotnicy, abyśmy mogli uniknąć wysyłania sygnałów w okolicy 9–11 kHz do głośnika niskotonowego.

Wybraliśmy więc tweeter Vifa, Altronics C3019. Jest to bardzo dobry głośnik wysokotonowy w akceptowalnej cenie, ale stawia przed konstruktorem wyzwanie w postaci znacznego piku impedancji przy około 1,75 kHz. Ten szczyt impedancji jest wynikiem rezonansu głośnika.

Głośnik wysokotonowy wykorzystuje ferrofluid (płyn o właściwościach ferromagnetycznych) w szczelinie układu magnetycznego.



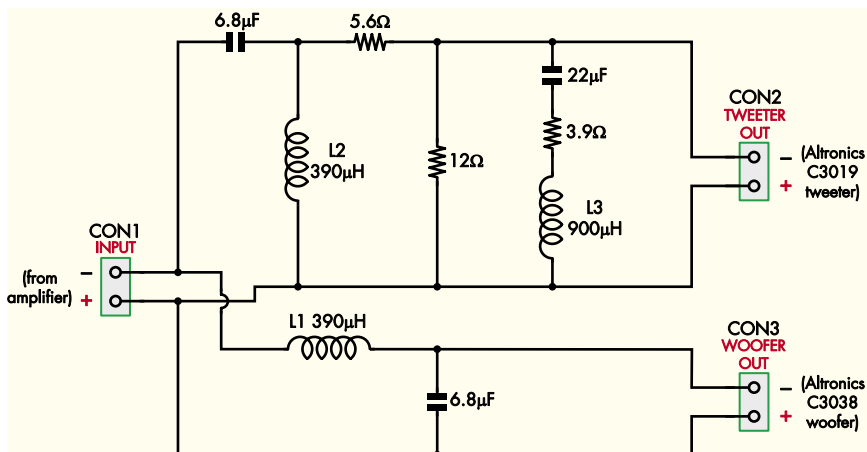
Rysunek 4. Impedancja głośnika wysokotonowego zmienia się wraz z częstotliwością, wpływając na działanie zwrotnicy. Niebieska linia pokazuje prostą zwrotnicę z obciążeniem rezystancyjnym 4 Ω . Czerwona krzywa pokazuje tę samą zwrotnicę z głośnikiem wysokotonowym Vifa jako obciążeniem. Zielona krzywa pokazuje skorygowaną odpowiedź naszej poprawionej zwrotnicy, z obwodem kompensacyjnym w celu zmniejszenia efektu rezonansu głośnika wysokotonowego.

Pomaga to w chłodzeniu cewki i zwykle tłumi rezonans przetwornika. Tak więc, w większości ferrofluidowych głośników wysokotonowych, impedancja przetwornika jest dość płaska pomimo rezonansu. Natomiast głośnik wysokotonowy C3019 ma charakterystykę pośrednią (między ferrofluidowcami, a głośnikami ze szczeliną powietrzną). Impedancja głośnika wynosi nominalnie 4 Ω , ale przy 1,75 kHz osiąga szczyt wynoszący 10 Ω .

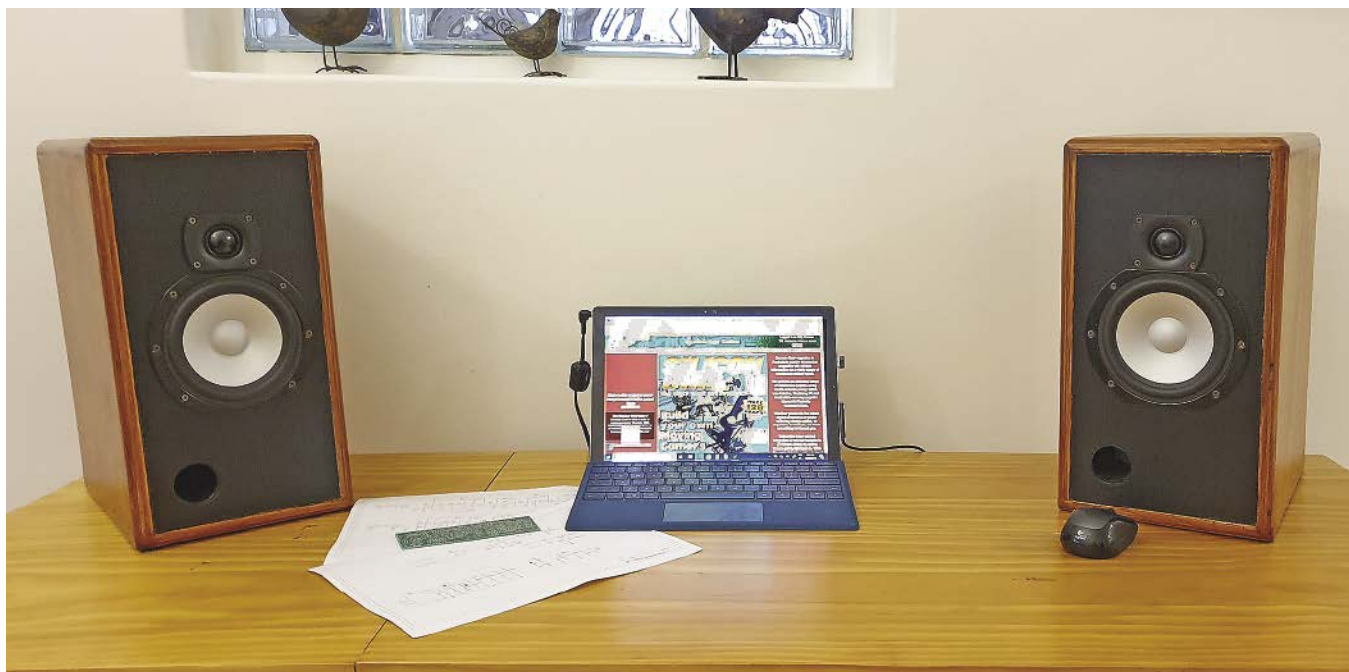
Musimy sobie z tym maksimum poradzić. Rysunek 4 pokazuje na niebiesko zachowanie idealnej zwrotnicy pierwszego rzędu, rzeczywistą charakterystykę na czerwono i skorygowaną na zielono. Korekta jest realizowana za pomocą obwodu tłumiącego LCR, składającego się (w naszym przypadku) z cewki o indukcyjności około 1 mH, kondensatora 22 μ F i rezystora 3,9 Ω .

Powiększa to wprawdzie koszty projektu, ale jest niezbędne do uzyskania dobrego dźwięku. Taki pik jak ten pokazany bez układu korekcyjnego jest odpowiedzialny za to, że wiele głośników wysokotonowych brzmi ostro i „męcząco”.

Wynikowy układ pasywnej zwrotnicy drugiego rzędu pokazano na Rysunku 5. Jest to dość skomplikowana zwrotnica jak na głośnik



Rysunek 5. Ostateczny schemat ideowy zwrotnicy, z dodatkowym filtrowaniem dla głośnika niskotonowego w celu skutecznego wycięcia sygnałów w rejonie załamania 9–11 kHz. Zawiera ona również obwód RLC (3,9 Ω /22 μ F/900 μ H) w celu wyrównania odpowiedzi głośnika wysokotonowego ze względu na rezonans pokazany na rysunku 4, oraz dzielnik rezystancyjny 5,6 Ω /12 Ω w celu dopasowania poziomów i impedancji obu przetworników do jednego źródła sygnału.



Nie musisz budować subwooferów w charakterze podstawek pod monitory; dwa główne głośniki są idealne do pracy na biurku z komputerem, odtwarzaczem MP3 itp. (choć kosztem nieco słabszego basu). Ponieważ mają własne zasilanie, podłącza się je prosto do praktycznie każdego źródła dźwięku, od „line out” po gniazda słuchawkowe...

dwudrożny, ale jest ona niezbędna dla uzyskania pożądanej jakości dźwięku.

Wszystkie trzy rezystory mogą być typu drutowego o mocy 5 W. Kondensatory również nie są zbyt trudne do zdobycia: kondensatory o pojemności 6,8 μF mogą być zarówno polipropylenowe metalizowane, jak i elektrolityczne bipolarne. Ja zdecydowałem się na te pierwsze, ale elektrolityczne też są w porządku. Ze względu na wysoką pojemność kondensatora 22 μF , musi on być elektrolityczny (bipolarny).

Od Red. EdW: zobacz uwagi na temat kondensatorów i cewek do zwrotnic, EdW 10/2022, str. 41.

Pozostaje nam tylko pytanie, skąd wziąć lub jak zrobić cewki powietrzne 390 μH i 900 μH o niskiej rezystancji dla prądu stałego, tak aby były jak najbardziej zbliżone do idealnych cewek.

Na szczęście okazuje się, że można po prostu kupić pełne szpule emaliowanego drutu miedzianego (DNE), a drut na szpuli będzie miał już mniej więcej właściwe wartości indukcyjności!

Od Red. EdW: Widzimy tu jeden problem – początek nawoju drutu na szpuli jest zwykle niedostępny.

Przetestowaliśmy szpule od Altronicsa (i są one podane na liście części). Nie jesteśmy pewni co do szpul od innych producentów. Musiałbyś sam zmierzyć ich indukcyjności.

Od Red. EdW: pomiary indukcyjności cewek i pojemności kondensatorów użytych do budowy zwrotnic należą do elementarnych czynności.

To naprawdę łut szczęścia, że 100 g szpula drutu DNE o średnicy 1 mm ma indukcyjność prawie dokładnie 390 μH . W rzeczywistości dla cewki L3 chcieliśmy zastosować indukcyjność 1 mH, ale 100-gramowa szpuka DNE o średnicy 0,8 mm ma indukcyjność 900 μH , i to jest wystarczająco blisko optimum.

Wynikająca z tego różnica to przesunięcie punktu podziału zwrotnicy z 3,0 kHz do 3,2 kHz. Wykorzystanie w ten sposób całych szpul zwalnia konstruktorów z pracy polegającej na żmudnym nawijaniu cewek wg żądania.

Trzy cewki są zamontowane na płycie drukowanej zwrotnicy prostopadle do siebie, wzdłuż osi X-Y-Z, tzn. jedna jest skierowana na północ/południe, jedna na wschód/zachód i jedna w górę/dół. Oznacza to,

że są one «ortogonalne», więc ich pola magnetyczne nie będą na siebie oddziaływać.

W przeciwnym razie uzyskalibyśmy niepożądany transformator powietrzny pomiędzy dwoma lub więcej cewkami i zwrotnica nie działałaby zgodnie z przeznaczeniem.

Wbudowany wzmacniacz

Używane przez nas gotowe moduły wzmacniaczy nie kosztują wiele, a mimo to zapewniają świetną wydajność. Dla zapalonych hobbystów myśl o zakupie gotowych modułów wzmacniaczy była trudną decyzją, z którą musieliśmy się zmierzyć... ale jesteśmy zadowoleni, że to zrobiliśmy.

Wzmacniacz dostarcza około 30 W RMS do dwóch głośników 8 Ω , co jest więcej niż wystarczające dla wszystkiego, co nie jest związane z dyskoteką.

W połączeniu z subwooferami monitory nigdy nie będą przetwarzać częstotliwości poniżej około 90 Hz, więc 30 W to naprawdę bardzo poważna moc. Wzmacniacz sterowany jest z liniowego wejścia stereo.

Jak wspomniano wcześniej, wzmacniacz korzysta z zewnętrznego zasilacza, który podłączany jest za pomocą wtyku DC 2,5/5,5 mm. Pozwala to na zachowanie prostoty i uniknięcie okablowania sieciowego wewnątrz głośnika.

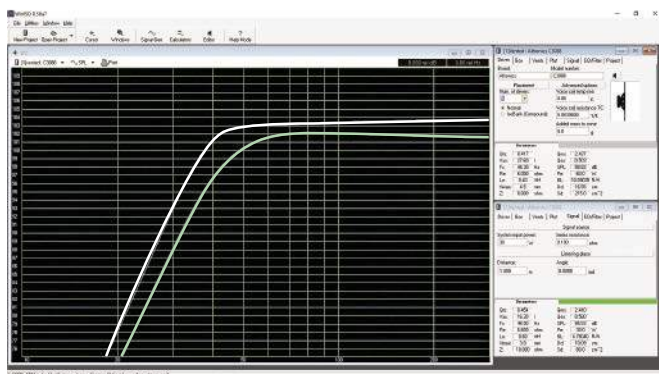
W specyfikacji podaliśmy moduły wzmacniaczy na układzie TDA7498, dostępne w serwisie eBay (i wielu innych). Teoretycznie są one w stanie dostarczyć 80 W do głośnika 8 Ω , ale my zasilamy je niższym napięciem niż maksymalne. Wybraliśmy ten moduł na TDA7498 po zakupie i przetestowaniu wielu innych wzmacniaczy.

Do tego projektu poważnymi kandydatami były dwie główne grupy wzmacniaczy: zawierające układ scalony TPA3116 lub TDA7498. Oba układy są wzmacniaczami klasy D i oba pracują z pojedynczym zasilaniem. Są bardzo wydajne, mają niewielki radiator w porównaniu do wzmacniaczy liniowych i są bardzo przystępne cenowo.

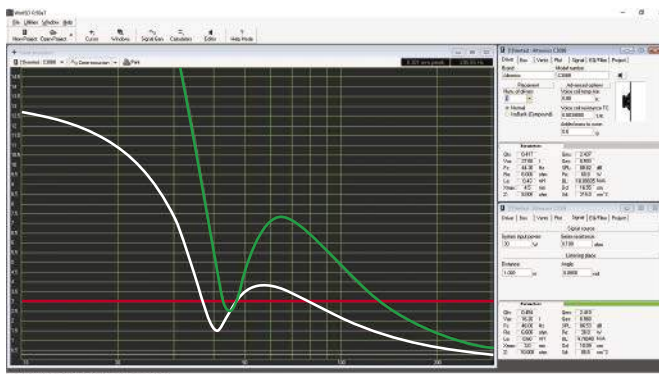
Rozważaliśmy zastosowanie wzmacniaczy liniowych, na przykład wzmacniaczy dyskretnych lub wzmacniaczy wykorzystujących świetne układy scalone LM3886. Zapewniłyby one nieco lepszą jakość, ale

wszystkie wymagają zasilaczy symetrycznych, a to prowadzi nas do umieszczenia transformatorów, prostowników i okablowania sieciowego wewnątrz głośnika. Byłyby one również droższe i generowałyby znacznie więcej ciepła wewnątrz obudowy.

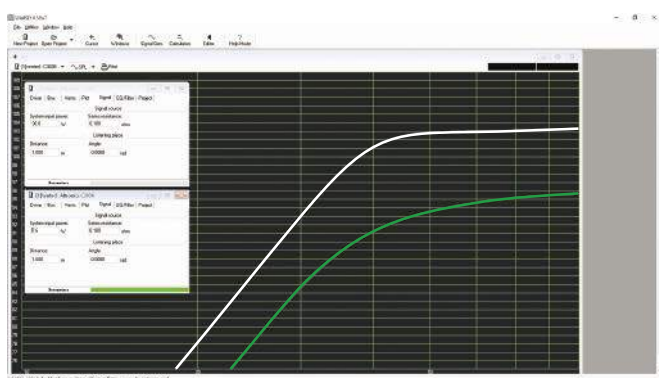
Patrząc na opcje klasy D z układami TPA3116 i TDA7498, kupiliśmy szereg wzmacniaczy do przetestowania. Znaleźliśmy kilka problemów z większością wzmacniaczy klasy D dostępnych obecnie na rynku.



Rysunek 6. Oczekiwane poziomy wyjściowe SPL dla głośników niskotonowych o średnicy 130 mm (zielony) w porównaniu do głośników o średnicy 200 mm (szary), oba wykresy przy mocy 30 W. Głośniki niskotonowe o większej średnicy generują nie tylko wyższy SPL, ale również mają dolny kraniec pasma -3 dB w przybliżeniu o 10 Hz niższy, około 35 Hz w porównaniu do 45 Hz (symulacja).



Rysunek 7. Symulowane wartości wychYLENIA membrany (w mm) dla głośników niskotonowych 130 mm (zielony) i 200 mm (szary). Głośniki niskotonowe 200 mm mają akceptowalne (<4,5 mm) wychYLENIA membrany aż do punktu -3 dB przy 35 Hz, podczas gdy głośniki niskotonowe 130 mm przekraczają maksymalny skok membrany i tym samym generują zniekształcenia przy znacznie wyższej częstotliwości przy tym samym poziomie mocy; tu około 100 Hz.



Rysunek 8. SPL w zależności od częstotliwości dla głośników niskotonowych 130 mm (zielony) i 200 mm (szary) przy najwyższym praktycznym poziomie mocy dla każdego z nich; odpowiednio 7,6 W i 30 W. Poprzez ograniczenie mocy basu głośnika 130 mm do 7,6 W, wychYLENIE membrany jest utrzymane w granicach rozsądku, ale maksymalny SPL jest o około 10 dB niższy w porównaniu do większych głośników 200 mm (symulacja).

Niektóre z nich są sprzedawane jako wzmacniacze „2.1 kanałowe”, z wyjściem na subwoofer i stereofonicznymi wyjściami na główne głośniki. Niestety, żaden z nich nie posiada filtrów na głównych wyjściach, co oznacza, że do podstawowych głośników wysyłany jest sygnał pełnego pasma, łącznie z częścią wysyłaną do subwoofera. Jest to wada, która czyni te urządzenia praktycznie bezużytecznymi.

Radiator wielu konstrukcji jest bardzo słaby. Bardzo często radiator mocowany jest za pomocą tylko jednej śrubki. Jest to tak delikatna konstrukcja, że nie możemy polecać jej do zastosowania w kolumnie.

Wydaje się dziełem przypadku, które wzmacniacze mają dobry kontakt pomiędzy układem scalonym wzmacniacza a radiatorem. Ale jest to coś, co możemy naprawić sami.

Również napięcie znamionowe kondensatorów elektrolitycznych w wielu z tych produktów jest bardzo zbliżone do napięcia pracy. To może nie brzmieć niepokojąco, ale tak jest. Niezawodność kondensatorów elektrolitycznych jest silnie uzależniona od tego, jak odległe od ich maksymalnych wartości znamionowych są warunki eksploatacji (obejmuje to temperaturę, napięcie i prąd tętniący).

Wylutowaliśmy z jednego ze wzmacniaczy kondensatory o napięciu znamionowym 25 V, przy napięciu zasilającym modułu 24 V, i przetestowaliśmy je na zasilaczu. Każdy z nich uległ katastrofalnemu uszkodzeniu przy napięciu 26–28 V. To zdecydowanie za mało, abyśmy mogli polecić ich użycie.

Wzmacniacze wykorzystujące TDA7498 mogą pracować przy napięciu do 32 V DC, a wybrany przez nas wzmacniacz ma solidną konstrukcję mechaniczną. Biorąc pod uwagę, że do zasilania wzmacniaczy podajemy 24 V DC z wtyczki, mamy spory margines napięcia pracy kondensatorów elektrolitycznych.

Jako bonus, polecany przez nas wzmacniacz nie zawiera regulatorów głośności i ma proste wejścia i wyjścia na złączach śrubowych. Dzięki temu jest bardzo przystępny cenowo. Powinien być w stanie znaleźć polecany wzmacniacz za około 9 USD (~AUD 14) za sztukę, co jest znacznie taniej niż kosztowałaby budowa wzmacniacza dyskretnego lub stosującego LM3886.

Wybraliśmy również zasilacz wtyczkowy 24 V 6 A z eBay'a za mniej niż AUD 35.

Integrując wzmacniacz, złącza wejściowe, gniazda wyjściowe do głośników i regulację głośności z aluminiowym panelem, możemy zbudować samodzielny wzmacniacz, gotowy do zainstalowania w tylnej części głośnika podstawkowego.

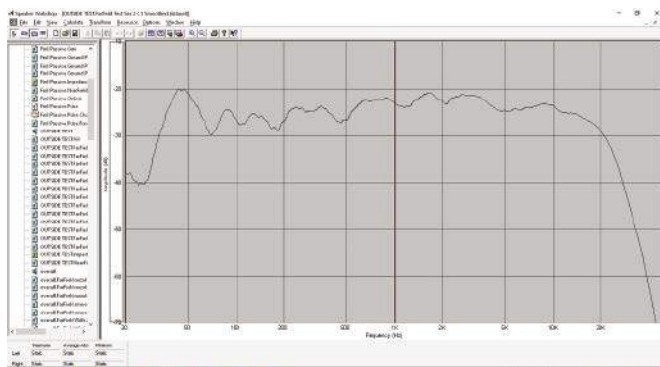
Konstrukcja subwoofera

Opcjonalne subwoofery zapewniają kilka korzyści. Ich większe, 200 mm (8 cali) przetworniki mogą przenosić znacznie więcej mocy ciągłej niż przetworniki w głośnikach półkowych, ponieważ mają cewki o średnicy 40 mm (1,5 cala). Dodatkowo, długość cewki i zakres wychYLENIA zawieszenia pozwala na większy skok membrany. Dzięki temu przetwornik ma liniowy skok membrany znacznie przekraczający ±4,5 mm.

To, w połączeniu z faktem, że membrany mają większą powierzchnię niż przetworniki basowe w kolumnach półkowych, oznacza, że subwoofery są znacznie lepiej przystosowane do emisji niskich częstotliwości przy wysokich poziomach mocy.

Aby zilustrować tę różnicę, na **rysunku 6** pokazano wynik symulacji w programie WinISD poziomu ciśnienia akustycznego (SPL) w zakresie częstotliwości poniżej 300 Hz, z subwoofera zasilanego mocą 30 W (szara krzywa) i z głośnika półkowego przy mocy 30 W (zielona krzywa).

Widać, że subwoofer zwiększa emisję niskich tonów o około 3–5 dB i rozszerza pasmo przenoszenia w dół o około 10 Hz, do mniej więcej 35 Hz.

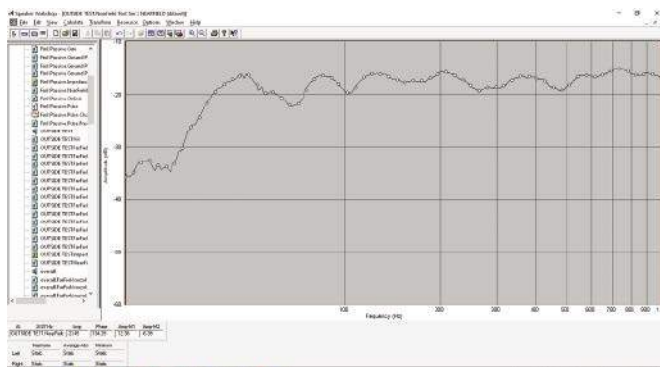


Rysunek 9. Pomiar „dalekiego pola” odpowiedzi systemu głośnikowego, dla jednego monitora i jednego subwoofera. Odpowiedź jest dość płaska od około 60 Hz do prawie 20 kHz, zmieniając się zaledwie o kilka dB. Szczyt przy 50 Hz został uznany za spowodowany odbiciami dźwięku od pobliskiej ściany.

Ale to nie wszystko. Rysunek 7 pokazuje obliczone wychylenia membrany dla obu głośników. Przy 30 W, przetwornik Altronics C3088 w subwooferze charakteryzuje się aż do około 35 Hz wychyleniem membrany znacznie poniżej swojego liniowego skoku 4,5 mm. Przy mocnymysterowaniu, przetwornik ten elegancko ogranicza wychylenie membrany bez jej uszkodzeń.

Jednak przy mocy 30 W, znacznie mniejszy przetwornik w głośniku półkowym przy około 38 Hz próbowałby wychylić membranę w granicach ± 7 mm, co znacznie przekracza jego możliwości. Głośnik po prostu nie jest w stanie tego zrobić i membrana osiąga krańce swojego mechanicznego skoku, powodując wręcz charczenie dźwięku.

Ponadto, gdy głośnik znajduje się w ekstremalnym zakresie wychyleń, cewka nie znajduje się całkowicie w polu magnetycznym „szczeliny”.



Rysunek 10. Pomiary w „bliskim polu” dają bardziej dokładny obraz odpowiedzi systemu w zakresie niskich tonów. Szczyt przy 50 Hz nie jest już tak zauważalny, a niskie tony rozciągają się aż do nieco poniżej 40 Hz.

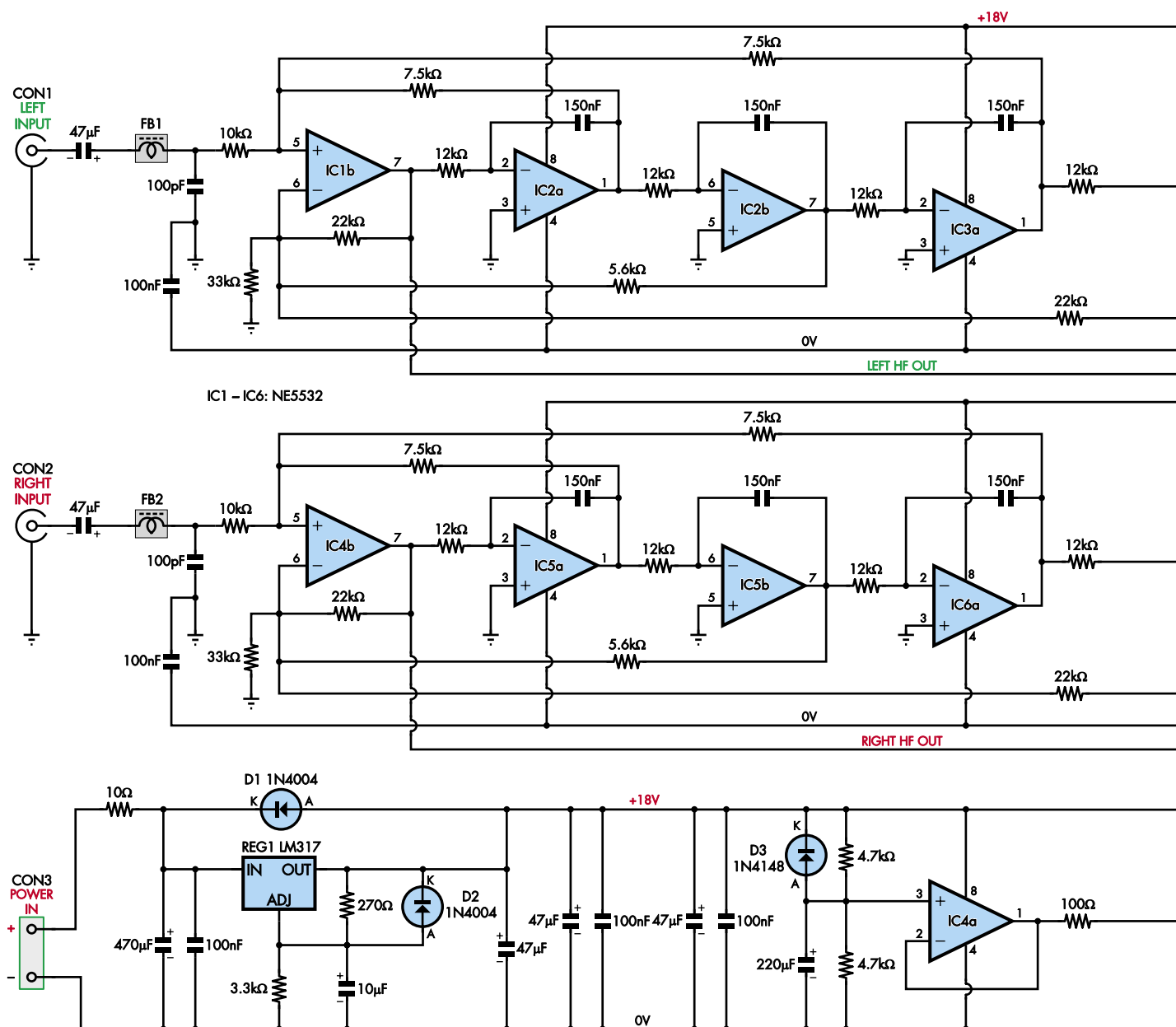
Tak więc nie tylko bas jest zniekształcony, ale zniekształcone są również wszystkie inne dźwięki emitowane z głośnika.

Oczywiście, przy zmniejszeniu głośności, monitor działa bardzo dobrze, ale musimy zdawać sobie sprawę, że prawa fizyki nakładają ograniczenia na to, czego możemy wymagać od głośnika. Dodanie subwooferów pozwala nam uniknąć wysyłania częstotliwości poniżej 90 Hz do głośników półkowych, co pozwala wyeliminować zniekształcenia opisane powyżej. Zamiast tego te najniższe dźwięki są odtwarzane przez subwoofery.

Dodatkową korzyścią jest znaczne zwiększenie mocy dostępnej dla monitorów do generowania średnich i wysokich częstotliwości, ponieważ cały sygnał niskotonowy został skierowany do oddzielnej wzmacniacza i głośnika.



...ale jeśli zbudujesz subwoofery, będą one świetnymi podstawkami dla głównych głośników. A ponieważ bas jest bezkierunkowy, można skierować membrany głośników tam, gdzie małe palce (czy ostre pazurki) nie zrobią krzywdy przetwornikom.



ACTIVE CROSSOVER FOR BOOKSHELF SPEAKERS

W idealnej sytuacji, monitory nie powinny reprodukować więcej niż około 7–10 W sygnałów o częstotliwości poniżej 100 Hz, co ogranicza wychylenie membrany do bardziej przystępnych 3–4 mm. Osiągalny SPL basu jest w tym przypadku oczywiście mniejszy. **Rysunek 8** pokazuje maksymalną praktyczną moc wyjściową w zakresie niskich częstotliwości, do 90 Hz, osiąganą przez przetworniki C3088 i C3038.

Aktywna zwrotnica, którą wykorzystujemy do rozdzielania sygnału pomiędzy subwoofery i głośniki podstawkowe, pozwala na wysterowanie monitorów z pełną mocą w całym ich zakresie odtwarzania, podnosząc osiągalny SPL do poziomu subwoofera.

Jeśli chodzi o obudowy subwoofertów, utrzymaliśmy ich szerokość i głębokość na takim samym poziomie jak monitorów. Dzięki temu subwoofery mogą być użyte lub wręcz „ukryte” jako podstawki pod głośniki. Dzięki temu mamy wygodną 35-litrową obudowę, w której możemy zamontować przetwornik Altronics C3088.

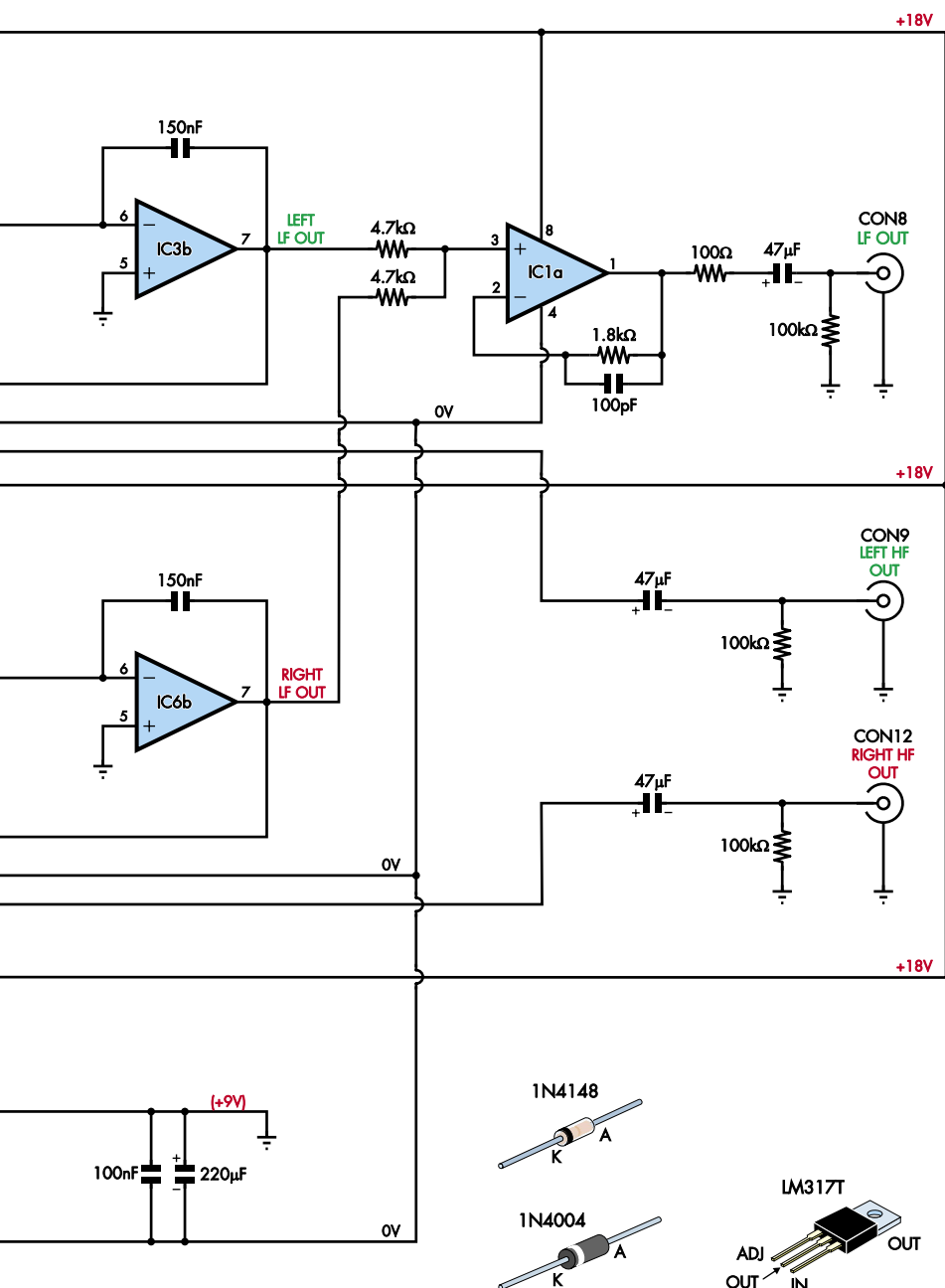
Być może zauważyliście pewien problem z tym związany: przetworniki niskotonowe o średnicy 200 mm raczej nie zmieszczą się w zwykły sposób w obudowie o szerokości 210 mm. Ale ponieważ jest to subwoofer i pracuje tylko do 90 Hz, jego propagacja dźwięku jest praktycznie dookólna, a ludzkie ucho nie jest zdolne do wycucia kierunku emisji tak niskich dźwięków.

Wykorzystujemy ten fakt i montujemy przetwornik z boku obudowy, a nie z przodu.

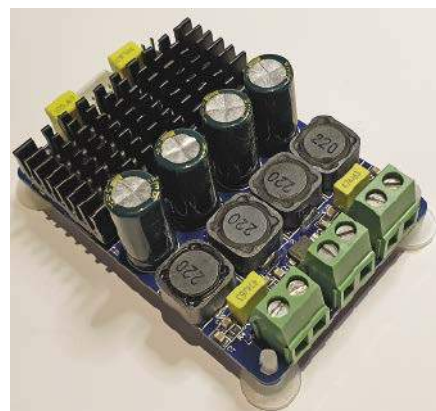
Podobnie, port We-Wy umieściliśmy z tyłu obudowy, gdyż jego dokładna lokalizacja nie jest krytyczna. Wszystkie te elementy można przenieść, jeśli wymaga tego np. umiejscowienie czy zastosowanie kolumn.

Ogólna wydajność

Pomiar odpowiedzi częstotliwościowej głośników, jeśli nie dysponujemy komorą bezchową, jest trudny. Jednakże spróbaliśmy tego, używając mikrofonu pomiarowego Behringer ECM8000,



Rysunek 11. Schemat ideowy aktywnej zwrotnicy, która służy do rozdzielania przychodzącego sygnału stereo, tak aby składowe o wysokiej częstotliwości mogły być doprowadzone do pary monitorów. Składowe o niskich częstotliwościach są miksowane do sygnału monofonicznego, buforowane przez IC1a, a następnie podawane do wzmacniacza, który może zasilac jeden lub dwa subwoofery. Układ pracuje z tym samym zasilaczem 24 V DC, co używany do zasilania wzmacniacza subwoofera. Napięcie zasilania jest redukowane regulowanym stabilizatorem do wartości 18 V, z generowaniem szyny wirtualnej masy o potencjale 9 V w celu napięciowego przesunięcia sygnału.



(Na zdjęciu powyżej): Wzmacniacz stereo 2×80 W klasy D, który kupiliśmy na eBayu za mniej niż 20\$ z wliczoną przesyłką. Nie mogliśmy zbudować takiego za podobną cenę, a jest doskonałej jakości!

niskoszumnego przedwzmacniacza mikrofonowego i oprogramowania Speaker Workshop PC.

Pomiary w bliskim polu mogą być wykonane ze znośną dokładnością do skromnej częstotliwości (powiedzmy około 1 kHz). Pomiary w polu dalekim są silnie zakłócone przez odbicia i rezonanse pomieszczenia, ale są bardziej reprezentatywne dla tego jak system głośnikowy brzmi w rzeczywistości.

Prezentowane tutaj pomiary są mieszanką obu tych metod. Najpierw przyjrzyjmy się pomiarom w polu dalekim pokazanym na **rysunku 9**. Zostały one wykonane na zewnątrz, z głośnikiem znajdującym się w odległości około 3 m od najbliższej przeszkody. Można zauważyć pik przy 50 Hz, który jest spowodowany odbiciem od otoczenia. Pomiary w polu bliskim (obok) dają lepszy wgląd w odpowiedź głośników w zakresie niskich częstotliwości.

Przeniesienie mikrofonu w miejsce bliższe obudów, około 50 cm od głośnika i znajdujące się w równej odległości od subwoofera

i monitora, daje odpowiedź w zakresie niskich tonów pokazaną na **rysunku 10**.

Zmierzona umowna dolna granica pasma przenoszenia, −3 dB, wynosi 34 Hz. Pozostaje mała nieliniowość w rejonie 50 Hz. Poza tym, odpowiedź jest, zgodnie z oczekiwaniami, bardzo płaska.

Wprawne oko zauważy, że drugi wykres jest usytuowany o kilka dB wyżej niż pierwszy. Jest to spowodowane tym, że mikrofon znajduje się bliżej głośnika.

Odpowiedź jest tak równomierna i dynamiczna, jak sugerują to wykresy. Jeśli zbudujecie te kolumny, myślimy, że będziecie zachwyceni ich dźwiękiem, a wasz portfel nie będzie zbyt chudy!

Konstrukcja aktywnej zwrotnicy

Jak wspomniano wcześniej, aktywna zwrotnica służy do rozdzielania przychodzącego stereofonicznego sygnału audio na trzy różne ścieżki:

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

Głośniki półkowe – do budowy jednej pary (2 sztuki)

Obudowy:

- 2 szt. 130 mm (5 cali) 40 W głośników niskotonowych z membraną aluminiową [Altronics C3038].
- 2 szt. głośników wysokotonowych 25 mm (1 cali) 100 W Vifa BC25SC55 [Altronics C3019]
- 1 zespół kompletnego wzmacniacza (patrz poniżej)
- 2 układy pasywnej zwrotnicy (patrz poniżej)
- 2 arkusze 600×1200 mm 15-milimetrowej sklejki szklanej
- 4 odcinki długości po 2 m listwy drewnianej „czwiercwałek” 15×15 mm lub 20×20 mm
- 80 szt. wkrętów do drewna 8G×25-28 mm samogwintujących z łbem stożkowym
- 20 szt. wkrętów do drewna 8G×15 mm samogwintujących z łbem stożkowym
- 16 wkrętów do drewna 8G×10-12 mm samogwintujących z łbem stożkowym
- 2 odcinki rury PVC o średnicy 40 mm i długości 105 mm
- 1 arkusz 80×40 mm blachy aluminiowej o grubości 1,5 mm
- 1 rolka cienkiej taśmy piankowej (np. taśma do uszczelniania drzwi)
- 1 paczka dużych zszywek (lub małe pudełko gwoździ 40 mm)
- 1 worek (w arkuszu) grubej waty Lincraft lub podobnej lekkiej akustycznej tłumiącej waty poliesterowej
- 4 arkusze papieru ściernego o gradacji 120
- 1 arkusz papieru ściernego o gradacji 240-400
- 1 mała puszka lakieru do drewna (typu jachtowego lub do okien z drewna egzotycznego)
- 1 mała puszka czarnej farby matowej lub satynowej
- 1 tubka 430-475 ml akrylowego kitu do szczelin
- 1 podwójny czerwono-czarny panel zaciskowy (do wtyków bananowych i widełek) [Altronics P9257A].
- 1 odcinek o długości 1m podwójnego kabla głośnikowego min. 2×2,5mm²
- 1 butelka 250 ml kleju do drewna typu Wikol

Dodatkowe części dla pary subwooferów

- 2 głośniki niskotonowe 200 mm (8 cali) 70 W z membranami polipropylenowymi [Altronics C3088].
- 1 płytka zespołu wzmacniacza subwoofera (patrz poniżej)
- 3 arkusze 600×1200 mm 15-milimetrowej sklejki szklanej
- 6 odcinków długości po 2 m listwy drewnianej „czwiercwałek” 15×15 mm lub 20×20 mm
- 2 odcinki rury PVC o średnicy 75 mm i długości 130 mm
- 100 szt. wkrętów do drewna 8G×25-28 mm samogwintujących z łbem stożkowym
- 16 szt. wkrętów do drewna 8G×15 mm samogwintujących z łbem stożkowym
- 8 szt. wkrętów do drewna 8G×10-12 mm samogwintujących z łbem stożkowym
- 1 arkusz 80×40 mm blachy aluminiowej o grubości 1,5 mm
- 6 arkuszy papieru ściernego o gradacji 120
- 1 arkusz papieru ściernego o gradacji 240-400
- 1 podwójny czerwono-czarny panel zaciskowy (do wtyków bananowych i widełek) [Altronics P9257A].
- 1 odcinek o długości 1 m podwójnego kabla głośnikowego min. 2×2,5mm²

Zespół wzmacniacza

- 1 arkusz 135×160 mm blachy aluminiowej o grubości 1,5 mm
- 1 wzmacniacz 100 W+100 W na układzie TDA7498, niebieska płytka PCB (dostępny na eBay)
- 1 zasilacz sieciowy 24 V 5-6 A typu pudełkowego z wtyczką DC o średnicy 2,5/5,5 mm
- 1 gniazdo wtyczkowe DC 2,5/5,5 mm do montażu na PCB [Altronics P0623]
- 1 czerwone gniazdo RCA do montażu na PCB [Jaycar PS0259]
- 1 czarne gniazdo RCA do montażu na PCB [Jaycar PS0496]
- 1 podwójny czerwono-czarny panel zaciskowy (do wtyków bananowych i widełek) [Altronics P9257A]
- 1 podwójny potencjometr logarytmiczny 10 kΩ [Altronics R2334, Jaycar RP3756]
- 1 żelka trójstykowa wtyczka do gniazda kołkowego raster 3,96 mm, oprawka i styki [Altronics P5643 + 3×P5640A, Jaycar HM3433]
- 1 gałka do potencjometru

- 8 szt. śrub M3×6
- 8 szt. podkładek ząbkowanych o średnicy 3 mm
- 4 szt. poliamidowych kołków dystansowych z gwintem M3 o długości od 10 mm do 25 mm
- 1 odcinek o długości 1 m pojedynczego przewodu ekranowanego
- 1 odcinek o długości 1 m podwójnego przewodu ekranowanego
- 1 odcinek o długości 1 m podwójnego kabla głośnikowego min. 2×2,5mm²
- 1 kawałek rurki termokurczliwej o średnicy 5 mm
- 1 mała tubka pasty termicznej
- 1 puszka czarnej farby w sprayu, odpowiedniej do aluminium

Zwrotnica pasywna

- 1 szt. dwustronna płytka PCB, kod 01101201, 137×100 mm
- 3 szt. 2-drożnych złączy śrubowych raster 5/5,08 mm do montażu na PCB (CON1-CON3)
- 1 cewka indukcyjna 900 μH z rdzeniem powietrznym (L1; 100 g DNE o średnicy 0,8 mm) [Altronics W0407].
- 2 cewki indukcyjne 390 μH z rdzeniem powietrznym (L2, L3; po 100 g DNE o średnicy 1,0 mm) [Altronics W0408]
- 1 szt. kondensator elektrolityczny bipolarny osiowy do zwrotnic 22 μF/100 V [Jaycar RY6912]
- 2 szt. kondensatorów foliowych osiowych do zwrotnic 6,8 μF/100 V [Jaycar RY6956 lub RY6906]
- 1 szt. rezystor drutowy 12 Ω 5 W 5%
- 1 szt. rezystor drutowy 5,6 Ω 5 W 5%
- 1 szt. rezystor drutowy 3,9 Ω 5 W 5%
- 4 szt. dużych plastikowych opasek kablowych

Zespół wzmacniacza subwoofera

- Wszystkie części wymienione powyżej dla zespołu wzmacniacza do monitora, oprócz blachy aluminiowej, plus:
- 1 arkusz 250×165 mm blachy aluminiowej o grubości 1,5 mm
 - 1 czerwone gniazdo RCA do montażu na panelu [Jaycar PS0259].
 - 1 czarne gniazdo RCA do montażu na panelu [Jaycar PS0496]
 - 1 dwustronna płytka PCB, kod 01101202, 132×45 mm
 - 6 szt. 2-drożnych złączy śrubowych raster 5/5,08 mm do montażu na PCB (CON4-CON9)
 - 6 szt. podstawek pod IC DIL-8 (dla IC1-IC6; opcjonalnie)
 - 2 szt. koralików ferrytowych (FB1, FB2)
 - 8 wkrętów M3×6
 - 8 podkładek ząbkowanych o średnicy 3 mm
 - 4 szt. poliamidowych kołków dystansowych z gwintem M3 o długości od 10 mm do 25 mm
 - 6 szt. podwójnych niskosumowalnych wzmacniaczy operacyjnych NE5532 (IC1-IC6)
 - 1 szt. regulowany stabilizator LM317 1,5 A (REG1)
 - 2 szt. diod 1N4004 400 V 1 A (D1, D2)
 - 1 szt. dioda 1N4148 (D3)

Kondensatory:

- 1 szt. kondensator elektrolityczny 470 μF 50 V 105 °C
- 2 szt. kondensatorów elektrolitycznych 220 μF 25 V
- 8 szt. kondensatorów elektrolitycznych 47 μF 35 V 105 °C
- 1 szt. kondensator elektrolityczny 10 μF 35 V
- 8 szt. kondensatorów foliowych 150 nF 63 V MKT
- 6 szt. kondensatorów ceramicznych 100 nF X7R MLC
- 3 szt. kondensatorów ceramicznych 100 pF NP0/COG (lub styroflexowych KSF)

Rezystory: (wszystkie 1/4 W 1% metalizowane)

- 3 szt. 100 kΩ 2 szt. 33 kΩ 4 szt. 22 kΩ 8 szt. 12 kΩ 2 szt. 10 kΩ
- 4 szt. 7,5 kΩ 2 szt. 5,6 kΩ 4 szt. 4,7 kΩ 1 szt. 3,3 kΩ 1 szt. 1,8 kΩ
- 1 szt. 270 Ω 2 szt. 100 Ω 1 szt. 10 Ω

sygnały lewy i prawy do monitorów, zawierające niewiele informacji o częstotliwości poniżej 90 Hz, plus trzeci sygnał monofoniczny dla subwooferów, który zawiera sygnały poniżej 90 Hz z obu kanałów (dźwięki basu w nagraniach są często monofoniczne, ponieważ ich emisja w stereo praktycznie nic nie wnosi).

Wzmacniacz subwooferów jest identyczny jak wzmacniacz monitorów, z wyjątkiem dodania tej aktywnej zwrotnicy, która jest zaprojektowana wg wymagań. Nie możemy nie podkreślić, jak ważne jest to dla osiągnięcia dobrej wydajności w systemie aktywnym, a także dla ochrony monitorów przed niepożądanymi sygnałami niskotonowymi.

Aktywna zwrotnica posiada filtr Linkwita-Rileya czwartego rzędu, o tłumieniu 24 dB na oktawę. Punkt podziału zwrotnicy znajduje się przy 90 Hz.

Zwrotnica czwartego rzędu, dająca bardzo strome nachylenie filtra, została wybrana po to, aby nawet gdy subwoofer znajduje się bardzo blisko słuchacza, nie można było go zlokalizować. Dzięki temu wydaje

się, że sygnały basowe pochodzą z tego samego miejsca, co pozostałe sygnały, czyli z monitorów.

Drugą zaletą jest to, że przy zwrotnicy czwartego rzędu, do monitorów wysyłany jest sygnał pozbawiony najniższych tonów, a to zapobiega wspomnianym wcześniej nadmiernym wychyleniom membrany, które mogą drastycznie zwiększyć zniekształcenia (i to nie tylko w obszarze basu).

Poziom sygnału o częstotliwości 90 Hz za filtrem dolnoprzepustowym wynosi zaledwie 25% poziomu sygnału niefiltrowanego. Przy 45 Hz do monitorów wysyłane jest zaledwie około 1% mocy sygnału. Jakość odtwarzania przez monitory jest więc znacznie lepsza, ponieważ membrana jest efektywnie nieruchoma, a nie wychylona w takt sygnału niskotonowego. Cewka membrany znajduje się więc zawsze całkowicie w szczeliny magnesu.

Zwrotnica jest zrealizowana jako „filtr o zmiennym poziomie tłumienia”, składający się w zasadzie ze czterech obwodów całkujących



Pasywna zwrotnica (w naturalnej wielkości) zostanie opisana w części 2, w przyszłym miesiącu.

połączonych szeregowo. Schemat ideowy zwrotnicy pokazano na rysunku 11.

Sygnały wejściowe są doprowadzane przez koralik ferrytowy, zbocznikowany do masy kondensatorem 100 pF, aby odfiltrować wszelkie ewentualnie odbierane sygnały RF. Następnie sygnał jest sprzężony zmiennoprądowo z obwodami całkującymi filtra aktywnego. Przesunięcie fazowe każdego obwodu jest ustalane przez wartości RC; w naszym przypadku 12 kΩ i 150 nF.

Zwrotnica lewego kanału jest zbudowana ze wzmacniaczy operacyjnych IC1b, IC2a, IC2b, IC3a i IC3b (na górze schematu), podczas gdy prawy kanał składa się z IC4b, IC5a, IC5b, IC6a i IC6b. Poza tym są one identyczne.

Jednym z nietypowych aspektów tego filtra jest to, że wykorzystuje on zagnieżdżone sprzężenie zwrotne. Drugi i czwarty stopień całkujący mają rezystory sprzężenia zwrotnego połączone z nieodwracającym wejściem pierwszego stopnia, podczas gdy trzeci i piąty stopień mają rezystory sprzężenia zwrotnego połączone z odwracającym wejściem pierwszego stopnia.

Wyjścia dolnozaporowe są w każdym kanale pobierane z wyjścia pierwszego stopnia. Wyjścia dolnoprzepustowe pochodzą

Narzędzia, których będziesz potrzebować:

- Pistolet do klejenia
- Rozgątniacz elektryczny
- Kłosek do szlifowania
- Zestaw dużych zacisków stolarskich
- Pistolet do zszywek (nie jest niezbędny, ale ułatwia budowę)
- Ciężkie rękawice (chronią ręce przed odpryskami podczas szlifowania)
- i okulary ochronne

z piątego stopnia. Są one miksowane 1:1 za pomocą pary rezystorów 4,7 kΩ, a następnie podawane do bufora IC1a, który wysyła zmiksowany sygnał doysterowania wzmacniaczy subwooferów.

Zazwyczaj wzmacniacze operacyjne w takim układzie zasilane są z szyny dodatniej i ujemnej („zasilanie symetryczne”), a poziomy sygnałów są odniesione do masy. Jednak w tym przypadku chcemy zasilać wzmacniacze operacyjne niesymetrycznie z zasilacza impulsowego DC, najlepiej 24–32 V.

Wejście zasilania 24–32 V jest filtrowane dolnoprzepustowo przez rezystor szeregowy 10 Ω i kondensator 470 µF, a następnie podawane do REG1, regulowanego stabilizatora LM317, aby uzyskać ładne, czyste napięcie wyjściowe 18 V DC do zasilania wszystkich wzmacniaczy operacyjnych. Dwa rezystory 4,7 kΩ na tej linii 18 V generują szynę wirtualnej masy na potencjale 9 V, która jest buforowana przez wzmacniacz operacyjny IC4a i filtr dolnoprzepustowy RC. Wirtualna masa jest wykorzystywana do zmiany poziomu odniesienia wszystkich sygnałów, tak aby pozostawały one w obrębie napięcia szyny zasilania wzmacniaczy operacyjnych, czyli 0 V i 18 V.

Sygnały są następnie ponownie sprzężone zmiennoprądowo z wyjściami, co usuwa napięcie podkładu 9 V DC i sprowadza poziom odniesienia sygnałów do masy.

Wnioski

Jeśli jesteś zainteresowany budową tych głośników (czy to jako samodzielnych głośników półkowych czy z subwooferami), już teraz możesz zacząć zbierać potrzebne części, wg załączonej listy.

W przyszłym miesiącu opiszemy jak wykonać oba zestawy obudów, wraz z wymaganą elektroniką. ■

Phil Prosser

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

REKLAMA

KEY PRODUCENT AUTOMATYKI GRZEWCZEJ
11-200 Bartoszyce ul. Bohaterów Warszawy 67 pwkey@onet.pl
tel. (89)7635050 fax (89)7635051

TANIE REGULATORY

DO KOTŁÓW WĘGLOWYCH I NA DREWNO

z wbudowanym termostatem pokojowym
zapewniającym komfort i oszczędność



REGULATORY DO KOTŁÓW Z PODAJNIKIEM

REGULATORY POGODOWE

- Prosta obsługa, bogate możliwości programowania
- Możliwość dopasowania do każdego kotła i rodzaju paliwa
- Wysoka jakość
- Gwarancja 24 miesiące

www.pwkey.pl



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

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

Wysokościomierz samochodowy z ekranem dotykowym

Ta zmodyfikowana wersja wysokościomierza Jima Rowe'a z ekranem dotykowym jest zoptymalizowana do użytku w samochodzie, ciężarówce lub innym pojeździe lądowym. Oczywiście ze względu na brak atestów nie nadaje się do użycia w szybowcu lub ultralekkim samolocie. Sprzęt został uproszczony i przystosowany do zasilania z pojazdu, a oprogramowanie zostało zaktualizowane w celu poprawy dokładności podczas typowej jazdy samochodem.

Jest to zmodyfikowana wersja projektu Touchscreen Altimeter and Weather Station (Stacja pogodowa z ekranem dotykowym i wysokościomierzem) z grudnia 2017 roku (siliconchip.com.au/Article/10898), lepiej dopasowana do użytku w samochodzie. Składa się ona z modułu ekranu dotykowego Micromite LCD BackPack, najlepiej w wersji V2 (kompletny zestaw do montażu zakupisz za AUD \$70 w sklepie

SC – siliconchip.com.au/Shop/20/4237), natomiast opis modułu i jego montażu znajdziesz w SC z maja 2017 (kopię tego artykułu można znaleźć w Internecie); oraz z płytki interfejsu zasilania z czujnikami temperatury, ciśnienia i wilgotności, do instalacji w samochodzie, zbudowanej wg poniższego opisu.

Być może zastanawiasz się, dlaczego chcę mieć wysokościomierz w moim samochodzie. Uważam, że to interesujące wiedzieć,

jak wysoko jestem podczas jazdy w górach, zwłaszcza podczas zatrzymywania się w punktach widokowych (niektóre mają podaną wysokość npm., ale nie wszystkie).

Ponadto, wydajność silnika zmniejsza się wraz z wysokością, więc ta informacja może zaspokoić nie tylko Twoją ciekawość.

Moc wolnossących silników benzynowych spada o około 3–4% na każde 300 m (1000 stóp) wzrostu wysokości; silniki



Oto wysokościomierz wbudowany w standardowe (DIN) wycięcie w desce rozdzielczej mojego samochodu. Ponieważ ma duży ekran, jest bardzo czytelny. Jak widać na ekranie, można zmienić zarówno tryb wyświetlania, jak i jednostki (np. stopy nad poziomem morza, jak pokazane tutaj [wg konwencji używanej w lotnictwie] na metry nad poziomem morza, co wszyscy znamy). Nawiasem mówiąc, QNH oznacza ciśnienie atmosferyczne skorygowane do średniego poziomu morza. Nie jest ono ani stałe, ani takie samo dla różnych miejsc – QNH można uzyskać od serwisów pogodowych.

z turbodoładowaniem są mniej wrażliwe, ale nadal mogą stracić trochę mocy z powodu rzadszego powietrza na większych wysokościach, w zależności od ich konstrukcji.

W pojazdach mechanicznych wysokościomierz może być zasilany z gniazda akcesoriów, więc nie ma potrzeby stosowania wewnętrznego akumulatora używanego w oryginalnym projekcie. Oznacza to, że możemy zmieścić całe urządzenie w jednej obudowie UB3 Jiffy, z wentylatorem chłodzącym do odprowadzania ciepła generowanego przez wyświetlacz, unikając konieczności montowania czujników w osobnej obudowie.

W tej konstrukcji zasilanie jest dostarczane kablem USB. Wiele nowoczesnych samochodów posiada gniazda ładowania USB. Jeśli Twój nie ma, możesz użyć ładowarki USB podłączonej do gniazda akcesoriów.

Można kupić gotowe tanie wysokościomierze, ale nie są one zbyt dokładne. Dzieje się tak, ponieważ przeliczają one zazwyczaj odczyt ciśnienia powietrza na wysokość w odniesieniu do „średniego ciśnienia na poziomie morza” (MSL – Mean Sea Level), czyli ciśnienia 1013,25 hPa. Ale ciśnienie na poziomie morza może się wahać (przy ekstremalnej

pogodzie) od 870 hPa do 1084,8 hPa, co daje zakres błędu 1770 m/5800 stóp.

Oczywiście, rzadko spotykamy tak skrajne warunki, ale widać, że wyliczanie odczytu wysokości przy pomocy MSL często będzie prowadzić do znacznych błędów wysokości, ze skrajną sytuacją odczytu poniżej poziomu morza.

Aby rozwiązać ten problem, zmodyfikowałem oprogramowanie wysokościomierza tak, że możesz ustawić lokalną wysokość, taką jak wysokość **npm**. Twojego podjazdu do garażu, lub punktu widokowego (jest zwykle podana), w celu wprowadzenia bardzo dokładnie ciśnienia odniesienia, Twojego lokalnego QNH (**Question Nil High** – pytanie o ciśnienie na zerowej wysokości, czyli lokalne ciśnienie na poziomie morza zmierzone na lądzie, co oznacza że jest ono podawane w odniesieniu do ciśnienia atmosferycznego w punkcie pomiaru. W praktyce jest to aktualne ciśnienie lokalne sprowadzone do poziomu morza, na przykład w Warszawie w celu otrzymania QNH, do lokalnego ciśnienia chwilowego należy dodać 13 hPa. Wysokość wyznaczoną według ciśnienia QNH podaje się jako „altitude”).

Oryginalne oprogramowanie wysokościomierza zapisywało ustawienie QNH,

gdy go wyłączałeś i ładowało je ponownie przy starcie.

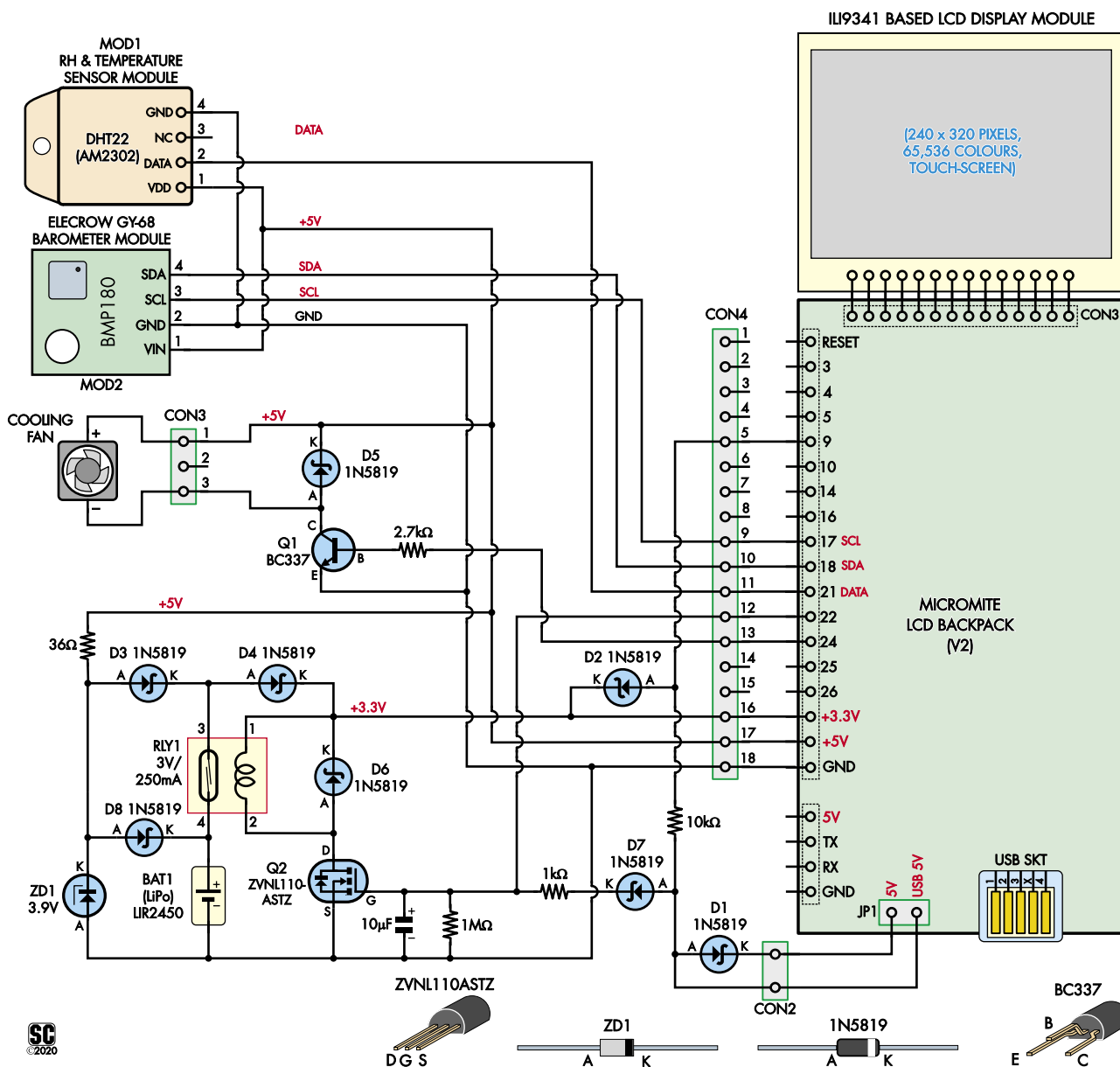
Jeśli pojedziesz gdzieś na piknik i wyłączysz wysokościomierz, przywróci on te same QNH i ustawienia, gdy włączysz go, aby odjechać.

Ale jeśli zostałeś na noc, gdy wyruszysz rano, QNH będzie prawdopodobnie znacząco różne, co prowadzi do błędów, które kumulują się z każdym postojem.

Aby rozwiązać ten problem, gdy wyłączasz zasilanie, oprogramowanie wysokościomierza pojazdu zapisuje wysokość nad poziomem morza, i używa tej wartości do obliczenia nowego QNH przy włączeniu zasilania.

Zakłada się, że pojazd nie zmienia wysokości, gdy go nie prowadzisz (miejmy nadzieję, że jest to bezpieczne założenie!). Tak więc urządzenie powinno pozostać dokładne przez całą podróż, pod warunkiem, że ustawisz na początku prawidłowo wysokość. Oszczędza to konieczności częstego sprawdzania aktualnego ciśnienia QNH w miejscu, w którym się znajdujesz (na przykład przez Internet) i aktualizowania urządzenia w celu utrzymania dokładności.

Wysokościomierz samochodowy jest tak skonstruowany, aby zmieścić się



Rysunek 1. Schemat ideowy wysokościomierza samochodowego jest wzorowany na układzie wysokościomierza dotykowego dla ultralekkich samolotów, ale został zoptymalizowany do użycia w pojazdach lądowych. Obejmuje to dodanie małego wentylatora sterowanego przez PWM, aby zapewnić czujnikom dostęp świeżego powietrza, oraz zapasowej baterii (BAT1) przełączanej przez RLY1, aby zapewnić zasilanie przez krótki czas po wyłączeniu, tak aby aktualna wysokość mogła być zapisana w pamięci flash.

w typowej kieszeni konsoli samochodowej (np. w Mazdzie 6). Kieszeń ta posiada gniazdo na akcesoria, które jest usytuowane, wraz z adapterem USB, po lewej stronie wysokościomierza.

Zmiany w schemacie

Zmodyfikowany schemat ideowy wysokościomierza pokazany jest na **rysunku 1**.

Oprócz zachowanego z poprzedniego projektu ekranu Micromite LCD Backpack, czujnika temperatury/wilgotności DHT22 i czujnika temperatury/ciśnienia BMP180, dodano następujące elementy: wentylator z regulacją obrotów PWM, mały akumulator Li-ion, przekaźnik sterowany MOSFET-em oraz kilka diod.

Obwód sterujący PWM wentylatora chłodzącego jest zaprojektowany tak, aby ograniczyć do minimum hałas, gdyż małe wentylatory chłodzące są notorycznie głośne. Zawiera on standardowy tranzystor NPN, Q1, sterowany z końcówki 24 Micromite przez rezystor 2,7 kΩ. Dioda Schottky'ego D5 zapobiega uszkodzeniu Q1 przez impulsy samoindukcji z wentylatora.

Oprogramowanie stosuje 20 Hz jako częstotliwość PWM, z wypełnieniem 50%. Daje to odpowiedni przepływ powietrza przy minimalnym hałasie.

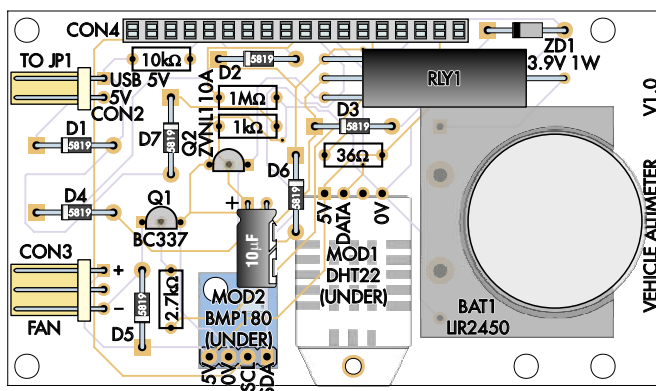
Aby urządzenie mogło zapisać wysokość przy wyłączaniu, musimy monitorować zasilanie 5 V i wykryć, kiedy zaczyna spadać. Ponieważ spada ono zbyt szybko, aby dać

oprogramowaniu wystarczająco dużo czasu na zapisanie ustawień, pastylkowy akumulator litowo-jonowy BAT1 zasilą obwód, gdy szyna zasilania 5 V wyłącza się.

Gdy skończymy zapisywać dane, zasilanie z akumulatora jest wyłączane programowo.

Jest jeszcze jedna korzyść z zastosowania takiego rozwiązania. Rozrusznik może poważnie wpływać na układ elektryczny pojazdu, a zasilanie 5 V może się wahać na tyle, by rozregulować wysokościomierz. Poprzez odizolowanie diody szyny 5 V od wejścia USB i użycie akumulatora litowo-jonowego do zapewnienia stabilnego zasilania 3,3 V, otrzymujemy niezawodny start układu.

Styki JP1 służą jako złącze dostępu do zasilania 5 V z gniazda USB oraz do doprowadzenia



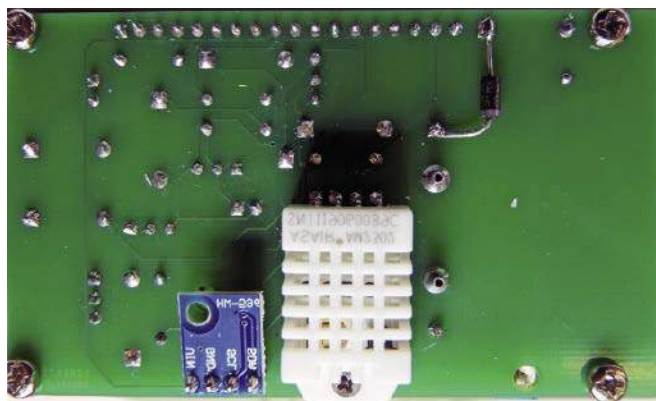
Rysunek 2. Aby ułatwić montaż, wszystkie elementy, które nie są częścią ekranu dotykowego Micromite LCD Backpack, są zamontowane na tej zbliżonej wielkością płytce drukowanej, ze zdjęciami zmontowanej płytki z przodu i z tyłu przedstawionymi po prawej. Tylko dwa czujniki i jedna dioda są zamontowane na spodzie PCB – wszystkie inne elementy zamontowano na wierzchu płytki, łącznie ze zwiernym cylindrycznym przekaźnikiem kontaktowym SPST (czarny element na górze po prawej stronie na górnym zdjęciu) i uchwytem na akumulator. Zauważ, że D8 jest zamontowana na spodzie PCB i jest przylutowana z anodą połączoną z katodą ZD1, a katodą z dodatnim zaciskiem BAT1.

napięcia 5 V z powrotem do ekranu LCD Micromite. Przepływ prądu możliwy jest tylko od zasilania USB 5 V do szyny zasilania 5 V przez diodę Schottky'ego D1. Napięcie USB +5 V trafia również na bramkę MOSFET-a Q1 poprzez kolejną diodę Schottky'ego (D7) i rezystor 1 kΩ. Dzięki temu Q2 włącza się, gdy tylko dostępne jest zasilanie USB, i zasila cewkę przekaźnika kontaktowego RLY1.

Kiedy linia 3,3 V jest zasilana z baterii BAT1 (tzn. USB 5 V jest wyłączone), linia 5 V jest na poziomie 3,3 V; szyna jest wtedy zasilana obydwiema przez regulator 3,3 V na płytce ekranu, z jego wyjścia do wejścia przez wewnętrzną diodę ochronną. Dioda D1 zapobiega wtedy zwarceniu napięcia 3,3 V do źródła 5 V USB.

Przekaźnik RLY1 włącza BAT1 do obwodu, gdy MOSFET Q2 jest włączony. BAT1 jest ładowany z szyny 5 V poprzez rezystor 36 Ω ograniczający prąd i diodę Schottky'ego D8. Dioda Zenera ZD1 ogranicza napięcie podawane na akumulator do bezpiecznego poziomu przy ładowaniu (około 3,6 V, biorąc pod uwagę napięcie przewodzenia diody D8).

BAT1 z kolei zasila szynę +3,3 V ekranu Backpack poprzez diodę Schottky'ego D4. Spadek napięcia na diodzie D4 redukuje napięcie 3,6 V z akumulatora do potrzebnego poziomu 3,3 V. Ta linia zasilania głównie



mikroprocesor PIC32 w Backpack-u, który ma zalecane maksymalne napięcie zasilania 3,6 V i absolutne maksimum 4,0 V.

Końcówka 9 Micromite przez rezystor 10 kΩ jest używana do wykrywania napięcia 5 V USB, aby określić, kiedy wyłączy się zewnętrzne zasilanie 5 V. Styk 22 Micromite jest wtedy ustawiany na niskim poziomie, aby wymusić obniżenie potencjału bramki Q2, wyłączenie RLY1 i wyłączenie układu.

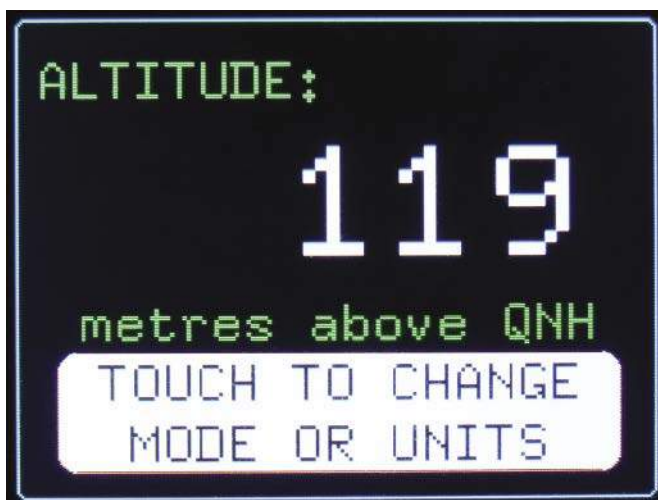
Używanie Weatherzone do uzyskania QNH

Weatherzone (weatherzone.com.au) jest darmową aplikacją mobilną do przeglądania prognoz pogody i związanych z nimi informacji. Zapewnia ona również prostą metodę uzyskania QNH.

W tym przykładzie, zrzut ekranu po lewej stronie pokazuje obserwacje na Terrey Hills; nie ma podanej wartości QNH, więc pole Pressure jest puste. Stuknięcie w ekran powoduje przejście do najbliższej lokalizacji z danymi, którą jest Sydney. Drugi zrzut ekranu pokazuje, że wskazuje on aktualną wartość QNH.

Jeśli chcesz uzyskać większą dokładność, użyj ekranu Obserwacji Pogody dla Twojego obszaru z Bureau of Meteorology. (www.bom.gov.au). BOM podaje QNH z rozdzielczością 0,1 hPa.





Ekran 1. Główny ekran po ustawieniu QNH. Pokazuje on wysokość nad QNH (efektywny poziom morza) nad poziomem morza wg QNH w metrach lub stopach.

Jedną z rzeczy, która nie jest pokazana na schemacie, jest to, że dodałem na przednim panelu do Backpack-a przełącznik ściemniania podświetlenia LCD. Łączy się on z potencjometrem nastawnym regulacji jasności Backpack-a (VR1), zwierając go, gdy przełącznik jest zamknięty i w ten sposób wybierając pomiędzy dwoma różnymi poziomami jasności: tym ustawionym przez VR1 i pełną jasnością.

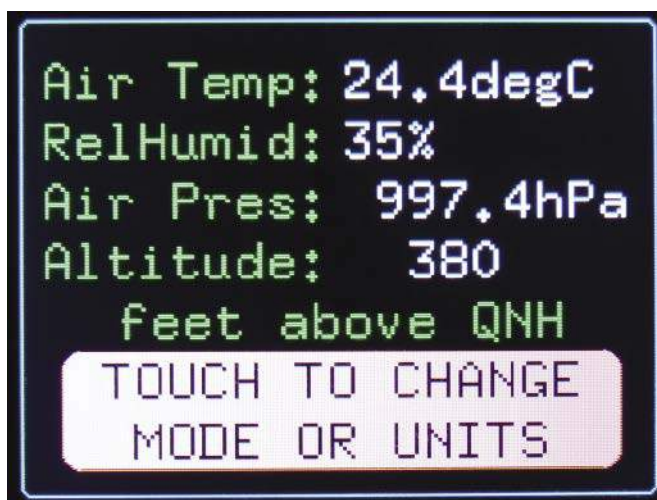
Jest to ważne, aby można było przełączyć podświetlenie na niską jasność nocą, aby nie pogorszyć sobie widzenia w ciemności.

Zmiany w oprogramowaniu

Oprogramowanie zostało zmienione w kilku miejscach, a niektóre z tych zmian zostały opisane powyżej. Wprowadzono również pewne usprawnienia w interfejsie użytkownika.

Ekran stacji pogodowej i wysokościomierza są podobne do oryginału. Pokazują one, do momentu wprowadzenia QNH lub dokładnej wysokości, wysokość wg MSL. Następnie pokazują wysokość wg QNH (Ekran 1 i Ekran 2).

Ekran zmiany trybu ma nowe opcje wyboru, które rozróżniają pomiędzy wprowadzeniem



Ekran 2. Rozszerzony ekran informacyjny po ustawieniu QNH, pokazujący wysokość w stopach wraz z temperaturą powietrza, wilgotnością względną i odczytami ciśnienia atmosferycznego.

QNH i aktualnej wysokości (Ekran 3). Bieżąca wartość QNH jest również pokazywana podczas wprowadzania aktualnej wysokości lub nowej wartości QNH (Ekran 4).

Jeżeli chcesz zmienić częstotliwość PWM wentylatora lub wypełnienie cyklu pracy, wyszukaj w kodzie źródłowym BASIC linię zaczynającą się od PWM i zmień wartości 20 (Hz) lub 50 (wypełnienie cyklu pracy) na odpowiednie dla Ciebie.

Zasilanie

Wysokościomierz pobiera około 90 mA przy napięciu 5 V. Może być zasilany z takiego banku energii USB (takiego jak Jaycar Cat MB3792), zapewniając czas pracy pomiędzy ładowaniami przekraczający 24 godziny, co czyni tę wersję użyteczną także poza jazdą mechaniczną.

Utrata zasilania USB 5 V jest wykrywana przez końcówkę 9 Micromite, z rezystorem 10 kΩ i diodą D2 zwierającą ten sygnał do szyny 3,3 V, ponieważ styk 9 Micromite nie akceptuje napięcia 5 V. Zasilanie Micromite jest podtrzymywane przez krótki czas po zaniku zasilania USB dzięki kondensatorowi 10 μF pomiędzy bramką i źródłem Q2, który powoli rozładowuje się przez równoległy rezystor 1 MΩ. W tym czasie Micromite pracuje zasilany z BAT1.

Zmiana poziomu na końcówce 9 Micromite wyzwała przerwanie programu, które powoduje, że Micromite zapisuje aktualne dane o wysokości. Następnie wyjście 22 Micromite jest przełączane w stan niski, wyłączając Q2 i zwalniając przekaźnik, a więc wyłączając wszystko. Dioda D6 tłumi wszelkie skoki napięcia na cewce przekaźnika.

W praktyce, Micromite działa przez około 200 ms po utracie zasilania 5 V. To daje Backpack-owi czas na wysłanie wiadomości

„Saved” do terminala podłączonego kablem USB, zanim zniknie zasilanie 3,3 V. Możesz zauważyć krótkie ściemnienie wyświetlacza, ponieważ podświetlenie ekranu, zanim się wyłączy, działa wtedy zasilane napięciem 3,3 V zamiast napięciem 5 V.

Zauważ, że wybór MOSFET-a Q2 nie jest krytyczny. Każdy N-kanalowy MOSFET o ciągłym prądzie drenu co najmniej 300 mA i maksymalnym napięciu wyzwalania bramki do 2,0 V (typowo zaprojektowany do zasilania z 3,3 V) powinien pracować równie dobrze jak ZVN110A.

Nie testowaliśmy jednak żadnych zamienników.

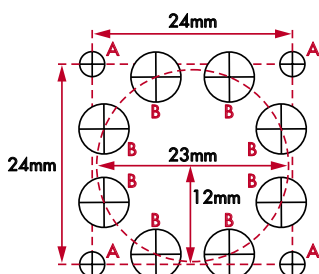
Budowa

Zaprojektowałem dwustronną płytkę drukowaną interfejsu, która mieści wszystkie elementy wysokościomierza samochodowego, jak pokazano na rysunku 2 i załączonych zdjęciach.

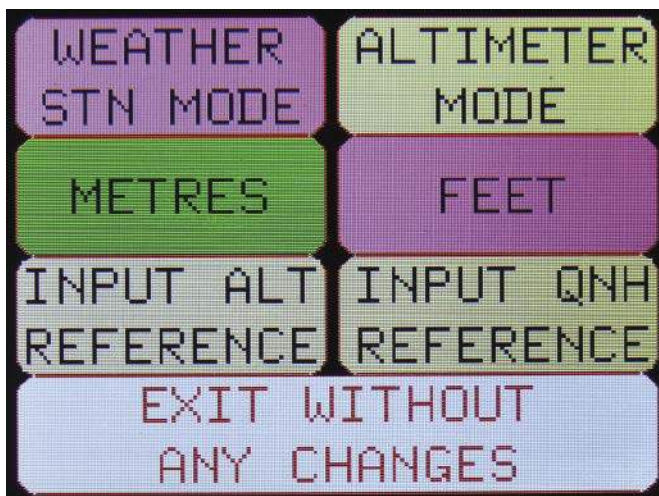
Dwa czujniki (BMP180 i DHT22) są zamontowane na spodzie. Dzięki temu czujniki znajdują się z dala od elementów wytwarzających ciepło, w dedykowanym strumieniu chłodnego powietrza pomiędzy jego wlotem i wylotem w obudowie. Płytką interfejsu jest podłączona bezpośrednio do LCD Backpack-a.

Aby uruchomić wysokościomierz, na początku możesz: wyłagać, pożyczyc lub zbudować Backpack-a. My jednak sugerujemy od razu zbudowanie wersji V2, chociaż oryginał też będzie działał. Nie zalecamy używania wersji V3, ponieważ oprogramowanie wysokościomierza nie jest przeznaczone do współpracy z większym ekranem, a wewnętrzna głębokość obudowy wersji V3 jest zmniejszona z powodu jej zagłębionego przedniego panelu.

Konstrukcja Backpack-a V2 jest w pełni opisana w Silicon Chip, z maja 2017,



Rysunek 3. Użyj tego schematu jako wzorca lub szablonu do wywiercenia ośmiu otworów wentylacyjnych na obu końcach obudowy oraz czterech otworów montażowych dla wentylatora na prawym krańcu. Otwory A: średnica 3 mm. Otwory B: średnica 6 mm, Uwaga: otwory A należy wiercić tylko po prawej stronie obudowy



Ekran 3. Ekran ustawień ma dwa przyciski na dole do kalibracji; jeden do wprowadzania aktualnie znanej wartości QNH, a drugi do wprowadzania aktualnej wysokości w stopach/metrach.



Ekran 4. Aktualna wartość QNH jest wyświetlana podczas wpisywania nowej wartości, aby przypomnieć, którą wartość aktualizujesz.

począwszy od strony 84 (siliconchip.com.au/Article/10652).

Ale biorąc pod uwagę jej względną prostotę i fakt, że dostępny jest zestaw montażowy, a w nadrukowanym na płytce opisie oznaczono umiejscowienie komponentów, tak naprawdę nie trzeba czytać tego artykułu. Po prostu zamontuj elementy tam, gdzie pokazano na PCB, i wszystko powinno działać.

Kiedy już zbudujesz i przetestujesz Backpack-a, podłącz przełącznik do potencjometru nastawnego VR1 tak, że kiedy przełącznik jest zwarty, VR1 też jest zwarty i ekran LCD pracuje z maksymalną jasnością. Kiedy jest wyłączony, jasność jest ustawiona przez VR1. Powinienesz ustawić ją na komfortowym poziomie do oglądania w nocy.

Zauważ, że istnieją dwie identyczne wersje 2,8-calowego ekranu dotykowego LCD 320×240, z których jedna wykorzystuje do regulacji podświetlenia zmianę prądu, a druga sterowanie napięciem.

Jeśli potencjometr nastawny 100 Ω dołączony do VR1 nie reguluje prawidłowo jasności podświetlenia, zastąp go potencjometrem 100 kΩ i podłącz jego wolny styk do masy. To powinno rozwiązać problem.

Teraz zmontuj opisaną tutaj płytkę interfejsu, lutując rezystory i diody na przedniej stronie.

Następnie dodaj pojemnik akumulatorka, złącza CON2-CON4 oraz przekaźnik RLY1. RLY1 jest w trochę dziwnej cylindrycznej obudowie, z trzema przewodami na jednym końcu

i jednym na drugim. Upewnij się, że jego oznaczenie typu jest skierowane do góry i przylutuj go jak pokazano na Rysunku 2 i zdjęciach.

Od spodu PCB przylutuj oba czujniki. Ostrożnie zagnij końcówki czujnika DHT22 względem jego korpusu, tak aby przeszły przez otwory montażowe w PCB.

Zamocuj czujnik za pomocą śrubki M2 i przylutuj końcówki, następnie przygotuj BMP180 do montażu, przylutowując do jego otworów dołączoną 4-stykową listwę kołkową. Przymocuj zespół do płytki drukowanej i przylutuj kołki do odpowiednich pól płytki. Sprawdź czy styk „SDA” przylutowany jest do kwadratowego pola.

Nie zapomnij zamontować diody D8, która też jest przylutowana na spodzie płytki, jak pokazano na zdjęciu na rysunku 2.

Pojedynczy kondensator jest typu elektrolitycznego. Montowany jest wygięty na boku. Upewnij się, że dłuższa (dodatnia) końcówka trafiła do kwadratowego pola, oznaczonego „+”. Przymocuj korpus kondensatora do płytki za pomocą odrobiny kleju silikonowego lub kawałka dwustronnej piankowej taśmy montażowej. Włutuj MOSFET Q2 i tranzystor BC337 Q1 tam, gdzie zaznaczono, i płytka jest gotowa.

Przygotowanie obudowy

Następnie przygotuj obudowę UB3 Jiffy. Wentylator chłodzący montuje się po prawej stronie, patrząc na obudowę od przodu (od strony pokrywki), tak daleko do tyłu jak to tylko możliwe.

Wywierć cztery otwory montażowe o średnicy 3 mm, każdy w rogu kwadratu 24×24 mm (lub po prostu zaznacz pozycje używając wentylatora, a następnie wywierć). Następnie wywierć kilka otworów wewnątrz kwadratu utworzonego przez mniejsze otwory w obudowie, aby umożliwić przepływ

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

- 1 zmontowany ekran dotykowy Micromite LCD Backpack (V1 lub V2) [SILICON CHIP nr kat. SC4024 lub SC4237].
- 1 moduł czujnika temperatury/wilgotności DHT22 (MOD1)
- 1 moduł czujnika temperatury/ciśnienia GY-68 BMP-180 (MOD2)
- 1 dwustronna płytka drukowana, nr katalog. 05105201, 86,5×49,5 mm
- 1 czarna lub szara obudowa UB3 Jiffy [Jaycar HB6013/HB6023].
- 1 miniaturowy przełącznik wychyłny 2-pozycyjny SPST/SPDT do montażu w obudowie [np. Jaycar ST0335]
- 1 cienki wentylator chłodzący 30 mm 12 V DC [Jaycar YX2501]
- 1 przekaźnik kontaktronowy cewka 3 V DC, 250 mA SPST [RS Cat 124-5129] (RLY1)
- 1 pojemnik akumulatora pastylkowego 2450 do montażu na PCB [element14 Cat 1216361] (BAT1)
- 1 akumulator LiR2450 Li-ion [element14 Cat 2009025] (BAT1)
- 1 komplet: wtyk kątowy 403-2 do druku męski + gniazdo 402-2 na przewód (ze stykami) (CON2)
- 1 komplet: wtyk kątowy 403-3 do druku męski + gniazdo 402-3 na przewód (ze stykami) (CON3)
- 1 szt. gniazdo jednorzędowe 20 styków, raster 2,54, proste do druku (CON4)
- 1 kabel USB 50 cm USB 2.0 A męski – 5-stykowy wtyk USB Mini-B [np. Jaycar WC7709]
- 1 zaciskowy przepust kablowy 6,2-7,4mm [Jaycar HP0718]
- 4 poliamidowe kołki dystansowe z gwintem M3 o długości 12 mm
- 4 śrubki M3×15/20

Półprzewodniki:

- 1 szt. tranzystor NPN BC337, TO-92 (Q1)
- 1 szt. MOSFET N-kanalowy lub podobny ZVN110ASTZ, TO-92 (Q2) [RS Cat 823-1833]
- 1 szt. dioda Zenera 3,9 V 1 W (ZD1) [np. 1N4730]
- 8 szt. diod Schottky'ego 1N5819 1A (D1-D8)

Kondensatory:

- 1 szt. 10 µF 16 V elektrolityczny

Rezystory: (wszystkie 1/4 W 1% metalizowane)

- 1 szt. 1 MΩ 1 szt. 10 kΩ 1 szt. 2,7 kΩ 1 szt. 1 kΩ 1 szt. 36 Ω

powietrza. Proponuję osiem otworów o średnicy 6 mm rozmieszczonych równo na obwodzie koła o średnicy 23 mm. Można użyć rysunku 3 jako szablonu do zaznaczenia tych otworów przed wierceniem.

Wywierć te same osiem otworów wlotowych na lewym końcu obudowy, naprzeciwko wentylatora, ale bez otworów montażowych.

Następnie zlokalizuj dogodne miejsce z tyłu obudowy, przez które będzie wychodził kabel USB. Wywierć otwór o średnicy 11,5 mm, w którym umieścisz zacisk do mocowania przewodu. My zlokalizowaliśmy go 20 mm od końca wentylatora, 10 mm od góry. Daje to wystarczającą długość kabla, aby wyciągnąć elektronikę z obudowy.

Wywierć otwór w przednim panelu, aby zamontować przełącznik ściemniacza, upewniając się, że przełącznik nie będzie przeszkadzał w podłączeniu wentylatora i Backpack-a.

Przytnij przewody wentylatora chłodzącego do około 150 mm i podłącz 3-stykowe gniazdo żeńskie typu 402-3, aby dopasować je do CON3 na płycie interfejsu. Następnie należy wykonać 2-przewodowy kabel łączący CON2 na płycie interfejsu z JP1 na Backpack-u.

Aby połączyć się z JP1, wytnij dwustykowy odcinek z resztek jednorzędowej listwy żeńskiej użytej do wykonania CON4, zagnij piny do korpusu, przylutuj do pinów przewody i obkurcz je odcinkiem rurki termokurczliwej do korpusu. Dzięki temu złącze jest na tyle krótkie, że zmieści się między złączem JP1 Backpack-a, a wyświetlaczem.

Dokładnie sprawdź połączenia. Jeśli zamienisz przewody, dioda D1 na płycie interfejsu odizoluje wszystko od wejścia USB 5 V.

Kabel USB ściśle przylega do końca pudełka. Ostrożnie usunęliśmy część plastikowego wzmocnienia przy mini złączu i delikatnie ogrzaliśmy kabel, aby ułożyć go w preferowanym przez nas kierunku.

Miniwtyczka USB może być włożona przez otwór wyjściowy w tylnej części pudełka, a kabel zabezpieczony zaciskiem do mocowania przewodu. Włóż akumulator LIR2450 do jego oprawki, zamontuj płytkę interfejsu na Backpack-u za pomocą 12 mm kołków dystansowych i śrubek M3×20. Budowa jest już zakończona.

Testowanie

Ładuj do Micromite poprawione oprogramowanie wysokościomierza o nazwie „Altimeter with power fail.bas” (dostępne do pobrania ze strony SILICON CHIP), oczywiście po uprzedniej kompilacji, i uruchom je. Jeśli nie wprowadzasz żadnych zmian do kodu źródłowego, możesz od razu załadować plik wykonywalny, dostępny tu: siliconchip.com.au/Shop/Download/5460/11782.

Przy pierwszym uruchomieniu, wyświetlacz powinien zainicjować się ekranem stacji pogodowej, używając MSL jako ciśnienie odniesienia.

Podłącz wysokościomierz do terminala takiego jak Teraterm lub MMEdit. Dioda LED na Backpack-u powinna migać dwa razy na sekundę, gdy Micromite wysyła wiadomość „pass” do terminala. Jeśli wysokościomierz nie uruchomi się, sprawdź połączenie z CON2 do JP1. Wentylator chłodzący powinien pracować, jeśli oprogramowanie zostało zainicjalizowane.

Sprawdź, czy akumulator jest ładowany. Napięcie na nim powinno dochodzić do 3,6 V. Spadek napięcia na rezystorze 36 Ω powinien wynosić około 0,9-1,1 V, gdy bateria jest naładowana. Możesz to sprawdzić na spodniej stronie płytki.

Sprawdź poprawność działania przycisków na ekranie dotykowym. Aby znaleźć QNH do wprowadzenia, najlepszą metodą jest użycie aplikacji takiej jak Weatherzone (patrz panel). Na ekranie aktualnej prognozy Weatherzone dla twojej lokalizacji znajduje się pole oznaczone jako „Pressure”. Jeśli wartość jest pusta, dotknij ekranu, aby przejść do najbliższej obserwacji QNH.

Kiedy dokonujesz zmiany, takiej jak wprowadzenie QNH lub wysokości odniesienia Alt (aktualna znana wysokość), możesz zauważyć, że odczyt wysokości ustala się do ostatecznej wartości w ciągu pięciu sekund.

Dzieje się tak, ponieważ ta wersja oprogramowania uśrednia odczyty, aby wyeliminować krótkotrwałe wahania i poprawić dokładność zapisanej wysokości przy wyłączeniu zasilania.

Przy podłączonym terminalu i monitorowaniu sygnału USB, terminal powinien

pokazywać „pass” raz na sekundę. Odłącz kabel od CON2. Terminal powinien wyświetlić komunikat „Saved”, wskazując, że aktualna wysokość została zapisana.

Zamontuj przedni panel do obudowy. Być może będziesz musiał zaopatrzyć się w dłuższe wkręty samogwintujące niż te dostarczone w zestawie, lub możesz nagwintować otwory montażowe i użyć wkrętów metrycznych. Wysokościomierz powinien być teraz gotowy do użycia.

Precyzja, dokładność i błędy

Pamiętaj, że wysokościomierz ciśnieniowy nie jest instrumentem o dokładności pomiarowej. Nawet jeśli może on wyświetlać wysokość z rozdzielczością do jednej stopy, prawdopodobnie będzie wyświetlał równie precyzyjnie błędną wysokość, ponieważ ma na to wpływ kilka czynników.

Jednym z nich jest dryft QNH. Biuro Meteorologii nieustannie zmienia QNH, a piloci muszą nieustannie korygować swoje wysokościomierze. Ponadto, QNH uzyskane z Weatherzone jest podawane z rozdzielczością 1 hPa. Na „dzień dobry” zostajesz z błędem ± 13 stóp/4 m.

Kolejny błąd wynika z różnicy temperatur. Jeśli zaparkujesz w słońcu i wyłączysz silnik, aktualna wysokość zostanie zapisana. Jednak po powrocie i ponownym uruchomieniu silnika, temperatura wewnątrz samochodu może być o 20°C wyższa niż temperatura otoczenia. Wysokościomierz użyje tej temperatury do obliczenia nowej wartości QNH. Ten błąd może wynosić do 6 m/20 stóp dla różnicy temperatur 20°C.

Te błędy są nie do zaakceptowania przy lądowaniach według wskazań przyrządów przy braku widoczności, ale nie stanowią większego problemu w przypadku podróży drogowych. Nie stresuj się. Wprowadź ponownie QNH i ciesz się tym wysokościomierzem! ■

Peter Bennett

REKLAMA

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

Już ponad rok publikujemy dla projektantów i programistów dla elektroniki. Odwiedź

ELPORTAL.pl

Analogowy automat perkusyjny

Prosty automat, opisany w artykule, może być zbudowany nawet przez niezbyt doświadczonego elektronika. Urządzenie odtwarza kultowe brzmienia starych maszyn perkusyjnych. Projekt z pewnością zainteresuje wielu Czytelników, choćby z racji panującej mody na muzyczne „analogi”.

Do czego to służy?

Automaty perkusyjne są to urządzenia imitujące dźwięki instrumentów perkusyjnych i odtwarzające rytmy zapisane w pamięci. Zaczęły się rozpowszechniać w latach 60. ubiegłego wieku i... szybko wywołały protesty perkusistów obawiających się bezrobocia. Czas jednak pokazał, że w większości stylów muzycznych żywego perkusistę maszyną zastąpić się nie da. Za to automaty walczyły przyczyniły się do powstania gatunków „pop elektroniczny” i „rock elektroniczny”, których nieodłączną cechą jest właśnie specyficzne brzmienie automatycznej perkusji.

W artykule opisano prosty automat perkusyjny z analogowym torem dźwiękowym. W krajowej prasie elektronicznej ostatnich 30 lat brakuje opisów tego typu urządzeń, choć jestem pewien, że spotkałyby się one z dużym zainteresowaniem.

Historia

Można wyróżnić trzy generacje automatów perkusyjnych:

- Pierwsza generacja (lata 60. i 70.) to automaty o stałym zestawie rytmów. Odtwarzały one modne w tamtym czasie rytmy latynoskie jak „bossa nova”, „rumba”, „cza-cza”. Brzmienia instrumentów (bongosy, klawesy, marakasy) powstawały w prostych układach tranzystorowych. Przykład: *Rhythm Ace FR6*, **fotografia 1a**.
- W drugiej generacji (wczesne lata 80.) była już możliwość wpisywania własnych rytmów do pamięci RAM. Zestaw brzmień uległ zmianie na rzecz dźwięków typowego rockowego zestawu

perkusyjnego (bęben taktowy, werbel, półkotły, talerze). Do ich symulacji stosowano układy analogowe ze wzmacniaczami operacyjnymi, niekiedy z regulacją parametrów jak wysokość tonu czy czas wybrzmiewania. Przykład: *Roland TR-808*, **fotografia 1b**.

- Trzecia generacja (od lat 80.) to urządzenia całkowicie cyfrowe. Układ syntezy odtwarza dźwięki instrumentów akustycznych zapisane w pamięci stałej. Rytmy są oczywiście programowane przez użytkownika. Jest wyświetlacz LCD, interfejs MIDI, czasem gniazdo dla dodatkowej karty pamięci. Przykład: *Boss DR-880*, **fotografia 1c**.

Dostępność pamięci o dużych pojemnościach sprawiła, że współczesny cyfrowy automat może odtwarzać z najwyższą jakością dźwięki setek instrumentów. Ten, zdawałoby się, bezkonkurencyjny system powinien „rozkładać na łopatki” i usuwać w zapomnienie dawne generacje maszyn perkusyjnych. Tymczasem w ostatnich trzydziestu latach bardzo modne – wręcz kultowe – stały się brzmienia starych analogowych automatów z lat 80., zwłaszcza TR-808 i TR-909. Urządzenia te osiągnęły niebotyczne ceny na giełdach, a ich dźwięki słychać w ogromnej części nagrań dzisiejszej muzyki pop.

I takie właśnie dźwięki wytwarza – analogowo, a jakże – automat opisany w artykule.



a)

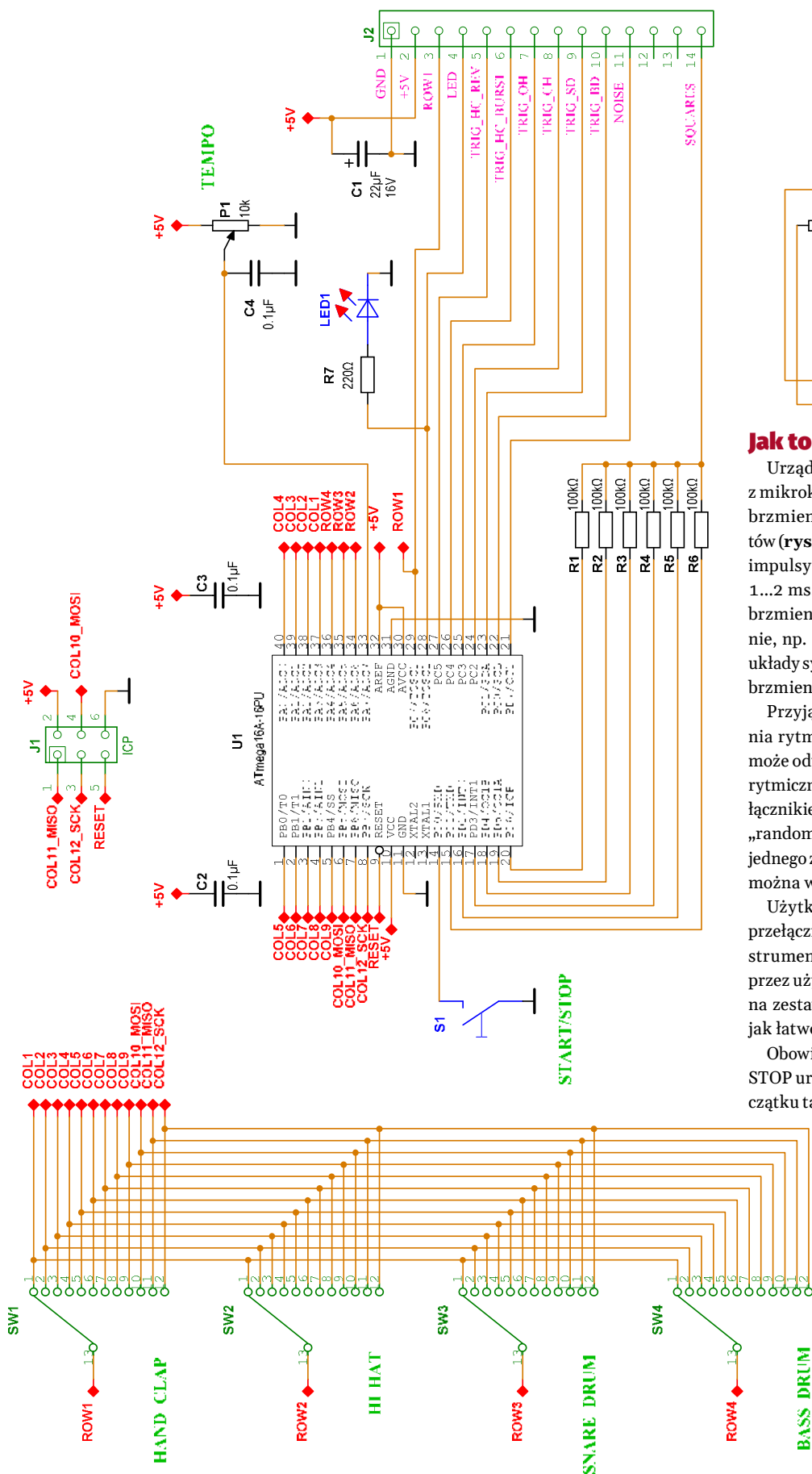


b)

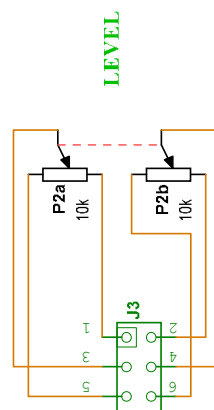


c)

Fotografia 1. Trzy generacje automatów perkusyjnych



Rysunek 1. Schemat ideowy płytki sterującej



Jak to działa?

Urządzenie składa się z płytki sterującej z mikrokontrolerem (**rysunek 1**) oraz płytki brzmieniowej z symulatorami instrumentów (**rysunek 2**). Z płytki sterującej wychodzą impulsy o amplitudzie 5 V i czasie trwania 1...2 ms, wyzwalające instrumenty na płycie brzmieniowej. Płytek można użyć niezależnie, np. dołączyć do płytki sterującej własne układy syntezy dźwięków albo sterować płytkę brzmieniową z bębnow-czujników.

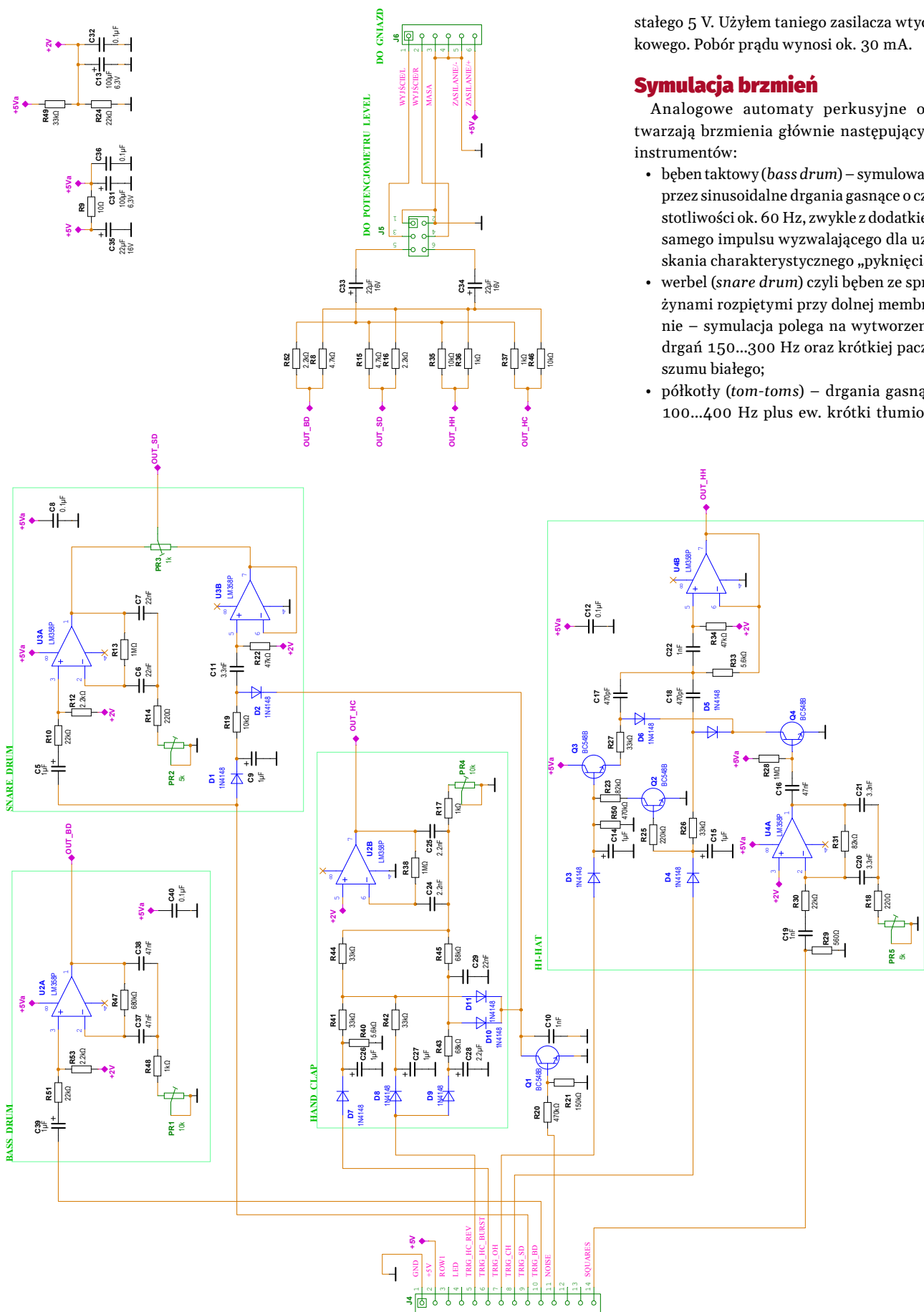
Przyjąłem nietypową koncepcję wytwarzania rytmów. Są 4 instrumenty. Każdy z nich może odtwarzać jeden z 10 wzorów (układów) rytmicznych, wybieranych odpowiednim przełącznikiem obrotowym. Można też wybrać tryb „random”, polegający na losowym wybieraniu jednego ze wzorów co 8 taktów. Przełącznikiem można wreszcie instrument wyłączyć.

Użytkownik, budując cały rytm, wybiera przełącznikami wzory rytmiczne każdego z instrumentów. Same wzory są niezmiennialne przez użytkownika, natomiast zabawa polega na zestawianiu ich kombinacji, których jest, jak łatwo policzyć, ponad 14 tysięcy.

Obowiązuje metrum 16/16. Przycisk START/STOP uruchamia odtwarzanie (zawsze od początku taktu) lub je zatrzymuje. W trakcie odtwarzania dioda LED sygnalizuje

odstęp czasu równe ćwierćnotom. Potencjometr TEMPO nastawia szybkość odtwarzania w zakresie 45...300 ćwierćnut na minutę, potencjometr LEVEL reguluje poziom sygnału na wyjściu. Wyjście jest dwukanałowe. Można je wykorzystać jako stereofoniczne, można też zewrzeć ze sobą oba kanały lub użyć tylko jednego z nich, uzyskując w ten sposób różne proporcje instrumentów.

Układ jest zasilany ze stabilizowanego napięcia



Rysunek 2. Schemat ideowy płytki brzmieniowej

stałego 5 V. Użyłem taniego zasilacza wtyczkowego. Pobór prądu wynosi ok. 30 mA.

Simulacja brzmień

Analogowe automaty perkusyjne odtworzą brzmienia głównie następujących instrumentów:

- bęben taktowy (*bass drum*) – symulowany przez sinusoidalne drgania gasnące o częstotliwości ok. 60 Hz, zwykle z dodatkiem samego impulsu wyzwalającego dla uzyskania charakterystycznego „pyknięcia”;
- werbel (*snare drum*) czyli bęben ze sprężynami rozpiętymi przy dolnej membranie – symulacja polega na wytworzeniu drgań 150...300 Hz oraz krótkiej paczki szumu białego;
- półkotły (*tom-toms*) – drgania gasnące 100...400 Hz plus ew. krótki tłumiony

szum imitujący odgłos uderzenia pałeczki o membranę;

- talerze (*cymbals*), przede wszystkim *hi-hat*, czyli zestaw dwóch talerzy złączonych ze sobą, o krótkim, ostrym brzmieniu (*closed hi-hat*); oba talerze można rozdzielić, uzyskując dźwięk dłuższy (*open hi-hat*) – dźwięki talerzy są imitowane przez gasnący „szum metaliczny”;
- różne tzw. „przeszkadzajki”: klawesy (*claves*) czy krowi dzwonek (*cowbell*) – szybko zanikające sygnały sinusoidalne, prostokątne, szumowe i ich mieszanki.

Większość instrumentów perkusyjnych zawiera membranę lub inny element, pobudzany uderzeniem do drgań gasnących, zwykle zbliżonych do sinusoidalnych. W automatach 1. generacji do symulacji tych instrumentów służyły obwody rezonansowe LC, wzbudzone krótkimi impulsami. W 2. generacji wyeliminowano kłopotliwe cewki, wprowadzając filtry aktywne RC. Stosowano też generatory o pracy ciągłej wraz ze sterowanymi wzmacniaczami zapewniającymi zanikanie amplitudy.

Innych sposobów używano do uzyskiwania dźwięków talerzy. Dźwięki te mają charakter szumowy, więc pierwotnie stosowano biały szum przepuszczony przez filtr górno- lub pasmowoprzepustowy. Tak uzyskany dźwięk miał jednak mało wspólnego z brzmieniem talerza i przypominał raczej syk powietrza wypuszczanego z dętki. W 1981 roku firma Roland zaczęła używać innego materiału dźwiękowego: mieszanki kilku przebiegów prostokątnych o różnych, nie powiązanych ze sobą, częstotliwościach z zakresu 100...1000 Hz. Dźwięk

taki, nieobrobiony, przypomina odgłos... uderzenia w pokrywkę garnka. Wystarczy jednak mocno odciąć filtrami wszystkie częstotliwości poniżej kilku kHz, pozostawiając tylko wysokie składowe harmoniczne prostokątów, aby uzyskać doskonały efekt jakby szumu, ale o silnym „blaszanym” zabarwieniu, jakie ma prawdziwy talerz. Poprawa brzmienia była rewelacyjna. Inni producenci, w tym nasza „Unitra”, otrzymywali metaliczny szum poprzez mieszanie 2...3 sygnałów prostokątnych w bramkach EXOR (spełniających rolę prymitywnych układów mnożących), przez co powstawały liczne częstotliwości nieharmoniczne.

Często używanym dźwiękiem jest odgłos klaśnięcia w dłonie (*hand clap*). Również w tym przypadku firma Roland nie zawiodła, uzyskując świetny efekt w bardzo prostym układzie. Trzykrotnie w odstępach 10 ms wyzwalała jest krótka paczka szumu białego. Po kolejnych 10 ms wyzwalał jest gasnący szum, trwający 0,7 sekundy, symulujący pogłos pomieszczenia. Wszystkie te sygnały przechodzą następnie przez filtr pasmowoprzepustowy 700...1000 Hz. Dźwięk analogowego *hand clap*u był tak dobry, że umieszczano go w pamięci automatów cyfrowych zamiast nagrania prawdziwych klaszczących dłoni. Czyżby symulacja była lepsza od oryginału?

Stosuje się też brzmienia specjalne. Tu realizowane są różne, często niekonwencjonalne, pomysły. Grupa *Depeche Mode* w roli dźwięku bębna taktowego użyła kiedyś czystych impulsów prostokątnych z generatora. Interesujące brzmienia powstają przez silne przesterowanie standardowych dźwięków perkusyjnych, wycięcie czy uwypuklenie niektórych pasm częstotliwości i tak dalej...

Ilekość studiuję schematy starych automatów analogowych firmy Roland, odczuwam podziw dla sztuki inżynierskiej łączącej prostotę ze znakomitymi efektami. Zdaniem wielu sztuka ta osiągnęła wyżyny w modelu TR-808. Istnieją opracowania naukowe analizujące użyte tam metody symulacji, a na temat powstania brzmień krążą nawet anegdoty. Podobno podczas prac nad symulacją talerzy niezdarzy inżynier rozlał herbatę na układ próbny. Wilgoć utworzyła nowe połączenie, które znacznie poprawiło brzmienie ☺.

Oczywiście w wersji końcowej połączenie to (po długich miesiącach prób!) skrupulatnie odtworzono. Inny mit głosi, że wyjątkowy dźwięk werbla automat zawdzięczał tranzystorom 2SC828R-NZ, użytym w generatorze białego szumu. Literki NZ oznaczały, że były to egzemplarze wadliwe – odrzuty z produkcji. Jednak te „buble”, w przeciwieństwie do pełnosprawnych 2SC828R, wytwarzały szum o świetnym brzmieniu. Do tego stopnia, że gdy poprawiła się technologia półprzewodnikowa i tranzystorów NZ zabrakło, zastępczego typu nie udało się znaleźć, a produkcję TR-808 rychło zakończono...

I po tym techniczno-historycznym wprowadzeniu wracamy do naszego projektu.

Rozwiązania układowe

Opisywany automat symuluje 4 instrumenty: *bass drum* (BD), *snare drum* (SD), *hi-hat* w wariantach *open* i *closed* (HH) oraz *hand clap* (HC). Układy są adaptacjami rozwiązań firmy Roland/Boss. Aby realizacja była oszczędna, opierałem się głównie na schematach automatów DR-110 i TR-606 – najprostszych, choć brzmieniowo niewiele ustępujących większym maszynom. Przyjrzyjmy się użytym rozwiązaniom układowym.

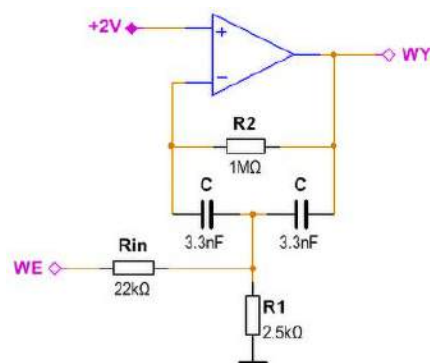
W torze symulacji każdego instrumentu natopkamy taki czy inny wariant filtru pasmowoprzepustowego, którego podstawowy układ z przykładowymi wartościami elementów jest przedstawiony na **rysunku 3**, a jego charakterystyka amplitudowa – na **rysunku 4**. Częstotliwość środkowa, wzmocnienie przy tej częstotliwości i dobroć („ostrość charakterystyki”) są dane wzorami:

$$f_0 = \frac{1}{2\pi\sqrt{[(R_{in} \parallel R_1)R_2]C}}$$

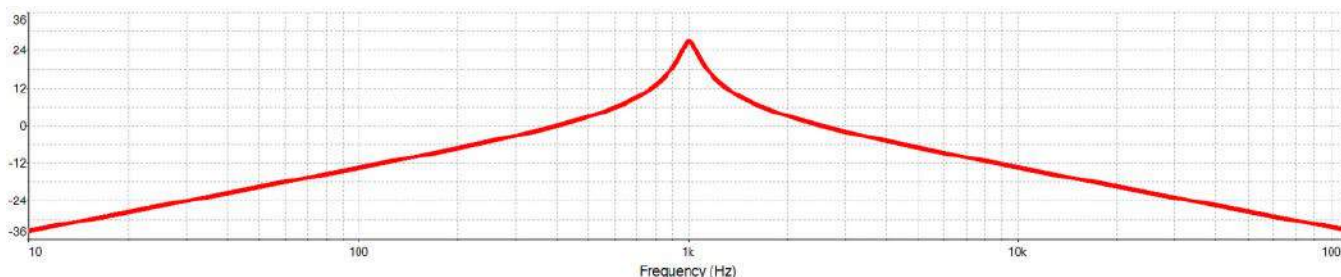
$$A_0 = -\frac{R_2}{2R_{in}}$$

$$Q = \frac{1}{2} \cdot \sqrt{\frac{R_2}{R_{in} \parallel R_1}}$$

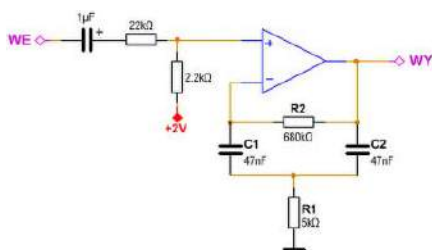
Układ pasmowoprzepustowy z rysunku 3 nadaje się do zawężenia pasma szumu



Rysunek 3. Filtr pasmowoprzepustowy



Rysunek 4. Charakterystyka filtru pasmowoprzepustowego



Rysunek 5. Symulator bębna

w instrumentach *hand clap* i *hi-hat*. Zawężenie to odzwierciedla rezonansowe właściwości symulowanego elementu bądź pomieszczenia odsłuchowego. Natomiast wariant tego samego filtru, z sygnałem wyzwalającym doprowadzonym do nieodwracającego wejścia wzmacniacza, stanowi symulator bębna (**rysunek 5**). Układ ten pracuje na granicy stabilności, a pobudzony na wejściu impulsem wytwarza drgania gasnące, modelując pracę napiętej membrany. Na **rysunku 6a** widzimy wejściowy impuls prostokątny i odpowiedź na wyjściu układu. Jak widać, zaraz na początku fala jest trochę zniekształcona, zatem brzmi ostrzej, co jest szczególnie istotne w symulacji bębna taktowego. Dla porównania przedstawiam przykładowy wykres fali prawdziwego bębna akustycznego (**rysunek 6b**).

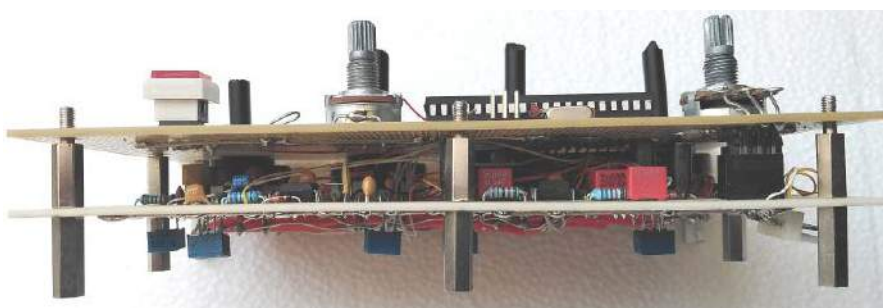
W symulatorze z rysunku 5 można przestrajac częstotliwość drgań. Jak? Spójrzmy na podane wcześniej wzory na f_0 i Q . Gdy nie ma rezystora R_{in} , wtedy wzory te przyjmują postać:

$$f_0 = \frac{1}{2\pi\sqrt{R_1 \cdot R_2 \cdot C}}$$

$$Q = \frac{1}{2} \cdot \sqrt{\frac{R_2}{R_1}}$$

f_0 wyznacza (w dobrym przybliżeniu) częstotliwość drgań. Natomiast dobroć Q określa w tym układzie jakby „ilość drgań” (ściślej – ile będzie okresów drgań, zanim ich amplituda nie zmaleje do 4% początkowej wartości). Widzimy, że okres drgań $1/f_0$ jest proporcjonalny do pierwiastka z iloczynu $R_1 \cdot R_2$, a dobroć – do pierwiastka ze stosunku R_2/R_1 . Wniosek jest taki, że przestrajanie okresu/częstotliwości najlepiej przeprowadzać rezystorem R_1 . Czas wybrzmiewania będzie wtedy niezmienny, bo jak skrócimy okres to jednocześnie tyle samo razy zwiększymy „ilość drgań” i vice versa. Niezależna regulacja zarówno częstotliwości jak i czasu wybrzmiewania byłaby w tym układzie trudna; należałoby raczej użyć generatora o pracy ciągłej, wzmacniacza sterowanego napięciem i generatora obwiedni, tak jak to uczyniono w modelu TR-909.

W automacie są używane bardzo proste bramki sygnałowe – patrz **rysunek 7**,



Fotografia 2. Połączone płytki układu

pochodzący z instrukcji serwisowej TR-808. Sygnał akustyczny dochodzi do bazy tranzystora, a przebieg sterujący (*envelope*; obwiednia) jest podawany przez rezystor i diodę do kolektora. Wadą układu są zniekształcenia sygnału oraz obecność na wyjściu składowej sterującej, którą trzeba odcinać filtrem górnoprzepustowym (co najmniej 1 kHz). Taki układ nadaje się więc zasadniczo tylko do bramkowania szumu. Jest za to, przynajmniej, wyjątkowo prosty. Podobnie jak układy obwiedniowe użyte do sterowania (**rysunek 8**).

Drgania gasnące bębna taktowego i werbla powstają w filtrach na U2a i U3a. W przypadku werbla mamy też dodatek bramkowanego szumu białego (Q1, U3b), wytwarzanego przez mikrokontroler.

Metaliczny szum dla *hi-hat*'u bierze początek w mikrokontrolerze, który generuje 6 przebiegów prostokątnych o częstotliwościach od 205 do 797 Hz – dokładnie takich, jak w kultowym TR-808. Ich mieszanka przechodzi przez filtr pasmowy ok. 7 kHz na U4a, po czym jest kluczowana z różnymi stałymi czasowymi w zależności od wariantu *closed* (D4) czy *open*



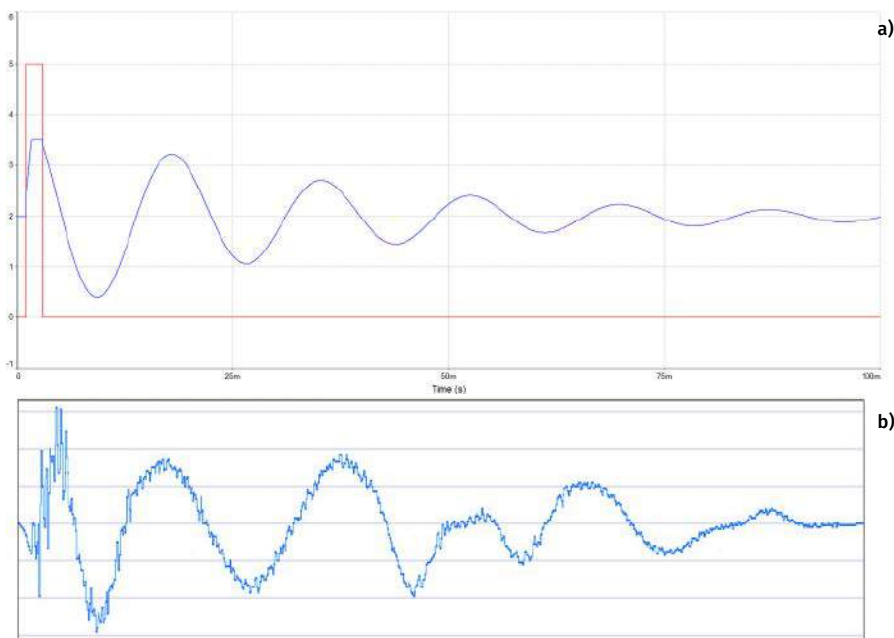
Fotografia 3. Układ w obudowie

(D3). Całość jest jeszcze dodatkowo filtrowana górnoprzepustowo w U4b.

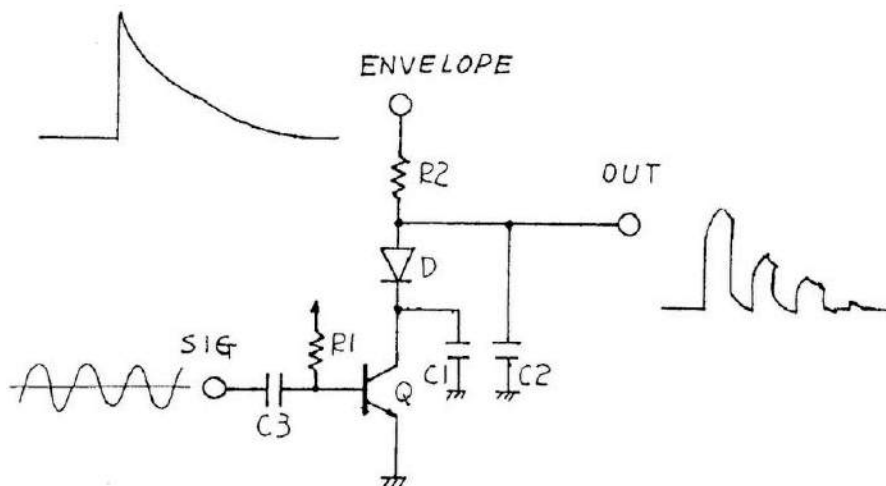
Sygnał *hand clap*'u, składający się z „kłaśnień” (D7) i „pogłosu” (D8, D9), jest przepuszczany przez pasmowy filtr ok. 1 kHz na U2b.

Ze względu na zasilanie pojedynczym napięciem stosowana jest „sztuczna masa” o potencjale +2 V.

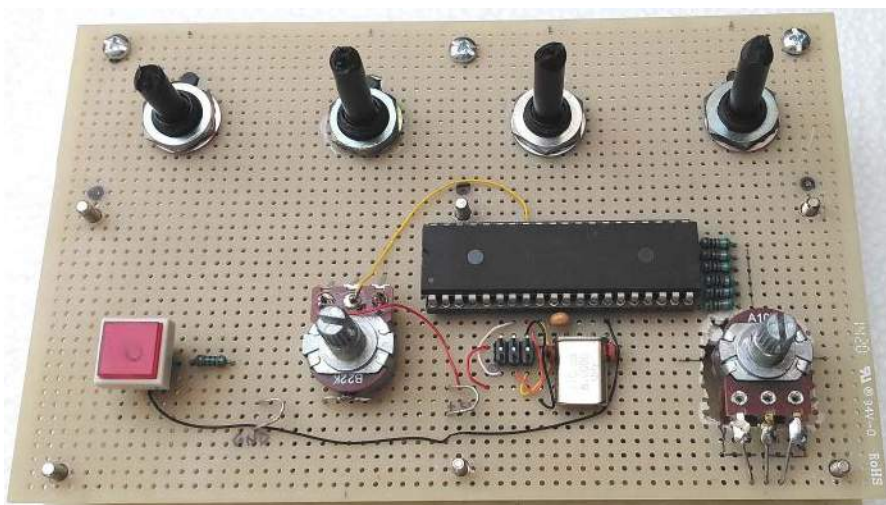
Zastrzeżenia może budzić użycie wzmacniaczy LM358, niezalecanych do aplikacji



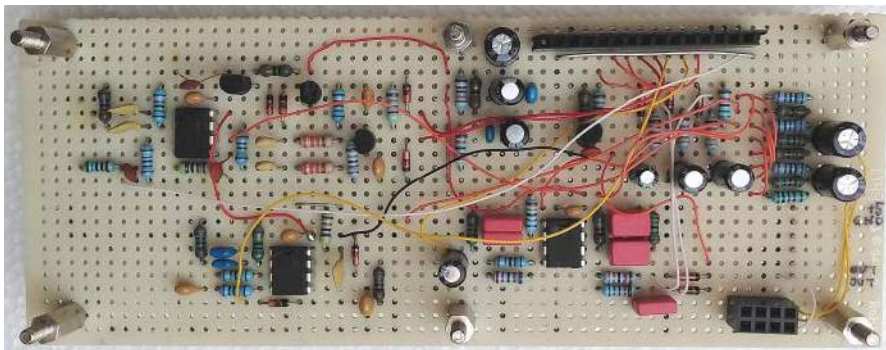
Rysunek 6. Odpowiedź symulatora na impuls oraz przebiegi prawdziwego bębna



Rysunek 7. Bramka sygnałowa



Fotografia 4. Płytkę sterującą



Fotografia 5. Płytkę brzmieniową – strona elementów

audio. Tutaj pracują one jednak w łatwych warunkach, z dużymi i niezbyt szybkimi sygnałami, a przy tym zadowolają się zasilaniem 5 V, są tanie i łatwo dostępne. Użycie lepszych typów jak RC4558 czy TL072 powodowałoby konieczność zwiększenia napięcia zasilania do przynajmniej 10 V. Byłoby to jednak w tak prostym układzie niecelowe. W syntezatorach firmy Korg z serii Volca są z powodzeniem używane wzmacniacze LM324

(poczwórna wersja LM358), również zasilane z 5 V.

Sterowanie

Sterownikiem automatu perkusyjnego może być prosty 8-bitowy mikrokontroler, np. AVR. Ma to tę zaletę, że można użyć układu w obudowie DIP z rozstawem wyprowadzeń 2,54 mm. W warunkach amatorskich bardzo ułatwi to montaż. W opisywanym urządzeniu

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

PŁYTKA STERUJĄCA

Rezystory:

R7: 220 Ω 5%
R1...R6: 100 kΩ 5%
P1: potencjometr obrotowy 10 kΩ, Ø 6 mm, liniowy
P2: potencjometr obrotowy 2×10 kΩ, Ø 6 mm, audio

Kondensatory:

C2...C4: 0,1 μF/25 V 20%, ceramiczny
C1: 22 μF/16 V 20%, elektrolityczny

Półprzewodniki:

U1: ATmega16-16PU, DIP40
LED1: dioda LED, czerwona (np. wbudowana w przycisk S1)

Inne:

S1: przycisk chwilowy, najlepiej z wbudowaną diodą LED, np. PB6133FBL-1
SW1...SW4: przełącznik obrotowy 12-pozycyjny, Ø 6 mm, np. CK1029
J1, J3: złącze męskie „goldpin”, 2×3
J2: złącze męskie „goldpin”, 1×14

PŁYTKA BRZMIENIOWA

Rezystory:

R9: 10 Ω 5%
R14, R18: 220 Ω 5%
R29: 560 Ω 5%
R17, R36, R37, R48: 1 kΩ 5%
R12, R16, R52, R53: 2,2 kΩ 5%
R8, R15: 4,7 kΩ 5%
R33, R40: 5,6 kΩ 5%
R19, R35, R46: 10 kΩ 5%
R10, R24, R30, R51: 22 kΩ 5%
R26, R27, R41, R42, R44, R49: 33 kΩ 5%
R22, R34: 47 kΩ 5%
R43, R45: 68 kΩ 5%
R23, R31: 82 kΩ 5%
R21: 150 kΩ 5%
R25: 220 kΩ 5%
R20, R50: 470 kΩ 5%
R47: 680 kΩ 5%
R13, R28, R38: 1 MΩ 5%
PR3: potencjometr montażowy 1 kΩ
PR2, PR5: potencjometr montażowy 5 kΩ
PR1, PR4: potencjometr montażowy 10 kΩ

Kondensatory:

C17, C18: 470 pF/25 V 5%, ceramiczny lub foliowy
C10, C19, C22: 1 nF/63 V 5%, foliowy
C24, C25: 2,2 nF/63 V 5%, foliowy
C11, C20, C21: 3,3 nF/63 V 5%, foliowy
C6, C7, C29: 22 nF/63 V 5%, foliowy
C16, C37, C38: 47 nF/63 V 5%, foliowy
C8, C12, C32, C36, C40: 0,1 μF/25 V 20%, ceramiczny
C5, C9, C14, C15, C26, C27, C39: 1 μF/50 V 20%, elektrolityczny
C28: 2,2 μF/50 V 20%, elektrolityczny
C33, C34, C35: 22 μF/16 V 20%, elektrolityczny
C13, C31: 100 μF/6,3 V 20%, elektrolityczny

Półprzewodniki:

U2...U4: LM358 DIP8
Q1...Q4: BC548B
D1...D11: 1N4148

Inne:

J4: złącze żeńskie „goldpin”, 1×14
J5: złącze żeńskie „goldpin”, 2×3
J6: złącze męskie KK, 1×6

Pozostałe:

złącze żeńskie KK, na kabel, 1×6 + styki
gniazdo Cinch/RCA, podwójne, na panel, np. CC-110
gniazdo zasilania 21/5.5, na panel obudowa, np. Kradex Z33 ABS
6 gatek Ø6 mm, np. PN-9D-6.4

używana jest spora ilość linii cyfrowych, co wymagało sięgnięcia po układ 40-nóżkowy. Zastosowałem ATmega16.

Mikrokontroler odczytuje położenia przełączników obrotowych, wykrywa wciśnięcia przycisku START/STOP, obsługuje LED, mierzy pozycję potencjometru TEMPO, wytwarza sygnały prostokątne szumu metalicznego i szum biały, wreszcie generuje impulsy wyzwalające dla instrumentów.

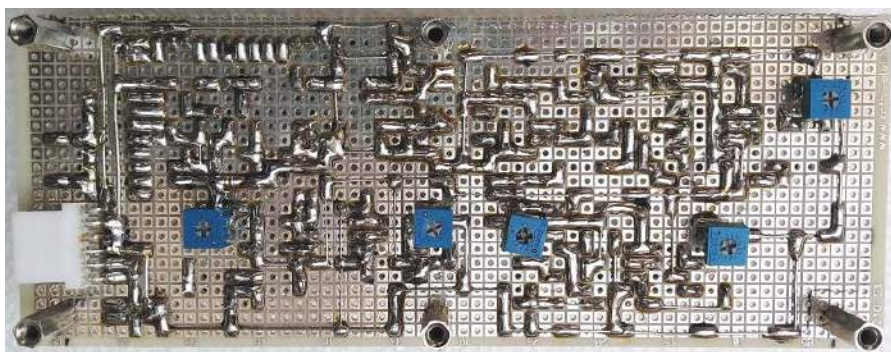
Są używane dwa timery. Timer 1 (16-bitowy) jest zwiększany co 128 μ s i służy do odmierzania czasu. Timer 2 (8-bitowy), pracujący w trybie CTC, zgłasza przerwania co 29,9 μ s, a w podprogramie ich obsługi generowane są przebiegi tworzące szum metaliczny, a także szum biały: pseudolosowy sygnał zerojedynkowy pochodzący z 16-bitowego rejestru LFSR ze sprzężeniem zwrotnym od 13. i 14. bitu. Drugi taki rejestr, służący do losowego wybierania wzorów rytmicznych, obsługiwany jest w programie głównym. Program główny odczytuje też stany przełączników i napięcie wejścia analogowego z potencjometru, a przede wszystkim odlicza szesnastki taktu i w odpowiednich chwilach generuje impulsy wyzwalające o odpowiedniej szerokości. Jest więc sporo odmierzania odcinków czasu, co jest przeprowadzane w oparciu o odczyty bieżących stanów Timera 1.

Program został napisany w języku C w darmowej wersji środowiska IAR Embedded Workbench for AVR.

Ustawienia *fuse bitów* mikrokontrolera ATmega16: High Byte = 11111111₂ = 0xFF, Low Byte = 11100100₂ = 0xE4.

Montaż i uruchomienie

Jak wspominałem, układ składa się z dwóch płytek. Na górze znajduje się płytka sterująca; na dole – płytka brzmieniowa (fotografia 2). Całość umieściłem w obudowie Kradex Z33 (fotografia 3).



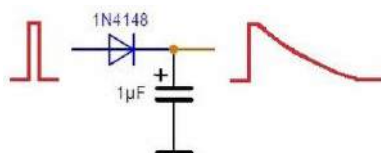
Fotografia 6. Płytki brzmieniowej – strona lutowania

Na zdjęciu płytki sterującej (fotografia 4) widać, że sterownik w modelu został wyposażony w rezonator kwarcowy. Jednak ostatecznie użyłem wewnętrznego generatora RC mikrokontrolera.

Perkusję zmontowałem na płytkach uniwersalnych z uwagi na możliwość łatwego dokonywania zmian układowych, zwłaszcza na płycie brzmieniowej (fotografia 5).

Obie płytki są połączone dwoma złączami typu „goldpin”: 14-stykowym, przenoszącym sygnały cyfrowe i zasilanie, oraz 6-stykowym, łączącym potencjometr głośności (znajdujący się na płycie sterującej) z płytą brzmieniową.

Przed montażem potencjometrów do sterujących na płycie brzmieniowej należy się zdecydować, z której strony je przylutujemy. Umieszczenie ich po stronie lutowania (fotografia 6) umożliwi nastawianie brzmień po tym, jak nastąpi połączenie obu płytek i do elementów na płycie brzmieniowej nie będzie już dostępu.



Rysunek 8. Układ obwiedniowy

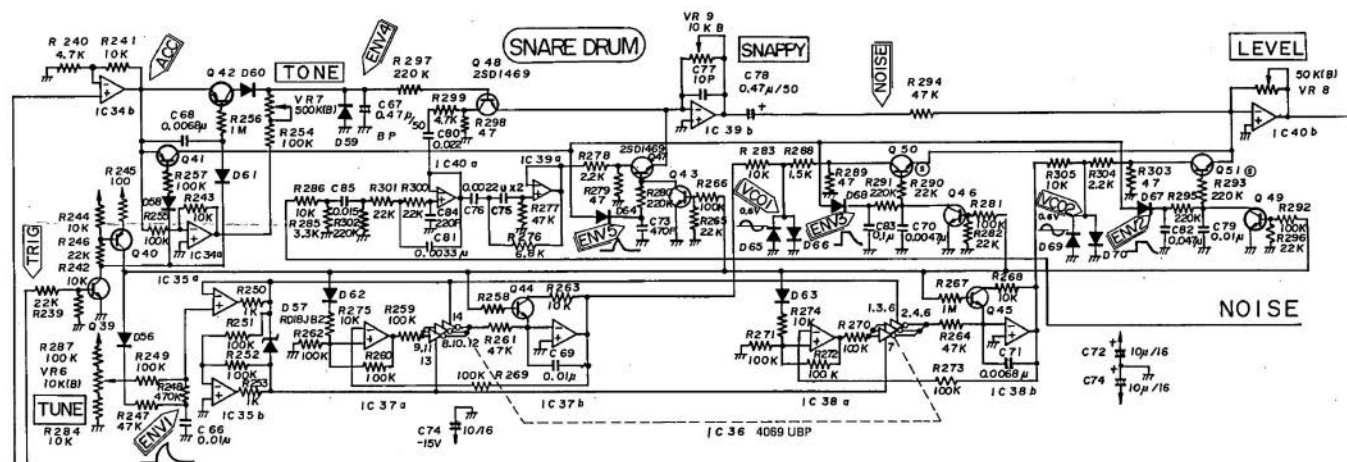
Gniazda wyjściowe (Cinch/RCA) i zasilania (5.5/2.1) są przymontowane do górnej części obudowy i dołączone do płytki brzmieniowej przewodami i 6-stykowym złączem typu KK.

Po zmontowaniu i złączeniu ze sobą obu płytek oraz podłączeniu gniazd, do układu doprowadzamy zasilanie 5 V i sprawdzamy działanie płytki sterującej. Z zaprogramowanym mikrokontrolerem powinna ruszyć od razu. Po wciśnięciu przycisku zacznie rozbłyskać LED z częstotliwością zależną od pozycji potencjometru. Dysponując oscyloskopem możemy zbadać obecność przebiegów prostokątnych na wyjściach PD1...PD6, szumu na PD7 oraz impulsów wyzwalających na PC0...PC5 (zależnie od ustawień przełączników obrotowych).

Następnie uruchamiamy płytkę brzmieniową. Zapewne, wskutek błędów montażowych, nie wszystkie instrumenty zabrzmią prawidłowo. Znowu pomocny będzie oscyloskop (w ostateczności multimetr, najlepiej z funkcją „bargraph”), którym prześledzimy propagację impulsów wyzwalających, wytwarzanie napięć obwiedniowych i sygnały wyjściowe.

Wskazówki dla eksperymentatorów

Co można zmienić? Na pewno sterownik. Dobry automat perkusyjny umożliwia



Rysunek 9. Układ zaawansowanego symulatora werbla

Quiz: cykl Audio Out

Skuteczność głośnika 80 dB/W/m można ocenić jako:

- ☐ niską
- ☐ wysoką
- ☐ bardzo wysoką

Przy skuteczności 80 dB/W/m wymagany jest wzmacniacz o mocy minimum:

- ☐ 5 W RMS
- ☐ 15 W RMS
- ☐ 25 W RMS

Głośniki wysokotonowe firmy Wavecor wykorzystują:

- ☐ duży pierścień ferrytowy
- ☐ małe magnesy neodymowe
- ☐ ani A, ani B, tylko inne rozwiązanie

Długość fali akustycznej dla dźwięku o częstotliwości 2,5 kHz wynosi:

- ☐ 98 mm
- ☐ 138 mm
- ☐ 182 mm

Charakterystyka częstotliwościowa impedancji głośnika pokazuje częstotliwość rezonansową w postaci:

- ☐ schodka
- ☐ dołka
- ☐ garbu

Częstotliwości rezonansowe głośników w miarę ich użytkowania spadają o:

- ☐ 5%
- ☐ 10%
- ☐ 15%

Zaletą chassis głośnika odlewane ciśnieniowo, w porównaniu do tłocznej stali, jest:

- ☐ niższy koszt produkcji
- ☐ brak efektu dzwonienia
- ☐ mniejsze zniekształcenia drugiej harmonicznej

Odpowiedzi częstotliwościowe głośników wysokotonowych i niskotonowych powinny pokrywać się:

- ☐ około dwie oktawy
- ☐ około jednej oktawy
- ☐ około trzy oktawy

W głośniku niskotonowym membrana z dopingiem to:

- ☐ membrana nasycona opatentowanym roztworem
- ☐ membrana pokryta specjalną powłoką
- ☐ membrana z mieszanki papieru i włókna szklanego

Membrana z dopingiem charakteryzuje się:

- ☐ mniejszymi zniekształceniami harmonicznymi
- ☐ wyższą częstotliwością rezonansową
- ☐ niższą częstotliwością rezonansową

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 03.02.2023.

zapamiętywanie rytmów programowanych przez użytkownika, a także ręczne wyzwalanie instrumentów. Wymaga to oczywiście zastosowania szeregu przycisków, kilkuset diod LED oraz wyświetlacza, no i napisania dość złożonego programu. Bardzo pożądane byłoby wejście MIDI. Wtedy automat dałoby się sterować z komputera, sekwencera itp. Można wprowadzić dynamiczne sterowanie instrumentami, zmieniając impulsom sterującym szerokość (co jest nietrudne) lub amplitudę (to jednak wymagałoby użycia wielokanałowego przetwornika C/A). Jako alternatywny sterownik można by było z pożytkiem wykorzystać Arduino lub płytkę ewaluacyjną jakiegoś mikrokontrolera.

Niejedno da się poprawić w części brzmieniowej. Według mojej oceny bardzo

dobrze brzmią *hi-hat* i *hand clap*, natomiast dźwięk werbla pozostawia trochę do życzenia. Może warto np. skrócić impuls wyzwalający poprzez zmniejszenie pojemności C5? Nienajlepiej brzmi „jednobitowy” szum biały wytwarzany przez mikrokontroler. Lepiej spisywałby się układ analogowy, wykorzystujący wzmocnione szumy tranzystora lub diody Zenera.

Oszczędna koncepcja urządzenia sprawiła, że na płycie czołowej nie przewidziałem regulatorów do nastawiania parametrów dźwięku. Ten stan łatwo zmienić, zastępując wszystkie 5 potencjometrów dostrojczych osiowymi.

Dalszy krok to eksperymenty z wartościami innych elementów. Można też dodać dalsze instrumenty, choćby zmodyfikowane wersje tu przedstawionych. Zachęcam

do prób z układami własnego pomysłu lub rozwiązaniami znalezionymi w Internecie. To świetna droga do lepszego poznania techniki analogowej! Ale uwaga: zaawansowany układ pojedynczego bębna potrafi być bardziej skomplikowany niż cała elektronika przedstawianego tu automatu. Niech zilustruje to przykład w postaci schematu symulatora werbla z TR-909 (**rysunek 9**).

Poeksperymentować na pewno warto. Produkty profesjonalne, mające trafiać do możliwie szerokiego kręgu nabywców, zwykle zawierają rozwiązania standardowe, by nie powiedzieć „oklepane”. Natomiast amatora w tworzeniu rzeczy oryginalnych ogranicza praktycznie tylko wyobraźnia. ■

Jarosław Ziembicki
j.ziembicki@wp.pl

Quiz: cykl Silniki krokowe w praktyce

Wraz ze zmniejszeniem napięcia zasilania silnika szczerkowego DC o połowę w stosunku do napięcia znamionowego, moc silnika i moment obrotowy zmniejsza się:

- ☐ dwa razy
- ☐ o 40%
- ☐ czterokrotnie

Prosty sterownik silnika krokowego można zrobić na timerach 555. Ile takich układów musi być w sterowniku:

- ☐ jeden
- ☐ dwa
- ☐ trzy

Czas T, po którym silnik osiąga 2/3 maksymalnego momentu obrotowego jest określony indukcyjną statą czasową według wzoru:

- ☐ $T = L \times C$
- ☐ $T = L/R$
- ☐ $T = R/L$

Czas potrzebny do osiągnięcia przez prąd uzwojenia stanu ustalonego w przybliżeniu wynosi:

- ☐ 2T
- ☐ 3T
- ☐ 5T

Jeśli silnik krokowy ma indukcyjność 3 mH, a rezystancja uzwojenia wynosi 54 Ω , to T wynosi:

- ☐ 55 μ s
- ☐ 110 μ s
- ☐ 150 μ s

Minimalny czas kroku oblicza się jako wielokrotność czasu T:

- ☐ 4T
- ☐ 3T
- ☐ 2T

Zastosowanie sterowania 2-fazowego dla dwukierunkowego, unipolarnego silnika krokowego, w porównaniu ze sterowaniem 1-fazowym daje:

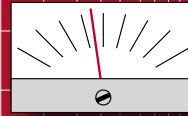
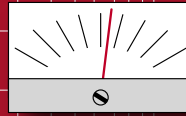
- ☐ większy moment obrotowy
- ☐ łatwiejszą zmianę kierunku
- ☐ większą szybkość obrotów

Czy zastosowanie sterownika na płycie Arduino pozwala bezpośrednio sterować silnikiem krokowym?

- ☐ tak
- ☐ nie – potrzebne są tranzystory przełączające prąd
- ☐ nie – potrzebny jest licznik dekadowy

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 10.02.2023.

AUDIO OUT



PE Mini-monitor – zwrotnica dla głośników Wavecor, część 1

LS3/5A to prawdopodobnie najlepszy głośnik – minimonitor, lecz jest bardzo drogi, a pozyskanie części może być kłopotliwe. Po latach spędzonych na próbach uzyskania jakości dźwięku jak najbardziej zbliżonej do LS3/5A, przedstawiamy mini-monitor w wersji PE. Wszystkie głośniki są dość niedoskonałe i mówi się, że trzeba się sporo męczyć ze zwrotnicą wykonując wiele iteracji, aż do uzyskania wymaganej „mieszanki elementów”. Najpierw jednak to podstawowe parametry techniczne muszą być właściwe.

Zdobycie części do głośników to zawsze problem, więc przetworniki, płytki ze zwrotnicą, cewki, kondensatory, a nawet zestawy obudów będą dostępne w sklepie PE. Ten mały głośnik został zaprojektowany z myślą o maksymalnej dokładności w relacji do swoich rozmiarów, więc maksymalna głośność jest z konieczności ograniczona. Jego niska skuteczność, wynosząca około 80 dB/W przy 1 m, oznacza, że wymagany jest wzmacniacz o mocy minimum 25 W RMS przy 8 Ω . Głośnik przeciąża się przy wzmacniaczach o mocy powyżej 60 W. Opisany w EPE w maju 2017 wzmacniacz

MX50 jest do niego idealny, ale proponuję jeszcze lepsze rozwiązania!

Korzystanie z płytki uniwersalnej zwrotnicy pasywnej

Zwrotnica PE Mini-Monitor jest zbudowana na płytce uniwersalnej zwrotnicy pasywnej opisaną w poprzednim odcinku. Numery komponentów odnoszą się tutaj do tych z ogólnego schematu obwodu (rysunek 9 w poprzednim odcinku). Oznacza to, że niektóre numery komponentów, które nie są używane, zostaną pominięte; na przykład C1.

Znalezienie głośników

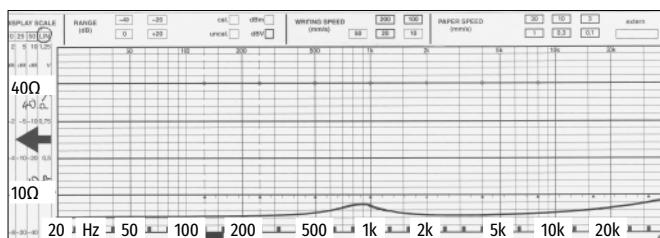
Projekt zwrotnicy nie może być rozpoczęty, dopóki nie zostaną wybrane głośniki. Oznacza to więc zwykle przeglądanie wielu kart katalogowych. B110A jest przez wielu uważany za najlepszy mały głośnik basowy, nawet po 40 latach, ponieważ technologia głośnikowa rozwija się dość powoli. Na rynku było kilku konkurentów, ale zbyt często znikali w momencie, gdy kończyłem projekt. Po otrzymaniu wielu próbek uznałem, że chiński producent Wavecor jest najbliższy oczekiwanej jakości basu, z kilkoma nowoczesnymi ulepszeniami technicznymi oraz ze sprawdzoną ciągłością



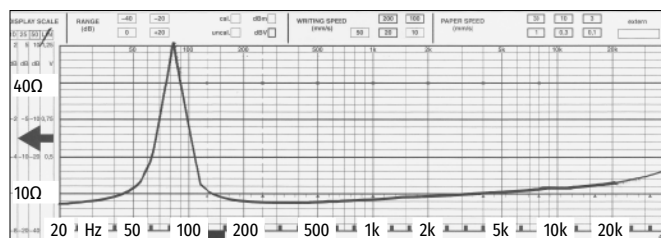
Rysunek 1. Mini-monitor PE może być wykorzystany do odsłuchu bliskiego pola. (Tak, 'podstawkami' są cegły!)



Rysunek 2. Głośnik wysokotonowy Wavecor TW022WA04 jest niewielki dzięki zastosowaniu pierścieniowego magnesu z metali ziem rzadkich.



Rysunek 3. Charakterystyka impedancji głośnika wysokotonowego Wavecor TW022WA04. 'Garb' pokazuje częstotliwość rezonansową. Zwróćmy uwagę, że skala pionowa jest liniowa i wskazuje omy (2 Ω /działkę).



Rysunek 4. Charakterystyka impedancji głośnika niskotonowego Wavecor WF120BD06. Częstotliwość rezonansowa została podniesiona przez efekt poduszki powietrznej pięciolitrowej obudowy zamkniętej (pokazanej na rysunku 1). Wzrost indukcyjności przy wysokich częstotliwościach jest znacznie mniejszy niż normalnie.

dostaw. Ich głośniki wysokotonowe są również dobre – wykorzystują małe magnesy neodymowe zamiast dużego pierścienia ferytowego stosowanego w T27. Wybrałem więc jeden z typów takiego głośnika.

Wybór częstotliwości zwrotnicy

Wielu konstruktorów PE będzie przywiązanych do swoich komputerów, więc uznałem, że idealny mini-monitor powinien nadawać się do użytku zarówno w towarzystwie monitora (bliskie pole) na biurku, jak i w optymalnym układzie wolnostojącym. Większość głośników dwudrożnych jest „fazowa” (rozczołkowana – głośnik wysokotonowy i niskotonowy brzmiały jakby były oddzielnymi źródłami), gdy słucha się ich z bliska. Często niewielka odległość, powiedzmy 1,5 m, jest potrzebna, aby wyjścia głośnika niskotonowego i wysokotonowego prawidłowo się zintegrowały. W naszej konstrukcji, głośniki są małe i zamontowane jak najbliżej siebie. Jeśli zastosuje się stosunkowo niski punkt zwrotnicy, centra akustyczne głośników mogą być oddalone od siebie o mniej niż długość fali. Zapobiega to wspomnianemu zjawisku braku integracji przy małej odległości. Typowe ustawienie biurka komputerowego pokazano na rysunku 1. Normalnie, przy 120 mm głośniku niskotonowym i 19 mm głośniku wysokotonowym, zastosowana zostałaby zwrotnica z zakresu

od 3 do 4 kHz. W tym projekcie docelowo zastosowano zwrotnicę pracującą w okolicach 2,5 kHz. Ramka głośnika wysokotonowego ma tylko 65 mm średnicy, pionowa odległość pomiędzy środkami przetworników wynosi 94 mm, a ich ramki prawie się stykają. Długość fali 2,5 kHz wynosi 138 mm (prędkość dźwięku (343 m/s) podzielona przez częstotliwość) jest znacznie większa od tej wartości, więc spójne czoło fali powinno zostać osiągnięte.

Częstotliwość rezonansowa

Większość 19 mm głośników wysokotonowych ma wysoką częstotliwość rezonansową wynoszącą około 1,2 do 1,8 kHz, co oznacza, że nie są w stanie obsłużyć zwrotnicy 2,5 kHz. Głośnik wysokotonowy 25 mm byłby lepszy, ale wtedy dyspersja przy wysokich częstotliwościach byłaby gorsza niż w T27. Okazało się, że Wavecor ma w swej ofercie idealny głośnik wysokotonowy, stanowiący kompromis pomiędzy dwoma wspomnianymi – 22 mm jednostkę o częstotliwości rezonansowej 800 Hz, oznaczoną jako 'TW022WA04', pokazaną na rysunku 2. Zawsze warto wykreślić charakterystyki impedancji interesujących nas głośników, ponieważ pokazują one wyraźnie częstotliwość rezonansową w postaci garbu. Mogą one również ujawnić przebiecia, tłumienie i indukcyjność. Charakterystyka głośnika wysokotonowego pokazana jest

na rysunku 3, ujawniając rezonans na poziomie 850 Hz. Warto zauważyć, że ten głośnik wysokotonowy jest dostępny tylko w wersji 4 Ω , prawdopodobnie po to, aby zachować lekkość cewki, co mogłoby powodować problemy z konstrukcją zwrotnicy. Na szczęście jest on bardziej czuły niż głośnik niskotonowy, więc można zastosować szeregową rezystancję, aby zapobiec zejściu impedancji całego systemu poniżej 8 Ω . Rysunek 4 pokazuje charakterystykę głośnika niskotonowego z 5-litrową skrzynią, dającą rezonans przy 82 Hz przy wysokim Q. Należy zauważyć, że częstotliwości rezonansowe generalnie spadają o 5% w miarę



Rysunek 7. Widok tyłu głośnika niskotonowego Wavecor WF120BD06; zauważ obszerne odpowietrzenie.



Rysunek 5. Zamontowany głośnik wysokotonowy Wavecor TW022WA04 do mini-monitora PE.



Rysunek 6. Głośnik niskotonowy Wavecor WF120BD06 do mini-monitora PE oraz front głośnika niskotonowego Wavecor zamontowany w skrzyni (proszę zwrócić uwagę na membranę z dopingiem).



użytkowania głośnika, co przypomina w pewnym sensie „docieranie się”.

Głośnik basowy

Najtrudniejszą sprawą do powielenia z B110A jest jej wysoka zgodność wynosząca 3 mm/N i niska częstotliwość rezonansowa o wartości 37 Hz. 4,75-calowa jednostka Wavecora WF120BD06 osiąga wyższe wartości w obu tych aspektach (1,5 mm/N i 50 Hz). Jej wysoka „jakość”, czy też „szczytowość” rezonansu, Q_{ts} na poziomie 0,4, oraz mniejsza średnica membrany (100 mm w porównaniu do 115 mm w B110A) pozwoliły na uzyskanie basu tylko nieznacznie gorszego od tego, który uzyskano w tej samej objętości obudowy, co LS3/5A. Mniejsza powierzchnia membrany oznacza, że rezonans swobodnego powietrza jest w mniejszym stopniu podnoszony przez „sprężyny powietrzne” z zamkniętego powietrza w obudowie. Element ten może pracować dobrze w dowolnej obudowie o objętości wewnętrznej od około 4 do 9 litrów. Częstotliwość rezonansowa będzie się wahać od 100 do 71 Hz (całkowite Q systemu = $Q_{tc} = 1$ do 0,6), więc standardowe skrzynki LS3/5A i PE o głębokości 9 cali będą idealne. Niewielkie podbicie basu z regulatorów barwy wzmocniacza również pomaga przy niskich głośnościach. W dużej skrzyni można zainstalować kanał reflex o wymiarach 4 na 1,4 cala, co daje możliwość uzyskania większego basu (ale gorszego jakościowo), dla tych, którzy lubią elektroniczną muzykę taneczną. Początkowo obliczyłem wielkość obudowy za pomocą darmowego programu kalkulacyjnego: <http://bit.ly/pe-feb20-a01>.

Przyjrzałem się też 5,5-calowym i 5,75-calowym głośnikom Wavecora, ale ich wyższy rezonans przy 57 Hz i większa powierzchnia membrany oznaczałyby konieczność zastosowania większej skrzyni. Przetworniki basowe są dostępne z membranami papierowymi lub z mieszanki papieru z włóknem szklanym. Ja preferuję tę z włókna szklanego, ponieważ nieco większa masa ruchoma daje niższą częstotliwość rezonansową. Wavecór udostępnia karty katalogowe w formacie PDF dla obydwu głośników, dostępne na stronie

internetowej Wavecór (<http://bit.ly/pe-feb20-a02> oraz <http://bit.ly/pe-feb20-a03>), a także do pobrania na stronie PE. **Rysunki od 5 do 7** przedstawiają omawiane głośniki.

Przerzucanie żelastwa

Układ napędzający głośniki Wavecór jest lepszy niż w B110A, co pomaga zrekompensować mniejszy i sztywniejszy zespół membrany. Cewka głosowa ma większą średnicę wynoszącą 32 mm, w porównaniu do 25 mm, chociaż jej długość jest taka sama i wynosi 12 mm. Liniowy skok membrany jest o 1 mm większy niż w przypadku B110A z powodu opatentowanego, specjalnie ukształtowanego nabiegunka o nazwie ‘Balanced Drive’, który daje bardziej symetryczną charakterystykę zależności prąd/przemieszczenie, co zaś skutkuje niższymi zniekształceniami drugiej harmonicznej. Charakterystyka basu pozostaje bardzo czysta, aż do momentu gwałtownego ograniczenia. Na stronie internetowej Wavecór (<http://bit.ly/pe-feb20-a04>) znajduje się dokument techniczny, który szczegółowo omawia tę koncepcję. Dostępny jest on również do pobrania na stronie PE. B110 ma bardziej stopniowaną charakterystykę przeciążenia. Przy średnich poziomach, Wavecór ma niższe zniekształcenia. B110 ma większą użyteczną moc maksymalną, ale kosztuje ponad dwa razy więcej. Rozwiązania redukujące modulację strumienia, składające się z aluminiowego pierścienia nabiegunka i miedzianej nakładki, których brakuje w B110A, również zmniejszają zniekształcenia. Cechy te znajdują odzwierciedlenie w krzywej impedancji jako zmniejszone narastanie przy wysokich częstotliwościach. Kolejnym plusem jest zastosowanie odlewanej ciśnieniowo chasis w przeciwieństwie do tłoczonej stali. Pozwala to uniknąć dzwonienia – ja sam często musiałem podpieścić magnes B110 rozpórką, aby uniknąć tego efektu. Dość dziwną cechą jednostki basowej Wavecór jest wentylowany pajak, który redukuje szumy wiatru i wspomaga chłodzenie cewki. Przez szczeliny można zobaczyć cewkę (patrz rysunek 7). Niestety, możliwe jest również przedostanie się do środka opiłków żelaza. Podłoga

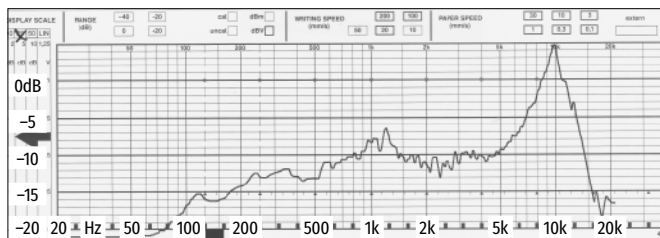
mojego warsztatu jest zaśmiecona ścinkami drutu stalowego z lutowania rezystorów, które stanowią prawdziwe zagrożenie dla magnesów głośnikowych. Ogólnie rzecz biorąc, zawieszenie Wavecór jest bardzo niskostratne (wysokie Q_m), co jest przydatną cechą dla głośnika w szczelnej obudowie.

Pożądane charakterystyki

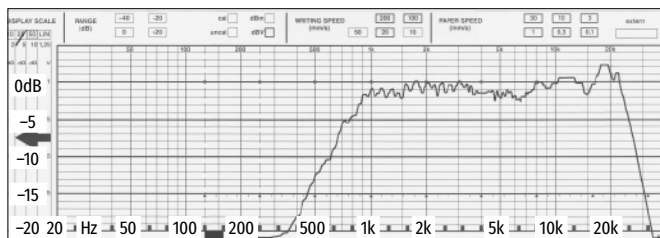
Aby zwrotnica była udana, odpowiedzi częstotliwościowe głośników powinny pokrywać się o około dwie oktawy, gdy są zamontowane w obudowie. Na **rysunku 8** widać, że jednostka basowa sięga do około 10 kHz, a głośnik wysokotonowy schodzi do 1 kHz, co widać na **rysunku 9**. Głośnik wysokotonowy ma najbardziej płaską krzywą odpowiedzi jaką kiedykolwiek mierzyłem. Głośnik niskotonowy jest dość dziwny, z potężnym rezonansem kopułki głośnikowej przy 10 kHz i zwykłym szczytem średnich częstotliwości przy około 1,2 kHz. Szczyt wysokich częstotliwości nie jest tak szkodliwy, jak się wydaje, ponieważ ostro opada poza osią. Filtry zwrotnicowe wykonują zarówno pracę związaną z łagodzeniem, jak i wyrównywaniem anomalii, więc powstały w ten sposób filtr niskotonowy jest dość złożony.

Skandal dopingowy

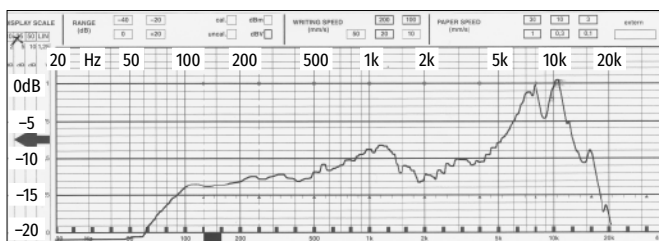
Membrana i kopułka przeciwpływowa są wykonane z twardej, sztywnej mieszanki papieru i włókna szklanego, co zapewnia twardą pracę aż do wyższych niż normalnie częstotliwości. Niestety, kiedy pojawia się załamanie, robi to z dużą intensywnością, powodując 15 dB wzrost przy 10 kHz. Stwierdziłem, że można to zjawisko zredukować do około 10 dB poprzez domieszkę do membrany plastifikowanej powłoki PVA, takiej jak Scola Master Medium (od dostawców z branży artystycznej). To również generalnie wygładziło odpowiedź i obniżyło podstawowy rezonans, dając lepszy bas, jak pokazano na **rysunku 10**. Niestety, efektywność jest również zmniejszona o kilka dB. Ta modyfikacja w postaci dopingu nie jest niezbędna, ale sprawia, że wyrównanie zwrotnicy jest bardziej efektywne. Mogę dostarczyć głośniki z dopingiem, jeśli ktoś się martwi



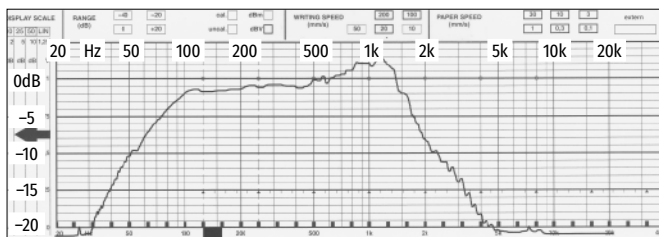
Rysunek 8. Odpowiedź częstotliwościowa głośnika niskotonowego WF120BD06 w 5-litrowej obudowie. Szczyt przy 10 kHz jest charakterystyczny dla sztywnej membrany o niskim tłumieniu. Zwrotnica została zaprojektowana tak, aby powstrzymać wzbudzenie tego rezonansu.



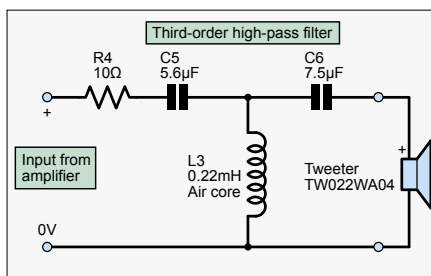
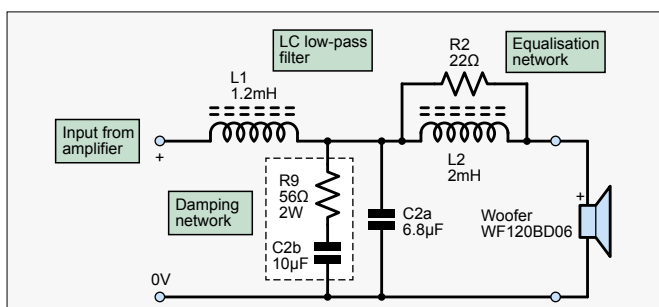
Rysunek 9. Odpowiedź częstotliwościowa głośnika wysokotonowego Wavecór TW022WA04 zamontowanego w obudowie. (Nie, nie wierzę, że charakterystyka jest aż tak płaska!)



Rysunek 10. Odpowiedź częstotliwościowa głośnika niskotonowego Wavecor WF120BD06 po dopingowaniu. Rezonans kopułki przeciwpyłowej spadł o 6 dB.



Rysunek 12. Zmniejszenie rozmiaru cewki do 1,2 mH spowodowało podniesienie częstotliwości, ale pojawia się „garb” przy 1,2 kHz.

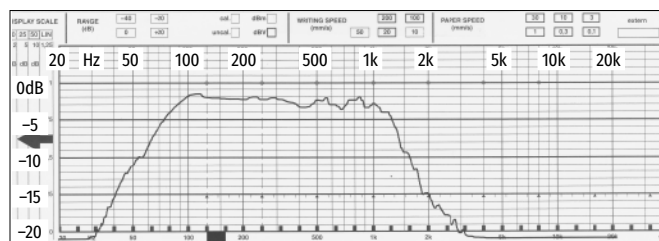


Rysunek 15. Wstępny układ górnoprzepustowy, standardowy filtr trzeciego rzędu, ale zasilający głośnik wysokotonowy 4 Ω.

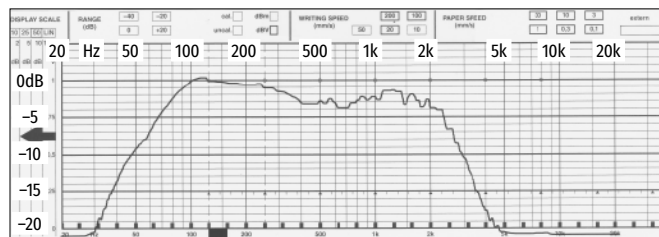
o jakość malowania. Uwaga – malować należy tylko przednią membranę i kopułkę przeciwpyłową. Doping nie może dostać się na elastyczną gumową otoczkę – jeśli się dostanie, należy go natychmiast wytrzeć wilgotną szmatką.

Projektowanie obwodu zwrotnicy

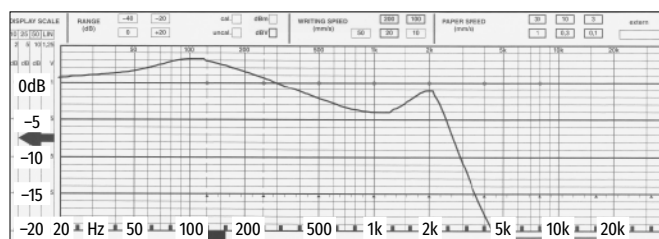
Redaktor wyraźnie poprosił mnie, abym nie pisał rozprawy na temat projektowania zwrotnic, więc po prostu przejdę pokrótce przez podstawową procedurę. Większość projektantów używa programów symulacyjnych, takich jak LTSpice lub LEAP do rozpoczęcia pracy. (Dokładny model elektryczny cewki jest



Rysunek 11. Pierwsza próba filtra dolnoprzepustowego drugiego rzędu z użyciem cewki 2,6 mH i kondensatora 15 μF. Wybrana częstotliwość jest za niska – na poziomie 1,2 kHz.



Rysunek 13. Dodatkowa cewka z równoległym rezystorem, plus zmniejszenie kondensatora do 6,8 μF zagina krzywą elektroakustyczną do wymaganego kształtu.



Rysunek 14. Układ dolnoprzepustowy zwrotnicy Wavecor (po lewej) i jego odpowiedź elektryczna (powyżej).

nawet dostępny w firmie Wavecor). Ja jestem doświadczonym, staroświeckim, analogowym konstruktorem, więc od razu przystępuję do pracy z moim ‘pudełkiem śmieci’ zawierającym części o standardowych wartościach i mając w głowie topologie układów, które już wcześniej stosowałem. Jednakże ta intuicyjna/iteratywna/empiryczna technika działa tylko jeśli masz sprzęt do wykreślania częstotliwości i charakterystyk impedancji, wraz z dość bezechowym pomieszczeniem testowym.

Pierwsza próba – sekcja dolnoprzepustowa

Lubię mieć rzeczy proste, ale nie tak proste, żeby nie działały. Z powodu szczytu na 10 kHz, filtr pierwszego rzędu składający się z samej cewki w szeregu nie zapewniłby wystarczającego tłumienia poza pasmem. Podstawowy filtr dolnoprzepustowy drugiego rzędu, składający się z cewki 2,6 mH i kondensatora 15 μF został przetestowany, dając charakterystykę elektroakustyczną pokazaną na rysunku 11. Uzyskałem częstotliwość odcięcia 1,2 kHz, a więc zbyt niską. Podniosłem częstotliwość poprzez zmniejszenie cewki do 1,2 mH, co dało w rezultacie charakterystykę z rysunku 12. Pojawił się garb przy 1,2 kHz, więc

zastosowałem jedną z moich starych sztuczek, polegającą na umieszczeniu kolejnej cewki w szeregu z głośnikiem niskotonowym, aby wyrównać ten garb. Dodałem równolegle rezystor, który mogłem dostroić, aby dostosować wymaganą wartość wyrównania. Zredukowałem kondensator do 6,8 μF, aby zwiększyć tłumienie, dając płaską odpowiedź do 2,5 kHz, jak pokazano na rysunku 13 w pobliżu wymaganej częstotliwości zwrotnicy. Wynikająca z tego odpowiedź elektryczna i obwód są zilustrowane na rysunku 14. To podejście jest „wsteczne” – większość konstruktorów projektuje najpierw charakterystykę elektryczną. Ja tego nie robię, ponieważ liczy się wyjście akustyczne. Cała procedura brzmi jak szybka praca, ale dojście do wyżej opisanego punktu zajęło około pięciu godzin.

Sekcja górnoprzepustowa

Mamy już działającą sekcję dolnoprzepustową, co zawsze jest zadaniem najtrudniejszym, więc teraz musimy zaprojektować sekcję górnoprzepustową. Częstotliwość zwrotnicy jest dość niska jak na 22-milimetrový głośnik wysokotonowy, więc konieczne będzie zastosowanie minimum filtra trzeciego rzędu, ponieważ zbyt duża ilość niskich częstotliwości na wejściu

mogłaby spowodować zniekształcenia i nawet uszkodzić głośnik wysokotonowy. Głośnik wysokotonowy ma cewkę $4\ \Omega$, więc cewka L3 w środku sekcji T jest mniejsza niż w przypadku konwencjonalnej konstrukcji $8\ \Omega$, a kondensator jest większy. Ponieważ zwrotnica jest projektowana nieco niżej niż standardowe $3\ \text{kHz}$ stosowane w większości głośników, wszystkie wartości zostały ponownie zwiększone. Patrz rysunek 15 dla obwodu wysokoprzepustowego. Rezystor wejściowy R4 ustawia tłumienie głośnika wysokotonowego i jego relatywny poziom w stosunku

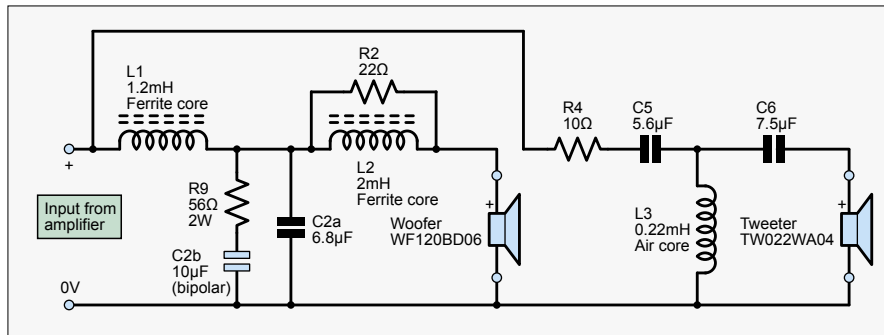
do głośnika niskotonowego. Jest to bardzo delikatna regulacja i zależy od akustyki pomieszczenia i ustawienia. Wartość może się wahać od 12 do $5,6\ \Omega$. W moim salonie optymalne było $10\ \Omega$, a w warsztacie $6,8\ \Omega$. Ostatecznie poszedłem na kompromis i wybrałem $7,5\ \Omega$ (użyłem $27\ \Omega$ połączonego równolegle $10\ \Omega$, co dało $7,3\ \Omega$). Jeśli głośnik niskotonowy nie był poddawany dopingowi, to rezystor trzeba będzie zmniejszyć o około 1 do $2\ \Omega$. Nawiasem mówiąc, piosenka zespołu 10cc *I'm Not in Love* jest dobrym testem na poziom głośnika wysokotonowego. Jeśli po lewej stronie słychać

wyraźnie brzdąkającą, zmikrofonowaną gitarę elektryczną, a wokal nie jest zbyt świszczący i wypluwany, wartość rezystora jest prawidłowa. Ponieważ rezystor o stosunkowo dużej wartości znajduje się w szeregu z wejściem głośnika wysokotonowego, to do zasilania obwodu głośnika wysokotonowego można użyć cienkiego, taniego przewodu, jeśli zastosowano bi-wiring.

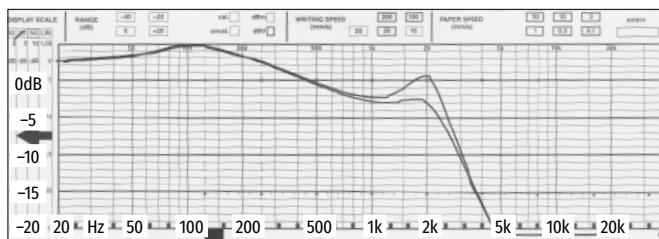
Podstawowy układ

Układ z rysunku 16 nie jest ostatecznym projektem zwrotnicy, ale podstawowe charakterystyki pokazane na rysunkach 16b do 16e są już zadowalające. W kolejnym odcinku sprawdzimy impedancję, aby upewnić się, że nie spadnie ona poniżej $6,7\ \Omega$, co jest niezbędne do zakwalifikowania kolumny jako $8\ \Omega$. Zamierzam również wprowadzić kilka drobnych zmian w zwrotnicy, aby poprawić subiektywne pasmo przenoszenia, tak zwane udźwięcznienie. ■

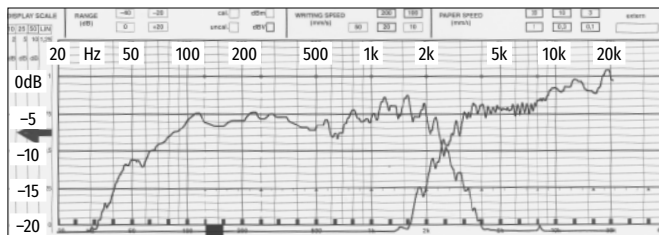
Jake Rothman



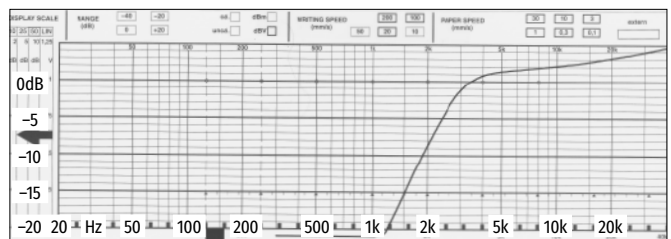
Rysunek 16a. Kompletny, prowizoryczny układ zwrotnicy uzyskany przez połączenie sekcji dolnoprzepustowej z rysunku 14 z górnoprzepustową z rysunku 15.



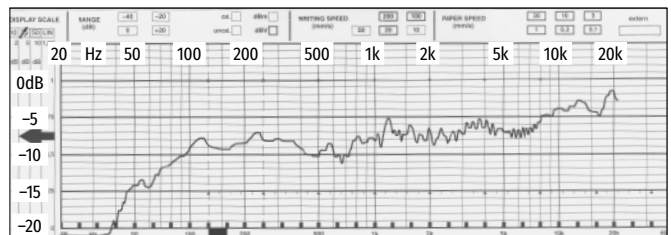
Rysunek 16b. Charakterystyka odpowiedzi elektrycznej układu dolnoprzepustowego pokazująca efekt dodania sieci tłumiącej zakreślonej linią przerywaną na rysunku 14.



Rysunek 16c. Wynikowe charakterystyki elektroakustyczne głośnika niskotonowego i wysokotonowego nałożone na siebie w celu pokazania działania zwrotnicy przy $2,5\ \text{kHz}$.



Rysunek 16d. Odpowiedź elektryczna filtra górnoprzepustowego.



Rysunek 16e. Ogólna odpowiedź układu z rezystorem tłumiącym.

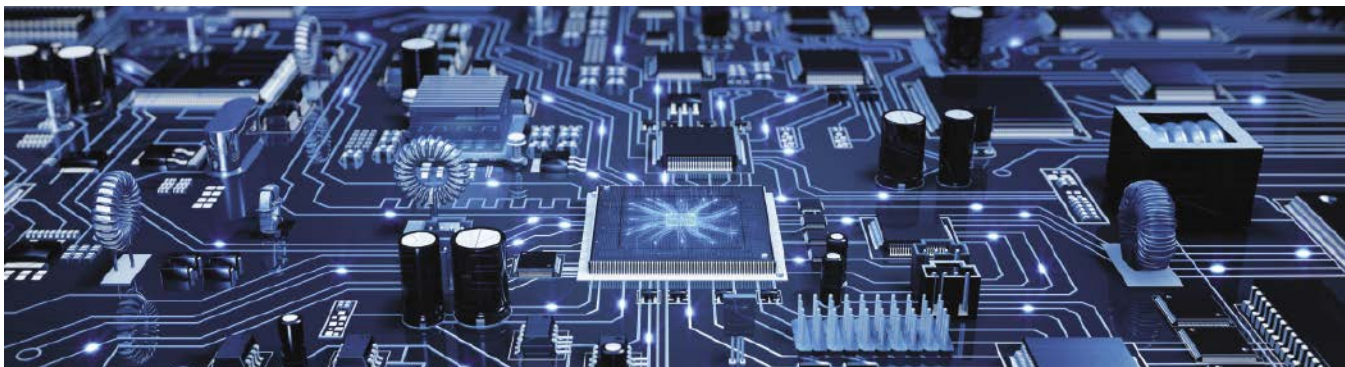
REKLAMA

Kursy w Ulubionym Kiosku

IT i Hi-tech • Muzyka i Dźwięk

Pełna oferta na stronie

www.ulubionykiosk.pl



Poziomy logiczne, część 2

W zeszłym miesiącu przyjrzelśmy się pewnym podstawowym pojęciom związanym z poziomami logicznymi (napięciami logicznymi), w tym (pokrótkę) pojęciu logiki dodatniej i ujemnej, zakresom napięć wejściowych i wyjściowych zera i jedynki dla bramek logicznych oraz marginesom szumów, które określają, jak duże przesunięcie napięcia może być tolerowane bez powodowania błędów. Przyjrzelśmy się podstawowym obwodom wewnętrznym bramek TTL i CMOS oraz historii technologii cyfrowych, która doprowadziła do obecnego stanu, w której większość używanych układów logicznych to układy CMOS, ale działające na różnych napięciach zasilania (jednak inne technologie, takie jak TTL, są nadal dostępne).

Szeroki wybór rodzajów układów logicznych i napięć roboczych prowadzi do potencjalnych problemów przy próbie połączenia układów scalonych z różnych rodzin, lub tych, które działają na różnych napięciach zasilających. Jest to częsty problem przy projektowaniu obwodów. Przykładowo, po wybraniu mikrokontrolera okazuje się, że układy peryferyjne najbardziej odpowiednie dla danego projektu pracują przy różnych napięciach zasilania i mają niekompatybilne poziomy logiczne. W zeszłym odcinku zakończyliśmy krótką dyskusję na temat kilku prostych technik łączenia starych układów 5 VLS TTL z HC CMOS; jednakże nie obejmuje to wszystkich sytuacji, które występują w przypadku nowocześniejszych układów. Jak wspomnieliśmy, najlepszym rozwiązaniem w przypadku łączenia obwodów o niekompatybilnych poziomach logicznych jest często użycie jednego z wielu dostępnych układów scalonych do translacji poziomów logicznych. Translacja poziomów logicznych jest częstym problemem w projektach komercyjnych, więc producenci

półprzewodników dostarczają wiele elementów, które spełniają taką rolę.

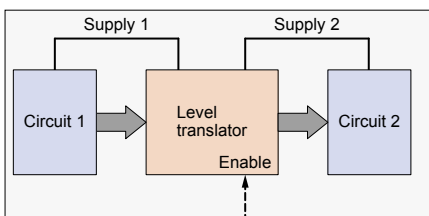
W tym odcinku zaczniemy od przyjrzenia się ogólnym formom dostępnych układów translacji, a następnie rozważymy kilka przykładowych elementów od różnych producentów. Układy te są jedynie przykładami – nie próbujemy ich polecać ponad inne. Często podobne układy są dostępne u innych producentów i każdy kto potrzebuje elementu do projektu powinien sprawdzić, co jest dostępne na rynku i wybrać najbardziej odpowiedni układ dla danego projektu. Ten artykuł dostarczy Ci pewnej wiedzy o tym, czego należy szukać w różnych sytuacjach.

Jak już wspomniano, istnieje wiele różnych typów obwodów translacji z poziomu logicznego, niektóre kluczowe przykłady są zilustrowane na rysunkach od 1 do 5. We wszystkich tych przypadkach istnieją dwa obwody działające na różnych zasilaniach (obwód 1 i obwód 2), gdzie każde z nich może mieć wyższe napięcie. Ogólnie rzecz biorąc, translacja z poziomu logicznego na niskim napięciu

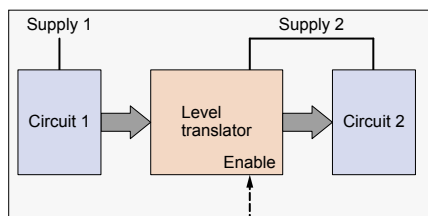
zasilania na wyższe nazywana jest „konwersją w górę”. W przypadku z wyższego na niższe napięcie nazywa się to „konwersją w dół”.

Translatory jednokierunkowe

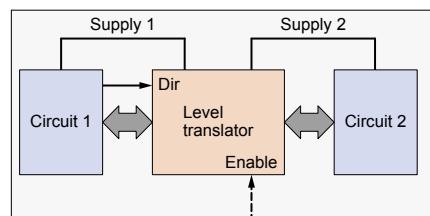
Rysunki 1 i 2 przedstawiają układy jednokierunkowe – jest to sytuacja, w której sygnał przemieszcza się tylko w jednym kierunku – w naszym przypadku z układu 1 do układu 2. Translatory jednokierunkowe mogą mieć podłączenia zasilania do obu źródeł (dual supply) lub tylko do zasilania związanego z wyjściem translatora (single supply). Translatory wymagające pojedynczego zasilania są wygodne w sytuacjach, gdy dostęp do zasilania obwodu źródłowego jest utrudniony. Zastosowanie podwójnego zasilania potencjalnie pozwala projektantom układu translatora łatwiej zapewnić optymalne parametry, takie jak niski pobór mocy w szerokim zakresie napięć zasilających. Jednokierunkowy układ translacji w swojej najbardziej podstawowej formie jest po prostu buforem (bramką logiczną z pojedynczym wyjściem, którego wartość jest



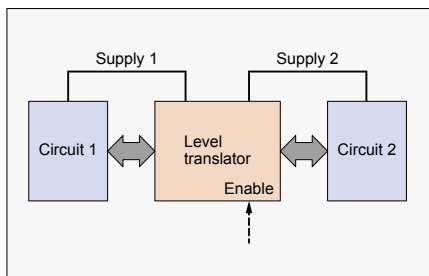
Rysunek 1. Jednokierunkowy translator poziomów z podwójnym zasilaniem.



Rysunek 2. Jednokierunkowy translator poziomów z pojedynczym zasilaniem.



Rysunek 3. Dwukierunkowy translator poziomów z kontrolą kierunku.

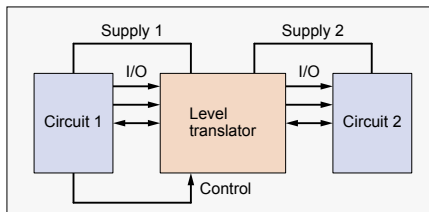


Rysunek 4. Dwukierunkowy translator poziomów z automatycznym wykrywaniem kierunku.

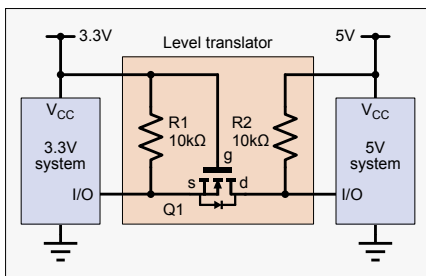
równa wartości jego pojedynczego wejścia). Układy translacyjne mogą mieć kilka równoległych buforów do translacji danych wielobitowych (lub wielu pojedynczych sygnałów). Dostępne są funkcje inne niż proste bufor, w tym bufor odwracający oraz bufor z wyjściami trójstanowymi (ułatwiający podłączenie wyjścia translatora do wspólnej szyny danych). Na rysunkach 1 i 2 pokazano możliwy sygnał „enable” – „uaktywnij”, który byłby wykorzystywany przez translatory trójstanowe. Dostępne są również translatory bramek logicznych o wielu wejściach (np. NAND lub NOR), które pozwalają np. na efektywne połączenie niestandardowego układu logicznego z translacją napięć za pomocą tego samego układu scalonego.

Translatory dwukierunkowe

Na rysunkach 3 i 4 przedstawiono dwukierunkowe translatory logiczne. Sygnały dwukierunkowe są powszechne w układach cyfrowych, w szczególności w kontekście współdzielonych magistral danych i szeregowych interfejsów peryferyjnych, takich jak I²C, które są wykorzystywane do komunikacji pomiędzy mikrokontrolerami i układami peryferyjnymi. Istnieją dwa rodzaje dwukierunkowych translatorów poziomów logicznych – takie, które posiadają kontrolę kierunku (rysunek 3) oraz takie, które automatycznie wykrywają kierunek sygnału (rysunek 4). Translatory poziomu z kontrolą kierunku mogą obsługiwać wiele bitów za pomocą jednego sterowania kierunkiem. Automatyczne wykrywanie kierunku odbywa się na zasadzie „per bit”. Oba typy (z kontrolą kierunku i bez) mogą mieć wejście zezwalające na „trzeci stan”, czyli odłączenie lub



Rysunek 5. Translator poziomów w zastosowaniu – przykład mieszanki translacji sygnału jednokierunkowego, dwukierunkowego i sygnału sterującego translatorem z układu 1.



Rysunek 6. Podstawowy dwukierunkowy translator poziomów z wykrywaniem kierunku.

ustawienie na wysoką impedancję wszystkich wyjść (w obu kierunkach). Wreszcie, jeśli chodzi o ogólne typy układów translatorów, istnieją układy scalone określane jako „translatory poziomu specyficznego dla danej aplikacji”. Elementy te mają na celu, jak sama nazwa wskazuje, translację poziomów w specyficznych sytuacjach, na przykład połączenie procesora z kartą SIM. Zazwyczaj takie sytuacje wymagają mieszanki jednokierunkowych i dwukierunkowych konwerterów poziomów i mogą również korzystać ze specyficznych sygnałów sterujących (np. w celu wdrożenia trybu wyłączenia) – patrz rysunek 5.

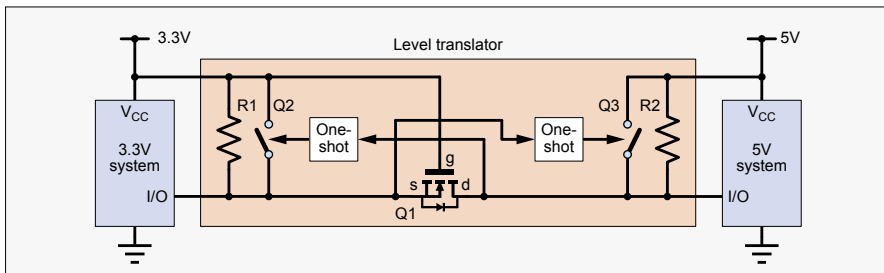
Automatyczne, dwukierunkowe translatory poziomów

Mimo stosowania różnych ulepszeń, translatory poziomów z rysunków 1 do 3 mogą być realizowane przez obwody podobne do standardowych bramek logicznych i buforów trójstanowych. Obwody stosowane do automatycznego wykrywania kierunku mogą być jednak nieco inne, dlatego przyjrzymy się bliżej powszechnie stosowanemu obwodowi tego typu. Podstawowa forma układu jest pokazana na rysunku 6. Przykład ten prezentuje układ z zasilaniem 3,3 V i 5,0 V, ale to samo podejście może być użyte dla innych napięć, zachowując ten sam względny kierunek dla napięcia od niskiego do wysokiego. Obwód może być zbudowany przy użyciu dyskretnych MOSFET-ów i rezystorów, ale powszechnie jest implementowany w układach scalonych do translacji poziomów, w których to mogą być zawarte także inne usprawnienia.

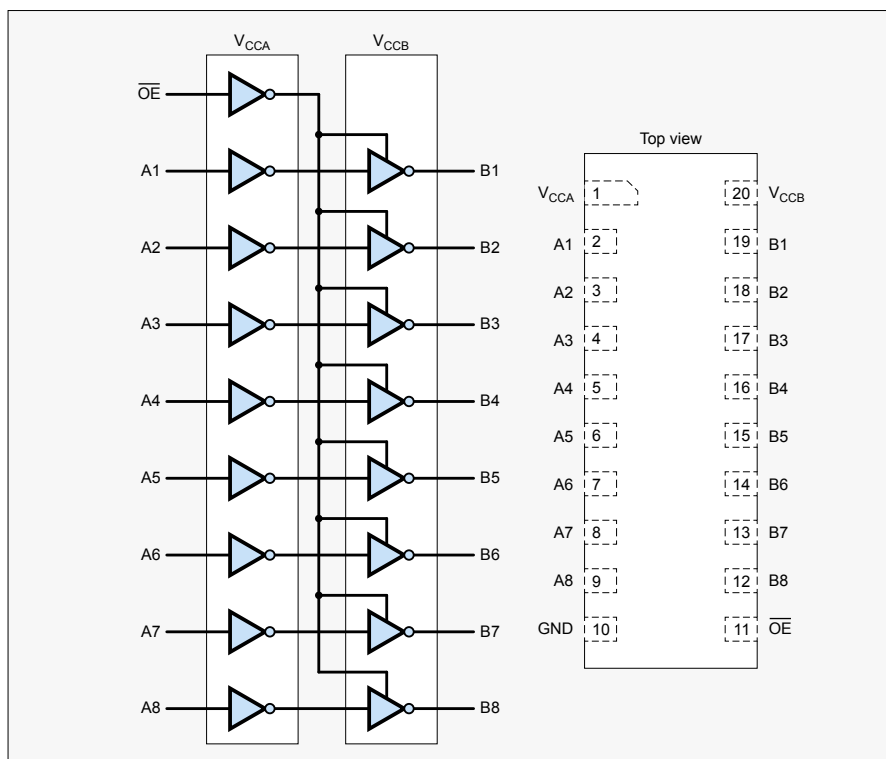
Obwód z rysunku 6 działa w następujący sposób przy przesyłaniu danych z układu

niższego napięcia do układu wyższego napięcia. Gdy układ niskiego napięcia podaje logiczną jedynkę (w tym przykładzie wynosi ona 3,3 V), napięcie bramka-źródło MOSFET-u jest równe zeru, a więc zostanie on wyłączony. W tym stanie wejście do układu wyższego napięcia jest podciągane do napięcia zasilania (w tym przypadku 5 V) przez R2, a zatem układ widzi logiczną 1 zgodnie z założeniami. Kiedy układ niskiego napięcia wprowadza logiczne 0 (0 V), napięcie bramka-źródło MOSFET jest równe niższemu napięciu zasilania (tutaj 3,3 V), więc MOSFET jest włączony, a wejście wyższego napięcia jest podciągane do 0 V przez niską rezystancję dren-źródło włączonego tranzystora. W ten sposób logiczne 0 jest widziane przez układ wyższego napięcia, zgodnie z założeniami.

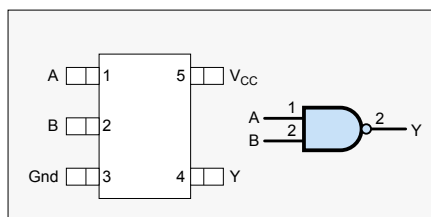
Układ na rysunku 6 działa w następujący sposób przy transmisji danych z układu o wyższym napięciu do układu o niższym napięciu. Jeśli układ wyższego napięcia podaje logiczną 1 (tutaj 5 V), to zakładając (biorąc pod uwagę, że strona wysoka nadaje), że układ niskiego napięcia nie wyprowadza aktywnie logicznego 0, MOSFET będzie wyłączony, więc wejście do układu niskiego napięcia będzie podciągnięte do napięcia zasilania (tutaj 3,3 V) przez R1, dając logiczną 1 zgodnie z założeniami. Kiedy układ wyższego napięcia podaje logiczne 0 (0 V), dioda podłoża MOSFET przewodzi (przez R1), ściągając wejście układu niskiego napięcia do około 0,7 V. W ten sposób powstaje dodatkowe napięcie bramka-źródło (w tym przypadku około 3,3–0,7=2,6 V), które jest wystarczające do włączenia MOSFET-a. Niska rezystancja dren-źródło włączonego tranzystora jest wtedy w stanie ściągnąć napięcie na wejściu układu niskiego napięcia nawet jeszcze niżej, co skutkuje logicznym 0, zgodnie z założeniami. Problem z obwodem na rysunku 6 polega na tym, że przejścia od niskiego do wysokiego poziomu logicznego mogą być powolne. Wejścia obwodów CMOS wraz z przewodami łączącymi układy stanowią obciążenie pojemnościowe dla napędzających je wyjść. Aby zmienić poziom logiczny, pojemność ta musi zostać naładowana lub rozładowana, a dzieje się to poprzez pewną rezystancję wyjściową,



Rysunek 7. Dwukierunkowy translator poziomów z przyspieszonym przetwarzaniem.



Rysunek 8. Wyprowadzenia pinów NLSV8T244 i schemat logiczny.



Rysunek 9. Przykładowe wyprowadzenia pinów i schemat logiczny elementu z rodziny LV1T – bramka NAND SN74LV1T00.

co prowadzi nas do pojęcia stałej czasowej RC, która określa jak szybko poziom logiczny może się zmienić. Rezystory podciągające na rysunku 6 są stosunkowo duże w porównaniu z efektywną rezystancją typowych wyjść bramek logicznych, co skutkuje powolnym przejściem, gdy rezystory te mają podciągnąć węzeł do poziomu logicznego 1.

Problem ten można pokonać stosując układ pokazany na **rysunku 7**. Tutaj przejście logiczne z 0 do 1 po stronie translatora otrzymującego sygnał wejściowy wyzwała krótki impuls (monostabilny, typowo rzędu kilkudziesięciu nanosekund). Impuls ten aktywuje przełącznik MOSFET, który powoduje tymczasowe zbocznikowanie rezystora podciągającego po stronie translatora wytwarzającego sygnał wyjściowy, zwiększając dostępny prąd i skracając czas potrzebny do naładowania pojemności węzła wyjściowego. Podczas używania tego obwodu translatora należy uważać, aby nie odwrócić kierunku sygnału, gdy impuls jest aktywny, w szczególności w celu wysterowania

logicznego 0, ponieważ spowoduje to powstanie sprzeczności sygnału i potencjalnie duży przepływ prądu.

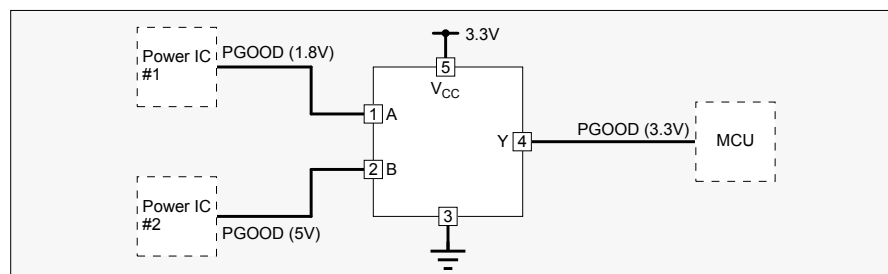
Translator jednokierunkowy NLSV8T244 z podwójnym zasilaniem

NLSV8T244 jest 8-bitowym, jednokierunkowym translatorem poziomów firmy ON Semiconductor o podwójnym zasilaniu. Schemat funkcjonalny i połączenia pinów przedstawiono na **rysunku 8**. Element posiada dwa źródła zasilania: V_{CCA} dla portu wejściowego (A), oraz V_{CCB} dla portu wyjściowego (B). Obie szyny zasilające mogą pracować w zakresie od 0,9 V do 4,5 V, co pozwala na bardzo elastyczną translację napięć logicznych. NLSV8T244 posiada wejście zezwolenia ($\overline{OE}$), które może być wykorzystane do ustawienia wyjść w stan wysokiej impedancji. Wyjścia przechodzą w stan

wysokiej impedancji również w przypadku braku zasilania na V_{CCB} . Jak wynika ze struktury układu na rysunku 8, wejście odnosi się do V_{CCB} , ale to i wszystkie inne wejścia są odporne na zbyt wysokie napięcia (w odniesieniu do zastosowanego napięcia zasilania) do maksymalnie 4,5 V. NLSV8T244 jest dostępny w kilku wariantach obudów, w tym SOIC-20 W o rozstawie pinów 1,27 mm (0,05 cala). Dostępne są różne układy typu '244' stworzone przez wielu producentów, o tej samej podstawowej strukturze obwodu i funkcji, ale różnej liczbie bitów i zakresach napięć zasilających. Podobną strukturę mają także inne jednokierunkowe translatory poziomów z podwójnym zasilaniem.

Jednokierunkowe translatory serii LV1T z zasilaniem pojedynczym

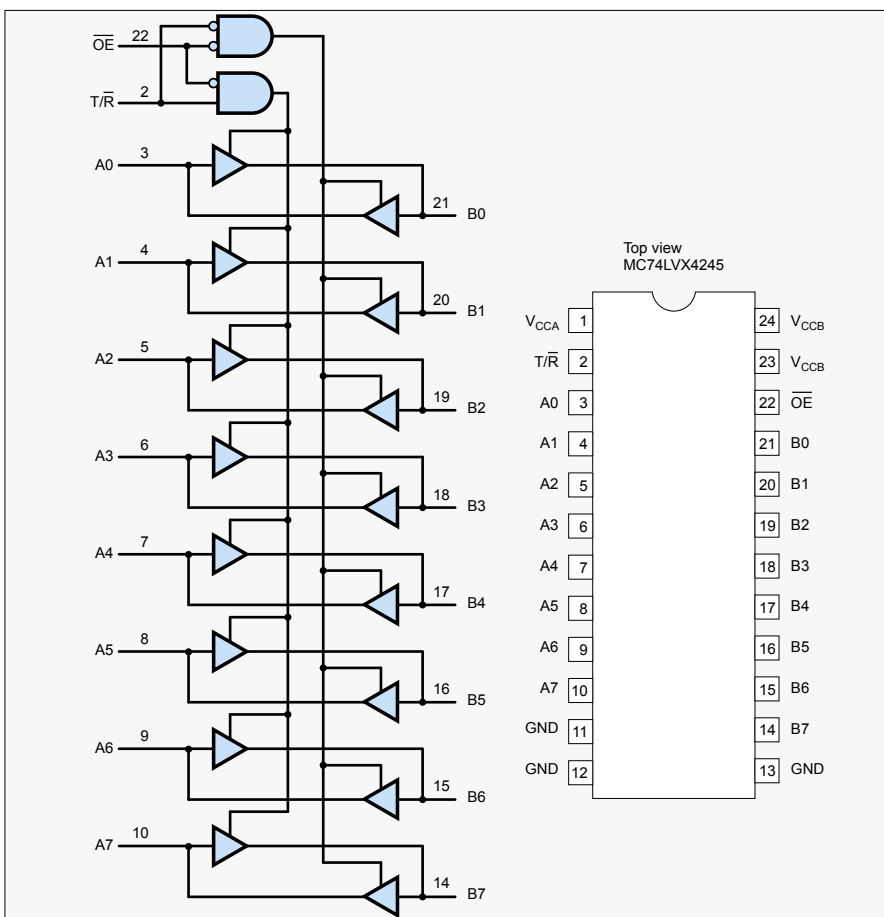
Seria LV1T układów CMOS firmy Texas Instruments stanowi przykład elementów do jednokierunkowej translacji poziomów logicznych z pojedynczym zasilaniem. Każdy układ z tej serii zapewnia pojedynczą funkcję logiczną wraz z translacją logiczną – w serii dostępnych jest dziewięć różnych bramek pojedynczych – patrz **tabela 1**. Opisy w arkuszach danych dla bramek NAND, NOR, AND i OR zawierają termin „pozytywny”, aby wskazać, że są przeznaczone do logiki dodatniej. Jak omówiono w poprzednim odcinku, oznacza to, że logiczna jedynka znajduje się na wyższym poziomie napięcia niż zero, co jest najczęściej spotykaną konwencją. Wszystkie te bramki, a także XOR, mają dwa wejścia. Możliwym problemem elementów z serii LV1T dla domowych konstruktorów jest ich mały rozmiar, choć jest to powszechne wyzwanie w przypadku układów, które są dostępne tylko w maleńkich obudowach do montażu powierzchniowego. Z punktu widzenia projektów komercyjnych, sensem produkcji takich układów scalonych z pojedynczą bramką jest oczywiście niewielka przestrzeń, jaką mogą zajmować na płytce w porównaniu ze starszymi seriami, gdzie zazwyczaj obudowa zawierała cztery bramki



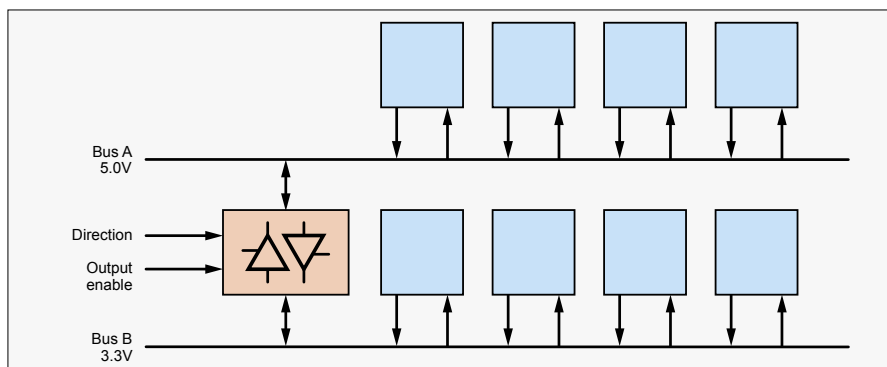
Rysunek 10. Przykładowy układ firmy Texas Instruments, w którym bramka AND serii LV1T dokonuje translacji dwóch różnych poziomów logicznych napięcia wejściowego (1,8 V i 5 V) na trzeci poziom napięcia wyjściowego (3,3 V).

Tabela 1. Elementy z serii LV1T – jednobramkowe translatory poziomów logicznych, zasilane z pojedynczego źródła.

Element	Funkcja
SN74LV1T 00	NAND
SN74LV1T 02	NOR
SN74LV1T 04	Inwerter
SN74LV1T 08	AND
SN74LV1T 32	OR
SN74LV1T 34	Bufor
SN74LV1T 86	XOR
SN74LV1T 125	Trzystanowy bufor – aktywacja: stan aktywny – poziom niski
SN74LV1T 126	Trzystanowy bufor – aktywacja: stan aktywny – poziom niski



Rysunek 11. Wyprowadzenia pinów i schemat logiczny układu MC74LVX4245.



Rysunek 12. Typowe zastosowanie MC74LVX4245 jako transceivera magistrali do połączenia dwóch magistral pracujących przy różnych napięciach.

dwuwęściowe. Seria LV1T jest dostępna w 5-pinowej obudowie SOT-23 o wymiarach zaledwie 2,9×1,6 mm lub jeszcze mniejszej SC70 o wymiarach 2,0×1,25 mm. Układ obudowy i schemat logiczny przedstawiono na **rysunku 9**. Wyjściowe poziomy logiczne układów serii LV1T są typowe dla układów CMOS i są związane z napięciem zasilania, które może być jednym ze standardowych napięć: 1,8 V, 2,5 V, 3,3 V lub 5 V. Układy te mogą wykonywać szereg translacji poziomów logicznych zarówno w górę jak i w dół.

LV1T – translacja w górę i w dół

Elementy z rodziny LV1T umożliwiają translację w górę, ponieważ są stworzone do pracy z poziomami logicznymi wejścia stanu wysokiego o wartości niższej niż typowe dla danego napięcia zasilania. Na przykład, przy napięciu zasilania 3,3 V typowa bramka CMOS miałaby poziom logiczny 1 na wejściu na poziomie co najmniej 2,3 V (około 0,7 napięcia zasilania). Jednak elementy z serii LV1T, pracując z napięciem 3,3 V, mają próg wejścia logicznego 1 na poziomie około 1,0 V. Pozwala to na translację sygnału z wyjścia logicznego 1,8 V do poziomów logicznych 3,3 V. Dostępne są następujące translacje w górę – drugie napięcie jest jednocześnie napięciem zasilania układu:

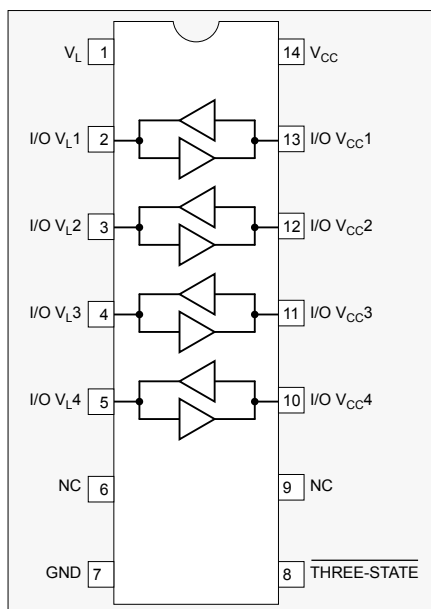
- 1,8 V do 2,5 V
- 1,8 V lub 2,5 V do 3,3 V
- 2,5 V lub 3,3 V do 5 V.

W zeszłym odcinku wspominaliśmy, że potencjalnym problemem z translacją w górę jest to, że wejście o niskim poziomie interpretowane może być jako słaby poziom logicznej jedynki lub pośredni poziom wejściowy przez układ pracujący na wyższym zasilaniu. Może to spowodować duży przepływ prądu, ponieważ w bramce mogą się włączyć zarówno tranzystory NMOS, jak i PMOS. Niski próg wyzwolenia układów LV1T pozwala uniknąć tego problemu, a pobór prądu jest na poziomie lub poniżej wartości określonej dla wszystkich wejść CMOS logicznej jedynki dla różnych translacji w górę.

Układy serii LV1T mają wejścia, które tolerują napięcie 5 V – na wejściach mogą występować napięcia wyższe od napięcia zasilania, do 5,5 V. Ułatwia to translację w dół, a dostępne konwersje poziomów są następujące, przy czym, ponownie, napięcie 'do' jest również napięciem zasilania:

- 2,5 V, 3,3 V, lub 5 V w dół do 1,8 V
- 3,3 V, do 5 V w dół do 2,5 V
- 5 V w dół do 3,3 V

Bramki dwuwęściowe w elementach serii LV1T są w stanie dokonać translacji dwóch osobnych poziomów logicznych na trzeci poziom wyjściowy (i może to być translacja w górę,

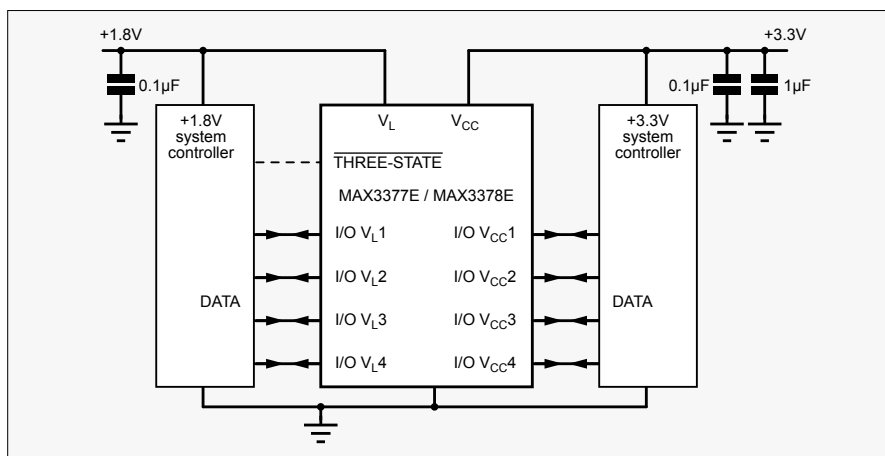


Rysunek 13. Wyprowadzenia pinów i schemat logiczny układów MAX3377E i MAX3378E (wersja w obudowie TSSOP-14).

albo w dół). Texas Instruments sugeruje, że przykładowe zastosowanie tego rozwiązania to systemy, w których mikrokontroler musi monitorować wiele układów scalonych zarządzających zasilaniem i musi wiedzieć, kiedy dwa źródła zasilania (inne niż jego własne) są włączone i gotowe do użycia. Jak wskazano w poprzednim odcinku, we współczesnych systemach powszechne jest występowanie wielu logicznych napięć zasilających. W takich sytuacjach często konieczne jest prawidłowe sekwencjonowanie ich włączania, lub aby główny procesor wiedział, kiedy wszystkie źródła są włączone przed rozpoczęciem wykonywania operacji. Przykładowy układ pokazany jest na rysunku 10.

MC74LVX4245 – dwukierunkowy translator poziomów

MC74LVX4245 jest dwukierunkowym, ośmiowejściowym translatorom poziomów z wyjściami trójstanowymi i o podwójnym zasilaniu. Schemat logiczny i wyprowadzenia pokazano na rysunku 11. Układ pochodzi z firmy ON Semiconductor i jest przeznaczony do pracy jako interfejs pomiędzy dwoma magistralami trójstanowymi pracującymi na logice 5 V i 3,3 V. Dostępne są różne układy serii „245” o tej samej podstawowej strukturze układu i spełnianej funkcji, ale różnej liczbie bitów i różnej translacji napięć. Układ MC74LVX4245 posiada dwa 8-bitowe porty wejściowe/wyjściowe trójstanowe, A i B, związane odpowiednio z zasilaniem V_{CA} (5 V) i V_{CB} (3,3 V). Wejście (transmit/receive – nadawanie/odbiór) stanowi sterowanie kierunkiem. Gdy jest w stanie wysokim,



Rysunek 14. Typowy układ aplikacyjny dla MAX3377E i MAX3378E.

dane są przesyłane z A do B; gdy jest w stanie niskim, przesyłane są z B do A. Wejście służy do włączania wyjścia (gdy jest w stanie niskim), bez względu na to, który port (A lub B) jest aktualnie wyjściem. Stan wysoki na wejściu wprowadza piny obu wejść/wyjść (I/O) A i B w stan wysokiej impedancji. Wejścia sterujące mogą pracować na poziomie niskiego lub wysokiego napięcia. ON Semiconductor zaleca włączanie zasilania V_{CCA} przed V_{CCB} , ponieważ w przypadku odwrotnej kolejności włączania mogą popłynąć wysokie prądy zasilające – więcej szczegółów w karcie katalogowej.

Typowe zastosowanie MC74LVX4245 pokazano na rysunku 12. Układ pełni rolę łącznika pomiędzy dwoma magistralami trójstanowymi pracującymi na napięciach 3,3 V i 5,0 V. Układy z jednej magistrali mogą odczytywać lub zapisywać dane z układów na drugiej. Zazwyczaj na magistrali o niższym napięciu znajdowałby się procesor wraz z układami peryferyjnymi podobnej generacji pracującymi na tym samym napięciu. Inne układy peryferyjne (być może starsze technologie) znajdowałyby się na magistrali o wyższym napięciu.

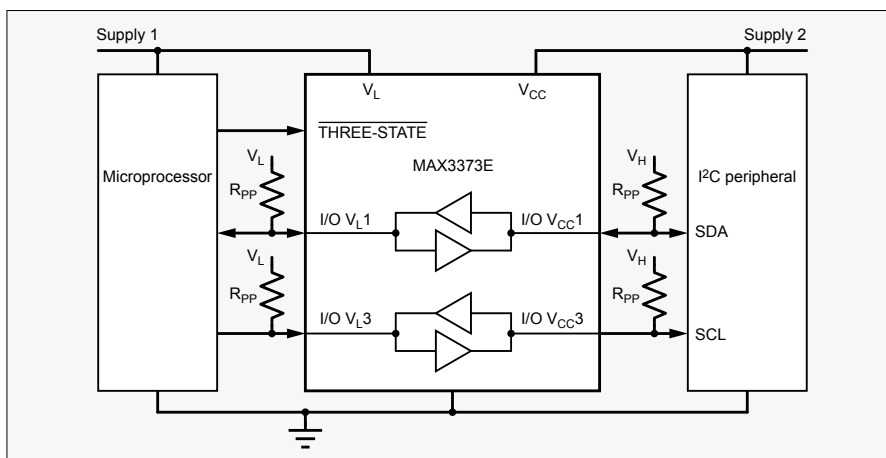
MAX3377E i MAX3378E – poczwórne, dwukierunkowe translatory poziomów z automatycznym wykrywaniem kierunku

Układy MAX3377E i MAX3378E firmy Maxim Integrated są czterowejściowymi, dwukierunkowymi translatorami poziomów z automatycznym wykrywaniem kierunku. Ich wyprowadzenia i schemat logiczny pokazano na rysunku 13. MAX3377E wykorzystuje układ podobny do rysunku 6 dla każdego z translatorów. Układ MAX3378E zawiera układy typu „one-shot” przyspieszające przełączanie układów logicznych, a każdy translator ma układ podobny do tego z rysunku 7. Dla obu układów strona niskiego napięcia (pin zasilający V_L) może mieć

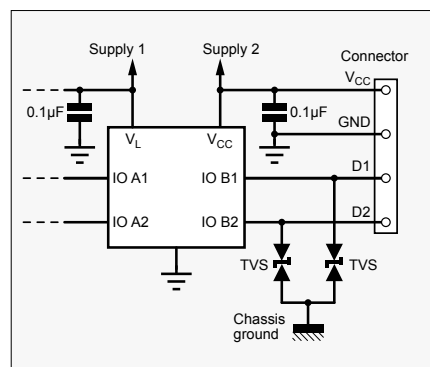
zakres od 1,2 V do 5,5 V, a strona wysokiego napięcia (pin V_{CC}) może mieć zakres od +1,65 V do +5,5 V. Istnieje wymóg, aby V_L było mniejsze od V_{CC} o co najmniej 0,3 V, ale nie dotyczy to stanów przejściowych podczas uruchamiania zasilania. Elementy te dostępne są w obudowach TSSOP, TDFN oraz w bardzo małej obudowie mikro (μ) chip-scale (UCSP), która posiada połączenia typu BGA pod obudową – patrz: https://en.wikipedia.org/wiki/Ball_grid_array. Oba układy posiadają aktywną w stanie niskim kontrolę trzeciego stanu (pin), który wprowadza wszystkie piny I/O w stan wysokiej impedancji i redukuje prąd zasilania układu do mniej niż 1 μ A. Dzieje się to poprzez odłączenie od zasilania wewnętrznych rezystorów podciągających 10 k Ω (patrz rysunki 6 i 7). Typowy układ aplikacyjny dla układów MAX3377E i MAX3378E pokazano na rysunku 14. Przykład ten pokazuje sterownik systemu 1,8 V (np. mikrokontroler) komunikujący się z systemem 3,3 V, ale inne kombinacje napięć miałyby podobną strukturę. Maxim produkuje również podwójną wersję MAX3378E, MAX3373E, która nadaje się do translacji poziomów w magistralach I²C. Przykładowy układ pokazano na rysunku 15, gdzie mikroprocesor na stosunkowo niskim napięciu zasilania (Supply 1) jest połączony z peryferyjnym układem scalonym I²C o wyższym napięciu zasilania (takim jak ADC lub DAC) poprzez układ MAX3378E. W zależności od konfiguracji magistrali, mogą być wymagane zewnętrzne rezystory podciągające (R_{pp}) (wewnętrzne mają wartość 10 k Ω) – szczegóły dotyczące wyboru wartości rezystorów znajdują się w dokumentacji firmy Maxim.

Użycie translatorów poziomów

Podobnie jak wszystkie cyfrowe układy scalone, translatory poziomów powinny posiadać kondensatory odsprężające zasilanie umieszczone blisko nich na płytce drukowanej. Układy scalone translacji poziomów o podwójnym



Rysunek 15. Translacja poziomów w magistrali I²C z wykorzystaniem układu MAX3373E.



Rysunek 16. Przykładowy układ translatora poziomów obsługującego sygnały spoza płytki poprzez złącze I/O. Zabezpieczenie stanowią diody TVS, należy również zwrócić uwagę na kondensatory odsprężające.

zasilaniu powinny być wyposażone w kondensatory odsprężające zasilanie na obu pinach zasilających. Kondensatory te zmniejszają szumy zasilania i zapobiegają powstawaniu zakłóceń, które mogą powodować nieprawidłowe działanie układów. Typowa wartość tych elementów to 0,1 μ F, ale karty katalogowe układów mogą zawierać inne zalecenia. Podobnie jak w przypadku wszystkich obwodów cyfrowych, zasilanie musi być podłączone za pomocą niskimpedancyjnych tras, zazwyczaj w postaci możliwie dużych pól miedzi znanych jako „poligony zasilania

i masy”. Jeśli translatory poziomów są używane do zapewnienia interfejsów do obwodów, które nie znajdują się na tej samej płytce, to zazwyczaj dobrym pomysłem jest zapewnienie jakiejś formy ochrony przed przejściowymi skokami napięcia. Jest to szczególnie ważne w przypadku, gdy złącza I/O będą podłączane i odłączane przez użytkowników. Ochrona jest zazwyczaj zapewniona przez dodanie diod spolaryzowanych zaporowo pomiędzy I/O a zasilaniem i masą, lub przez zastosowanie diod TVS (dwukierunkowe diody lawinowe) podłączonych pomiędzy I/O a masą (patrz

rysunek 16). Elementy zabezpieczające powinny być umieszczone blisko złącza I/O, a droga pomiędzy układem translatora a złączem I/O powinna być jak najkrótsza, aby zmniejszyć zarówno emisję, jak i podatność na zakłócenia radiowe. ■

Ian Bell

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, luty 2020 (www.epemag3.com)

REKLAMA

Sięgnij po archiwalne wydania ELEKTRONIKI dla WSZYSTKICH

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

Prześlijka
GRATIS

Zamów wygodnie na
www.UlubionyKiosk.pl

Miejsca dla specjalistów

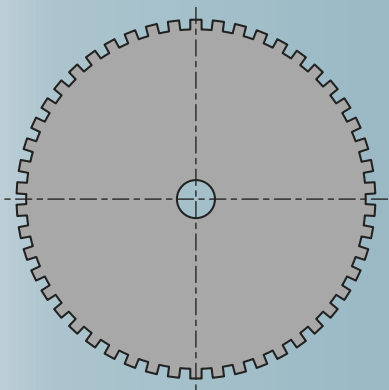
Elektronika dla wszystkich

Elektronika dla wszystkich

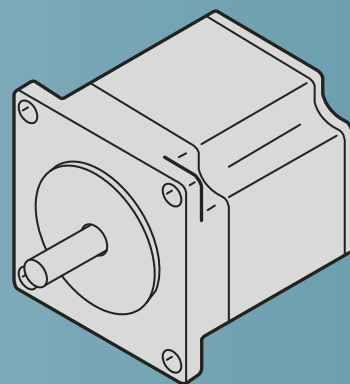
Miejsca dla specjalistów

Elektronika dla wszystkich

Elektronika dla wszystkich



Silniki krokowe w praktyce



Część 3: Podstawowe sterowniki silników krokowych

W pierwszych dwóch częściach niniejszej serii omówiliśmy różne typy silników krokowych oraz kombinacje ich uzwojeń. W części trzeciej przechodzimy do elektroniki wymaganej do sterowania silnikami krokowymi. Najpierw przypomnienie – w naszej serii omawiamy tylko sterowniki silników krokowych z magnesem trwałym oraz silników krokowych hybrydowych, ponieważ są to najpowszechniej dostępne w sprzedaży silniki, do tego w wielu odmianach. Możliwy jest zakup silników wielofazowych (np. 5-fazowych), ale są one mniej popularne, podobnie jak ich sterowniki, dlatego też skupimy się na systemach 2-fazowych bipolarnych i 4-fazowych unipolarnych.

Wydańność silnika krokowego w zakresie momentu obrotowego, prędkości i dokładności zależy w dużym stopniu od wyboru sterownika silnika krokowego, trybu pracy oraz napięcia sterowania. Wiemy już, że silników krokowych nie można wprawić w ruch wyłącznie poprzez podłączenie do baterii, jak to ma miejsce w przypadku silnika szczotkowego DC. Uzwojenia w silniku krokowym muszą być zasilane w określonej sekwencji aby stopniowo obracać wirnik. Silniki unipolarne, które mają pięć lub sześć przewodów (patrz rysunek 8 w pierwszej części naszej serii) mogą się obracać poprzez wielokrotne zasilanie każdej z czterech faz w odpowiedniej kolejności. Sterowniki silników bipolarnych, które zostaną opisane w kolejnym odcinku, mają dodatkową komplikację. Wymagają one mianowicie, aby sterownik odwrócił prąd w każdej fazie w ramach sekwencji.

Niskoobrotowe silniki szczotkowe czy silniki krokowe?

Powszechnym błędnym przekonaniem jest, iż sterownik silnika krokowego to wszystko, co jest potrzebne, aby wprawić silnik krokowy w obrót. Sterownik silnika krokowego to przede wszystkim elektronika mocy, która przełącza duże prądy pomiędzy fazami silnika wraz z logiką, która ustala sekwencję przełączania. Niezbędny jest także kontroler, który generuje (co najmniej) impulsy zegarowe taktujące

kroki obrotów silnika. Większość kontrolerów dostarcza również sygnał kierunkowy, który określa kierunek obrotu silnika. Inne, zaawansowane sygnały stanu kontrolera i interfejsy komunikacyjne dostępne w bardziej zaawansowanych sterownikach komercyjnych zostaną opisane w późniejszych artykułach z tej serii.

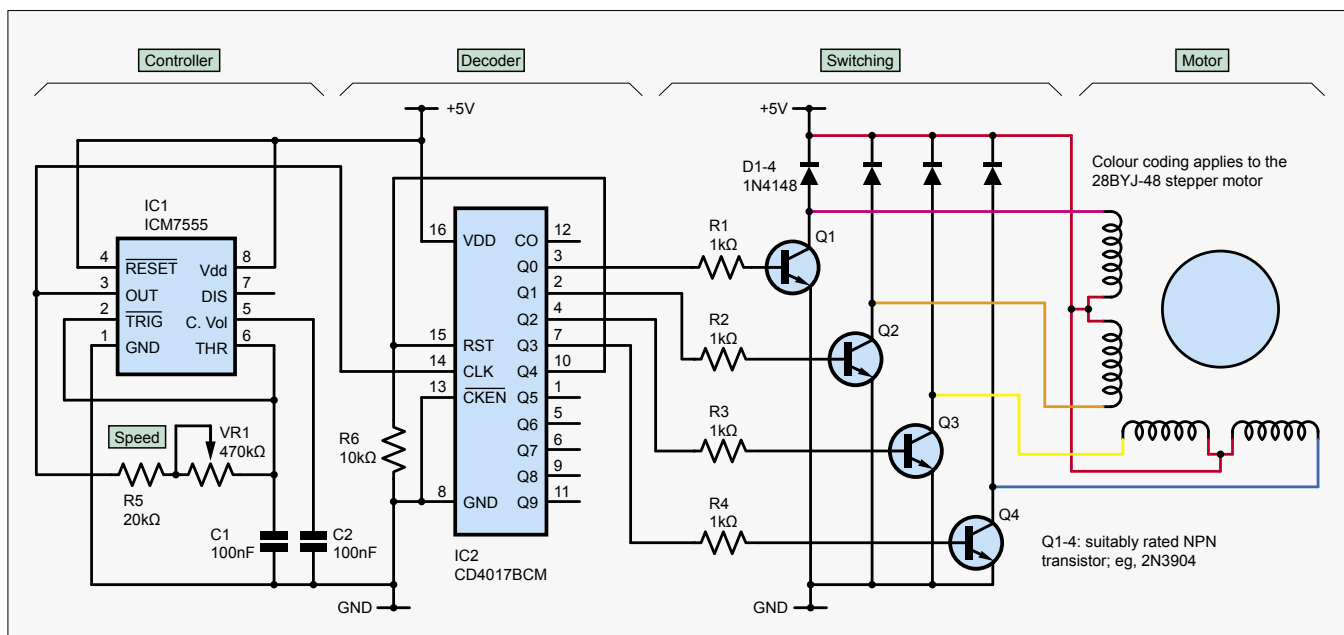
Kontrola pozycji silnika krokowego będzie zazwyczaj wymagać mikrokontrolera, ponieważ kontroler musi „liczyć” liczbę kroków od znanej pozycji wyjściowej. Wiele zastosowań silnika krokowego nie wymaga jednak kontroli pozycji – po prostu używany jest on jako silnik o zmiennej prędkości, zwłaszcza niskiej. Tani silnik DC z typowym zakresem prędkości obrotowej 2000 do 10000 obrotów na minutę potrzebuje przekładni, aby uzyskać niskie prędkości obrotowe. Można wprawdzie obniżyć napięcie zasilania silnika prądu stałego, co mogłoby dać obniżenie prędkości być może w stosunku 6:1, ale ma to niepożądaną konsekwencję w postaci zmniejszenia mocy silnika, a więc i momentu obrotowego, czterokrotnie z każdym zmniejszeniem napięcia o połowę w stosunku do napięcia nominalnego. Prędkość silnika szczotkowego DC również zmienia się w zależności od obciążenia wału silnika. W przeciwieństwie do silników DC, silniki krokowe nie mają takich ograniczeń momentu i prędkości, chociaż mają inne przypadłości, które wymagają uwzględnienia na etapie projektowania.

Wysokoobrotowe silniki szczotkowe DC są z natury głośnie. Może to być również problemem silników krokowych ze względu na wibracje powodowane przez dyskretny ruch krokowy. Wibracje mogą być zmniejszone poprzez mechaniczne tłumienie, ale zmniejszenie wielkości kroku przez techniki zwane „półkrok” i „mikrokrok” jest również bardzo skuteczne. Mikrokroki zostaną omówione w następnym odcinku, kiedy przyjrzymy się sterownikom silników bipolarnych.

Sterownik/kontroler unipolarnego, jednofazowego i jednokierunkowego silnika krokowego L/R

Zanim zagłębimy się w bardziej zaawansowane techniki sterowania, zacznijmy od być może najbardziej podstawowego obwodu sterownika/kontrolera dla unipolarnych silników krokowych. Rysunek 21 pokazuje schemat obwodu takiego sterownika, który wykorzystuje elementy, które być może już posiadasz. Jeśli nie, to są one łatwo dostępne za niewielką cenę. Obwód posiada trzy odrębne sekcje, które zostały podsumowane w tabeli 4.

Jest to stosunkowo prosty układ, więc ma pewne ograniczenia: jest jednokierunkowy (brak możliwości ruchu w drugą stronę) oraz działa tylko w trybie pełnokrokowym. Sterowniki typu L/R, zwane również „sterownikami stałonapięciowymi”, podają stałe napięcie do każdego uzwojenia, a wynikowy



Rysunek 21. Jednokierunkowy sterownik/kontroler unipolarnego, jednofazowego silnika krokowego

prąd uzwojenia określa moment obrotowy silnika. Ponieważ tylko jedno uzwojenie jest zasilane w danym momencie (jedna faza), moment obrotowy będzie niższy niż w przypadku silników bipolarnych o podobnej wielkości lub 2-fazowego sterownika unipolarnego, który zasila dwa uzwojenia jednocześnie. Mimo to, dla zastosowań wymagających niskiego momentu obrotowego może to być więcej niż wystarczające. Zrozumienie jak ten obwód działa z silnikami unipolarnymi zapewni dobrą podstawę do zrozumienia zalet sterowników silników bipolarnych w następnym odcinku.

Prąd w danym uzwojeniu fazowym zależy od napięcia (V) na uzwojeniu, indukcyjności uzwojenia (L) i rezystancji uzwojenia (R). Z prawa Ohma wynika, że R określa końcowy prąd uzwojenia (I) – stosując typowy wzór $I = V/R$. 'Końcowy' to ważny element rozważań analizy silników krokowych. Prąd końcowy nie jest osiągany natychmiast (lub nie do momentu zatrzymania kroczenia), ponieważ indukcyjność uzwojenia silnika determinuje szybkość zmian prądu. Jest to ważne; im dłuższy czas konieczny na zmianę, tym wolniejsza szybkość kroku i niższa maksymalna możliwa prędkość silnika krokowego. Gdy napięcie kroku zmienia się szybciej niż prąd, wtedy bardzo prawdopodobne jest, że silnik się zatrzyma. Niektóre sterowniki L/R mogą automatycznie zwiększać napięcie napędowe wraz ze wzrostem częstotliwości kroku, aby skompensować straty wynikające z indukcyjności. (Sterowniki bipolarne mogą również częściowo złagodzić efekt indukcyjności poprzez zastosowanie wyższego napięcia napędowego bez uszkodzenia silnika – kolejny odcinek wyjaśni tę kwestię).

Tabela 4. Główne funkcje w układzie jednokierunkowego, unipolarnego sterownika krokowego

Funkcja	Implementacja
Kontroler – generator impulsów	Timer 555 skonfigurowany jako astabilny multiwibrator, który generuje impulsy zegarowe.
Dekoder sterownika	Licznik dekadowy 4017 tworzy czterokrokową sekwencję.
Przełączanie sterownika	Sterowniki tranzystorowe przełączają zasilanie do uzwojeń.

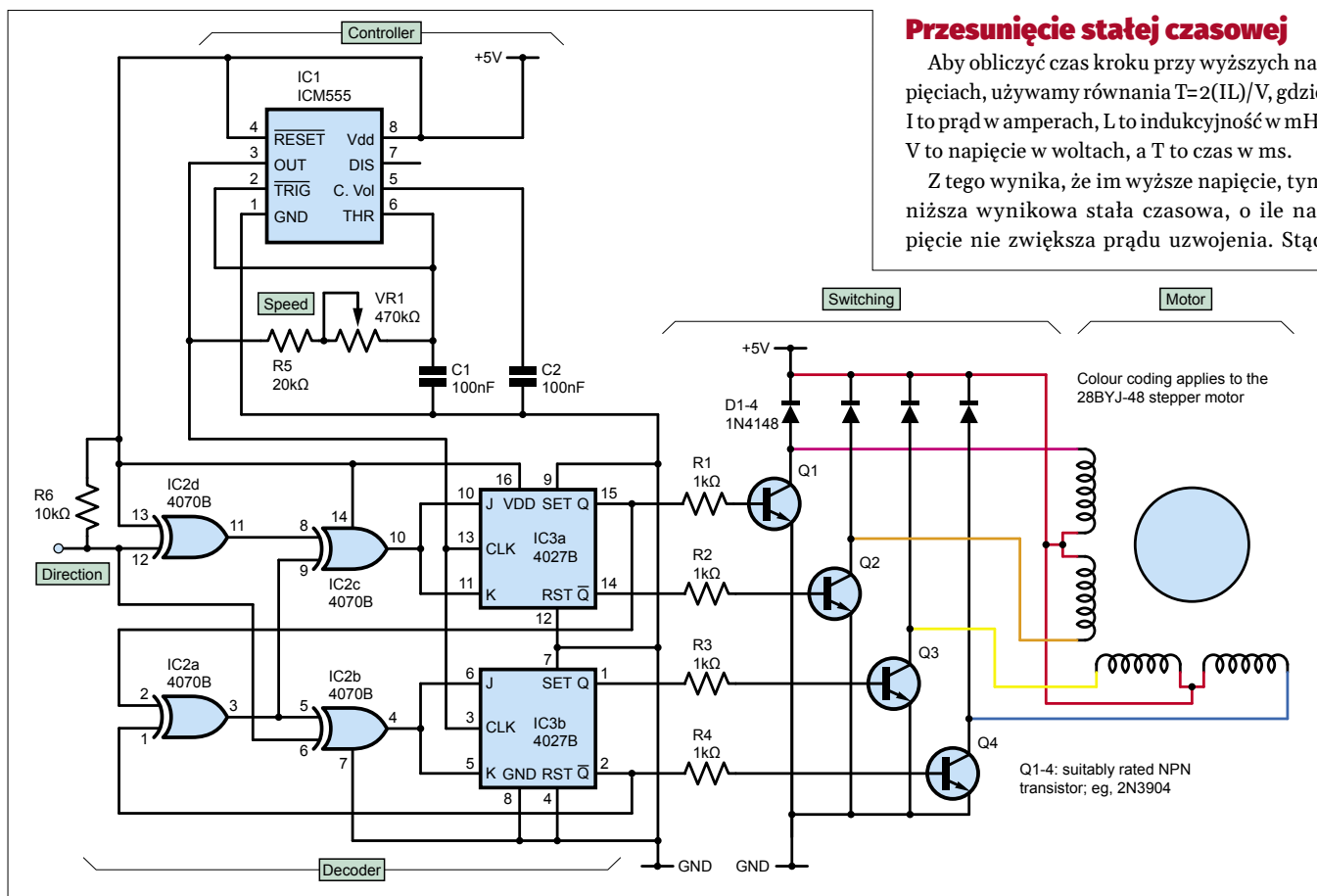
Stałe czasowe

Podobnie jak układy rezystor-kondensator posiadają stałą czasową ($T=RC$), tak obwody rezystor-cewka również wykazują ten efekt w oparciu o wzór na indukcyjną stałą czasową, $T=L/R$, gdzie T jest czasem w sekundach do osiągnięcia 2/3 maksymalnego momentu obrotowego silnika. Ponadto, dobrym przybliżeniem jest, iż czas potrzebny

do osiągnięcia przez prąd uzwojenia stanu ustalonego (maksymalnego momentu obrotowego silnika) wynosi 5T. Na przykład, jeśli nasz silnik krokowy ma indukcyjność 3 mH, a rezystancja uzwojenia wynosi 54 Ω , to $T = 3 \times 10^{-3} / 54 = 55 \mu s$. Każdy krok składa się z pojawienia się i zaniku pola magnetycznego, więc minimalny czas dla kroku jest dwukrotnie większy od tej liczby – wynosi więc 110 μs .



Rysunek 22. Przebiegi 1-4 dla wyjść od Q0 do Q3 układu 4017.



Rysunek 23. Dwukierunkowy sterownik unipolarnego, 2-fazowego silnika krokowego.

Tabela 5. Sekwencja sterowania 2-fazowego uzwojenia unipolarnego pokazująca jak pary uzwojeń są zasilane w 4-krokowej sekwencji. Ciemniejsze pole oznacza uzwojenie zasilane, jaśniejsze pole oznacza uzwojenie niezasilane.

Uzwojenie	Krok 1	Krok 2	Krok 3	Krok 4
1 Różowe (QA)				
2 Pomarańczowe (QA)				
3 Żółte (QB)				
4 Niebieskie (QB)				

„inteligentne” sterowniki unipolarne wykorzystują przebieg sinusoidalny i zmniejszają napięcie napędowe w miarę spadku częstotliwości kroku. Ten typ obwodu sterownika jest dość skomplikowany, wymaga strojenia i nie będzie dalej omawiany, gdyż sterowniki bipolarne umożliwiają osiągnięcie lepszych rezultatów przy mniejszym nakładzie pracy.

Budowa sterownika

Załóżmy, że używamy silnika krokowego, który może osiągnąć pożądaną częstotliwość kroku przy swoim napięciu nominalnym i możemy przystąpić do budowy obwodu z rysunku 21. Przykładowo w pierwszym odcinku opisano bardzo tani silnik krokowy z reduktorem typu 28BYJ-48. Jest on unipolarny więc możemy wykorzystać go jako podstawę do testowania naszego obwodu sterownika krokowego i równie dobrze mógłby on znaleźć zastosowanie w wielu praktycznych aplikacjach. Silnik 28BYJ-48 można kupić wraz ze sterownikiem na płycie drukowanej za około 1,5 funta w serwisie eBay (patrz pozycja 123922567974) z darmową przesyłką. Płyta sterownika zawiera tylko trzecią sekcję przełączania ze sterownikami tranzystorów (patrz tabela 4), nie posiada dekodera sekwencji i jest oparta na układzie ULN2003. Chociaż jest ona użyteczna, nadal



Rysunek 24. Przebiegi 2-fazowe dla dwukierunkowego sterownika unipolarnego silnika krokowego.

Tabela 6. Sekwencja kodowa dla sytuacji, gdy wejście sterujące kierunkiem jest zwarte do masy. Ciemniejsze pole oznacza uzwojenie zasilane, jaśniejsze pole oznacza uzwojenie niezasilane.

Uzwojenie	Krok 1	Krok 2	Krok 3	Krok 4
1 Różowe (QA)				
2 Pomarańczowe (QA)				
3 Żółte (QB)				
4 Niebieskie (QB)				



Rysunek 25. Przebiegi 2-fazowe odwrócone dla dwukierunkowego sterownika unipolarnego silnika krokowego. Zauważ, że ścieżki 3 i 4 są odwrócone lub przesunięte w fazie o 180° w stosunku do tych z rysunku 24, aby silnik obracał się w przeciwnym kierunku.

potrzebujemy dwóch pozostałych sekcji, więc do płytki ULN2003 wrócimy później. 28BYJ-48 wymaga około 2048 kroków/obrót wału, co jest oparte na silniku mającym 32 pełne kroki na obrót wirnika, który z kolei napędza wał wyjściowy poprzez przekładnię redukcyjną 64:1. (Uwaga, przekładnia ma w rzeczywistości 63,68:1, więc liczba kroków wynosi 2037, ale nie jest to istotne dla naszych testów).

Kontroler z timerem 555

Sterownik na rysunku 21 jest CMOS-ową wersją timera 555 skonfigurowaną jako astabilny multiwibrator wytwarzający na wyjściu sygnał prostokątny o 50% cyklu pracy. Eksperymentuj z wartością kondensatora C1 – elementem o pojemności 100 nF podłączonym do pinu 6 układu 7555, aby zmienić główny zakres częstotliwości zapewniany przez potencjometr prędkości o wartości 470 kΩ lub nastawę wstępną. Jeśli nie potrzebujesz możliwości zmiany prędkości, możesz pominąć potencjometr i umieścić zwykły rezystor między pinami 2 i 3 układu scalonego 555. Rezystor ten powinien mieć wartość co najmniej 20 kΩ. Taka wartość z kondensatorem czasowym 100 nF daje częstotliwość około 400 Hz, co daje prędkość obrotową wału około 12 obrotów na minutę.

Dekoder

Sekcja dekodera obwodu jest oparta na moim ulubionym, starym liczniku dekadowym 4017B. Ten układ będzie sekwencyjnie ustawiał stan wysoki kolejno na jednym z jego 10 pinów wyjściowych (Q0 do Q9). Ponieważ potrzebujemy tylko czterech wyjść, piąte (Q4, pin 10) jest podłączone do pinu reset, aby uczynić go licznikiem do czterech. Rysunek 22 pokazuje przebiegi dla wyjść od Q0 do Q3, gdzie widać powtarzające się poszczególne kroki. Wyjścia Q0 do Q3 są następnie podłączone do sekcji przełączania prądu składającej się z czterech tranzystorów NPN. Użyłem tranzystorów 2N3904 dla silnika 28BYJ-48, ale możesz użyć dowolnego typu tranzystora NPN pod warunkiem, że jego prąd kolektora jest odpowiednio dobrany. W układzie występują także cztery diody, których rolą jest stłumienie wszelkich skoków wysokiego napięcia, które mogą wystąpić po odłączeniu uzwojenia, spowodowanych jego indukcyjnością – są one określane jako „diody flyback”. W tym miejscu można zastąpić tranzystory za pomocą sterownika ULN2003 (patrz rysunek 26). ULN2003 ma siedem par tranzystorów Darlingtona, z których potrzebujemy tylko cztery, i zawiera już diody

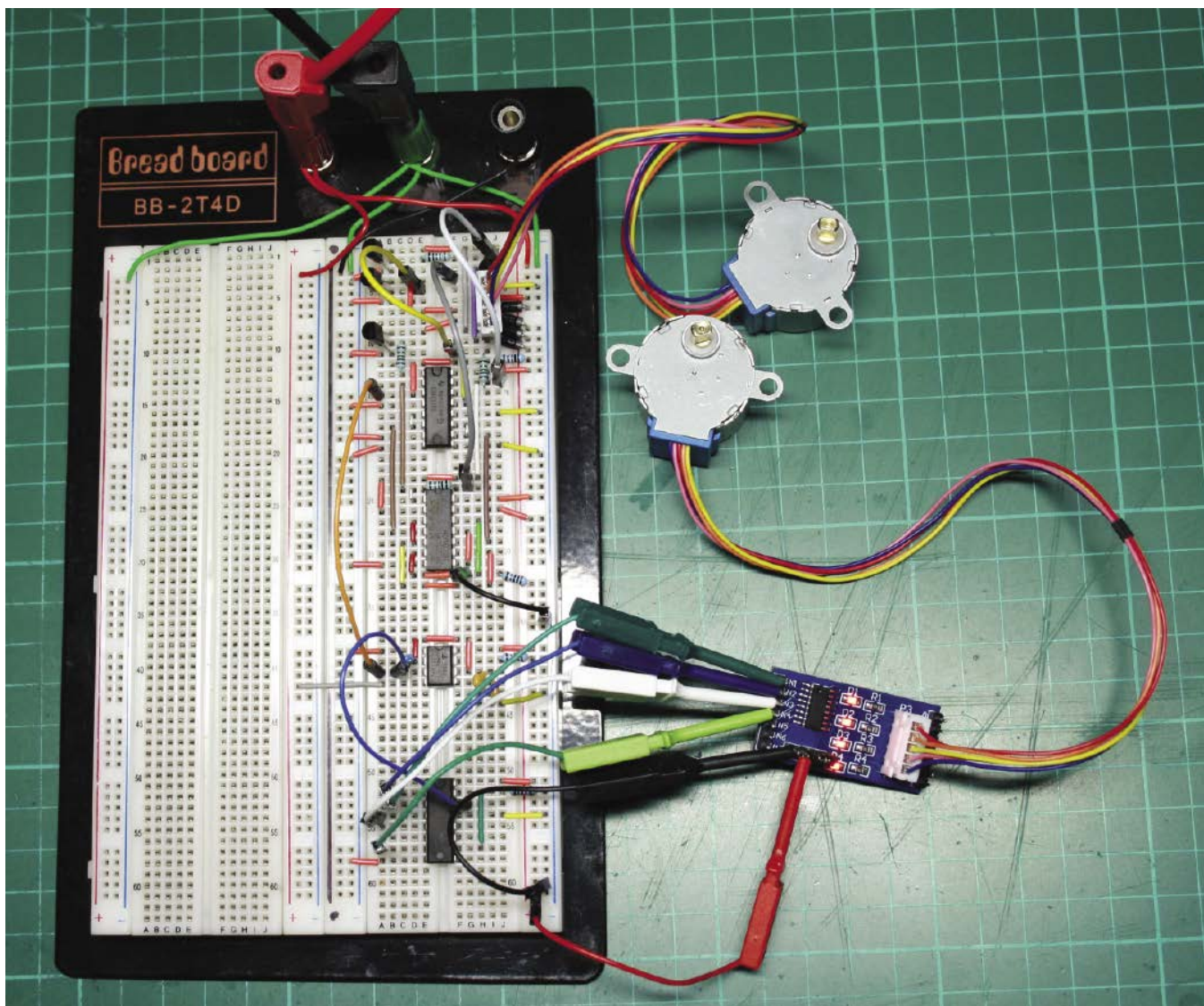
flyback. Ważne jest, aby silnik był podłączony do tranzystorów sterujących w odpowiedniej kolejności, w przeciwnym razie silnik może się nie obracać lub mieć niewystarczający moment obrotowy. W przypadku 28BYJ-48, tranzystory od Q1 do Q4 są podłączone do przewodów silnika odpowiednio: różowego, pomarańczowego, żółtego i niebieskiego, przy czym czerwony przewód idzie do plusa zasilania. Zasilanie powinno znajdować się w zakresie 5–6 V, aby uzyskać najlepsze parametry silnika. Jeśli używasz silnika krokowego z wyższym napięciem cewki, to ten układ będzie dobrze działał również dla 12 V. Taka konfiguracja jest nazywana sekwencją typu „wave” i jako że będzie napędzać unipolarny silnik krokowy, możemy ją usprawnić. Odwrócenie kierunku pracy silnika i uzyskanie większego momentu obrotowego poprzez przejście do kontroli 2-fazowej jest następnym logicznym usprawnieniem do wykonania.

Sterownik/kontroler unipolarnego, dwufazowego, dwukierunkowego silnika krokowego L/R

Jeżeli zasilimy dwa uzwojenia jednocześnie w odpowiedniej kolejności to możemy zwiększyć dostępny moment obrotowy z silnika. Jeśli równocześnie odwrócimy sekwencję, możemy sprawić, że silnik zmieni kierunek. Licznik dekadowy 4017B liczy tylko w jednym kierunku, więc potrzebujemy innego podejścia. Tabela 5 pokazuje, jak każde uzwojenie naszego unipolarnego silnika powinno być zasilane dla sterowania 2-fazowego.

Obwód z rysunku 23 jest podobny do naszego obwodu jednokierunkowego z rysunku 21, z tym że sekcja dekodera 4017B została zastąpiona dwoma nowymi układami logicznymi. Nie przypisuję sobie żadnych zasług w kwestii tego obwodu – podobne pojawiają się w wielu książkach i na stronach internetowych dotyczących silników krokowych i biorąc pod uwagę jego prostotę, działa on dość dobrze. Dekoder lub obwód translatora jest oparty na układzie poczwórnego XORA 4070B i podwójnym przerzutniku JK 4027B. Zamiast pokazanego typu CMOS można użyć równoważnych układów TTL, choć numeracja pinów w układach będzie się różnić.

Wejście sterowania kierunkiem jest po prostu zwierane do masy, aby silnik zmienił kierunek. Tabela 5 odnosi się do wyjść Q i, gdzie IC3a to „A”, a IC3b to „B”. Przyłożenie oscyloskopu do tych czterech wyjść daje przebiegi pokazane na rysunku 24. Możemy sprawdzić poprawność



Rysunek 26. Płytki prototypowa zawierająca sterowniki silników krokowych oparte na układach 4017B i 4027B. Zaczynając od dołu, 4017B jest podłączony do płytki sterownika dostarczonej z silnikami krokowymi. Powyżej 4017B znajduje się multiwibrator astabilny 7555 wytwarzający sygnał zegarowy. Górna połowa płytki prototypowej to układy 4070B i 4027B dla dwukierunkowego sterownika napędzającego tranzystory dyskretny znajdujące się na samej górze.

działania naszego układu porównując tabelę 5 z rysunkiem 24. Wybierzmy punkt początkowy sekwencji jako punkt wyzwania 'T', który widać na górze, pośrodku ekranu nad przebiegiem 1 na rysunku 24. Jest to początek kroku 1 i patrząc na kolejne przebiegi widzimy, że przebieg 1 ma stan wysoki, 2 i 3 zaś niski, a 4 ma ponownie stan wysoki, co odpowiada stanom przedstawionym w tabeli 5.

Krok 2 rozpoczyna się nieco ponad jedną podziałkę później, gdy przebiegi 3 i 4 zmieniają stan. Patrząc na kolejne przebiegi widzimy, że przebieg 1 jest nadal wysoki, przebieg 2 niski, 3 jest wysoki, a 4 jest niski, co pasuje do stanów kroku 2 z tabeli 5.

Analogicznie krok 3 i krok 4 mogą być analizowane przed rozpoczęciem od nowa całej sekwencji, tak jak przebieg 1 przechodzi w stan wysoki około 5 podziałek dalej od naszego punktu startowego „T”.

Zwarcie do masy wejścia kierunku odwraca sekwencję kodu (patrz tabela 6 a także rysunek 25). Jest to łatwiejsze do zinterpretowania, jeśli wyjścia przerzutnika JK są interpretowane w systemie szesnastkowym (traktując każde „włączone” uzwojenie jako „1”, każde „wyłączone” jako „0”, a QA jest LSB). Tabela 5 ma sekwencję numeracji 9-5-6-A, natomiast tabela 6 ma sekwencję 5-9-A-6, co pokazuje, że sekwencja kodów rzeczywiście się odwróciła.

Na rysunku 26 pokazano płytkę prototypową, która posłużyła do wykonania opisanych powyżej układów. Próba zatrzymania silników poprzez uchwycenie ich wałów pokazuje, że sterownik 2-fazowy daje większy moment obrotowy niż sterownik 1-fazowy. Jednakże, dzięki zintegrowanej przekładni, nawet sterownik jednofazowy byłby odpowiedni do wielu zastosowań, w których obciążenie wału nie jest zbyt duże.

Wykorzystanie mikrokontrolera

Wielu czytelników korzysta z mikrokontrolerów, takich jak na przykład Arduino. Można je wykorzystać do tworzenia impulsów krokowych i sekwencji zastępując logikę dyskretną, ale nadal niezbędne będą opisane wcześniej sterowniki tranzystorowe (takie jak ULN2003) do przełączania uzwojeń.

W kolejnym odcinku

W części czwartej przejdziemy do ważnego tematu, jakim jest konstrukcja sterowników bipolarnych silników krokowych, wraz z omówieniem elementów dostępnych komercyjnie. ■

Paul Cooper

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

Patronat EdW nad szkołami i uczelnianymi Kołami Naukowymi rozkwiata i daje redakcji EdW impulsy zachęcające do wspierania edukacji szkolnej i uczelnianej. Działa sprzężenie zwrotne. Dostajemy mnóstwo listów od uczniów, nauczycieli i studentów. Dla nich jest ta rubryka.



A. Historia, czyli lektura lekka na rozgrzewkę

To jakiś cud. Gdy siadłem do pisania tego tekstu w przeddzień Wigilii i napisałem datę przypisywaną wynalazkowi tranzystora (23.12.1947 r.), nagle uświadomiłem sobie, że stało się to równo 75 lat temu. Był wieczór. Jacyś kretyni odpalali pierwsze fajerwerki. (Kretyni, bo noc sylwestrową spędzam zwykle w szafie z moim psem umierającym ze strachu.) Nie były to fajerwerki dla uczczenia jubileuszu 75-lecia epokowego wynalazku tranzystora. Nie zauważyłem, żeby świat pamiętał o tej dacie. A wyobraźmy sobie na chwilę, jak wyglądałby świat, gdyby nie wynaleziono tranzystora. Trochę jak świat Indian, którzy nie wynaleźli koła. Nie byłoby komputerów osobistych, telefonów komórkowych (nie mówiąc o smartfonach), GPS, internetu, aparatów słuchowych, itd.

Tranzystor to największy wynalazek XX wieku, który zapoczątkował nową epokę w rozwoju cywilizacji. Pomyślałem więc, że jubileusz 75-lecia obliguje mnie do skupienia się w historycznej części tego wykładu na opisanu narodzin wynalazku tranzystora, na przedstawieniu nie tylko suchych faktów, ale też emocji, ambicji, bólu porażek, euforii sukcesów, czyli żywych ludzi, którzy dokonali tego wspaniałego aktu twórczego. Sam odczuwam dreszcz emocji, gdy o tym piszę. Przecież ja miałem wtedy 4 lata. To wszystko stało się za mojego życia. Niewiarygodne. Pozwólcie, że na chwilę wzruszę się wspomnieniami. Nie pamiętam co robiłem 23 grudnia 1947 roku, ale 4 lata później zrobiłem mój pierwszy odbiornik kryształkowy. Chodziłem do 3-ej klasy i zaczęła się botanika, której nie cierpiełem, ale w pracowni przyrodniczej leżał kawałek galeny, który był mi potrzebny dla zrobienia detektora. Nie było wtedy Wikipedii, w której mógłbym wyczytać, że „galena jest minerałem pospolitym i szeroko rozpowszechnionym”. Gdybym to wiedział, to szukałbym w polu, ale nie wiedziałem i okrucieństwo galeny przemieścił się ze szkoły do mojego domu. Radio zagrało, a ja – słowo daję, – nie wyrosłem na złodziejaska. Jestem pewien, że Czytelnicy EdW doskonale wiedzą, jak silne emocje towarzyszą pasji do tworzenia czegoś, co ma zagrać, zadziałać. Zatem nietrudno będzie nam wczuć się w atmosferę towarzyszącą zespołowi Bell Labs w dniach narodzin tranzystora. Łatwo zrozumiemy W. Brattaina, który Wigilię 1947 r. spędził w laboratorium ciesząc się jak dziecko, że jego konstrukcja na drucikach i sprężynie zaczęła generować sygnał. Dojdziemy do tego, ale zaczniemy od prehistorii.

Prehistoria

Zastanawiam się jak daleko sięgnąć w zakamarki historii, bo wynalazek tranzystora nie wziął się przecież z niczego. Droga do każdego odkrycia wiedzie przez lata, a nawet stulecia nawarstwiających się obserwacji, doświadczeń, pomysłów, których kulminacją jest „wielkie odkrycie” lub „przełomowy wynalazek”. Sięgnijmy do początku XX wieku, czyli do początku dwóch ścieżek, które doprowadziły do wynalazku tranzystora. Jedna ścieżka to lampy elektronowe, a druga – detektor kryształkowy. W 1904 roku J.A. Fleming zbudował pierwszą lampę elektronową – diodę, a w 1906 r. Lee De Forest wprowadził do diody próżniowej siatkę i uzyskał pierwszą lampę wzmacniającą – triodę. Na początku lat 1900-tych pojawiły się też odbiorniki kryształkowe, w których głównym elementem

był detektor, czyli dioda wykorzystująca zjawisko jednokierunkowego przepływu prądu na styku metalu z półprzewodnikiem (igły metalowej z kryształem galeny). Dodajmy, że zjawisko niesymetrii przepływu prądu przez styk metalu z półprzewodnikiem (m-s) w zależności od kierunku polaryzacji, odkrył F. Braun w 1874 roku. W kolejnych latach powstała teoria działania złącza metal – półprzewodnik (W. H. Schottky) oraz rozwinęła się masowa produkcja germanowych diod ostrzowych. Już w latach dwudziestych rosyjski wynalazca Oleg Łosiew obserwował na styku m-s efekty świecenia (protoplasta diod świecących) oraz charakterystykę prądowo-napięciową z ujemnym zboczem (protoplasta diody tunelowej, wynaleziona w 1957 r. przez Leo Esaki, za co japoński fizyk otrzymał Nagrodę Nobla). Zatem ścieżka zapoczątkowana detektorem kryształkowym doprowadziła nas w latach czterdziestych do dojrzałej teorii złącza m-s i szerokich zastosowań diod ostrzowych, nie tylko jako elementów prostowniczych, ale również w wielu innych aplikacjach, na przykład w mieszaczach stacji radiolokacyjnych. Natomiast sam detektor kryształkowy przepadł w wyniku rozwoju lamp elektronowych. Zastąpiła go dioda próżniowa, która wraz ze wzmacniaczem na triodach elektronowych pozwoliła skonstruować odbiornik radiowy nieporównywalnie lepszy od radia kryształkowego. Jednocześnie fizycy zajmujący się złączem m-s spełniającym funkcję diody poddawani byli silnej pokusie, żeby powtórzyć wyczyn Lee De Foresta, który do diody próżniowej wprowadził siatkę i uzyskał triodę, czyli element wzmacniający. Gdyby się udało dla diody m-s wymyślić odpowiednik siatki w lampie, to mielibyśmy wzmacniacz na bazie ciała stałego, wolny od wielu niedostatków lamp elektronowych. Próbowano na różne sposoby wprowadzać siatkę między metal i półprzewodnik. Aż znalazł się fizyk, który odszedł od koncepcji dosłownie narzucającej się z triody próżniowej i wymyślił półprzewodnikowy element wzmacniający z elektrodą sterującą prądem, która spełniała rolę siatki w lampie, ale konstrukcyjnie nic wspólnego z konstrukcją triody próżniowej nie miała. Ten element to protoplasta współczesnych tranzystorów polowych z izolowaną bramką, a jego wynalazca nazywał się Julius Edgar Lilienfeld (patent w Kanadzie 1925 r. i w USA 1926 r.) – **rysunek 1**. Opatentowanego elementu Lilienfeld nie wykonał praktycznie. W świetle obecnej wiedzy jest pewne, że nie mógł zadziałać. Patent spoczął na półkach urzędów patentowych na dwadzieścia lat. Zainteresował się nim dopiero po wojnie William Shockley, a mimo że był geniuszem i wiedział niemal wszystko, to jednak nie wiedział, że ten patent nie może zadziałać, nie dlatego, że idea jest błędna, ale dlatego, że technologia musi dojrzeć do realizacji tej idei. Tego W. Shockley nie wiedział i czekało go bolesne rozczarowanie.

William Shockley i jego drużyna

Bell Laboratories to jedyny takiej klasy ośrodek naukowy na świecie – wylęgarnia patentów technologicznych, która szczyty się 15-ma laureatami Nagrody Nobla, przyznanej za 9 przełomowych, wręcz rewolucyjnych technologii. W szczytowym okresie świetności w Bell Labs zatrudniano blisko 6000 pracowników, w tym 2000 z tytułem naukowym doktora. Fantastyczne sukcesy były możliwe przede wszystkim dzięki agresywnej i skutecznej rekrutacji młodych, świetnie zapowiadających się naukowców. W 1936 roku w Bell Labs podjęto decyzję o powołaniu zespołu badawczego, którego celem jest zastąpienie lampy elektronowej elementem na ciele stałym, a ściślej półprzewodnikowym. 26-letni, wybitnie uzdolniony fizyk William Shockley bez wahania przyjął propozycję pracy w tym zespole. We współpracy ze starszym o 10 lat Walterem Brattainem przez trzy lata W. Shockley sprawdził eksperymentalnie wiele pomysłów na tranzystor polowy, podobny do patentu J.E.

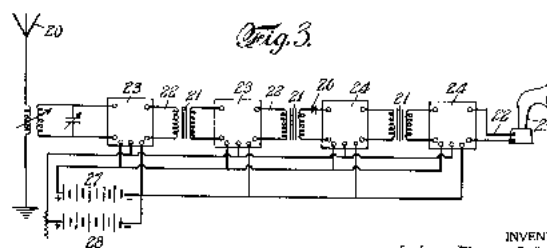
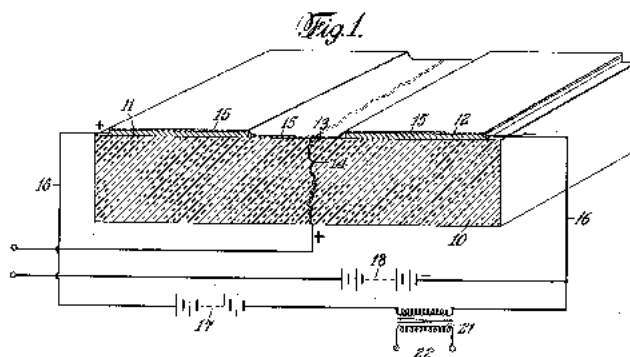
Jan. 28, 1930.

J. E. LILIENTELD

1,745,175

METHOD AND APPARATUS FOR CONTROLLING ELECTRIC CURRENTS

Filed Oct. 8, 1926



INVENTOR
Julius Edgar Lilienfeld
BY *Frederick H. ...*
ATTORNEY

Rysunek 1. Patent J.E. Lilienfelda (1925 r. – Kanada, 1926 r. – USA), protoplasta współczesnych tranzystorów polowych z izolowaną bramką



Julius Edgar Lilienfeld – twórca patentu pierwszego tranzystora polowego

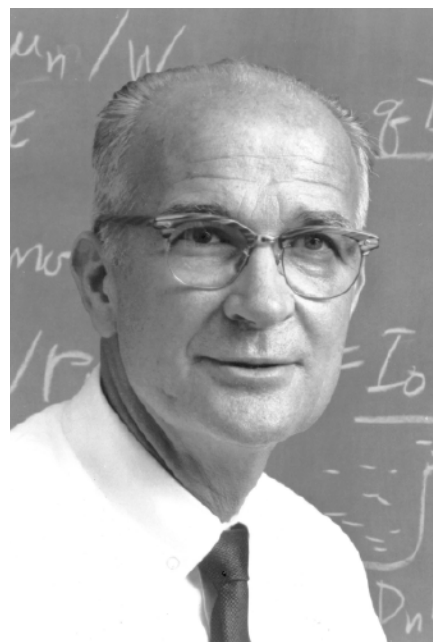
Lilienfelda, z zastosowaniem tlenku miedzi jako półprzewodnika. Wszystkie próbne prototypy takiego tranzystora nie dały pozytywnych rezultatów. W tych latach W. Shockley opublikował kilka fundamentalnych prac na temat półprzewodników, ale nie przybliżył się do osiągnięcia celu, tj. opracowania elementu półprzewodnikowego zdolnego zamienić lampę elektronową. Wybuch wojny przerwał ten kierunek prac i W. Shockley został skierowany do zadań związanych ze zwalczaniem niemieckich łodzi podwodnych i w dziedzinie radiolokacji. Nie będzie przesady w twierdzeniu, że wniósł istotny wkład w zwycięstwo Ameryki. Na zlecenie rządu USA opracował raport, w którym wyliczył koszty inwazji na Wyspy Japońskie – 5 do 10 mln zabitych Japończyków i 1,7 do 4 mln strat osobowych armii amerykańskiej, w tym 400.000 do 800.000 zabitych. Ten raport ostatecznie przesądził o użyciu bomby atomowej.

W 1945 roku, zaraz po zakończeniu wojny, w Bell Labs powrócono do projektu opracowania półprzewodnikowego odpowiednika lampy elektronowej i zaproszono ponownie W. Shockleya, tym razem jako szefa zespołu badawczego (wraz z chemikiem Stanleyem Morganem). Poza W. Shockleyem

w zespole znaleźli się fizycy John Bardeen, Walter Brattain, Gerald Pearson, chemik Robert Gibney, elektronik Hilbert Moore oraz kilku pracowników inżynierskich i technicznych. Już na samym początku W. Shockley podjął bardzo szczęśliwą decyzję, że ograniczają pole eksperymentów do dwóch materiałów półprzewodnikowych – germanu i krzemu. Oczywiście, przyjęto założenie, że celem jest opracowanie tranzystora polowego, działającego według idei z patentu J.E. Lilienfelda. Ta idea sprowadza się do sterowania przewodnością przypowierzchniowej warstwy półprzewodnika przez oddziaływanie polem elektrycznym prostopadłym do tej powierzchni. Początkowo atmosfera w zespole była bardzo dobra. Niemal codziennie zbierano się dla przedyskutowania bieżących postępów prac. Jednak stopniowo narastała frustracja i rozczarowanie – oczekiwanych rezultatów nie było. Pole elektryczne działające prostopadle do powierzchni półprzewodnika nie powodowało żadnych zmian przewodności warstwy przypowierzchniowej półprzewodnika. W 1946 r. fizyk teoretyk John Bardeen (dwukrotny laureat Nagrody Nobla, poza nagrodą za tranzystor w 1956 r. otrzymał też Nagrodę Nobla za teorię nadprzewodnictwa w 1972 r.) wysunął hipotezę, że elektrony przyciągane przez pole elektryczne z głębi półprzewodnika do jego powierzchni nie biorą udziału w przewodzeniu prądu, bo są uwięzione w pułapkach, które nazwał stanami powierzchniowymi. Istotnie, powierzchnia kryształu półprzewodnika jest obszarem „brutalnego” urwania ciągłości sieci krystalicznej i nieciągłości strukturalnej na styku półprzewodnika z warstwą izolatora nałożoną na półprzewodnik. Jest to więc miejsce z dużą liczbą defektowych stanów energetycznych, stanowiących swego rodzaju „dołki”, do których wpadają elektrony. John Bardeen całkowicie pogrążył się w pracach nad nowym, niejako pobocznym kierunkiem badań i zaczął tworzyć teorię stanów powierzchniowych. Jego pionierskie prace zapoczątkowały rozwój nowej dziedziny nauki – fizyki powierzchni ciała stałego. W. Shockley uznał, że zespół powinien zaniechać wykonywania trochę bezładnych eksperymentów z kolejnymi prototypami półprzewodnikowego elementu wzmacniającego i wesprzeć prace J. Bardeena nad stanami powierzchniowymi. W. Shockley rozumował logicznie, że dopóki nie zostanie usunięta przeszkoda, jaką okazały się stany powierzchniowe, dopóty nie ma szans na osiągnięcie zasadniczego celu i wykonanie zadania postawionego przed zespołem. Przy wiodącej roli fizyka eksperymentatora W. Brattaina przystąpiono do realizacji serii eksperymentów służących badaniom stanów powierzchniowych. Sam W. Shockley uznając wiodącą rolę J. Bardeena w badaniach nad stanami powierzchniowymi, coraz mniej uczestniczył w bieżących pracach zespołu, oddając się pracy „dla siebie”, tj. tworzeniu teorii półprzewodników, która niebawem miała mu zapewnić bezsprzeczną pozycję światowego guru w tej dziedzinie.



Wynalazcy tranzystora: John Bardeen (z lewej), William Shockley (w środku), Walter Brattain (z prawej)



William Shockley

Tranzystor ostrzowy

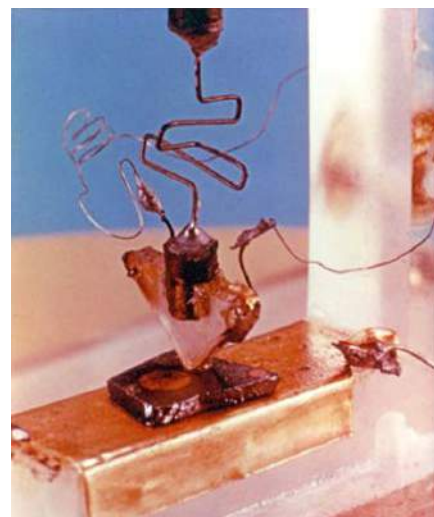
W toku wielu eksperymentów prowadzących do opanowania problemu stanów powierzchniowych zdarzyło się, że na płytce germanowej wytworzono warstwę dielektryczną tlenku germanu, by przez ten dielektryk oddziaływać polem elektrycznym na półprzewodnik. Po przyłożeniu dwóch elektrod metalowych okazało się, że płynie przez nie prąd i obserwuje się nieznaczne wzmocnienie. Ta przypadkowa obserwacja skierowała uwagę eksperymentatorów na zupełnie nowy trop prowadzący do zbudowania elementu wzmacniającego, działającego na zupełnie innej zasadzie, niż zakładano w projekcie pracy badawczej. Zrezygnowano z wykonywania warstwy tlenku germanu, która i tak nie spełniała roli izolatora (może została przypadkowo „zmyta”) i wykonano konstrukcję

jak na fotografii (rysunek 2), którą przekrojowo pokazuje rysunek 3. Rolę elektrod ostrzowych spełniała złota folia nałożona na klin plastikowy i przecięta żyłką na jego czubku. W ten sposób uzyskano względnie mały odstęp między elektrodami (ok. 100 μm). Uzyskano wyraźne wzmocnienie sygnału, dochodzące do 100 razy. Stało się to 16 grudnia 1947 r. Historyczna demonstracja wzmacniacza półprzewodnikowego odbyła się 23 grudnia 1947 r. Na rysunku 4 widzimy notatkę W. Brattaina z tego pokazu. Nazwę **tranzystor** wymyślił członek zespołu inżynier John Pierce; jak sam wyjaśniał przez analogię z lampą, dla której najważniejszym parametrem jest transkonduktancja, w tranzystorze właściwe będzie pojęcie transresistance i brzmienie powinno być podobne do słów thermistor, varistor. 13.02.1948 inny członek zespołu Shockleya John N. Shive zbudował tranzystor ostrzowy z kontaktami z brązu na przeciwnych ścianach płytki z germanu, dowodząc tym samym, że dziury mogą dyfundować poprzez płytkę Ge, a nie tylko wzdłuż powierzchni płytki, jak sądzono wcześniej. Ten eksperyment umocnił W. Shockleya w słuszności jego idei tranzystora złączowego (nazywanego wówczas również tranzystorem warstwowym), który w zmodyfikowanych wariantach technologicznych jest do dzisiaj produkowany i szeroko stosowany jako tranzystor bipolarny. Był to też pierwszy rodzaj tranzystora, którego sposób działania został w pełni wyjaśniony i opisany w ramach teorii Shockleya wyłożonej w jego słynnej opasłej książce (558 stron) *Electrons and Holes in Semiconductors*, wydanej w roku 1950. Shockley zazdrośnie nie włączał do prac nad tranzystorem złączowym żadnego z członków swojego zespołu. Jak sam określił, w jego zespole panowała atmosfera „współpracy i rywalizacji”. Z czasem pozostała tylko rywalizacja. W wynalazku tranzystora ostrzowego W. Shockley bezpośrednio nie uczestniczył, co zresztą bardzo ciężko przeżywał. Trzeba przyznać, że W. Shockley nie miał zwyczaju przypisywać sobie współautorstwa prac, w których nie uczestniczył. Tym razem jednak był przekonany, że powinien opatentować ten wynalazek jednoosobowo, bo on kierował zespołem i formułował zadania dla pozostałych. Wstrząśnięty tym W. Brattain powiedział: „Oh, hell, Shockley, there's enough glory in this for everybody”.

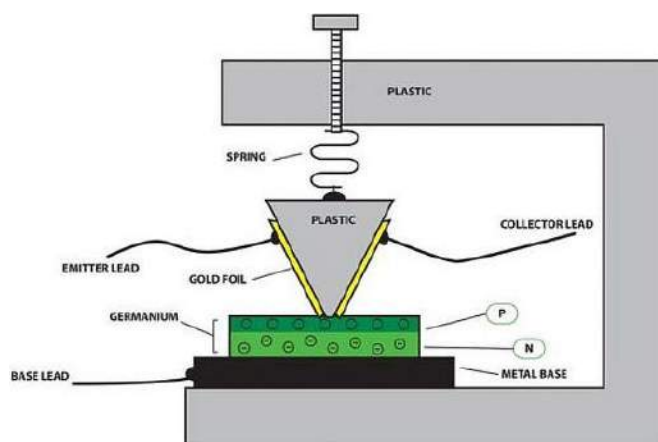
Ostatecznie patent na tranzystor ostrzowy został zgłoszony przez dwóch autorów – J. Bardeena i W. Brattaina (rysunek 5). Ze względu na najwyższą rangę wynalazku tranzystora kierownictwo Bell Labs zdecydowało, że oficjalnie jest trzech wynalazców – W. Shockley, W. Brattain, J. Bardeen i nakazało, żeby zawsze występowali w trójkę na wszystkich fotografiach. W wywiadach prasowych w imieniu całej trójki występował W. Shockley, co bardzo źle znosili W. Brattain i J. Bardeen. Tym bardziej, że z czasem, ze względu na rosnącą popularność W. Shockleya jako wynalazcy tranzystora złączowego, twórcy teorii elementów półprzewodnikowych i autora wiekopomnej książki, wszystkie zasługi, w tym również wynalazek pierwszego tranzystora, zaczęto przypisywać W. Shockleyowi, a nazwiska W. Brattaina, J. Bardeen schodziły w cień. W poczuciu doznanej krzywdy i nieprzyjaźni do W. Shockleya obaj odeszli z jego zespołu w roku 1951, przy czym J. Bardeen odszedł całkiem z Bell Labs. Pokłócenie z W. Shockleyem spotkali się z nim po kilku latach podczas wręczania Nagrody Nobla w 1956 r.

Epilog

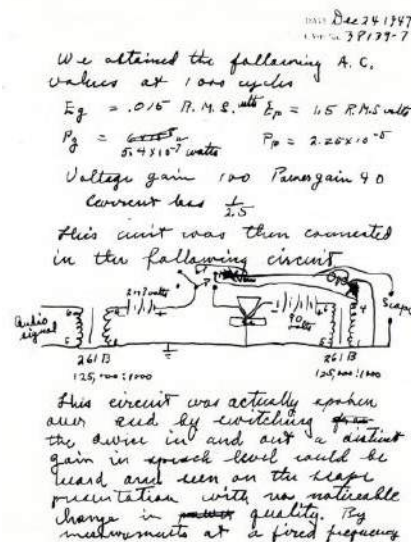
Ostatecznie okazało się, że tranzystor ostrzowy był ślepą uliczką. Spory na temat niejasnej teorii jego działania stały się bezprzedmiotowe po kilku latach, gdy rynek zdobyły bez reszty tranzystory złączowe, według koncepcji i patentu W. Shockleya, których działanie zrozumiałe i przekonująco wyjaśniała teoria Shockleya. Z biegiem lat doskonalono technologię, już w 1954 r. pojawiły się tranzystory krzemowe, które wyparły z rynku tranzystory germanowe, ale tranzystory złączowe, znane obecnie jako bipolarne, niezmiennie działają na zasadach opisanych ponad 70 lat temu przez W. Shockleya. A dlaczego zmieniono im nazwę na bipolarne? Bo pojawiły się unipolarne, najpierw złączowe, a później tranzystor MOS. Ale to już odrębna historia. A jak rozwijała się dalej kariera naukowa W. Shockleya? W powojennym



Rysunek 2. Pierwszy tranzystor ostrzowy, zademonstrowany 23 grudnia 1947 r. przez W. Brattaina w Bell Labs. © 2006–2007 Alcatel-Lucent. All rights reserved (za zgodą Computer History Museum)



Rysunek 3. Przekrój fizyczny konstrukcji pierwszego tranzystora ostrzowego (za zgodą Computer History Museum)

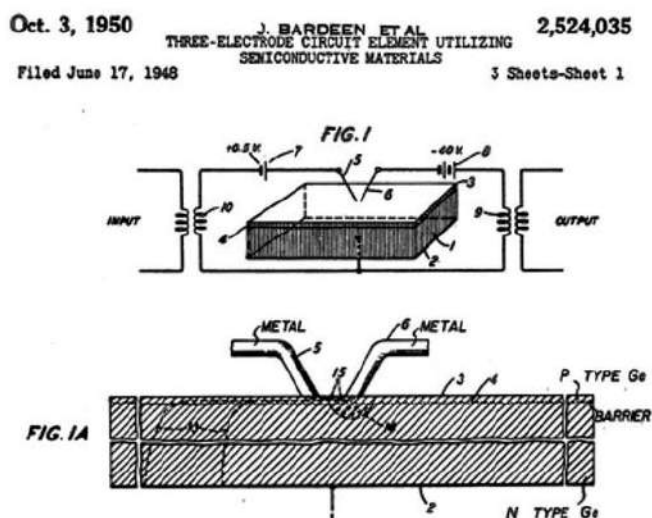


Rysunek 4. Notatka W. Brattaina z pokazu działającego tranzystora w dniu 23.12.1947 r. Courtesy Lucent Technologies 1997 (za zgodą Computer History Museum)



„Zdradziecka ósemka” – założyciele firmy Fairchild Semiconductor, którzy odeszli z firmy Shockley Semiconductor Laboratory (za zgodą Computer History Museum)

Życiorys W. Shockleya można wyodrębnić cztery okresy. Okres pierwszy to 10 lat pracy w Bell Labs (lata 1945–1955). Były to lata najbardziej płodne naukowo, które wyniosły W. Shockleya na szczyty sławy. Oczywiście, największym jego osiągnięciem był wynalazek tranzystora złączowego. W niektórych źródłach można przeczytać opowieść o tym, że W. Shockley dokonał tego wynalazku niejako „na złość” twórcom tranzystora ostrzowego – W. Brattainowi i J. Bardeenowi. Faktem jest, że był niepokieszony, gdy dowiedział się o przypadkowym wynalazku o niezwykle doniosłości, w jego laboratorium, a jego przy tym nie było. Od pokazu tego wynalazku (23.12.1947 r.) upłynęło zaledwie kilka tygodni, gdy zespół Shockleya zebrał się na naradę, by ustalić dalsze etapy badań nad doskonaleniem tranzystora ostrzowego. Było to 23.01.1948 r. W. Shockley w pewnym momencie podszedł do tablicy i osłupiałym kolegom przedstawił szczegółowo koncepcję całkiem innego tranzystora, właśnie tranzystora złączowego. To się nazywa „wywrócić stolik”. Nie mógł tego wymyślić w ciągu kilku tygodni. Od dłuższego czasu niezbyt się interesował eksperymentami kolegów w laboratorium, bo pracował nad teorią złącza p-n, z której zrodziła się koncepcja tranzystora złączowego. Złączem p-n był zafascynowany już od 1946 roku, gdy po przyznaniu patentu ujawniono dotychczas utajnione odkrycie w Bell Labs złącza p-n jeszcze w 1940 roku. Zgłoszenie patentowe W. Shockleya na wynalazek tranzystora złączowego nastąpiło 26.06.1948 r. Kolejne lata pracy W. Shockleya były niezwykle płodne. W sumie miał on około 90 patentów, z których większość dotyczyła półprzewodników i była dokonana podczas jego pracy w Bell Labs. Poza tranzystorem złączowym do największych jego osiągnięć należą patenty na tranzystor polowy złączowy (zgłoszenie 24.08.1951 r.) i tzw. dioda Shockleya, tj. struktura czterowarstwowa n-p-n-p (współcześnie znana jako diak, triak i tyrystor). Widząc ogromny potencjał komercyjny w diodzie przełączającej n-p-n-p Shockley w 1955 r. zdecydował się „pójść na swoje”, tj. powołać firmę do produkcji tych elementów. Tak się rozpoczął drugi okres jego powojennej kariery. Opuścił Bell Labs i utworzył firmę Shockley Semiconductors Laboratory, przekształconą w 1958 roku w firmę Shockley Transistor Corporation z solidnym wkładem finansowym Arnolda O. Beckmanna jako inwestora strategicznego. Początki firmy Shockleya były obiecujące. Słynął z umiejętności rekrutowania do pracy utalentowanych współpracowników. I tym razem zebrał bardzo silny zespół na czele z Gordonem Moore’em i Robertem Noyce’em, choć odmówił mu parę osób z Bell Labs, na których mu bardzo zależało (nie chcieli się przeprowadzać z New Jersey na drugą stronę Stanów do Kalifornii). Jednak grupa najlepszych specjalistów odeszła od niego we wrześniu 1957 r. Była to słynna „zdradziecka ósemka”, która założyła firmę Fairchild Semiconductors. Jako przyczynę odejścia podawano apodyktyczny charakter W. Shockleya, jego nieumiejętność zarządzania i brak zgodności co do kierunku rozwoju firmy – Shockleya interesowały wyłącznie elementy czterowarstwowe n-p-n-p, a zespół chciał pracować nad tranzystorami krzemowymi. Był to cios dla W. Shockleya, ale historycy teraz przypisują mu zasługę w powstaniu Doliny Krzemowej. Gordon Moore



Rysunek 5. Patent J. Bardeena i W. Brattaina na tranzystor ostrzowy. US Patent Office (za zgodą Computer History Museum)

(prowodny „zdrady”, ten sam od prawa Moore’a) ukuł określenie „creative fission process”, czyli „proces kreatywnego rozszczepienia”. Przecież Fairchild utworzony przez „zdradziecką ósemkę” okazał się kuźnią kadr dla wielu startupów. Przez odwieranie (spin-off) od Fairchilda powstały w Dolinie Krzemowej takie kolosy jak Intel, National Semiconductor, AMD i dziesiątki innych firm (rysunek 6). W istocie W. Shockley, może w sposób niezamierzony, zapoczątkował proces powstawania Doliny Krzemowej (ta nazwa pojawiła się później – w roku 1971). Biznesowy etap życia W. Shockleya skończył się niepowodzeniem. Inwestor, A. O. Beckmann, wycofał się sprzedając z dużą stratą firmę Shockley Semiconductor Laboratory i w życiu W. Shockleya rozpoczął się kolejny etap (1963–1975), gdy pracował na stanowisku profesora w Stanford University. W tym okresie jego aktywność na polu półprzewodników stopniowo gasła. Niestety, pojawiła się nowa pasja badawcza, co najmniej kontrowersyjna dla współczesnego społeczeństwa. W roku 1975, w wieku 65 lat W. Shockley przeszedł na emeryturę i zajął się eugeniką, całkowicie oddał się poszukiwaniu dowodów na niższą wartość intelektualną ras kolorowych, zagrażającą ludzkości degeneracją wskutek mieszania się ras. Przez jakiś czas renoma jego nazwiska sprawiała, że był zapraszany na uczelnie, gdzie wygłaszał swoje teorie, narażając się nawet na fizyczne napaści. Na jednym z wykładów posłużył się prostym dowodem na słuszność swoich eugenicznych teorii. Przywołał przykład własnej rodziny, mówiąc, że z jego związku małżeńskiego z kobietą znacznie ustępującą mu intelektualnie narodził się syn, który wprawdzie uzyskał stopień doktora fizyki, ale nigdy nie osiągnie w nauce poziomu swojego ojca. Użył przy tym określenia, że nastąpiła regresja zdolności intelektualnych. Tego było za wiele i był to jego ostatni wykład (na ile mnie wiadomo). W języku rosyjskim

Kamienie milowe w historii tranzystorów według Computer History Museum (www.computerhistory.org)

1833: Michael Faraday opisał „niezwykły przypadek” obserwacji wzrostu przewodnictwa elektrycznego w kryształach siarczku srebra wraz ze wzrostem temperatury. Faradaya zdziwiło to zjawisko, bo było odwrotne do obserwowanych dla miedzi i innych metali. Można uznać, że jest to **pierwsza historycznie obserwacja właściwości elektrycznych półprzewodników**.

1874: 24-letni fizyk niemiecki Ferdinand Braun zaobserwował, że przez styk drutu metalowego z kryształem galeny (siarczek ołowiu) prąd płynie tylko w jednym kierunku. Był to **pierwszy opis diody metal-półprzewodnik**.

1901: Jagadis Chandra Bose, profesor fizyki w Kalkucie (Indie) i pionier radiotechniki, zaprezentował **detekcję milimetrycznych fal elektromagnetycznych przy użyciu kontaktu ostrza metalowego z kryształem galeny**. W latach 1902–1906 amerykański pionier radiotechniki Greenleaf W. Pickard przetestował właściwości prostownicze styków metalu z tysiącami różnych minerałów. W grupie z najlepszymi rezultatami znalazły się kryształy krzemu z firmy Westinghouse. W 1906 roku zgłosił patent na krzemowy detektor ostrzowy. Pickard z dwoma partnerami założył firmę Wireless Specialty Apparatus Company produkującą detektory kryształkowe nazywane „cat’s whisker” (wąs kota). Była to prawdopodobnie **pierwsza firma produkująca krzemowe przyrządy półprzewodnikowe**.

1926: Polsko-amerykański fizyk Julius E. Lilienfeld złożył w USA patent (rok wcześniej w Kanadzie) pod tytułem „Method and Apparatus for Controlling Electric Currents”, w którym zaproponował strukturę trójelektrodową z użyciem siarczku miedzi jako półprzewodnika. Był to **prototyp dzisiejszych tranzystorów polowych**.

1931: Alan Wilson przedstawił **podstawową teorię właściwości półprzewodników**, opracowaną na gruncie fizyki kwantowej. W dwóch artykułach pod tytułem „The Theory of Electronic Semi-Conductors” zaproponował model wyjaśniający wpływ atomów domieszek (zanieczyszczeń) na właściwości materiałów półprzewodnikowych. Słowo półprzewodnik po raz pierwszy pojawiło się w roku 1911 w literaturze niemieckiej (niem. halbleiter), jako materiał, którego przewodnictwo elektryczne zawiera się między przewodnikami (metalami) i izolatorami. Dopiero siedem lat po publikacjach Wilsona, tj. w roku 1938 pojawiło się wyjaśnienie właściwości prostowniczych złącza metal-półprzewodnik. Teorię złącza m-s opracowali niezależnie – Walter Schottky (Niemcy), Boris Davydov (ZSRR) i Nevill Matt (Anglia).

1940: Russell Ohl i Jack Scaff z Bell Labs **odkrywają właściwości prostownicze i fotoelektryczne krzemowego złącza p-n**. Wprowadzili określenia: typ n (negative) dla obszaru krzemu domieszkowanego fosforem i wykazującego nadmiar elektronów oraz typ p (positive) dla obszaru krzemu domieszkowanego borem i wykazującego braki elektronów (później nazwane „dziurami”).

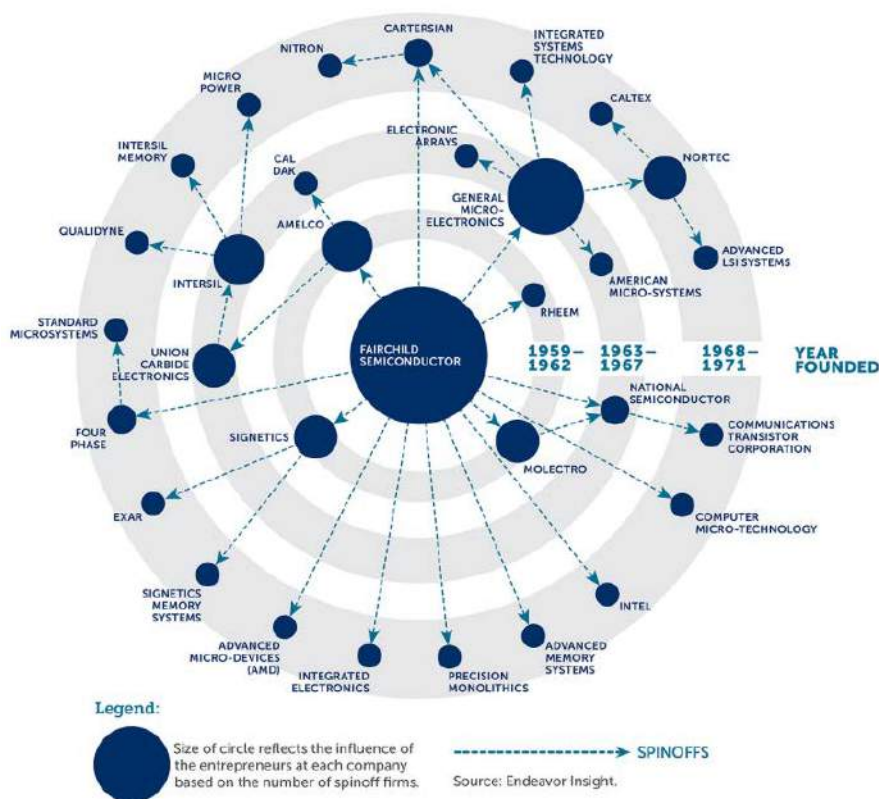
1941: Opracowanie metod wytwarzania kryształów Ge i Si o bardzo wysokiej czystości. Druga wojna światowa stymulowała gwałtowny rozwój radiolokacji, a w radarach stosowano diody m-s germanowe i krzemowe do detekcji i przetwarzania sygnałów mikrofalowych na częstotliwościach przekraczających możliwości diod na lampach próżniowych. Dla produkcji diod m-s niezbędne były płytki Ge, Si o bardzo wysokiej czystości, aby można było swobodnie kształtować ich przewodnictwo typu p i n przez kontrolowane domieszkowanie. W okresie drugiej wojny światowej osiągnięto duże postępy w technologii wytwarzania bardzo czystych kryształów Ge i Si.

1947: Wynalazek tranzystora ostrzowego. Zespół kierowany przez W. Shockleya w Bell Labs pracował nad wynalazkiem półprzewodnikowego elementu wzmacniającego. W toku eksperymentów, gdy do powierzchni germanu docięnięto zamocowane blisko siebie dwa kontakty ze złota, nieoczekiwanie zaobserwowano efekt wzmocnienia sygnału, dochodzącego do 100 razy. Stało się to 16.12.1947 r., jednak za datę tego wynalazku przyjęto 23.12.1947 r., gdy W. Brattain zademonstrował wynalazek kolegom i przełożonym. Autorami zgłoszenia patentowego są W. Brattain i J. Bardeen. Natomiast Nagrodę Nobla za wynalazek tranzystora w 1956 roku otrzymali: W. Brattain, J. Bardeen i W. Shockley. Tranzystor ostrzowy pracował zawodnie i choć w niewielkich ilościach był produkowany przez kilka lat, to nie odegrał istotnej roli w rozwoju techniki.

1948: Koncepcja tranzystora złączeniowego. Zaledwie miesiąc po demonstracji wynalazku tranzystora ostrzowego, 23 stycznia 1948 r. William Shockley przedstawił koncepcję zupełnie innej struktury tranzystora, której działanie było oparte na opracowanej przez Shockleya teorii działania złącza p-n. Tranzystorowa struktura n-p-n lub p-n-p stanowi podstawę budowy współczesnych tranzystorów bipolarnych. W. Shockley zgłosił patent na ten tranzystor w czerwcu 1948 roku, a w 1949 roku opublikował szczegółową teorię jego działania. Po upływie kolejnych dwóch lat, w roku 1951, w Bell Labs powstała technologia umożliwiająca produkcję tego tranzystora w ilościach liczących się na rynku.

jest zgrabne słowo, które fonetycznie brzmi „nierukopożatnyj”, czyli ktoś, komu nie podaje się ręki. Takiego statusu społecznego pod koniec życia dorobił się W. Shockley. Zmarł w hospicjum na raka prostaty w roku 1989, pracując niemal do ostatniego dnia. Był geniuszem, człowiekiem o największych zasługach dla rozwoju elektroniki półprzewodnikowej. Dodajmy jeszcze jedną uwagę. Skąd się wzięła tak wielka fascynacja Shockleya strukturami czterowarstwowymi n-p-n-p. Otóż od samego początku jego pracy w Bell Labs, jeszcze przed wojną, w poszukiwaniu półprzewodnikowego zamiennika lampy chodziło nie tyle o wzmacniacz, ile o klucz, czyli element dwustanowy, który mógłby zastąpić przekaźniki w telefonii. Nie zapominajmy, że pełna nazwa Bell Labs to Bell Telephone Laboratories, firmy badawczo-rozwojowej należącej wówczas do AT&T (American Telephone & Telegraph Company). (Od 2016 roku właścicielem Bell Labs jest Nokia.) W ogóle tranzystory w USA były potrzebne do komputerów, komunikacji i zastosowań militarnych, czyli głównie w roli przełączników, a nie wzmacniaczy liniowych. Tranzystorem jako elementem wzmacniającym zainteresowali się Japończycy. W 1952 roku firma Sony kupiła od Bell Labs licencję na tranzystor złączowy za 25.000 USD. W tej samej cenie tę licencję kupiło wówczas również 40 innych firm, ale dla Sony był to początek światowej ekspansji w sprzęcie powszechnego użytku. Parę lat później tranzystorowe odbiorniki radiowe produkcji Sony zalały rynek amerykański, a w 1959 roku Japonia stała się największym producentem tranzystorów na świecie. A struktury n-p-n-p też z czasem zajęły ważne miejsce w elektronice, głównie jako półprzewodnikowe elementy mocy. A propos, może to byłby dobry temat następnego wykładu? Biorę pod uwagę jeszcze optoelektronikę. Co wybrać na temat kolejnego wykładu? Proszę o głosy w tej sprawie.

THE CREATION OF SILICON VALLEY: GROWTH OF THE LOCAL COMPUTER CHIP INDUSTRY



Rysunek 6. Powstanie Doliny Krzemowej. Grafika pokazująca pączkowanie firm powoływanych do życia przez pracowników Fairchild Semiconductor (za zgodą Computer History Museum)

Quiz: Historia tranzystora

Właściwości prostownicze złącza metalu z półprzewodnikiem zostały odkryte:

- ☐ w drugiej połowie XIX wieku
- ☐ w XX wieku przed 1-szą wojną światową
- ☐ w latach 1920–1930

Zgłoszenie patentowe J. E. Lilienfelda na tranzystor polowy miało miejsce w Kanadzie, w roku:

- ☐ 1916
- ☐ 1925
- ☐ 1930

Wynalazek pierwszego tranzystora miał miejsce w firmie:

- ☐ Bell Telephone Laboratories
- ☐ Fairchild Semiconductor
- ☐ Telefunken

Pierwszy tranzystor ostrzowy zademonstrowano:

- ☐ 3 lipca 1946 roku
- ☐ 20 grudnia 1946 roku
- ☐ 23 grudnia 1947 roku

Patent na pierwszy tranzystor ostrzowy należy do:

- ☐ W. Shockleya
- ☐ W. Brattaina i J. Bardeena
- ☐ W. Shockleya, W. Brattaina i J. Bardeena

Patent na pierwszy tranzystor złączowy, nazywany również warstwowym należy do:

- ☐ W. Shockleya
- ☐ W. Brattaina i J. Bardeena
- ☐ W. Shockleya, W. Brattaina i J. Bardeena

Nagrodę Nobla za wynalazek tranzystora przyznano w roku:

- ☐ 1952
- ☐ 1954
- ☐ 1956

Laureatami Nagrody Nobla za wynalazek tranzystora są:

- ☐ W. Brattain i J. Bardeen
- ☐ W. Shockley
- ☐ W. Shockley, W. Brattain i J. Bardeen

Firma Sony kupiła licencję na tranzystor złączowy za kwotę:

- ☐ 1 USD
- ☐ 25.000 USD
- ☐ 1 mln USD

Nazwa Dolina Krzemowa (Silicon Valley) jest używana od roku:

- ☐ 1949
- ☐ 1967
- ☐ 1971

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 17.02.2023.

B. Meritum, czyli sedno tematu

Tranzystory to temat na kilka książek, więc w tym wykładzie pojęcie sedna tematu ograniczymy do absolutnie minimalnego zakresu, stanowiącego jakby uzupełnienie opowieści historycznej w części A. Po pierwsze, była to opowieść o wynalazku tranzystora bipolarnego, więc nie będziemy się zajmować tranzystorami unipolarnymi. Po drugie, była to opowieść o dokonaniach fizyków, więc nie dotkniemy tematyki pracy tranzystora w układach elektronicznych. Tymi zagadnieniami zajmujemy się w EdW bez przerwy (np. w EdW 7, 9, 10/2022 był cykl artykułów „Zrozumieć tranzystory bipolarnie”, prezentujących pracę tranzystora w poszczególnych konfiguracjach obwodowych). Pozostają nam zagadnienia fizycznego opisu działania tranzystora bipolarnego. To też tematyka bardzo obszerna i być może wykraczająca poza potrzeby elektronika hobbisty. Zatem ograniczymy tematykę tego wykładu do elementarnego opisu zjawisk fizycznych, na poziomie „sosu fizycznego”, pozwalającego lepiej zrozumieć zasadę działania tranzystora. Może to się przydać tym Czytelnikom, którym nie wystarcza traktowanie tranzystora jako czwórnika, czyli czarnej skrzynki z trzema nóżkami.

Posłużę się fragmentami z podręcznika mojego autorstwa, wydanego prawie pół wieku temu, więc zapewne już zapomnianego.

Tranzystor jako wzmacniacz

Wiemy, że w układach cyfrowych tranzystor pracuje jako klucz, czyli sterowany element dwustanowy, który można przełączać ze stanu przewodzenia do nieprzewodzenia i odwrotnie. Natomiast w układach analogowych, często nazywanych liniowymi, tranzystor spełnia rolę wzmacniacza.

Zadajmy sobie – pozornie tylko trywialne – pytanie co to oznacza, że tranzystor wzmacnia.

Co to jest wzmacniacz?

Często rozumie się w ten sposób: „Wzmacniacz – jak wskazuje nazwa – jest przyrządem służącym do wzmacniania, czyli do powiększania czegoś”. Przykładowo lupa lub lornetka teatralna powiększają obraz, dźwignia mechaniczna zwiększa siłę, transformator zwiększa napięcie (lub prąd). Czy te przyrządy są wzmacniaczami?

Nie. Nie są to wzmacniacze, gdyż: „Wzmacniacz jest przyrządem umożliwiającym sterowanie większej mocy mniejszą”. Aby nastąpił efekt wzmocnienia, są konieczne dwie rzeczy: źródło energii i przyrząd do sterowania przepływu tej energii – wzmacniacz. Jeżeli na przykład strumień wody z węża ogrodniczego skierujemy na łopatkę turbiny, która obracając się będzie wykonywała jakąś pracę, to przekręcając kran wodny (jest to czynność niewymagająca dużych energii) można powodować znaczne zmiany energii obracającej się turbiny. W tym przykładzie kran jest przyrządem służącym do sterowania dużej mocy za pomocą małej, czyli jest wzmacniaczem. Wzmacniaczem w ogólnym sensie jest również wyłącznik sieci oświetleniowej, wyłącznik radiowy itp. Transformator natomiast, nawet 1000-krotnie podwyższający napięcie lub prąd, nie jest wzmacniaczem, gdyż moc wydzielana na jego wyjściu może być co najwyżej równa mocy dostarczanej na wejście.

Tranzystor jest wzmacniaczem stosowanym zarówno do liniowego (wprost proporcjonalnego) zwiększania mocy sygnału, jak również do nieliniowego, przy czym często dyskretnego (skokowego, kluczującego) sterowania mocy. Stąd generalnie można wymienić dwa rodzaje zastosowań tranzystorów:

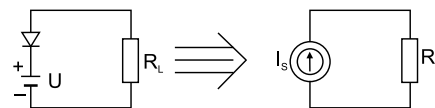
- liniowe,
- nieliniowe (analogowe nieliniowe oraz impulsowe – głównie cyfrowe).

Spróbujmy wymyślić tranzystor

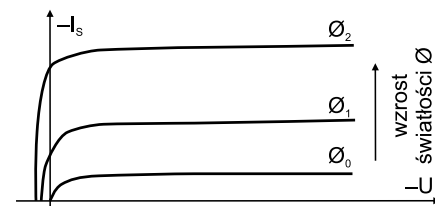
Na ogół odkrywanie rzeczy znanych jest zajęciem mało pożytecznym, ale w tym przypadku może być uzasadnione ze względów dydaktycznych.

Tranzystor powinien mieć wejście i wyjście. Do wejścia doprowadzimy sygnał sterujący, wyjście zaś włączymy w obwód sterowany. Najpierw rozpatrzmy właściwości obwodu sterowanego (wyjściowego). Z uwagi na dążenie do uzyskania jak największego wzmocnienia napięciowego obwód wyjściowy powinien mieć właściwości źródła prądowego. Źródło prądowe można obciążyć bardzo dużą rezystancją, wówczas niewielkie przyrosty prądu wywołują bardzo duże przyrosty napięcia na rezystancji obciążenia. Złącze p-n, spolaryzowane w kierunku zaporowym, ma właściwości źródła o stałej wydajności prądowej, gdyż w szerokim zakresie zmian napięcia polaryzacji lub rezystancji obciążenia prąd w obwodzie jest wielkością stałą, określoną prądem nasycenia złącza p-n oznaczanym jako I_s (rysunek 7). Jeżeli potrafimy sterować prąd I_s , to uzyskamy wzmacniacz o bardzo dużym wzmocnieniu napięciowym. (Chcielibyśmy, aby wzmacniacz działał tylko w kierunku od wejścia do wyjścia. Prąd I_s powinien więc zależeć tylko od sygnału wejściowego. Zauważmy, że brak zależności prądu I_s od napięcia wyjściowego jest cechą korzystną nie tylko ze względu na możliwość uzyskania dużego wzmocnienia napięciowego, lecz również dlatego, że zostaje wyeliminowane niepożądane oddziaływanie sygnału wyjściowego na wejściowy). Prąd I_s jest prądem nośników mniejszościowych, zależnym wprost proporcjonalnie od koncentracji tych nośników. Z kolei koncentrację nośników mniejszościowych można zmieniać różnymi sposobami. Między innymi oświetlenie złącza lub wzrost temperatury powoduje zwiększenie koncentracji nośników mniejszościowych – zwiększenie prądu I_s (rysunek 8).

Dodatkową liczbę nośników mniejszościowych można również wprowadzić do bazy złącza p-n z obwodu zewnętrznego, jeżeli kontakt bazy z obwodem zewnętrznym będzie miał właściwości wstrzykiwania. Takie właściwości ma złącze p-n spolaryzowane w kierunku przewodzenia. Zatem kontakt metalu z bazą złącza należy



Rysunek 7. Złącze p-n spolaryzowane w kierunku zaporowym jako źródło o stałej wydajności prądowej. R_L – rezystancja obciążenia; I_s – wsteczny prąd nasycenia złącza p-n



Rysunek 8. Zmiany charakterystyki $I_s(U)$ spowodowane zmianami koncentracji nośników mniejszościowych wskutek oświetlenia złącza

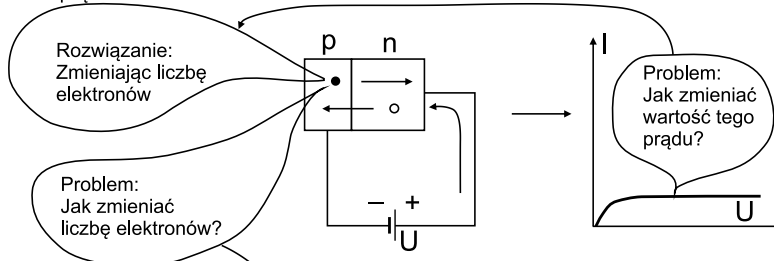
zamienić na złącze n-p. W ten sposób powstaje struktura dwuzłączowa n-p-n. Polaryzacja złącza pierwszego w kierunku przewodzenia powoduje wstrzykiwanie elektronów z obszaru N do obszaru P będącego wspólną bazą obu złączy. Elektrony dostarczane do obszaru P – jako nośniki mniejszościowe – biorą udział w prądzie I_s drugiego złącza spolaryzowanego w kierunku zaporowym (oczywiście pod warunkiem, że baza będzie cienka). W ten sposób obwód wyjściowy ma właściwości sterowanego źródła prądowego, gdyż wszelkie zmiany prądu płynącego przez pierwsze złącze w kierunku przewodzenia powodują proporcjonalne zmiany prądu I_s drugiego złącza. Kolejne fazy rozumowania prowadzącego do utworzenia tranzystora, czyli struktury mającej pożądane właściwości wzmacniacza sygnałów elektrycznych, przedstawiono na **rysunku 9**.

Analogicznie rozumując można by zaproponować również tranzystor w postaci struktury p-n-p, w którym pierwsze złącze spolaryzowane w kierunku przewodzenia wstrzykuje dziury do obszaru N, skąd są one odbierane przez drugie złącze spolaryzowane w kierunku zaporowym. Oczywiście źródła polaryzujące oba złącza miałyby odwrotną biegunowość niż w przypadku struktury n-p-n. Trzy kolejne warstwy tranzystora n-p-n lub p-n-p są nazywane zgodnie z ich funkcjami:

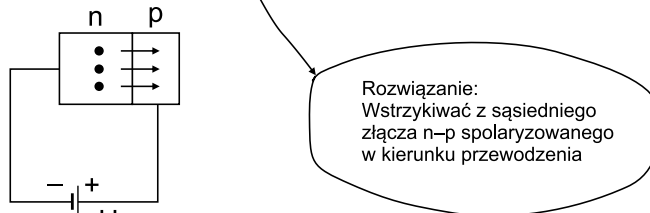
- emiter (pierwsza warstwa, która dostarcza nośników mniejszościowych do drugiej warstwy),
- baza (druga warstwa),
- kolektor (warstwa zbierająca nośniki wstrzykiwane z emitera do bazy).

Na **rysunku 10** przedstawiono obie struktury tranzystora bipolarnego wraz z symbolami tych tranzystorów. Symbole są określone z zachowaniem zasady, że strzałka oznaczająca emiter ma zwrot zgodny z kierunkiem

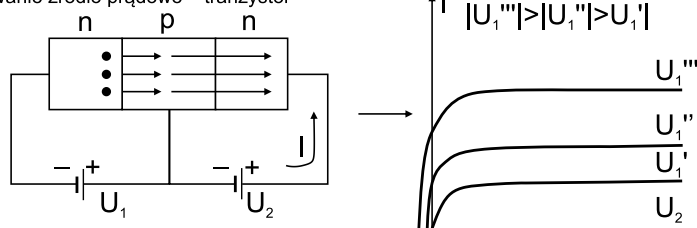
a) Źródło prądowe niesterowane



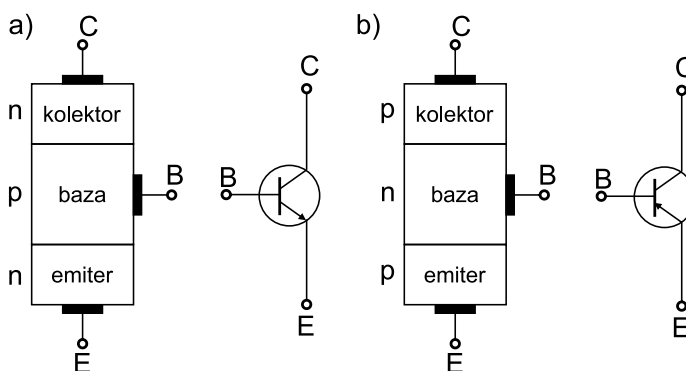
b) Sposób sterowania



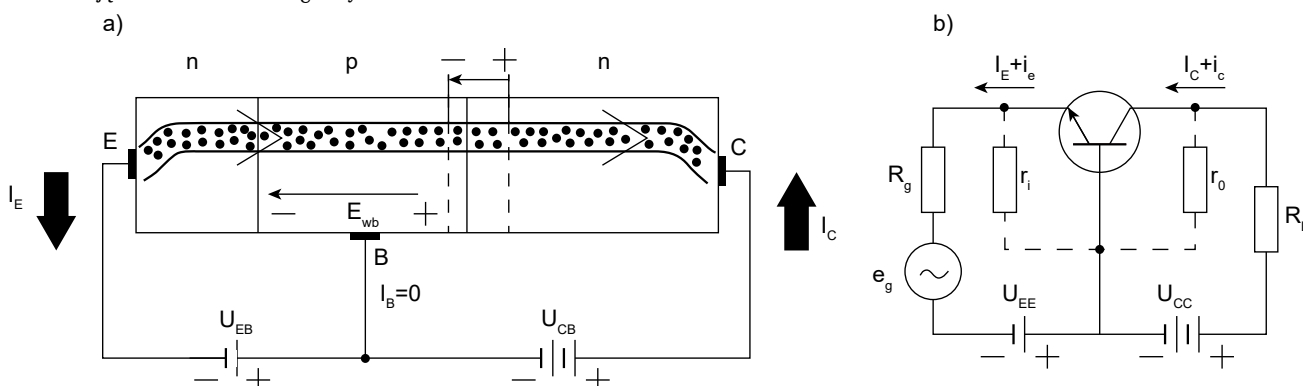
c) Sterowanie źródła prądowe – tranzystor



Rysunek 9. Szkic kolejnych faz rozumowania prowadzącego do utworzenia tranzystora



Rysunek 10. Dwie struktury tranzystora bipolarnego i ich symbole elektryczne: a) tranzystor n-p-n; b) tranzystor p-n-p



Rysunek 11. Najbardziej uproszczona ilustracja zjawisk zachodzących w tranzystorze: a) obraz fizyczny przepływu prądu; b) układ włączenia tranzystora pracującego jako wzmacniacz

przepływu prądu według konwencji przyjętej w elektrotechnice (od plusa do minusa). Jest to zwrot zgodny z kierunkiem przepływu ładunków dodatnich (dziur), a przeciwny do kierunku przepływu ładunków ujemnych (elektronów).

Tranzystor (TRANSfer resISTOR) to element transformujący rezystancję.

Rozpatrzmy przepływ prądów w strukturze $n-p-n$ przy polaryzacji złącza $E-B$ w kierunku przewodzenia, a złącza $B-C$ w kierunku zaporowym.

Przy takiej polaryzacji tranzystor spełnia funkcję elementu czynnego, tj. może służyć do liniowego wzmacniania sygnałów elektrycznych. Najbardziej uproszczony obraz zjawisk zachodzących w tranzystorze przedstawiono na **rysunku 11**. Przyjmujemy, że napięcia polaryzacji U_{EB} , U_{CB} odkładają się wyłącznie na odpowiednich warstwach zaporowych. Wskutek polaryzacji złącza $E-B$ w kierunku przewodzenia z emitera do bazy są wstrzykiwane elektrony. W bazie istnieje tzw. *wbudowane pole elektryczne* E_{wb} , spowodowane nierównomiernym rozkładem koncentracji domieszek.

Ponieważ koncentracja domieszki akceptorowej w bazie maleje w kierunku od emitera do kolektora, to pole elektryczne E_{wb} , przeciwdziałające dyfuzji dziur, jest skierowane od potencjału dodatniego przy kolektorze do potencjału ujemnego przy emiterze. Elektrony wstrzykiwane z emitera do bazy są unoszone przez E_{wb} w kierunku kolektora. Istnienie pola wbudowanego w bazie nie jest warunkiem koniecznym dla pracy tranzystora. W starych tranzystorach stopowych nie było pola E_{wb} , a transport nośników w bazie (od emitera do kolektora) odbywał się wskutek dyfuzji. Po przejściu przez bazę elektrony dostają się do warstwy zaporowej złącza $B-C$, w której istnieje silne pole elektryczne „wymiatające” te elektrony dalej do obwodu kolektora. Strumień elektronów wstrzykiwanych z emitera do bazy tworzy prąd emitera w obwodzie wejściowym, a strumień elektronów odbieranych przez kolektor jest równy strumieniowi elektronów wstrzykiwanych przez emiter, czyli prąd kolektora nie zależy od napięcia U_{CB} , lecz jest funkcją napięcia U_{EB} . Zgodnie z konwencją przyjętą w elektrotechnice prąd jest skierowany przeciwnie do strumienia elektronów. W tak uproszczonym modelu tranzystora prąd wyjściowy I_C jest równy prądowi wejściowemu I_E , $I_C = I_E$, czyli *współczynnik wzmocnienia prądowego* α_N definiowany jako I_C/I_E jest równy jedności, przy czym α_N jest współczynnikiem wzmocnienia dla prądu stałego. Współczynnik wzmocnienia dla małych przyrostów prądu $\alpha = \Delta I_C/\Delta I_E$ jest też równy jedności.

W tym miejscu trudno oprzeć się odruchowi zwątpienia w jakąkolwiek użyteczność elementu, który bez wzmocnienia przenosi prąd od wejścia do wyjścia. Można jednak łatwo wykazać, że moc wydzielana w obwodzie wyjściowym jest większa niż moc dostarczona do wejścia tranzystora.

Zgodnie z układem przedstawionym na rysunku 11 tranzystor jest polaryzowany z baterii U_{EE} i U_{CC} , a ponadto w obwodzie wejściowym jest włączone źródło e_g małego sygnału sinusoidalnego. Baterie U_{EE} i U_{CC} powodują przepływ prądów stałych I_C i I_E , natomiast ze źródła e_g płynie w obwodzie wejściowym prąd sinusoidalny i_e o amplitudzie I_{em} , który powoduje przepływ prądu sinusoidalnego i_c o amplitudzie I_{cm} w obwodzie wyjściowym. Obliczamy moc sygnału sinusoidalnego na wejściu i wyjściu tranzystora.

Moc dostarczana do wejścia tranzystora

$$P_i = I_{em}^2 r_i \quad (1)$$

przy czym r_i – rezystancja wejściowa tranzystora.

Moc odbierana w obciążeniu R_L

$$P_o = I_{cm}^2 R_L \quad (2)$$

Maksimum mocy w obciążeniu uzyskuje się przy spełnieniu warunku dopasowania

$$r_o = R_L; \text{ przy czym } r_o \text{ – rezystancja wyjściowa tranzystora} \quad (3)$$

Uwzględniając (1) do (3) można wyrazić wzmocnienie mocy w postaci

$$k_p = I_{cm}^2 R_L / I_{em}^2 r_i = I_{cm}^2 r_o / I_{em}^2 r_i \quad (4a)$$

Biorąc pod uwagę, że $\alpha = I_{cm}/I_{em}$, otrzymujemy

$$k_p = \alpha^2 r_o / r_i \quad (4b)$$

Quiz: Tranzystor bipolarny. Podstawy działania

Nazwę tranzystor można wywodzić od transformowania rezystancji:

- ☐ z małej na dużą
- ☐ z dużej na małą
- ☐ ze zmiennej na stałą

Gdy tranzystor pracuje jako wzmacniacz, to jego złącza E (emiter) – B (baza) i C (kolektor) – B (baza) są spolaryzowane następująco:

- ☐ E-B w kierunku przewodzenia
- ☐ C-B w kierunku przewodzenia
- ☐ E-B w kierunku zaporowym
- ☐ C-B w kierunku przewodzenia
- ☐ E-B w kierunku przewodzenia
- ☐ C-B w kierunku zaporowym

W tranzystorze $n-p-n$ z emitera do bazy są wstrzykiwane:

- ☐ dziury
- ☐ elektrony
- ☐ zarówno elektrony jak i dziury

Nośniki wstrzykiwane z emitera do bazy odgrywają w bazie rolę nośników:

- ☐ samoistnych
- ☐ mniejszościowych
- ☐ większościowych

Transport nośników w bazie od emitera do kolektora odbywa się wskutek:

- ☐ unoszenia
- ☐ dyfuzji
- ☐ unoszenia i dyfuzji

Charakterystyki prądowo-napięciowe złącza B-C spolaryzowanego zaporowo są:

- ☐ płaskie
- ☐ liniowe
- ☐ wykładnicze

Wzmocnienie prądowe α , definiowane jako stosunek prądu kolektora I_C do prądu emitera I_E wynosi:

- ☐ prawie 1
- ☐ nieco ponad 1
- ☐ kilkadziesiąt do kilkaset razy

Wzmocnienie prądowe β , zdefiniowane jako stosunek prądu kolektora I_C do prądu bazy I_B wynosi:

- ☐ prawie 1
- ☐ nieco ponad 1
- ☐ kilkadziesiąt to kilkaset razy

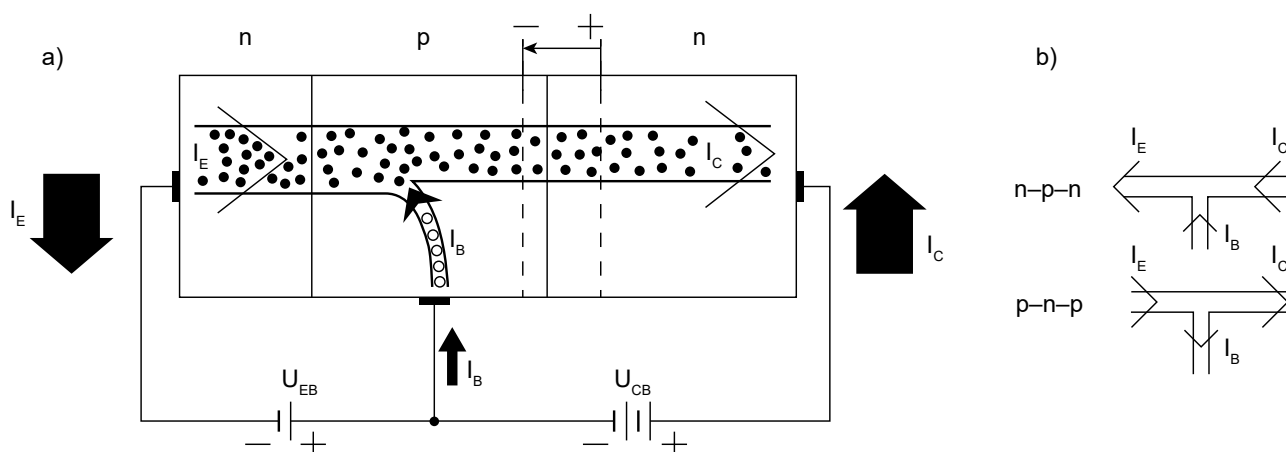
Prąd zerowy kolektora I_{CBO} płynie przy polaryzacji złącza C-B:

- ☐ napięciem zerowym
- ☐ w kierunku przewodzenia
- ☐ w kierunku zaporowym

Pod wpływem rosnącej temperatury prąd zerowy kolektora I_{CBO} :

- ☐ nie zmienia się
- ☐ rośnie
- ☐ maleje

Rozwiązanie znajdziesz na www.elportal.pl/quizy od dnia 24.02.2023.



Rysunek 12. Dokładniejsza ilustracja zjawisk zachodzących w tranzystorze z uwzględnieniem prądu spowodowanego rekombinacją nośników w bazie: a) obraz strumieni nośników; b) schematyczny rozptył prądu w tranzystorach n-p-n i p-n-p

Ponieważ w rozpatrywanym modelu uproszczonym $\alpha = 1$, więc

$$k_p = r_o / r_i \quad (5)$$

Stosunek r_o / r_i wynosi kilka tysięcy, gdyż r_i jest małą rezystancją przyrostową złącza E-B spolaryzowanego w kierunku przewodzenia (przykładowo w temperaturze pokojowej przy prądzie emitera $I_e = 1$ mA, $r_i = 25 \Omega$), r_o zaś jest bardzo dużą rezystancją przyrostową złącza B-C spolaryzowanego w kierunku zaporowym (w przypadku idealnego źródła prądowego rezystancja r_o byłaby nieskończenie wielka, w rzeczywistości jest rzędu kilkuset kiloomów).

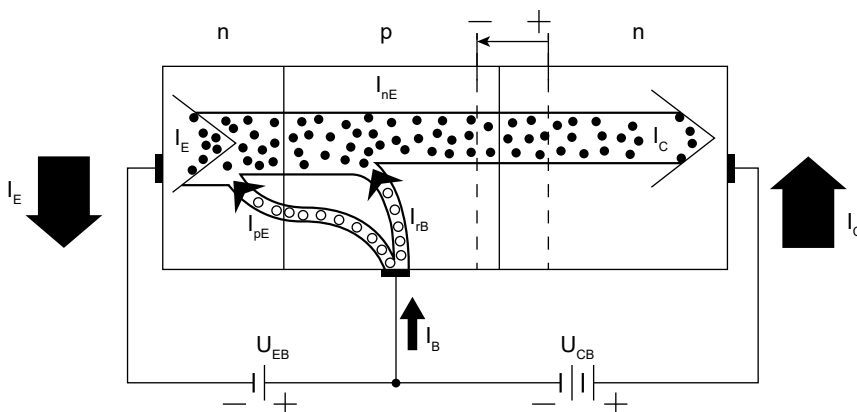
Tranzystor jest rzeczywiście „elementem transformującym rezystancję” i wzmacniaczem mocy.

Dokładniejszy opis zjawisk w tranzystorze

Uwzględnijmy teraz, że przez elektrodę bazy płynie niewielki prąd, wynikający z rekombinacji elektronów i dziur w obszarze bazy. Jednym z podstawowych założeń wprowadzanych w analizie tranzystora jest zasada obojętności elektrycznej całego obszaru bazy. Stąd wynika, że liczby elektronów i dziur nadmiarowych w bazie są sobie równe. Jeżeli na przykład z emitera w pewnej chwili wpływa do bazy 100 elektronów, to ładunek ujemny, jaki tworzą te elektrony, przyciąga z najbliższego sąsiedztwa 100 dziur. Niedomiar tych stu dziur „w najbliższym sąsiedztwie” jest uzupełniany przez przyływ dziur z następnych obszarów bazy, aż ostatecznie wpływa 100 dziur z obwodu zewnętrznego przez elektrodę bazy (jest to oczywiście jednoznaczne z usunięciem 100 elektronów z warstwy bazy do obwodu zewnętrznego, gdyż w obwodzie zewnętrznym płynie tylko prąd elektronowy). Cały proces równoważenia się ładunków nośników nadmiarowych przebiega w czasie $\tau \approx 10^{-11} \dots 10^{-13}$ s, czyli – praktycznie biorąc – jest to proces natychmiastowy. Dlatego można twierdzić, że zawsze istnieje równowaga ładunku elektronów i dziur nadmiarowych w obszarze bazy.

Skorzystajmy teraz z zasady obojętności elektrycznej bazy w celu wyjaśnienia rozptyłu prądów, przedstawionego na **rysunku 12**. Strumień elektronów wpływających z emitera do bazy niech przykładowo wynosi 100 elektronów na sekundę, a szybkość rekombinacji par elektron-dziura niech będzie 1 para na sekundę. W pierwszej sekundzie z emitera wpływa do bazy 100 elektronów i z uwagi na zasadę obojętności elektrycznej bazy w tym samym czasie przez elektrodę bazy wypływa 100 elektronów do obwodu zewnętrznego, co oznacza inaczej, że do obszaru bazy wpływa 100 dziur. W ten sposób w chwili włączenia tranzystora prąd bazy jest równy prądowi emitera. Jest to stan nieustalony. Obecnie interesuje nas stan ustalony, który w rozpatrywanym przykładzie oznacza, że w jednej sekundzie wpływa do bazy 100 elektronów z emitera, wypływa 99 elektronów do kolektora, a jeden elektron rekombinuje z dziurą. Oznacza to ubytek ładunku jednej dziury, która musi być dostarczona do bazy z obwodu zewnętrznego.

Zatem w stanie ustalonym liczba elektronów odbieranych w jednostce czasu przez kolektor jest mniejsza niż liczba elektronów wstrzykiwanych do bazy z emitera, a więc prąd kolektora jest mniejszy niż prąd emitera. Różnica tych dwóch prądów jest spowodowana rekombinacją elektronów z dziurami. A ponieważ baza musi być obojętna elektrycznie, z zewnętrznego obwodu bazy wpływa strumień dziur uzupełniających „straty” ładunku



Rysunek 13. Dalsze uściślenie obrazu zjawisk w tranzystorze; uwzględniono prąd dyfuzji dziur z bazy do emitera I_{pE}

dodatniego, spowodowane rekombinacją. Ten strumień nośników tworzy prąd bazy I_B . Stąd można zapisać podstawowe równanie prądów w tranzystorze słuszne również dla małych przyrządów prądu:

$$I_E = I_B + I_C \quad \Delta I_E = \Delta I_B + \Delta I_C \quad (6)$$

Bilans prądów w tranzystorze wynika konsekwentnie z zasady obojętności elektrycznej bazy i obowiązuje zarówno dla tranzystorów $n-p-n$ jak również $p-n-p$ (rysunek 12b), zarówno dla stanu ustalonego jak i nieustalonego.

Tranzystor jest tym lepszy (tym większe ma wzmocnienie), im mniej nośników rekombinuje w bazie. W dobrym tranzystorze

$$I_C \leq I_E; \quad I_B \ll I_C; \quad I_B \ll I_E$$

Zatem współczynnik wzmocnienia prądowego $\alpha = \Delta I_C / \Delta I_E$ jest nieco mniejszy niż jedność. Najczęściej $\alpha \approx 0,980 \dots 0,995$.

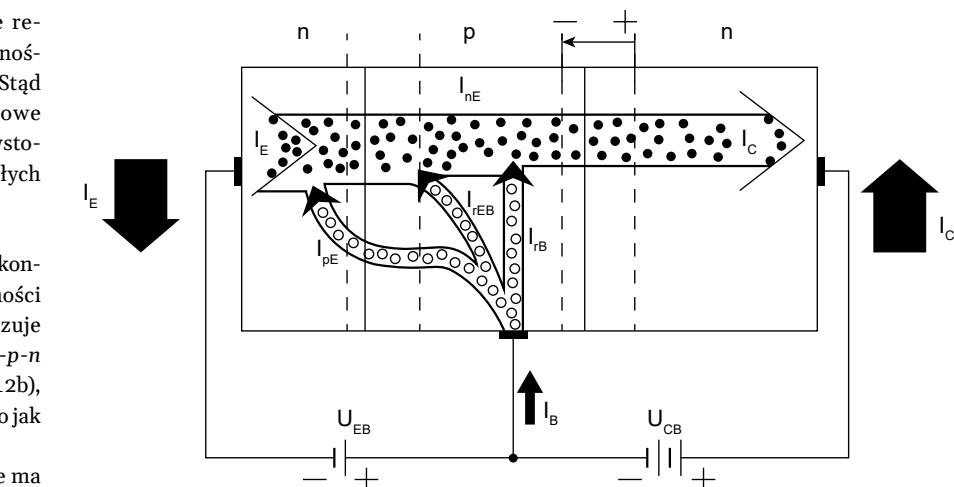
Współczynnik wzmocnienia prądowego tranzystora można również zdefiniować jako stosunek prądu kolektora do prądu bazy:

$$\beta_N = I_C / I_B \quad \beta = \Delta I_C / \Delta I_B \quad (7)$$

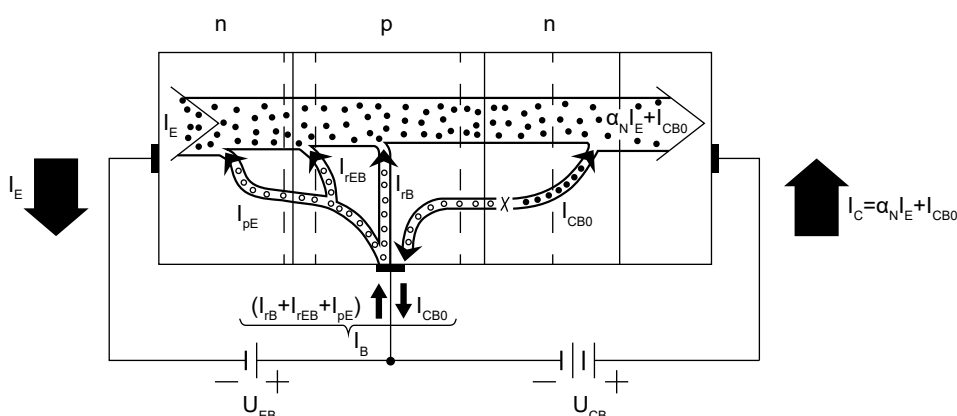
Ta definicja wzmocnienia prądowego ma zastosowanie w przypadku takiego układu włączenia tranzystora, w którym prądem wejściowym jest prąd bazy. Trzy warianty układu włączenia tranzystora były omawiane w cyklu artykułów „Zrozumieć tranzystory bipolarne” w EdW 7, 9, 10/2022. Zauważmy, że istnieje bezpośredni związek między α i β , gdyż uwzględniając (6):

$$\alpha = \beta / (1 + \beta) \quad \beta = \alpha / (1 - \alpha)$$

Dalsze uściślenia opisu zjawisk zachodzących w tranzystorze, przedstawione na **ryśunkach 13, 14** wynikają z uwzględnienia dodatkowych składowych prądów emitera i bazy. Na rysunku 13 przedstawiono składową I_{pE} prądu dyfuzji dziur z bazy do emitera, gdzie rekombinują one z elektronami. Składowa I_{pE} cyркуluje w obwodzie wejściowym dając jednakowy wkład do prądu emitera i prądu bazy. (Ponieważ obszar emitera jest znacznie silniej domieszkowany niż obszar bazy, to strumień dyfuzji dziur z bazy do emitera jest znacznie mniejszy niż strumień dyfuzji elektronów z emitera do bazy). Na rysunku 14 uwzględniono również składową prądu rekombinacji w obszarze warstwy zaporowej złącza emiter–baza.



Rysunek 14. Kolejne uściśnienie obrazu zjawisk w tranzystorze; uwzględniono prąd rekombinacji I_{reB} w warstwie zaporowej złącza E–B



Rysunek 15. Najbardziej dokładny obraz zjawisk zachodzących w tranzystorze; uwzględniono prąd zerowy kolektora I_{CBO}

Pełny obraz rozplywu prądu w tranzystorze przedstawiono na **ryśunku 15**, na którym uwzględniono również tzw. *prąd zerowy* I_{CBO} .

W tranzystorze krzemowym jest prąd nośników mniejszościowych, generowanych w obszarze warstwy zaporowej złącza C–B spolaryzowanego w kierunku zaporowym. Para elektron-dziura powstająca w warstwie zaporowej jest natychmiast „wymiatana”, przy czym elektron podąża do kolektora, a dziura do bazy. Prąd I_{CBO} dodaje się do prądu kolektora, a odejmuje od prądu bazy. Zatem efektywnie w obwodzie kolektora płynie prąd

$$I_C = \alpha_N I_E + I_{CBO} \quad (9)$$

a w obwodzie bazy

$$I_B = I_{rB} + I_{reB} + I_{pE} - I_{CBO} \quad (10)$$

przy czym: I_{rB} – prąd rekombinacji w bazie; I_{reB} – prąd rekombinacji w warstwie zaporowej złącza emiter – baza; I_{pE} – prąd dyfuzji dziur z bazy do emitera.

Uwzględniając zależności (9), (10) należy skorygować wzory na wzmocnienie dla prądu stałego:

$$\beta_N = (I_C - I_{CBO}) / I_E \quad \beta_N = (I_C - I_{CBO}) / (I_B + I_{CBO}) \quad (11)$$

Współczynniki wzmocnienia prądowego alfa, beta dla małych sygnałów pozostają bez zmian.

Tyle „sosu fizycznego” chyba wystarczy jak na jeden wykład. ■

Wiesław Marciniak

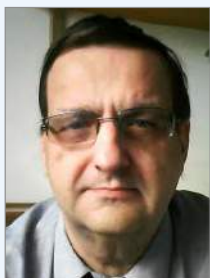
Uczmy się na cudzych błędach

Celem tej rubryki jest kształtowanie u Czytelników EdW umiejętności krytycznego czytania schematów i opisów projektów autorskich. Wszyscy jesteśmy omylni. Konstruktorzy projektów elektronicznych też. W projektach publikowanych w Internecie, ale też w artykułach drukowanych zdarzają się błędy różnej wagi, w tym też takie, które sprawiają, że układ nie może działać prawidłowo. Uczmy się wykrywać te błędy na przykładach projektów sprawdzonych w naszym redakcyjnym Pokoju Nauczycielskim.

Pamiętajmy! Nie oceniamy Autorów, tylko uczymy się na cudzych błędach.

Zapraszamy Czytelników do współpracy z naszym Pokojem Nauczycielskim. Jeśli natrafiłicie w Internecie lub źródłach drukowanych na opis projektu z poważnymi Waszymi zdaniem błędami, to przysyłajcie takie opisy do naszej redakcji (redakcja@elportal.pl w tytule wiadomości: Pokój Nauczycielski) wraz z Waszymi uwagami.

Projekt sprawdza i poprawia Paweł Sujko



Mgr inż. elektronik po Politechnice Warszawskiej, specjalność aparatura elektroniczna. Od 1992 roku pracownik Polskiego Radia SA jako inżynier serwisowy. Największa forma w jakiej „maczałem” palce, to nadajnik w Solcu Kujawskim (1 MW), najmniejsza, to pendrive (<1 W).



Prosty, ale uniwersalny muzyczny sygnalizator dźwiękowy

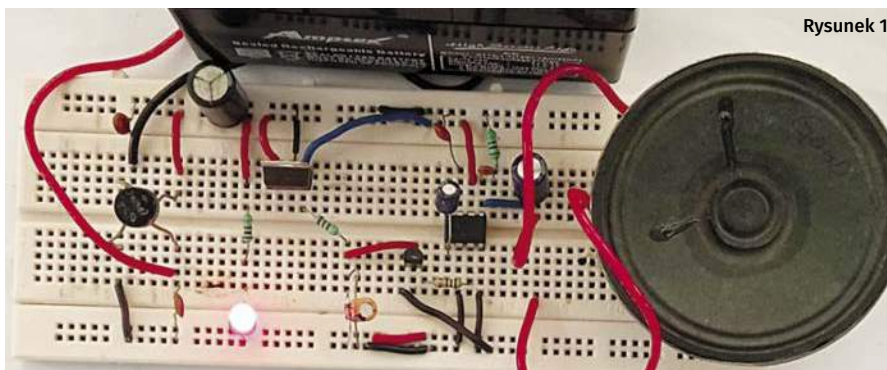
Ten prosty sygnalizator muzyczny może być używany w samochodzie, skuterze, rowerze lub motocyklu, a nawet jako dzwonek do drzwi czy ogólnie jako sygnalizator pozytywna.

Red. EdW: Zgodnie z obowiązującymi w Polsce przepisami dotyczącymi parametrów technicznych samochodów dopuszczonych do ruchu drogowego, sygnalizator dźwiękowy (popularnie klakson) może wydawać wyłącznie dźwięk ciągły (niemuzyczny) i nieprzeźwiący. Sygnały zmiennotonowe są zastrzeżone wyłącznie dla pojazdów uprzywilejowanych. Używanie sygnałów niedopuszczonych homologacją grozi zatrzymaniem dowodu rejestracyjnego.

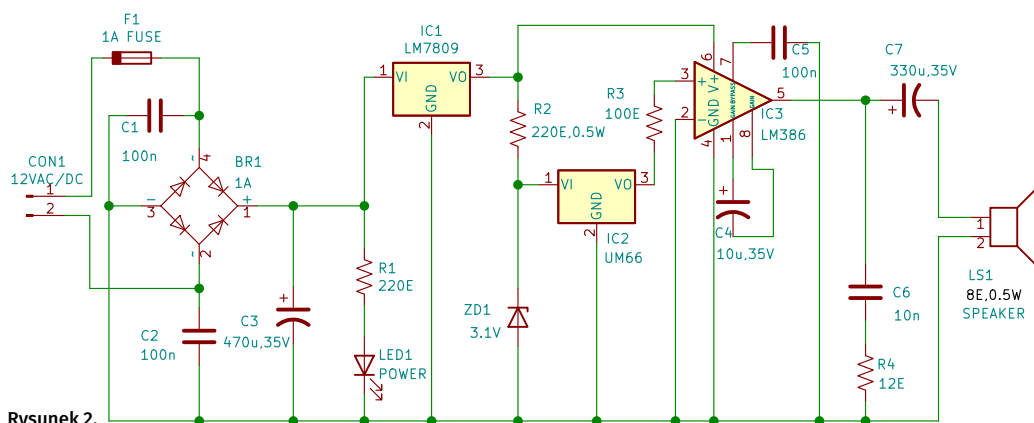
Jako źródło dźwięku sygnalizator wykorzystuje scalony generator melodii UM66T, którego sygnał wyjściowy jest wzmacniany przez wzmacniacz scalony LM386. Głośności dźwięku nie można jednak zmieniać, ponieważ nie przewidziano wymaganego do tego potencjometru, aby obwód był prosty. Prototyp autora pokazano na rysunku 1.

Schemat ideowy pokazano na rysunku 2. Układ składa się z mostka prostowniczego (BR1), stabilizatora napięcia typu 7809 (IC1), generatora melodii UM66T (IC2), wzmacniacza audio małej mocy LM386 (IC3) i kilku innych elementów. Kondensatory podłączone do zacisków zasilających służą do minimalizacji wszelkich zakłóceń szumowych.

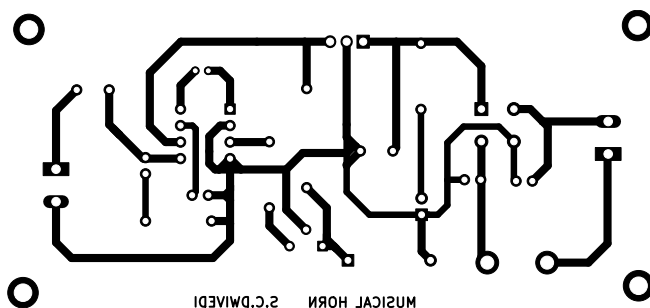
Sercem układu jest generator melodii UM66T. Ma on wbudowany generator jednostek rytmicznych i tonów. Układ jest zamknięty w trójnóżkowej obudowie tranzystorowej typu TO-92 i jest dostępny w wielu wersjach odtwarzających różne melodie. Układy scalone tej serii, są przeznaczone do użytku w dzwonekach, telefonach komórkowych i zabawkach. UM66T ma wbudowaną pamięć ROM dla melodii (maksymalnie 62 nuty/pauzy w jednym lub kilku utworach) i pobiera bardzo małą moc z zasilania. Sygnał melodyczny jest dostępny na końcówce 3 układu i jest wzmacniany



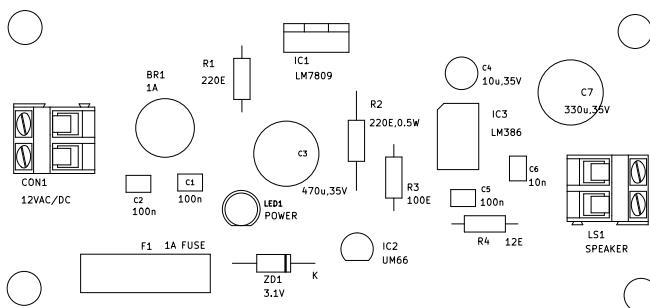
Rysunek 1.



Rysunek 2.



Rysunek 3.



Rysunek 4.

przez LM386 dlaysterowania głośnika. Układ może też odtwarzać melodie nawet bez wzmacniacza audio LM386 (wtedy używa się brzęczyka piezoelektrycznego), ale głośność będzie bardzo niska. Seria UM66xx została zaprojektowana specjalnie do odtwarzania melodii przy minimalnej ilości elementów zewnętrznych. Ostatnie dwie cyfry numeru układu identyfikują zawartość pamięci melodii.

Red. EdW: Litera L lub S, na końcu oznaczenia układu informują o sposobie pracy: S – odtworzenie jednokrotne, L – odtwarzanie ciągłe do wyłączenia zasilania.

Poniższa lista podaje przykładowe melodie odtwarzane przez układ o konkretnym numerze (repertuar obejmuje melodie zagraniczne):

- UM66T01 – „Jingle bells”, „Santa is coming to town” i „We wish you a merry Christmas”
- UM66T02 – „Jingle bells”
- UM66T04 – „Jingle bells”, „Rudolph the red nosed reindeer” i „Joy to the world”
- UM66T05 – „Home sweet home”
- UM66T06 – „Let me call you Sweetheart”
- UM66T08 – „Happy Birthday to you”

Red. EdW: Ciekawostki z cyklu „Nie samą elektroniką człowiek żyje”:

„Jingle bells” pierwotnie była znana pod tytułem „The One Horse Open Sleigh” (Jednokonne sanie) i została skomponowana w USA przez Jamesa Lorda Pierponta w 1857 roku, do śpiewania na Święto Dziękczynienia.

„Rudolph the Red-Nosed Reindeer” został skomponowany w 1949 roku przez Johna Davida Marxa, w odniesieniu do postaci bajkowego renifera, stworzonego przez Roberta Lewisa Maya w 1939 roku.

„We wish you a merry Christmass”, to angielska kolęda, w obecnej formie znana od 1935 roku, z aranżacji Arthura Warrella, brytyjskiego kompozytora, dyrygenta i organisty. Sama melodia powstała gdzieś w połowie XIX wieku (pełne pochodzenie nie jest znane).

„Santa is coming to town”, amerykańska piosenka świąteczna z 1934 roku skomponowana przez Johna F. Cootsa i Jamesa L. Gillespiego.

„Home sweet, home” powstała w 1823 roku, tekst stworzył amerykański aktor i dramaturg John Howard Payne, a muzykę angielski kompozytor Sir Henry Bishop.

„Let me call you Sweetheart”, to amerykańska melodia popularna skomponowana w 1910 roku przez amerykańskiego kompozytora Leo Friedmana do tekstu Beth Whitson, znana np. z wykonania przez Binga Crosby’ego w 1934 roku.

„Happy Birthday to you” – melodia urodzinowa pochodzi z piosenki „Good Morning to All” autorstwa siostr Patty i Mildred Hill, powstałej w 1910 roku... dla dzieci w przedszkolu, gdzie obie pracowały. Wersja urodzinowa pojawiła się w 1912 roku. Piosenka ta bardzo długo miała zastrzeżone prawa autorskie (miały być ważne do 2030 roku ale sądownie skrócono je do 2015 roku).

Układ scalony generatora melodii UM66 używany w opisywanym układzie można zamienić na dowolny inny z tej serii, aby zmienić odtwarzaną melodię/melodie. Układ może pracować zarówno na 12 V DC (Red. EdW: w przypadku napięcia stałego, to lepiej je doprowadzić już za mostkiem prostowniczym, by pominąć spadek napięcia na nim dla prawidłowej pracy stabilizatora 9 V), jak i na 12 V AC. Jeśli chcesz korzystać z sieci AC, użyj transformatora obniżającego napięcie z 230 V AC na 12 V, o wydajności 500 mA (Red. EdW: w tym miejscu nasuwa się pytanie o sensowność stosowania bezpiecznika F1 na prąd 1 A? Kto kogo będzie zabezpieczał?). Obniżone napięcie przemiennie może być podłączone bezpośrednio do złącza CON1.

Napięcie 12 V AC jest prostowane za pomocą prostownika mostkowego BR1 i filtrowane przez kondensatory C1, C2 i C3, aby usunąć wszelkie tętnienia napięcia. Wyprostowane i wygładzone filtrem napięcie jest następnie stabilizowane przez IC1 na poziomie 9 V. Jeśli chcesz użyć 12 V DC, podłącz go bezpośrednio przez CON1. LM386 to wzmacniacz mocy przeznaczony do użytku w niskonapięciowych aplikacjach konsumenckich. Jego wzmocnienie jest wewnętrznie ustawione na 20, ale dodanie zewnętrznych: rezystora i kondensatora między pinami 1 i 8 zwiększy wzmocnienie do dowolnej wartości ale nie większej niż 200. Wejścia wzmacniacza są uziemione, podczas gdy wyjście automatycznie ustawia składową stałą na poziomie połowy napięcia zasilania.

Napięcie robocze obwodu wynosi 12 V AC lub DC. UM66 pracuje między 1,5 V a 3,3 V DC, więc zastosowano tutaj diodę Zenera, aby wytworzyć 3 V potrzebne do jego zasilania. Wzmacniacz mocy LM386 działa na napięcia od 4 V do 18 V; tutaj zasilany jest napięciem 9 V stabilizowanym przez IC1 (7809) wystarczające dla uzyskania wystarczająco głośnego dźwięku. Sygnał z UM66 jest podawany na wejście wzmacniacza mocy LM386. LM386 ma kondensator 10 µF podłączony między pinami 1 i 8, który służy do zwiększenia wzmocnienia wzmacniacza (do 200). Przycisk/przełącznik chwilowy (nie pokazany na schemacie obwodu) służy do włączania obwodu, wytwarzając w ten sposób głośną melodię.

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

Półprzewodniki:

- IC1 – stabilizator napięcia 7809, 9 V, TO-220
- IC2 – generator melodii UM66T, obudowa TO-92
- IC3 – LM386 wzmacniacz mocy
- BR1 – mostek prostowniczy, 50 V, 1 A
- ZD1 – BZX55C 3 V, 0,5 W dioda Zenera
- LED1 – dioda LED, 5r, 5 mm

Rezystory:

- (wszystkie 1/4 W, ±5% węglowe), chyba, że zaznaczono inaczej
- R1 – 220 Ω (lepiej 1 lub 1,5 kΩ)
- R2 – 220 Ω, 0,5 W (lepiej 1 kΩ)
- R3 – 100 Ω
- R4 – 12 Ω

Kondensatory:

- C1, C2, C5 – 100 nF, ceramiczne
- C3 – 470 µF, 35 V elektrolityczny
- C4 – 10 µF, 35 V elektrolityczny
- C6 – 10 nF, ceramiczny
- C7 – 330 µF, 35 V elektrolityczny

Inne:

- CON1 – złącze 2-stykowe
- LS1 – głośnik 8-omowy, 0,5 W
- BATT1 – akumulator 12 V (lub 12 V AC)
- Bezpiecznik F1 – 1 A z uchwytem

Działanie układu jest proste. Po podłączeniu 12 V AC/DC do CON1, obwód jest gotowy do użycia. Po włączeniu zasilania zaświeci się dioda LED1 i jednocześnie z głośnika będzie słyszeć muzykę. Projekt jednostronnej płytki drukowanej pokazano na rysunku 3, a rozmieszczenie elementów na niej, na rysunku 4. Po zmontowaniu układu na płytce umieścić go w odpowiedniej obudowie. Zamocuj diodę LED1 z przodu obudowy, a głośnik z tyłu. Układ jest teraz gotowy do użycia. ■

S. C. Dwivedi

Uwagi i poprawki

Jeżeli transformator ma prąd maksymalny 500 mA, to bezpiecznik F1, o wartości 1 A nie ma w tym miejscu sensu.

Przy zasilaniu 12 V AC, na kondensatorze C3 będzie występować napięcie ok. 15,6 V, co oznacza, że przez diodę LED1 popłynie prąd, ok. $(15,6 \text{ V} - 2 \text{ V}) / 220 \Omega = 61,8 \text{ mA}$. Nie wiem ile ta dioda pożyje przy takim traktowaniu. Rezystor R1 powinien mieć wartość 1–1,5 k Ω .

Układ pozytywny (IC2), pobiera prąd nie większy niż 100 μA , więc mało sensownym jest przepuszczenie przez diodę ZD1 27 mA prądu, wystarczy tam 6 mA (w tym 5 mA dla poprawnej pracy diody Zenera). R2 powinien mieć wartość 1 k Ω .

Karta katalogowa układu LM386 podaje, że absolutnie maksymalna wartość wejściowa sygnału powinna mieścić się w zakresie $\pm 0,4 \text{ V}$, a w tym układzie dostaje on z UM66T sygnał o amplitudzie prawie 3 V. Co prawda producent wzmacniacza nie pisze nic o tym, czym grozi takie przesterowanie wejścia ale widać jest tam coś na rzeczy. Pomiędzy nimi powinien być potencjometr albo dzielnik rezystorowy (np. 8,2 k Ω na 1,2 k Ω) obniżający amplitudę sygnału z UM66T.

Podobnie nie ma sensu zwiększenie wzmocnienia wzmacniacza do 200 V/V przez dołączenie kondensatora C4 pomiędzy końcówki 1 i 8 wzmacniacza, bo to tylko da totalne jego przesterowanie.

Zalecana przez producenta wzmacniacza wartość C6 wynosi 47 nF.

Brak filtracji napięcia zasilania UM66T może powodować zniekształcenia dźwięku przez skoki napięcia zasilania powodowane przez przesterowany wzmacniacz LM386. Generator RC wbudowany w UM66 ma częstotliwość mocno zależną od napięcia zasilania, co pozwala też na dostrojenie go do jakiejś znanej tonacji, by nie męczył co bardziej muzycznych uszu graniem gdzieś pomiędzy tonami skali temperowanej. On i tak, ze względu na uproszczony dzielnik częstotliwości ma spore odchylenia tonów od wartości poprawnych. ■

Zasilacz z cyfrowym panelowym miernikiem napięcia i prądu

Artykuł prezentuje zasilacz regulowany z cyfrowym woltomierzem i amperomierzem panelowym. Jego zaletą jest to, że nie potrzebujesz używać osobnych mierników do pomiarów napięcia wyjściowego i prądu obciążenia.

Schemat ideowy zasilacza z cyfrowym miernikiem panelowym V/A przedstawiono na rysunku 2.

Zasilacz składa się z: transformatora (X1) obniżającego napięcie z 230 V na 15 V AC, prostownika w układzie Graetza (diody D1 do D4), trójkątówkowego stabilizatora napięcia (IC1), typu LM317, cyfrowego miernika panelowego (DPM1) i kilku innych elementów.

LM317 to popularny nastawny stabilizator napięcia, który jest wyposażony w: wewnętrzne ograniczenie prądu, zabezpieczenie przed przegrzaniem i zabezpieczenie pracy stopnia wyjściowego w obszarze bezpiecznym. Zapewnia więc regulowane i stabilizowane napięcie na wyjściu, niezależnie od wahań napięcia wejściowego i prądu obciążenia. Opisany układ przetwarza niestabilizowane napięcie stałe z kondensatora filtrującego C₁, na stabilizowane napięcie o nastawianej wartości.

Dioda świecąca LED1 wskazuje obecność napięcia po prostowniku, a przed stabilizatorem.

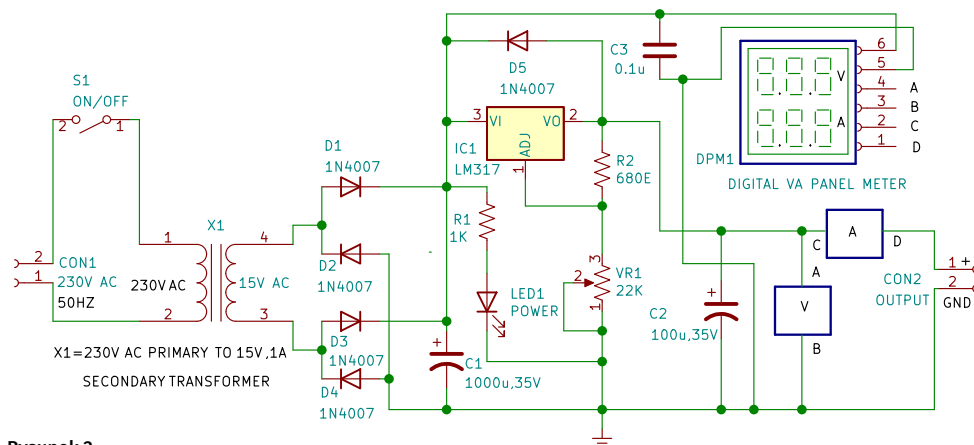
Potencjometr V_{R1} jest podłączony do końcówki ADJ stabilizatora IC, w celu regulacji napięcia wyjściowego. Kondensator C₃ zabezpiecza zasilanie panela pomiarowego przed zakłóceniami impulsowymi. Dioda D₅ podłączona między końcówki Vo i Vi stabilizatora napięcia zabezpiecza go w sytuacji zwarcia na wejściu, tj. gdyby napięcie na C₁ stało się mniejsze niż na C₂ i miałby się ten ostatni rozładować przez wyjście stabilizatora. Do stabilizatora LM317 wymagany jest odpowiedni radiator.

Aby można było ustawić żądane napięcie na wyjściu, między końcówkami: Vo i Adj włączone są dwa rezystory: R2 (680 Ω) i nastawny VR1 (22 k Ω !).

Red. EdW: Wewnętrzne obwody LM317 utrzymują napięcie odniesienia U_{ref}, wynoszące typowo 1,25 V, pomiędzy końcówkami Vo i Adj, o ile tylko zapewniony jest warunek, że $U_{Vi} - U_{Vo} > 3 \text{ V}$. Dzięki tej stabilizacji przez rezystor R1 płynie stały prąd $I_{ref} = U_{ref} / R_2$. Układ LM317 jest tak skonstruowany, że z końcówki Adj wypływa



Rysunek 1.



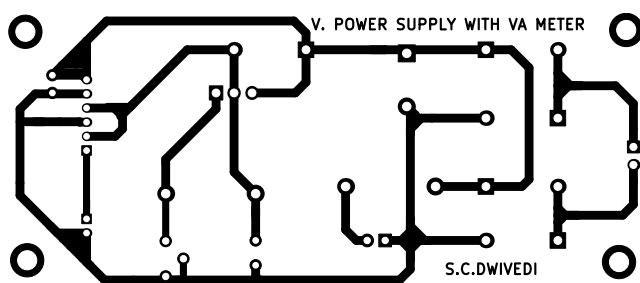
Rysunek 2.



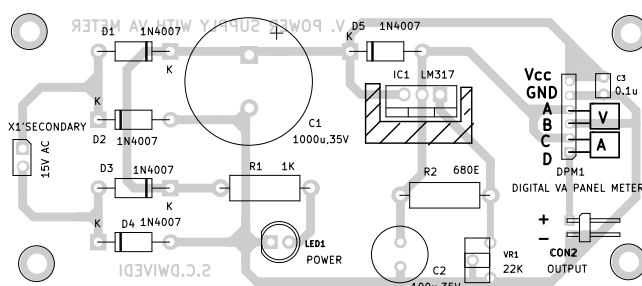
Rysunek 3.



Rysunek 4.



Rysunek 5.



Rysunek 6.

praktycznie niezmienny prąd polaryzacji wewnętrznej źródła odniesienia, wynoszący typowo 50 μA (maksymalnie 100 μA). Cała reszta prądu potrzebna do pracy tego układu wypływa przez końcówkę V_o .

Suma prądów I_{ref} i I_{adj} przepływa przez rezystor nastawny V_{R1} do masy.

Tak więc napięcie wyjściowe całego stabilizatora składa się z sumy: spadku napięcia na V_{R1} i spadku napięcia na R_2 równego U_{REF} .

Wynika z tego, że minimalne napięcie wyjściowe stabilizatora jest równe napięciu U_{REF} przy $V_{R1}=0$.

Dodatkowo, ponieważ prąd I_{adj} zmienia się nieznacznie (do 5 μA) przy zmianach prądu obciążenia oraz istnieje pewien rozrzut jego wartości nominalnej pomiędzy egzemplarzami układu LM317 (od 50 μA do 100 μA), to dobrze jest wybrać wartość I_{REF} odpowiednio dużą by prąd I_{adj} był możliwie mały na tle całkowitego prądu

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

Półprzewodniki:

IC1 – LM317, TO-220 – stabilizator
D1-D5 – 1N4007 – diody prostownicze
LED1 – 5mm LED

Rezystory: ($\pm 5\%$ węglowe)

R1 – 1 k Ω , 0,5 W
R2 – 680 Ω (lepiej 150 Ω 1% metalizowany)
VR1 – 22 k Ω (lepiej 2,2 k Ω)

Kondensatory:

C1 – 1000 μF , 35 V elektrolityczny
C2 – 100 μF , 35 V elektrolityczny
C3 – 0,1 μF ceramiczny

Inne:

X1 – 230 V/15 V, 1 A transformator sieciowy
DPM1 – Dwukolorowy, 3-cyfrowy wyświetlacz panelowy
S1 – Wtycznik zasilania na 230 V AC
CON1 – wtyczka sieciowa 230 V
CON2 – złącze 2-stykowe
- Radiator dla LM317

plynącego przez R_2 . Z powyższego wynika, że dobór wartości elementów podany przez autora jest daleki od optymalnego, a zakres regulacji napięcia zaczyna się od 1,25 V, a nie od 0 V.

VR1 można regulować od pozycji minimalnej do maksymalnej, aby uzyskać napięcie od 0 V (patrz uwagi wyżej) do 15 V na obciążeniu.

Red. EdW: Przyjmując typową wartość $V_{R1} = 2200 \Omega$ i typową tolerancję tej rezystancji wynoszącą $\pm 20\%$, to obliczenia należy rozpocząć od zapewnienia uzyskania 15 V na wyjściu dla najgorszego przypadku, tj. $V_{R1} - 20\%$ czyli 1760 Ω , $U_{\text{REFmin}} = 1,2 \text{ V}$ oraz $I_{\text{ADJmin}} = 50 \mu\text{A}$ i dla tej wartości wyliczyć maksymalną R_2 .

Jest to typowa wartość z szeregu 1%. Przyjmijmy jednak wartość o stopień mniejszą czyli 150 Ω , tol. 1%, co zapewni nam możliwość ustawienia 15 V na wyjściu nawet w najgorszym przypadku odchyłek wartości użytych elementów.

W przypadku realnych elementów U_{WYmax} będzie większe od 15 V ale z tym zawsze można walczyć dobierając odpowiedni rezystor stały lub potencjometr montażowy dołączany równolegle do V_{R1} dla ograniczenia jego zakresu regulacji.

Gdy obwód jest włączony, cyfrowy miernik panelowy pokazuje również 000 (patrz uwagi wyżej) dla napięcia i 000 dla prądu (przy braku obciążenia). Miernik panelowy może natomiast mierzyć napięcie do 100 V i prąd do 10 A. Niektóre ważne parametry miernika panelowego są wymienione w tabeli.

Konstrukcja i testowanie

Rysunek 5 pokazuje projekt ścieżek płytki drukowanej w skali 1:1, rysunek 6 pokazuje rozmieszczenie elementów na niej. Obwód, jako że jest prosty, może być również zmontowany na płytce stykowej (to jest rozwiązanie raczej dla małych prądów obciążenia) lub płytce uniwersalnej.

Po sprawdzeniu poprawności wykonania połączeń można włączyć układ podłączając napięcie 230 V do uzwojenia pierwotnego transformatora X1. Aby zmierzyć napięcie na wyjściu zasilacza podłącz przewody A i B miernika panelowego do punktów A i B na płytce drukowanej (jeżeli jeszcze nie połączyliśmy przewodów amperomierza, to na złączu CON2 nie będzie napięcia). Miernik pokaże aktualną wartość napięcia, która będzie się zmieniała wraz ze zmianą VR1. Zakładając, że przewody C i D nie są połączone, odczyt odpowiadający prądowi wyniesie 0,00. Odczyty 2 V i 0 A na mierniku panelowym pokazano na rysunku 3.

Jeśli chcesz zmierzyć prąd obciążenia, to podłącz przewody C i D miernika panelowego do punktów C i D płytki drukowanej, szeregowo z obciążeniem, tak jak pokazano na schemacie. Odczyty prądu 3 V i 0,04 (lub 40 mA) na mierniku panelowym pokazano na rysunku 4.

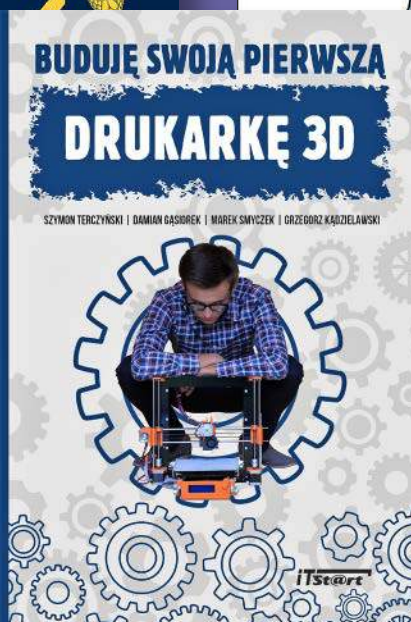
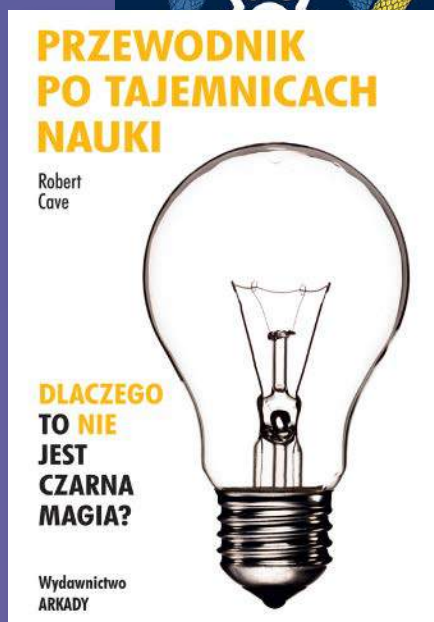
W zależności od zasilacza i zastosowanego obciążenia można mierzyć prąd od 0 do 1,5 A. Jednak obecny obwód jest przeznaczony tylko dla napięcia do 15 V i prądu do 1 A. ■

Sani Theo

Specyfikacja panelu pomiarowego

sklep AVT kod handlowy 03314
Napięcie zasilania: 4,5 do 30 V DC
Prąd zasilania: 20 mA
Kolor wyświetlacza: dwukolorowy (czerwony i niebieski)
Wymiary: 48×29×21 mm
Okres odświeżania: około 500 ms
Dokładność pomiaru: 1%
Prąd roboczy: 20 mA
Zakres pomiarowy: DC 0–100 V, 0–10 A
Temperatura pracy: od 10 do 65°C

KSIĄŻKI W ULUBIONYM KIOSKU Z RABATEM DO 30%



Zobacz pełną ofertę! PONAD 500 TYTUŁÓW

Zamów wygodnie na www.UlubionyKiosk.pl

Projekty układów zwrotnic do zestawów głośnikowych

Z uwagi na duże zainteresowanie projektowaniem i samodzielnym wykonaniem nagłośnienia do użytku domowego, w poniższym artykule przedstawiono dwa projekty pasywnych zestawów głośnikowych o wysokiej jakości odtwarzania dźwięku, jakie każdy czytelnik może wykonać we własnym zakresie przy użyciu elementów elektronicznych, które są bez problemu dostępne na rynku. Obydwa zestawy głośnikowe wykorzystują nieznacznie tylko zmodyfikowane obudowy popularnych zestawów głośnikowych firmy Tonsil, typu „Altus 300” oraz „Zeus”, co w znaczącym stopniu obniża koszt wykonania tych urządzeń.

1.1. Zestaw głośnikowy w obudowie typu „Altus 300”

Pierwszą propozycję stanowi zestaw głośnikowy w obudowie typu „Altus 300”. Taki zestaw można sobie wykonać tanim kosztem we własnym zakresie. Charakteryzuje się on wyrównaną charakterystyką poziomu ciśnienia akustycznego w funkcji częstotliwości, wyrównaną charakterystyką opóźnienia grupowego w funkcji częstotliwości oraz wyrównaną charakterystyką modułu impedancji w funkcji częstotliwości poza obszarem rezonansów głośnika niskotonowego w obudowie basrefleks, co jest szczególnie korzystne z punktu widzenia użytkowników wzmacniaczy lampowych. Zwrotnica tego zestawu głośnikowego była wielokrotnie optymalizowana pod kątem uzyskania jak najlepszych rezultatów dźwiękowych i parametrów elektrycznych. Zestaw jest trójdrożny. Częstotliwości podziału to 500 Hz oraz 5 kHz. Impedancja znamionowa to 8 Ω .

Aby wykonać taki zestaw będziemy potrzebować następujących elementów:

1. Dwie obudowy od zestawu głośnikowego „Altus 300” wraz z przyłączami, gniazdami maskownic, maskownicami, cokołami i rurami basrefleks, wykonane z płyty MDF o grubości 18 mm. Można je zamawiać w najrozmaitszych wykonaniach przy zastosowaniu

różnych rodzajów oklein. Najbardziej popularne wykonanie to front i tył w okleinie typu czarna morka, natomiast ścianki boczne w okleinie typu czarny jesion. Istnieje jednak możliwość zastosowania różnych kombinacji rodzajów oklein i jeśli komuś zależy na bardziej klasycznym wyglądzie zestawu, można np. zastosować kombinację – front i tył w okleinie typu czarna morka, natomiast ścianki boczne w okleinie typu orzech. Producent przewiduje również możliwość wykonania obudów do tego zestawu głośnikowego w odbiciu lustrzanym.

2. Dwa głośniki niskotonowe GDN 30/60/3 w wykonaniu ośmioomowym z kobaltowym obwodem magnetycznym (najlepiej ze starej produkcji z lat siedemdziesiątych dwudziestego wieku, po regeneracji, ponieważ wersja współczesna nie posiada na rdzeniu obwodu

magnetycznego miedzianego pierścienia Faradaya i jej zastosowanie wymaga przerobienia obwodu kompensacyjnego Zobła w proponowanej zwrotnicy). Istnieje także możliwość zastosowania tego głośnika w wersji produkowanej współcześnie. W tym artykule znajdzie się komentarz dotyczący modyfikacji układu zwrotnicy celem umożliwienia jej współpracy z głośnikiem w tym wykonaniu.

3. Dwa głośniki średnionowe GDM 18/80 w wykonaniu ośmioomowym.
4. Dwa głośniki wysokotonowe GDWK 9/80/1 w wykonaniu ośmioomowym.
5. Pierścienie dekoracyjne do głośników niskotonowych i średnionowych wykonane w Tonsilu na zamówienie bez fazowania otworów montażowych (otwory z fazowaniami dla śrub z łbem stożkowym „wyrabiają się” podczas montażu i demontażu co wygląda nieestetycznie), przeznaczone



Rysunek 1.2. Głośnik średnionowy Tonsil typu GDM 18/80 Źródło: <https://www.skleptonsil.pl/>



Rysunek 1.1. Głośnik niskotonowy Tonsil typu GDN 30/60/3 Źródło: <https://www.skleptonsil.pl/>



Rysunek 1.3. Głośnik wysokotonowy Tonsil typu GDWK 9/80/1 Źródło: <https://www.skleptonsil.pl/>



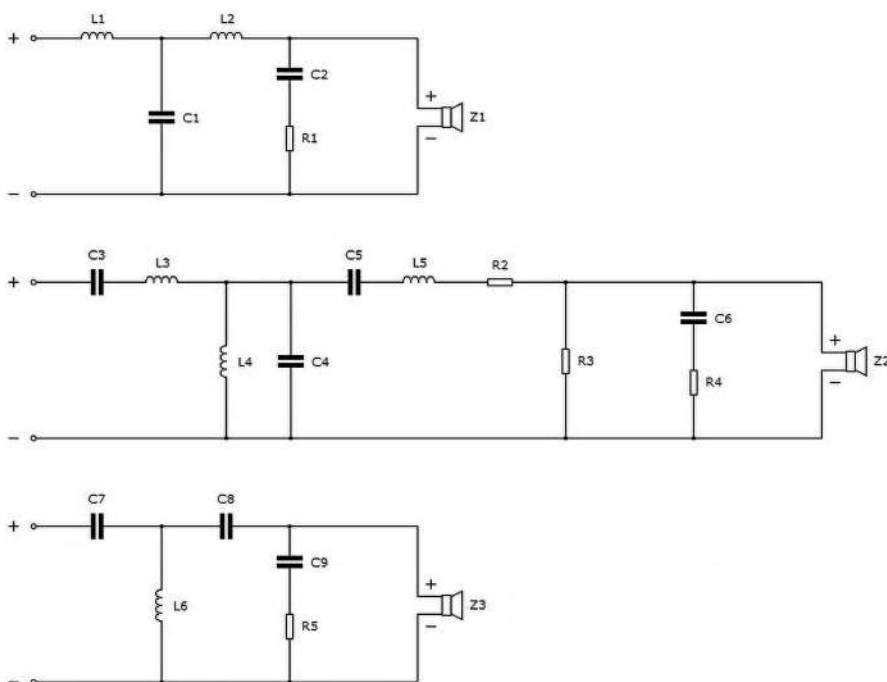
Rysunek 1.4. Śruba z gwintem metrycznym M4 z łbem walcowym pod płaski śrubokręt wykonana zgodnie z normą DIN84



Rysunek 1.5. Pierścień dekoracyjny głośnika wysokotonowego Tonsil typu GDWK 9/80/1

Tabela 1.1. Wykaz elementów składowych zwrotnicy zestawu głośnikowego w obudowie typu „Altus 300”

R1	8,2 Ω/20 W
R2	2,2 Ω/20 W
R3	15,0 Ω/20 W
R4	8,2 Ω/20 W
R5	8,2 Ω/20 W
L1	3,58 mH/0,820 Ω/Ø 1,2 mm
L2	1,15 mH/0,445 Ω/Ø 1,2 mm
L3	0,35 mH/0,270 Ω/Ø 1,0 mm
L4	1,60 mH/0,650 Ω/Ø 1,0 mm
L5	0,12 mH/0,140 Ω/Ø 1,0 mm
L6	0,16 mH/0,240 Ω/Ø 0,8 mm
C1	56,0 µF/400 V/MKP
C2	18,0 µF/400 V/MKP
C3	27,0 µF/400 V/MKP
C4	6,8 µF/400 V/MKP
C5	68,0 µF/400 V/MKP
C6	6,8 µF/400 V/MKP
C7	2,7 µF/400 V/MKP
C8	8,2 µF/400 V/MKP
C9	1,0 µF/400 V/MKP
Z1	GDN 30/60/3
Z2	GDM 18/80
Z3	GDWT 9/80/1



Rysunek 1.6. Schemat ideowy zwrotnicy zestawu głośnikowego w obudowie typu „Altus 300”



Rysunek 1.7. Nakrętki pazurkowe z gwintem metrycznym wewnętrznym M4



Rysunek 1.10. Wygląd zewnętrzny gotowego zestawu głośnikowego w obudowie typu „Altus 300”

do przykręcania wkrętami M4 z łbem walcowym pod płaski śrubokręt (wykonane zgodnie z normą DIN84).

6. Pierścienie dekoracyjne do głośników wysokotonowych wykonane w Tonsilu na zamówienie bez fazowania otworów montażowych (nie są one produkowane seryjnie ale można je wykonać w produkcji jednostkowej).
7. Watolina do wytłumienia.
8. Elementy zwrotnicy.

9. Nakrętki pazurkowe.

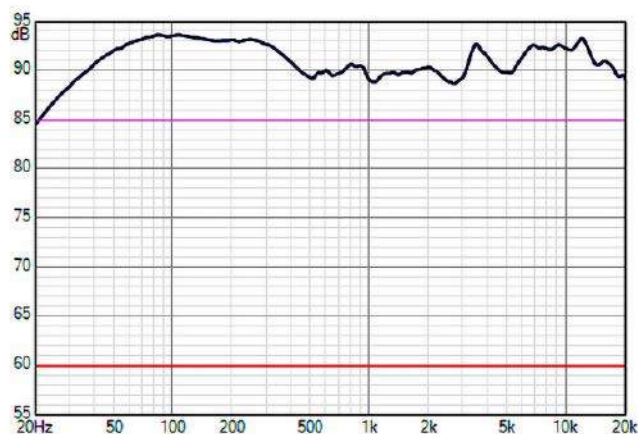
Obudowy tego zestawu głośnikowego różnią się od fabrycznie produkowanych. Musimy wziąć to pod uwagę podczas składania zamówienia. Przede wszystkim w fabrycznym zestawie głośnikowym występuje głośnik wysokotonowy tubowy typu GDWT 9/100, którego średnica montażowa wynosi $\varnothing 78$ mm. Proponowany zestaw głośnikowy



Rysunek 1.8. Zwrotnica zestawu głośnikowego w obudowie typu „Altus 300” – tor niskotonowy



Rysunek 1.9. Zwrotnica zestawu głośnikowego w obudowie typu „Altus 300” – tor średniotonowy i wysokotonowy



Rysunek 1.11. Wypadkowa charakterystyka poziomu ciśnienia akustycznego w funkcji częstotliwości zestawu głośnikowego w obudowie typu „Altus 300” (wszystkie głośniki podłączone synfazowo)



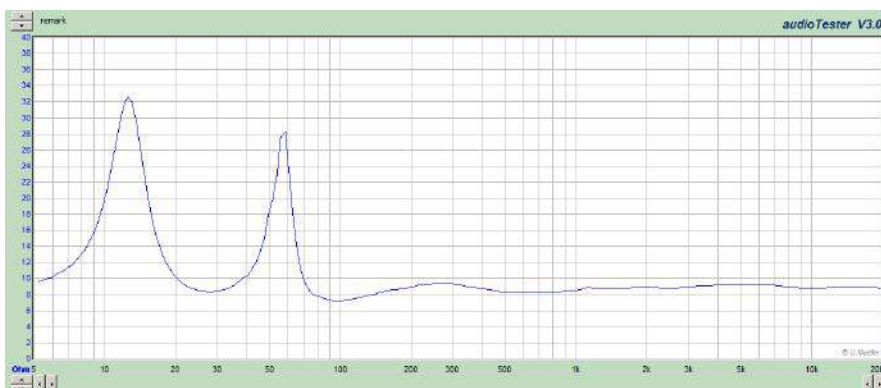
Rysunek 1.12. Wypadkowa charakterystyka poziomu ciśnienia akustycznego w funkcji częstotliwości zestawu głośnikowego w obudowie typu „Altus 300” (głośnik średniotonowy podłączony afazowo)

wykorzystuje głośnik wysokotonowy kopułkowy typu GDWK 9/80/1, którego średnica montażowa wynosi $\varnothing 80$ mm. A zatem otwór pod ten głośnik powinien zostać wyfrezowany na obrabiarce CNC właśnie na wymiar $\varnothing 80$ mm. Oprócz tego w przypadku gdy chcemy zamocować głośniki oraz przyłączyć przy pomocy nakrętek pazurkowych z gwintem metrycznym wewnętrznym M4, musimy zlecić wywiercenie na obrabiarce CNC wszystkich otworów montażowych na wymiar $\varnothing 5,0$ mm, $\varnothing 5,5$ mm lub $\varnothing 6,0$ mm, w zależności od tego jakim rodzajem nakrętek pazurkowych akurat dysponujemy. Jeśli zamierzamy natomiast zamocować głośniki oraz przyłączyć przy pomocy wkrętów do drewna, możemy z tej modyfikacji zrezygnować. Nie jest to jednak wskazane ze względu na to, że każdorazowy demontaż zestawu głośnikowego (np. podczas zaistnienia konieczności przeprowadzenia napraw lub regeneracji głośników) może sprawić, że otwory w płycie MDF się „wyrabiają”. Na ściankach przyklejamy od wewnątrz butaprenem lub przymocowujemy zszywaczem tapicerskim watolinę.

Następnie przystępujemy do montażu zwrotnicy według schematu zaproponowanego na rysunku 1.6. Zwrotnicę najlepiej umieścić w każdej obudowie na dwóch osobnych płytach (czyli na zestaw wyjdą cztery płyty tekstolitowe bądź pilśniowe). Osobno tor niskotonowy i osobno średniotonowy i wysokotonowy. Rozmieszczenie elementów pokazano na rysunkach 1.8. oraz 1.9.

Zwrotnice przykręcamy za pośrednictwem nóżek dystansowych wkrętami do drewna do tylnej i dolnej ścianki zestawu. Pod każdą płytą zwrotnicy podkładamy piankę poliuretanową.

Połączenia pomiędzy głośnikami a zwrotnicą realizujemy dwużyłowymi przewodami miedzianymi o następujących przekrojach:



Rysunek 1.13. Wypadkowa charakterystyka modułu impedancji w funkcji częstotliwości zestawu głośnikowego w obudowie typu „Altus 300”

- 2,5 mm² – tor niskotonowy,
- 1,5 mm² – tor średniotonowy,
- 0,5 mm² – tor wysokotonowy.

Połączenie pomiędzy zwrotnicami a przyłączem – 2,5 mm².

Zastosowanie głośnika GDN 30/60/3 z bieżącej produkcji wymaga zwiększenia pojemności kondensatora C2 w układzie kompensacyjnym Zobla z 18 μ F do 33 μ F. Głośniki w wersji z bieżącej produkcji nie posiadają bowiem na rdzeniu obwodu magnetycznego miedzianego pierścienia Faradaya.

1.2. Zestaw głośnikowy w obudowie typu „Zeus”

Drugą propozycję stanowi zestaw głośnikowy w obudowie typu „Zeus”. Taki zestaw można sobie także wykonać tanim kosztem we własnym zakresie. Charakteryzuje się on wyrównaną charakterystyką poziomu ciśnienia akustycznego w funkcji częstotliwości, wyrównaną charakterystyką opóźnienia grupowego w funkcji częstotliwości oraz wyrównaną charakterystyką modułu impedancji w funkcji częstotliwości poza obszarem rezonansów głośnika niskotonowego w obudowie basrefleks,

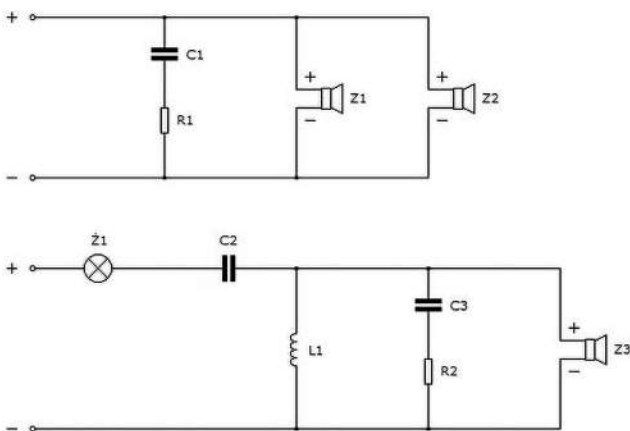
a także bardzo dużą efektywnością, co jest szczególnie korzystne z punktu widzenia użytkowników wzmacniaczy lampowych. Zwrotnica tego zestawu głośnikowego stanowi modyfikację fabrycznego układu, który został uzupełniony o dwa odpowiednio dobrane obwody kompensacyjne Zobla celem wyrównania wypadkowej charakterystyki modułu impedancji w funkcji częstotliwości. Zastosowano także inne niż w oryginale głośniki niskotonowe. Zestaw jest dwudrożny. Częstotliwość podziału to 4 kHz. Impedancja znamionowa to 4 Ω .

Aby wykonać taki zestaw będziemy potrzebowali następujących elementów:

1. Dwie obudowy od zestawu głośnikowego typu „Zeus” wraz z przyłączami, gniazdami maskownic, maskownicami i rurami basrefleks, wykonane z płyty MDF o grubości 18 mm. Można je zamawiać w najrozmaitszych wykonaniach przy zastosowaniu różnych rodzajów oklein. Najbardziej popularne wykonanie to front i tył w okleinie typu czarna morka, natomiast ścianki boczne w okleinie typu czarny jesion.



Rysunek 1.14. Głośnik niskotonowy Tonsil typu GDN 30/80/2 Źródło: <https://www.skleptonsil.pl/>



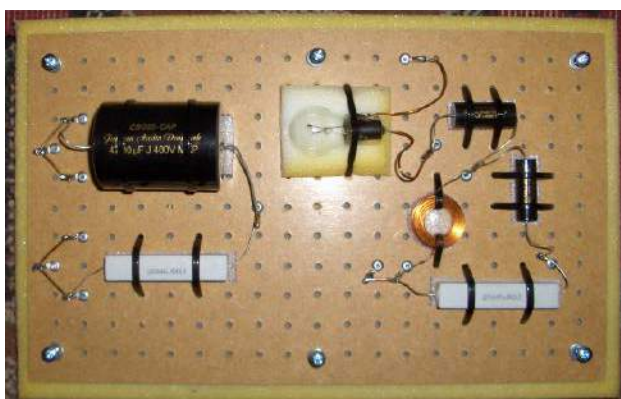
Rysunek 1.16. Schemat ideowy zwrotnicy zestawu głośnikowego w obudowie typu „Zeus”



Rysunek 1.18. Żarówka samochodowa od kierunkowskazów 12 V / 21 W Źródło: <https://www.skleptonsil.pl/>



Rysunek 1.15. Głośnik wysokotonowy Tonsil typu GDWT 12-19/150 Źródło: <https://www.skleptonsil.pl/>



Rysunek 1.17. Zwrotnica zestawu głośnikowego w obudowie typu „Zeus”



Rysunek 1.19. Wygląd zewnętrzny gotowego zestawu głośnikowego w obudowie typu „Zeus”

Tab. 1.2. Wykaz elementów składowych zwrotnicy zestawu głośnikowego w obudowie typu „Zeus”

R1	6,8 Ω /20 W
R2	6,8 Ω /20 W
L1	0,20 mH/0,350 Ω /Ø 0,7 mm
C1	47,0 μ F/400 V/MKP
C2	1,5 μ F/400 V/MKP
C3	1,0 μ F/400 V/MKP
Ż1	12 V/21 W
Z1	GDN 30/80/2
Z2	GDN 30/80/2
Z3	GDWT 12-19/150

Istnieje jednak możliwość zastosowania różnych kombinacji rodzajów oklein i jeśli komuś zależy na bardziej klasycznym wyglądzie zestawu, można np. zastosować kombinację – front i tył w okleinie typu czarna morka natomiast ścianki boczne w okleinie typu orzech.

2. Cztery głośniki GDN 30/80/2 w wykonaniu ośmioomowym.
3. Dwa głośniki GDWT 12-19/150 w wykonaniu ośmioomowym.
4. Watolina do wytłumienia.

5. Elementy zwrotnicy.

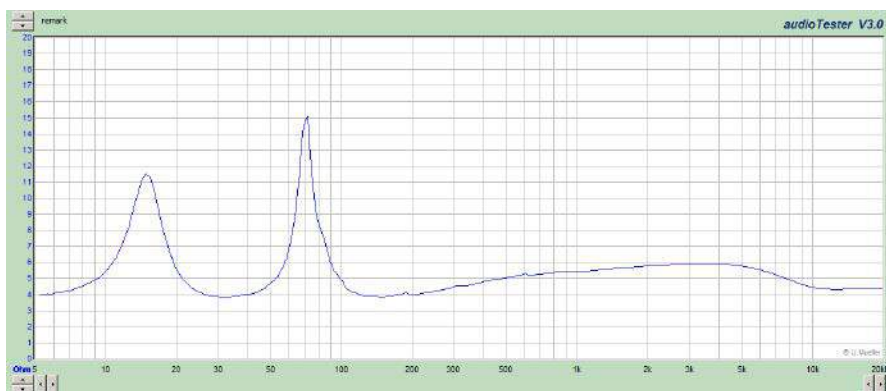
O b u d o w y tego zestawu głośnikowego w zasadzie nie różnią się od fabrycznie produkowanych. Możemy jedynie opcjonalnie zlecić wywiercenie na obrabiarce CNC wszystkich otworów montażowych o średnicach $\varnothing$ 5,0 mm, $\varnothing$ 5,5 mm lub $\varnothing$ 6,0 mm w przypadku gdy chcemy zamocować wszystkie elementy przy pomocy nakrętek



Rysunek 1.20. Wypadkowa charakterystyka poziomu ciśnienia akustycznego w funkcji częstotliwości zestawu głośnikowego w obudowie typu „Zeus” (wszystkie głośniki podłączone synfazowo)



Rysunek 1.21. Wypadkowa charakterystyka poziomu ciśnienia akustycznego w funkcji częstotliwości zestawu głośnikowego w obudowie typu „Zeus” (głośnik wysokotonowy podłączony afazowo)



Rysunek 1.22. Wypadkowa charakterystyka modułu impedancji w funkcji częstotliwości zestawu głośnikowego w obudowie typu „Zeus”

pazurkowych. Jeśli zamierzamy natomiast zamocować głośniki wysokotonowe oraz przyłączyć przy pomocy wkrętów do drewna, możemy z tej modyfikacji zrezygnować. Nie jest to jednak wskazane ze względu na to, że każdorazowy demontaż zestawu głośnikowego (np. podczas zaistnienia konieczności przeprowadzenia napraw lub regeneracji głośników) może sprawić, że otwory w płycie MDF się „wyrobią”. Na ściankach przyklejamy od wewnątrz butaprenem lub przymocowujemy zszywaczem tapicerskim watolinę.

Następnie przystępujemy do montażu zwrotnicy według schematu zaproponowanego na rysunku 1.16. Zwrotnicę najlepiej umieścić w każdej obudowie na jednej osobnej płycie (czyli na zestaw wyjdą dwie płyty tekstolitowe

bądź pilśniowe). Rozmieszczenie elementów pokazano na rysunku 1.17.

Opcjonalnie przed filtrem toru wysokotonowego można zamontować żarówkę samochodową od kierunkowskazów w celu zabezpieczenia głośnika wysokotonowego przed przeciążeniem.

Zwrotnicę przykręcamy za pośrednictwem nóżek dystansowych wkrętami do drewna do tylnej ścianki zestawu. Pod każdą płytę zwrotnicy podkładamy piankę poliuretanową.

Połączenia pomiędzy głośnikami a zwrotnicą realizujemy dwużyłowymi przewodami miedzianymi o następujących przekrojach:

- 2,5 mm² – tor niskotonowy,
- 1,0 mm² – tor wysokotonowy.



Rysunek 1.23. Okładka książki pt. „Wprowadzenie do projektowania układów zwrotnic zestawów głośnikowych. Poradnik praktyczny.”

Połączenie pomiędzy zwrotnicą a przyłączem – 2,5 mm².

Książka o zwrotnicach do zestawów głośnikowych

Zapraszam do zapoznania się z moją najnowszą książką pt. „Wprowadzenie do projektowania układów zwrotnic zestawów głośnikowych. Poradnik praktyczny.”:

- <https://bit.ly/3Zl38cy>
- <https://bit.ly/3VYCHql>
- <https://bit.ly/3ikG6SL>

oraz:

- <https://bit.ly/3jPXCyn>
- <https://bit.ly/3jWhWhH>

mgr inż. Tomasz Łysek

REKLAMA

przejrysz i kupisz na
www.ulubionykiosk.pl

m.technik

Ciekawi świata są zawsze młodzi

Wykrywanie obiektów z użyciem modułu Lidar

Radar – RAdio Detection And Ranging. To system wykrywania obecności i odległości obiektów z użyciem fal radiowych. Działanie radaru oparte jest na kalkulacji czasu echa, fali odbitej od napotkanych przeszkód. Mierząc to opóźnienie, układ oblicza w jakiej odległości jest obiekt. Mierząc przesunięcie częstotliwości Dopplera można obliczyć prędkość poruszania się obiektu w kierunku fali padającej; o czym wie każdy kierowca. **Lidar to Light Detection And Ranging** i działa na tej samej zasadzie. Różnica jest w tym, iż zamiast fali radiowej użyta jest fala świetlna lasera.

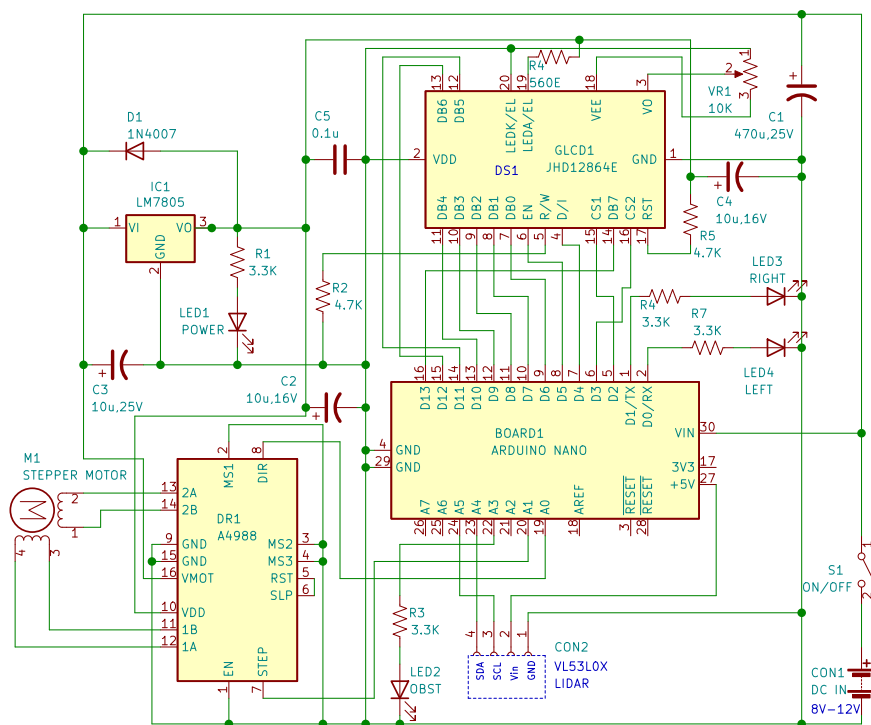
Projekt bieżącego DIY wykorzystuje koncepcję znajdowania obiektów na ww. zasadzie. To prosty a zarazem wydajny projekt, dzięki wykorzystaniu modułu lidar VL53L0X. Nadajnik na tym module emituje impulsy światła laserowego i oblicza czas echa jeśli wykryje impuls powrotny. Układ skanuje przestrzeń w stosunkowo szerokim zakresie kąta, a wynik przedstawia na graficznym wyświetlaczu LCD (GLCD Graphic Liquid Crystal Display).

Opis układu

Schemat ideowy układu wykrywania obiektów z użyciem modułu lidar pokazuje rysunek 1.

Zastosowany tu wyświetlacz graficzny ma rozdzielczość 128×64 piksele. Ekran ten jest podzielony na dwie połowy po 64×64 piksele. Aktywacja obu części realizowana jest za pośrednictwem wyprowadzeń CS1 i CS2 (piny 15 i 16). Komunikacja między wyświetlaczem a mikrokontrolerem jest ośmio-bitowa. Pełnymi bajtami przesyłane są dane i komendy, a rozróżnia je aktywacja wyprowadzenia statusu RS. W układzie zastosowano także potencjometr montażowy VR1, którym można ustawić kontrast wyświetlacza.

Moduł lidar zamontowano na osi silnika krokowego (co widać na rysunku 2). Zastosowano mały silnik NEMA17 o kroku



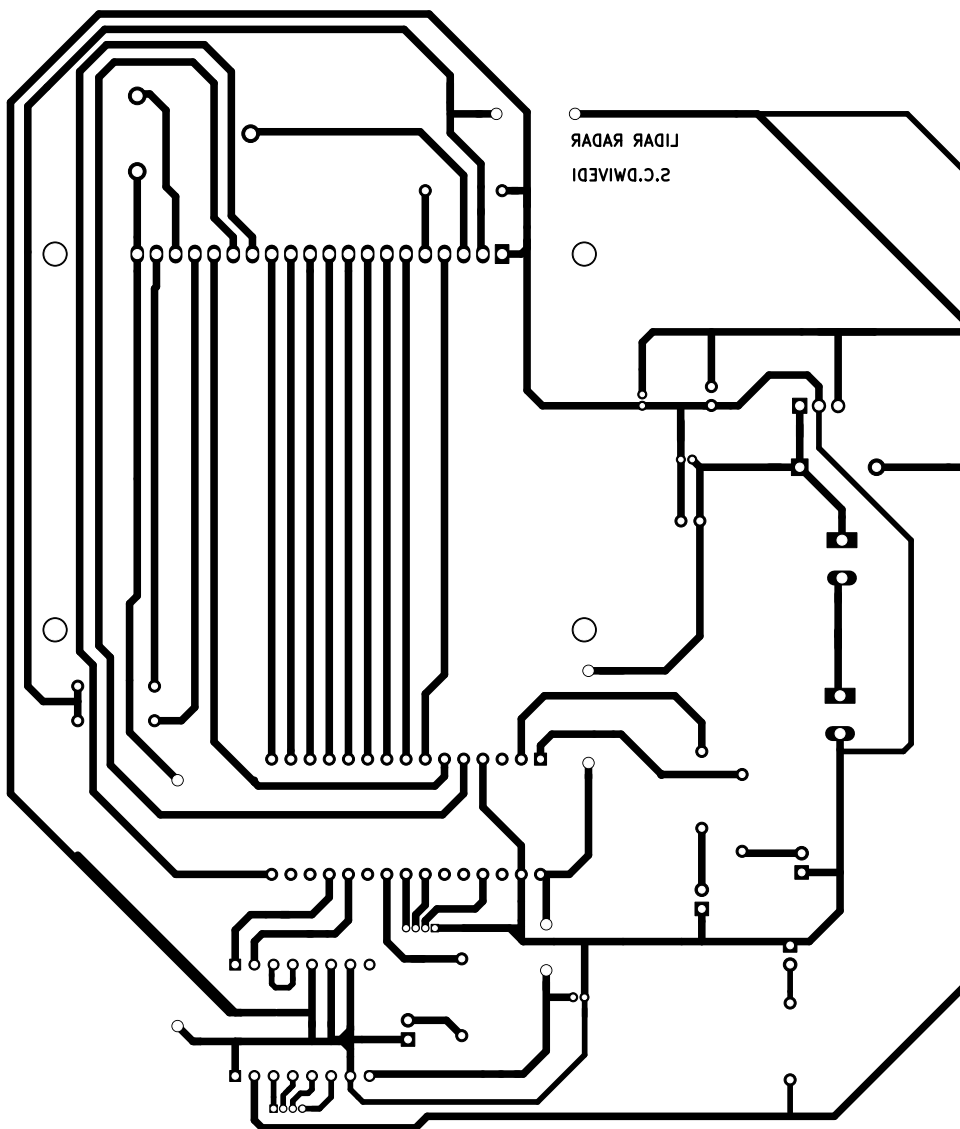
Rysunek 1. Schemat ideowy



Rysunek 2. Moduł LIDAR umocowany na osi silnika krokowego



Rysunek 3. Ekran prototypu wykonanego przez autora pokazuje odległość i kąt skanowania Lidaru



Rysunek 4. Płytkę PCB od strony druku w skali 1:1

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

Półprzewodniki:

IC1: LM7805 stabilizator napięcia 5 V
BOARD1: Arduino Nano
DR1: A4988 driver silnika krokowego
LIDAR: VL53L0X moduł LIDAR
D1: 1N4007 dioda prostownicza
LED1...LED4: diody LED 5 mm

Rezystory: (wszystkie 0,25 W ±5%)

R1, R3, R6, R7: 3,3 kΩ
R2, R5: 4,7 kΩ
R4: 560 Ω
VR1: 10 kΩ potencjometr montażowy

Kondensatory:

C1: 470 µF/25 V elektrolityczny
C2...C4: 10 µF/16 V elektrolityczny
C5: 0,1 µF ceramiczny

Inne:

CON1: złącze 2-pinowe
CON2: 4-ro pinowe żeńskie złącze typu Berg
GLCD1: 128×64 JHD12864E moduł wyświetlacza GLCD
M1: silnik krokowy (4-wire)
S1: przetątnik SPST
Ponadto: zasilacz 8 V do 12 VDC oraz radiator na IC1

1,8 stopnia i prądzie cewki 1 A (z możliwością ustawienia ograniczenia prądowego). Jako driver silnika krokowego pracuje układ scalony A4988, który może wystawiać cewki silnika prądem 2 A. Mózgiem układu jest Arduino Nano z mikrokontrolerem ATmega328 w wersji obudowy płaskiego montażu. Dzięki zastosowaniu złączy typu Berg connector łatwy jest montaż Arduino na PCB urządzenia.

Montaż i instalacja oprogramowania

Montaż elektryczny należy wykonać zgodnie ze schematem, zaś „szkic” oprogramowania należy wgrać do mikrokontrolera w tradycyjny dla Arduino sposób. Część prac mechanicznych sprowadza się do poprawnego montażu modułu lidar na osi silnika oraz montażu całości tak, aby lidar widział przewidziany zakres monitoringu. Płytkę PCB wykonano z drukiem jednostronnym, której

projekt pokazano na **rysunku 4**. Można wprost skorzystać z tego projektu, gdyż wymiary (w papierowej wersji EdW) powinny odpowiadać rzeczywistej skali 1:1. Na **rysunku 5** pokazano schemat ułożenia elementów na PCB.

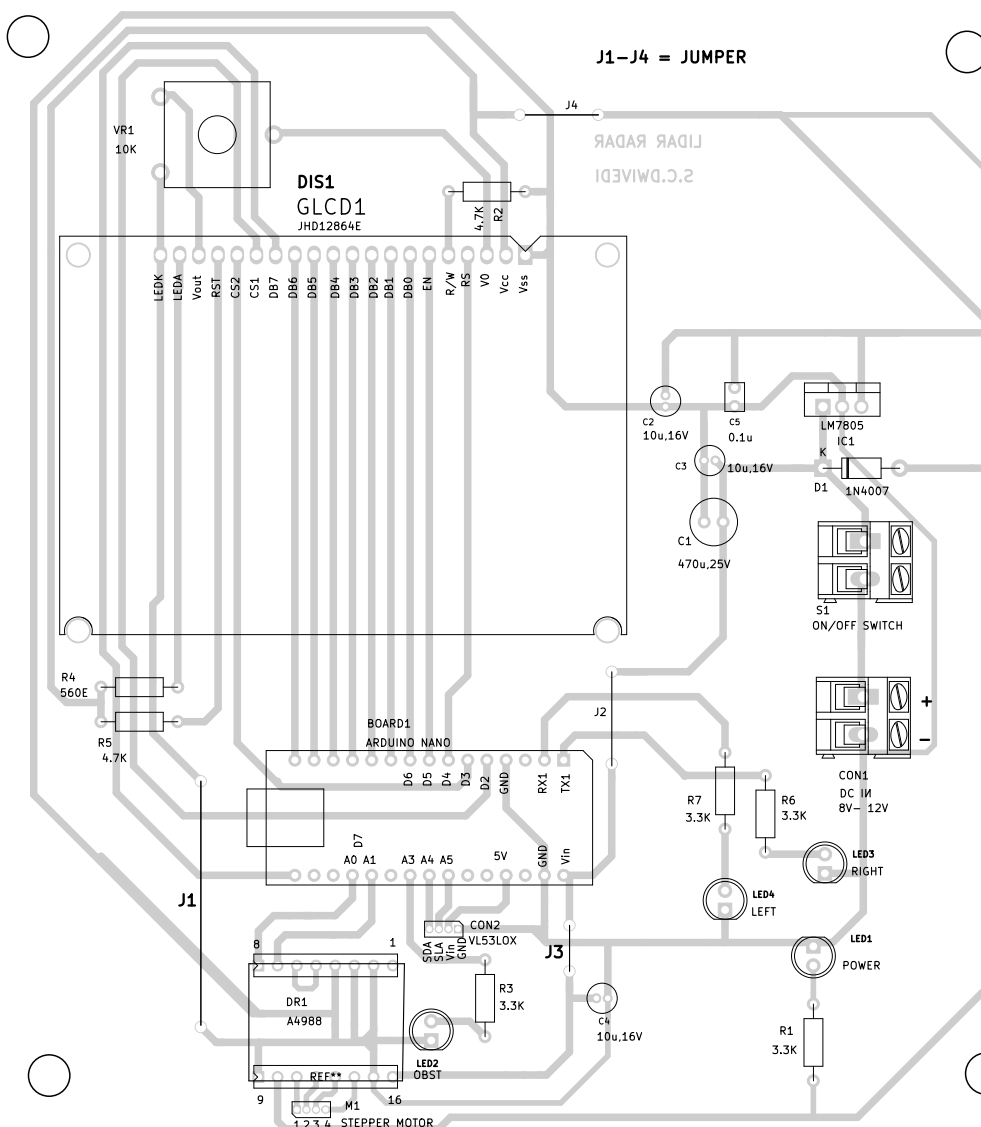
Po włączeniu zasilania (wyłącznikiem S1), na wyświetlaczu powinien pokazać się napis LIDAR RADAR, a silnik krokowy powinien ustawić się do „zerowej” pozycji. Następnie układ powinien przejść do skanowania przestrzeni celem wykrycia obiektów w zakresie pola widzenia. Silnik powinien obrócić swoją oś w obu kierunkach w zakresie kąta ok. 60 stopni. Lidar powinien widzieć obiekty w zakresie odległości do 1 metra. Mikrokontroler na Atmega przelicza odległość na centymetry i wyświetla tę informację na wyświetlaczu graficznym (do 99 cm). Przykładowy wynik na **rysunku 3** jest 40 cm, wraz z „promieniami” wskazującymi na katowy zakres obserwacji. Jeśli obiekt jest dalej niż 1 metr, na wyświetlaczu powinien pojawić się napis OUT.

Poza wyświetlaczem GLCD na płytce zamontowano cztery diody LED. LED1 sygnalizuje obecność zasilania. LED2 to OBST (Obstacle), co oznacza iż lidar dostrzegł obiekt w założonym polu widzenia. Obszar obserwacji podzielono na dwie części i świecenie LED3 oznacza obiekt w części lewej, zaś LED4 w prawej części pola widzenia lidar. Jeśli obiekt

znajduje się pośrodku, powinny świecić obie diody LED3 i LED4.

Są sytuacje, kiedy obserwacja światłem w zakresie widzialnym nie zdaje egzaminu. To przede wszystkim obiekty przezroczyste, których lidar nie zobaczy. Chcąc zmienić kierunek obrotu silnika krokowego wystarczy odwrotnie podłączyć jego cewki do drivera A4988. Wydajność prądowa drivera jest na tyle duża, że cewki silnika mogą ulec uszkodzeniu. Dlatego w razie takiej potrzeby można ustawić ograniczenie prądowe potencjometrem montażowym (szczegóły w tym zakresie można znaleźć w karcie katalogowej układu scalonego A4988). W zakresie oprogramowania też można dokonać kilku „regulacji”. W szczególności w „szkicu” można modyfikować takie dane jak: `const int STEPPER_DELAY = 100; // ustawienie opóźnienia skutkującego prędkością obrotu osi silnika krokowego`

**Kod źródłowy
tego projektu jest dostępny
do pobrania ze strony
<https://bit.ly/3Mf2V4P>**



Rysunek 5. Rozmieszczenie elementów na płycie PCB

```
const int
stepsPerRevolution =
200; // ustawienie
kąta obrotu jednego kroku
(wartość 200 odpowiada
kątowni 1,8 stopnia)
int ONESTEP=2; // liczba
impulsów wysyłanych
z drivere A4988
```

W zakresie oprogramowania należy dodatkowo ściągnąć kilka plików: VL53L0X.cpp, VL53L0X.h, GLCDfont.h i licence.txt.

Foldery te można znaleźć na stronie: <https://github.com/pololu/vl53l0x-arduino>.

Tak wykonany moduł obserwacji lidarowej może być użyteczny w projektach różnego rodzaju robotów, które powinny widzieć przestrzeń w której się poruszają. ■

Fayaz Hassan

Ten artykuł był wcześniej opublikowany na łamach „EFY”, czerwiec 2022 (efymag.com)

REKLAMA

świat radio

Magazyn wszystkich użytkowników eteru
KRÓTKOFALARSTWO CB RADIOTECHNIKA

przejrzyj i kup na
www.ulubionykiosk.pl

WYDANIE 11-12/22





Smart Glasses na wzór Google Glass z wykorzystaniem wyświetlacza OLED

Google Glass zostały stworzone kilka lat temu. Wyglądają podobnie do zwykłych okularów, lecz wykonane są ze specjalnego materiału i mają interfejs użytkownika UI (User Interface). Google Glass realizują kilka standardowych funkcji typowych dla smartfonów.

Wśród nich wyświetlanie czasu i daty, odczytanie wiadomości, aczkolwiek możesz też czytać książkę itp. Obraz tworzony przed oczami może być nałożony na inny obraz lub robić wrażenie zawieszenia w powietrzu. Obraz ten nie przesłania widoku realnego świata, który widzisz mając „na nosie” zwykłe okulary. Google Glass stały się rewelacyjnym wynalazkiem i jednym z najbardziej zaawansowanych technicznie gadżetów dostępnych w obecnym czasie. Jest to jednak gadżet drogi, więc wiele osób nie ma możliwości doświadczenia wrażeń jakich dostarcza. Bieżący DIY jest próbą wykonania

substytutu, który nazwiemy Smart Glasses (jako że obecnie, wszystko musi być smart – przypis Redakcji). Rysunek 1 pokazuje zdjęcie oryginalnych Google Glass-es, a rysunek 2 prototypu autora Smart Glasses.

Wykaz potrzebnych podzespołów:

- Raspberry Pi Zero W
- Przezroczysty wyświetlacz OLED o przekątnej 29 mm
- Eye glass – zwykłe okulary

Jak widać, układ bazuje na mikrokomputerku Raspberry Pi, a „image” jest tworzony na „Transparent OLED” czyli przezroczystym wyświetlaczu w technologii OLED.

Zatem, choć urządzenie jest skomplikowane, schemat jest zaskakująco prosty, co widać na rysunku 3.

Okulary są niejako „stelażem”, na którym montujemy wyświetlacz. Można w tym celu usunąć oryginalne szkło lub soczewkę korekcyjną okularów. Schemat pokazuje interfejs między wyświetlaczem i Raspberry Pi, aczkolwiek jest jeszcze potrzebna bateria w celu zasilania gadżetu.

Oprogramowanie

Jako że wybrano mikrokomputer Raspberry Pi, naturalnym środowiskiem programistycznym jest Python, i w tym języku napisano oprogramowanie naszych Smart Glasses. Zainstalowano najnowszą wersję systemu operacyjnego Raspbian wraz z Python-em w wersji 3. Jeśli nie dysponujesz najnowszą wersją, zainstaluj najpierw Pythona wraz z jego „środowiskiem” IDE. Z bogatej biblioteki należy ściągnąć składniki potrzebne naszym „okularom”. W tym zakresie można kierować się poniższymi wskazówkami, które reprezentują listę

-minimum. Umożliwia bowiem, wyświetlanie tylko daty/czasu i informacji tekstowych na tle obrazu przed naszymi oczami.

W celu zainstalowania modułów Pythona, należy uruchomić oprogramowanie pod Linux-em i wpisać następujące komendy:

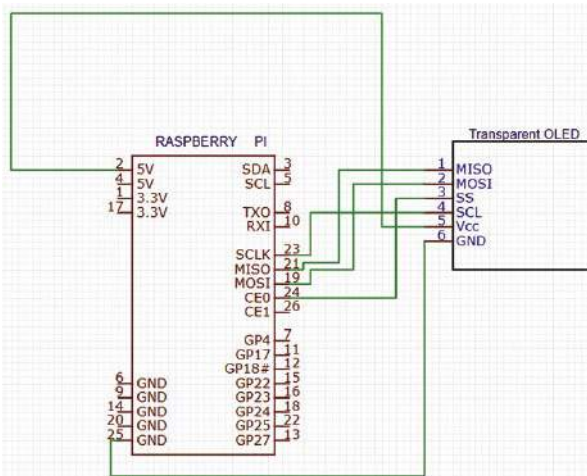
```
wget http://www.airspayce.com/mikem/bcm2835/bcm2835-1.71.tar.gz
tar zxvf bcm2835-1.71.tar.gz
cd bcm2835-1.71/
sudo ./configure && sudo make &&
sudo make check && sudo make
install
```

```
sudo apt-get install wiringpi
```

Jeśli system OS Raspberry Pi jest starszy niż maj 2019, możesz dodatkowo potrzebować:



Rysunek 1. Oryginalne Google Glass



Rysunek 3. Schemat ideowy



Rysunek 2. Prototyp „Smart Okularów” autora

```
File Edit Format Run Options Window Help
import datetime
picdir = os.path.join(os.path.dirname(os.path.realpath(__file__)), 'pic')
libdir = os.path.join(os.path.dirname(os.path.realpath(__file__)), 'lib')
if os.path.exists(libdir):
    sys.path.append(libdir)

import logging
import time
import traceback
from waveshare_OLED import OLED_1in5
from PIL import Image, ImageDraw, ImageFont
logging.basicConfig(level=logging.DEBUG)

try:
    disp = OLED_1in5.OLED_1in5()
    x = datetime.datetime.now()
```

Rysunek 4. Zrzut ekranu fragmentu programu

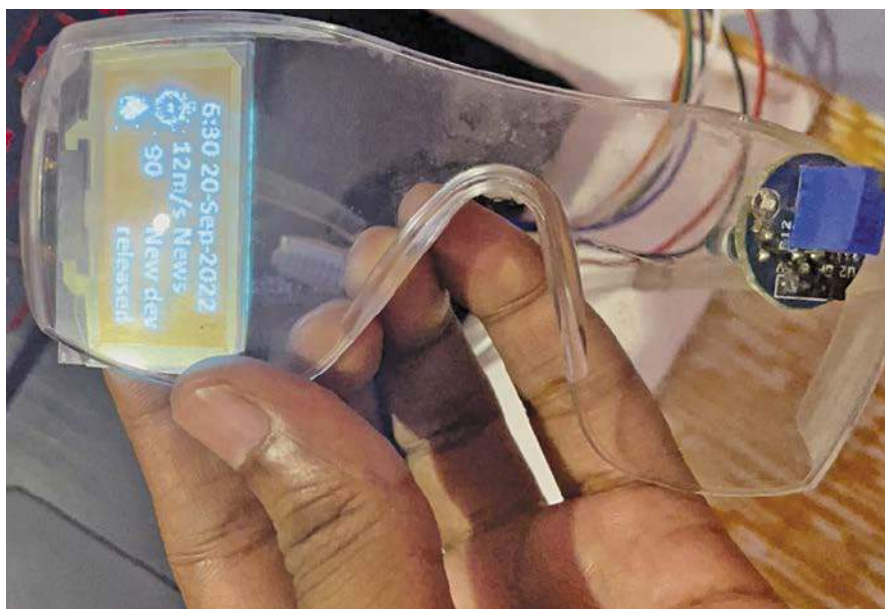
```
File Edit Format Run Options Window Help
draw = ImageDraw.Draw(image1)
font = ImageFont.truetype(os.path.join(picdir, 'Font.ttc'), 12)
font1 = ImageFont.truetype(os.path.join(picdir, 'Font.ttc'), 18)
font2 = ImageFont.truetype(os.path.join(picdir, 'Font.ttc'), 24)
logging.info ("****draw line")
draw.line([(0,0),(127,0)], fill = 15)
draw.line([(0,0),(0,127)], fill = 15)
draw.line([(0,127),(127,127)], fill = 15)
draw.line([(127,0),(127,127)], fill = 15)
logging.info ("****draw text")
draw.text((20,0), str(x), font = font1, fill = 15)

image1 = image1.rotate(0)
disp.ShowImage(disp.getbuffer(image1))
time.sleep(3)

logging.info ("****draw rectangle")
image1 = Image.new('L', (disp.width, disp.height), 0)
draw = ImageDraw.Draw(image1)
for i in range(0, 16):
    draw.rectangle([(0, 8*i), (128, 8*(i+1))], fill = i)
disp.ShowImage(disp.getbuffer(image1))
time.sleep(3)

logging.info ("****draw image")
Himage2 = Image.new('L', (disp.width, disp.height), 0) # 0: clear the frame
bmp = Image.open(os.path.join(picdir, '1in5.bmp'))
Himage2.paste(bmp, (0,0))
Himage2=Himage2.rotate(0)
```

Rysunek 5. Inny fragment oprogramowania



Rysunek 6. Gotowy prototyp poddany testom

wget https://project-downloads.drogon.net/wiringpi-latest.deb
sudo dpkg -i wiringpi-latest.deb
Poprawność instalacji sprawdzisz prostą komendą:

gpio -v

Teraz zainstaluj następujące moduły Pythona
sudo apt-get install python3-pil
sudo apt-get install python3-numpy
sudo pip3 install RPi.GPIO
sudo pip3 install spidev

Biblioteka dla przeźroczystego wyświetlacza OLED zawiera:
sudo apt-get install p7zip-full
sudo wget https://www.waveshare.com/w/upload/2/2c/OLED_Module_Code.7z
7z x OLED_Module_Code.7z -o./OLED_Module_Code
cd OLED_Module_Code/RaspberryPi

Fragmenty kodu programu pokazują zrzut ekranu komputera na rysunkach 4 i 5. Wpierw trzeba zaimportować moduły biblioteki Pythona. W prototypie ściągnięto bibliotekę daty-czasu i obsługę zainstalowanego wyświetlacza OLED. Chcąc rozbudować funkcjonalność „smart okularów” można ściągnąć wiele dodatkowych modułów oprogramowania, które dostępne są w środowisku Pythona.

Uwaga: Kod źródłowy można ściągnąć spod następującego adresu:
<https://www.electronicsforu.com/electronics-projects/smart-google-glasses-with-transparent-oled-display>

Konstrukcja mechaniczna i przetestowanie działania Smart-okularów

Po wykonaniu wyżej wymienionych czynności nasze okulary powinny być gotowe do przetestowania. Po podłączeniu zasilania i uruchomieniu kodu programu, na wyświetlaczu powinny pojawić się informacje o bieżącym czasie i dacie wraz z innym tekstem powiązanym z zainstalowanymi modułami oprogramowania. Obraz przed oczami może przypominać sceny z filmów science fiction, które pewnie wielokrotnie widziałeś. Tekst może być nałożony na inny obraz lub być zawieszony – „pływać w powietrzu”. Na rysunku 6 jest zdjęcie prototypu okularów wykonanych przez autora i poddanych testom. To jest pierwsza wersja projektu. Można ją rozbudować dodając funkcje sterujące za pomocą ruchów oczami lub np. funkcje wyświetlania „powiadomień” przychodzących na twój telefon. ■

Ashwini Kumar Sinha

Ten artykuł był wcześniej opublikowany na łamach „EFY”, listopad 2022 (efymag.com)

Sterowanie komputerem ruchem oczu

Pełny tytuł tego DIY to „sterowanie komputerem ruchem oczu – urządzenie dla osób niepełnosprawnych”. Do-pisek wydaje się równie logiczny jak i zbędny, gdyż tylko w takim zastosowaniu opisane tu urządzenie ma sens. Ruch oczu ma bowiem zastąpić ruch myszką.

Jest na tym Świecie zapewne tysiące (lub więcej) osób, których niepełnosprawność uniemożliwia jednak interfejs z komputerem przy pomocy tradycyjnych narzędzi „IN”,

z których tzw. myszka znalazła największą popularność.

Opisany tu DIY umożliwia pracę na kom-puterze nawet takim osobom. Możesz być

sparaliżowany „od pasa, a nawet powyżej pasa”. Wystarczy ruch głową i oczami. I takie osoby w pełni docenią ten wynalazek.

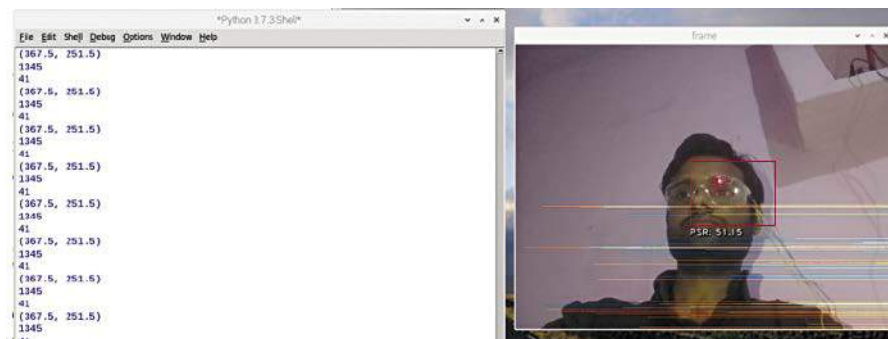
Zatem, urządzenie tu opisane zastąpi tra-dycyjną myszkę i działa nawet przy słabym oświetleniu bądź w pełnej ciemności. Musi być więc możliwość przesuwu kursora po całym ekranie monitora i kliknięcie na określonej pozycji. Mrugnięcie okiem powinno odpow-iać kliknięciu lewym przyciskiem myszy.

Współrzędne położenia kursora mieszczą się w zakresie $(x, y) = (0, 0)$ do $(1080, 720)$; tj. lewy-dolny i prawy-górny róg ekranu (współ-rzędna pozioma, pionowa). Działanie układu (a przede wszystkim programu) składa się z trzech części.

1. Rozpoznanie przez system operacyjny mrugnięcia prawym okiem.
2. Detekcja ruchu gałki ocznej przez system rozpoznawania obrazu „image processing”.
3. Przesłanie powyższych informacji do kom-puteru pod kontrolę graficznego interfejsu GUI Graphic User Interface.

Rozwiązanie głównych problemów

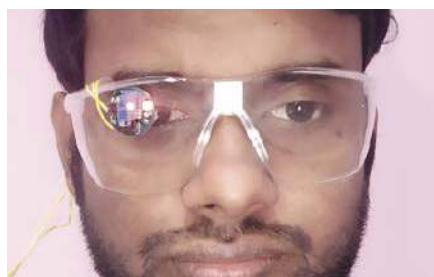
Realizacja wyżej nakreślonego zada-nia napotyka na kilka istotnych trudności.



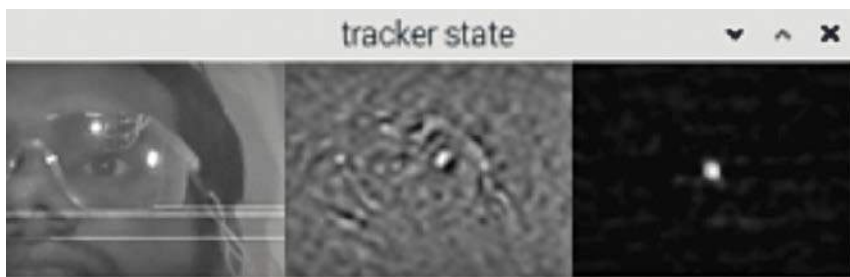
Rysunek 1. Współrzędne odczytane na podstawie ruchu gałek ocznych autora



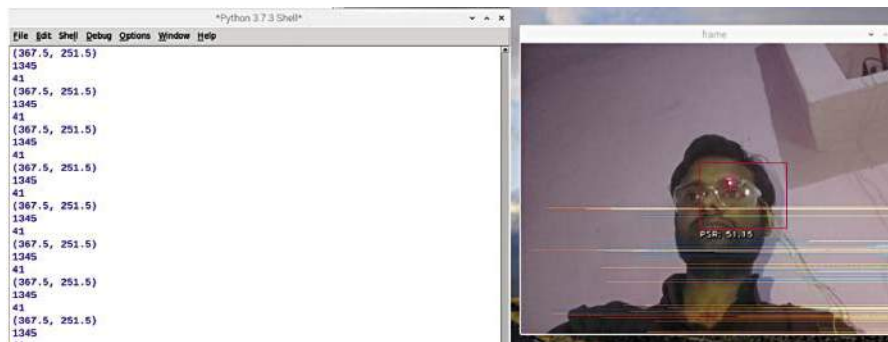
Rysunek 2. Autor testuje działanie czujnika zamknięcia powieki prawego oka



Rysunek 3. Zdjęcie twarzy autora w okularach



Rysunek 4. Obraz gałki ocznej autora



Rysunek 5. Autor testuje urządzenie

Przezwyciężenie ich jest konieczne dla po-prawnej funkcjonalności urządzenia i jest zaspokojone głównie na drodze programowej

Główne problemy to:

Rozróżnienie między odruchowym a ce-lowym mrugnięciem powieki

Naturalnym jest, iż człowiek mruga oczami co kilka sekund i trudno powstrzymać ten odruch. Jest to pewnie niezbędne do nawilżenia gałki ocznej oraz usuwania na bieżąco drob-nego pyłu czy kurzu. Ale to sprawa medyczna, która nas nie interesuje. Dla poprawy działania urządzenia trzeba odróżnić

```
# Python 2/3 compatibility
from __future__ import print_function
import sys
PY3 = sys.version_info[0] == 3

if PY3:
    xrange = range
from pynput.mouse import Button, Controller
from nmap import nmap

import numpy as np
import cv2 as cv
from common import draw_str, RectSelector
import video
from gpiozero import Button as butt
from signal import pause
import mouse
import time

prex=0
prey=0
button = butt(27)
mouse = Controller()

mapfnx = nmap(290, 400, 10, 1905, normfn=int)
mapfny = nmap(250, 300, 10, 1059, normfn=int)
```

Rysunek 6. Zrzut ekranu fragmentu programu

Rysunek 7. Fragment programu przestania współzrzednych do oprogramowania GUI komputera

```
class MOSSE:
    def __init__(self, frame, rect):
        x1, y1, x2, y2 = rect
        w, h = map(cv.getOptimalDFTSize, [x2-x1, y2-y1])
        x1, y1 = (x1+x2-w)//2, (y1+y2-h)//2
        self.pos = x, y = x1+0.5*(w-1), y1+0.5*(h-1)
        self.size = w, h
        img = cv.getRectSubPix(frame, (w, h), (x, y))

        self.win = cv.createHanningWindow((w, h), cv.CV_32F)
        g = np.zeros((h, w), np.float32)
        g[h//2, w//2] = 1
        g = cv.GaussianBlur(g, (-1, -1), 2.0)
        g /= g.max()

        self.G = cv.dft(g, flags=cv.DFT_COMPLEX_OUTPUT)
        self.H1 = np.zeros_like(self.G)
        self.H2 = np.zeros_like(self.G)
        for _i in xrange(128):
            a = self.preprocess(rnd_warp(img))
            A = cv.dft(a, flags=cv.DFT_COMPLEX_OUTPUT)
            self.H1 += cv.mulSpectrums(self.G, A, 0, conjB=True)
            self.H2 += cv.mulSpectrums(A, A, 0, conjB=True)
        self.update_kernel()
        self.update(frame)
```

Rysunek 8. Inny fragment programu

ruchy mimowolne od celowych. Nasuwają się dwie koncepcje rozwiązania tego problemu.

Mruganie mimowolne cechuje krótkie mrugnięcia obiema powiekami jednocześnie. Można zatem zastosować detektory na obu oczach i ignorować takie ruchy, co na drodze programowej nie jest trudne. Jeśli zostanie zidentyfikowane mrugnięcie tylko jednym okiem, program wykona funkcję, jakoby został naciśnięty lewy lub prawy klawisz myszy. Rozwiązanie takie byłoby zapewne efektywne, aczkolwiek wymaga detektorów na obu oczach, co komplikuje i podraża cały hardware. Ponadto, uniemożliwiłoby to korzystanie ze „zdalnej myszy” przez osoby mające jakiś problem z jednym okiem. Lepszym wydaje się detekcja czasu mrugnięcia powieką. To nie powinno być trudne a zarazem skuteczne. Dlatego iż naturalne mrugnięcia mimowolne są bardzo krótkie. Rozwiązanie w oparciu o tę koncepcję jest też tańsze, gdyż w całości można je zrealizować na drodze programowej.

Można zdecydować, jeśli mrugnięcie będzie krótsze niż jedna sekunda, należy je zignorować. Jeśli będzie dłuższe, uznać za intencjonalne. Aczkolwiek przyjęcie progu nawet na poziomie czterech sekund jest do przyjęcia

```
def update(self, frame, rate = 0.125):
    (x, y), (w, h) = self.pos, self.size
    self.last_img = img = cv.getRectSubPix(frame, (w, h), (x, y))
    img = self.preprocess(img)
    self.last_resp, (dx, dy), self.psr = self.correlate(img)
    self.good = self.psr > 8.0
    if not self.good:
        return

    self.last_img = img = cv.getRectSubPix(frame, (w, h), self.pos)
    img = self.preprocess(img)
    A = cv.dft(img, flags=cv.DFT_COMPLEX_OUTPUT)
    H1 = cv.mulSpectrums(self.G, A, 0, conjB=True)
    H2 = cv.mulSpectrums(A, A, 0, conjB=True)
    self.H1 = self.H1 * (1.0-rate) + H1 * rate
    self.H2 = self.H2 * (1.0-rate) + H2 * rate
    self.update_kernel()

@property
def state_vis(self):
    f = cv.idft(self.H, flags=cv.DFT_SCALE | cv.DFT_REAL_OUTPUT)
    h, w = f.shape
    f = np.roll(f, -h//2, 0)
    f = np.roll(f, -w//2, 1)
    kernel = np.uint8((f-f.min()) / f.ptp())*255)
    resp = self.last_resp
    resp = np.uint8(np.clip(resp/resp.max(), 0, 1)*255)
    vis = np.hstack([self.last_img, kernel, resp])
    return vis

def draw_state(self, vis):
    (x, y), (w, h) = self.pos, self.size
    x1, y1, x2, y2 = int(x-0.5*w), int(y-0.5*h), int(x+0.5*w), int(y+0.5*h)
    cv.rectangle(vis, (x1, y1), (x2, y2), (0, 0, 255))
    if self.good:
        cv.circle(vis, (int(x), int(y)), 2, (0, 0, 255), -1)
        #print(x,y)
        self.pos = x, y
        print(self.pos)
        #mouse.move(x+dx,y+dy)
        xmove=mapfnx(x)
        ymove= mapfny(y)
        print(xmove)
        print(ymove)
        mouse.position = (xmove,ymove)

        #mouse.position = (x,y)
```

Rysunek 9. Kod programu zastępujący działanie tradycyjnej myszy

i na pewno nie wystąpi pomyłka niecelowego uruchomienia kursora myszy.

Detekcja ruchu gałki ocznej z użyciem zobrazenia „image processisng”

Drugim problemem do rozwiązania, trudniejszym od kliknięcia myszą, jest przesuw kursora po ekranie monitora/komputera. Trzeba w jakiś sposób powiązać położenie gałki ocznej z obiektem na ekranie. To musi być „połączenie sztywne”, gdyż jego „zerwanie” może być irytujące lub wręcz uniemożliwiające obsługę komputera wg nakreślonej koncepcji. Ale samo zobrazenie, odwzorowanie położenia oka jest mało precyzyjne względem wymaganej rozdzielczości. Proces ten komplikuje się, szczególnie jeśli założymy, że układ ma również działać w warunkach słabego oświetlenia lub nawet w ciemności. Dla przezwyciężenia tych trudności założono punktowe, słabe źródło światła na sensorze mocowanym do okularów.

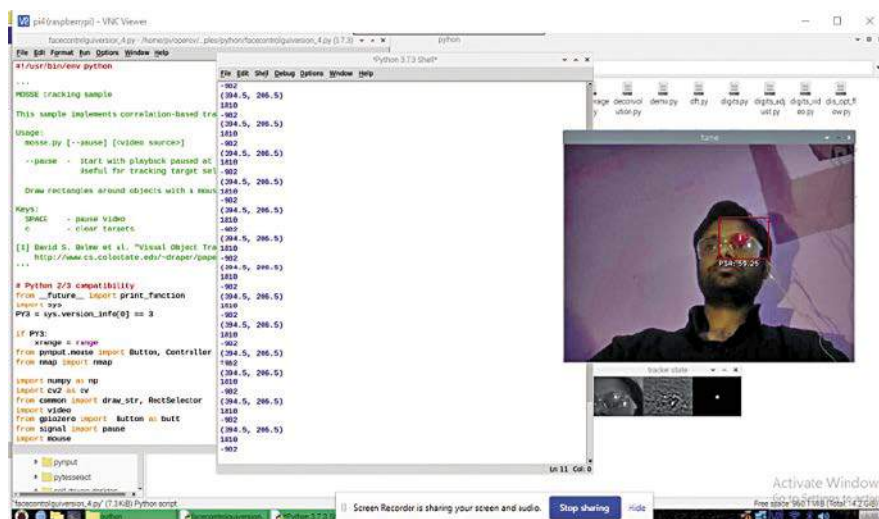
Przesłanie do komputera informacji o ruchu i mrugnięciu okiem

Nawet tradycyjna obsługa komputera jest graficzna. Odbywa się poprzez rozmieszczenie

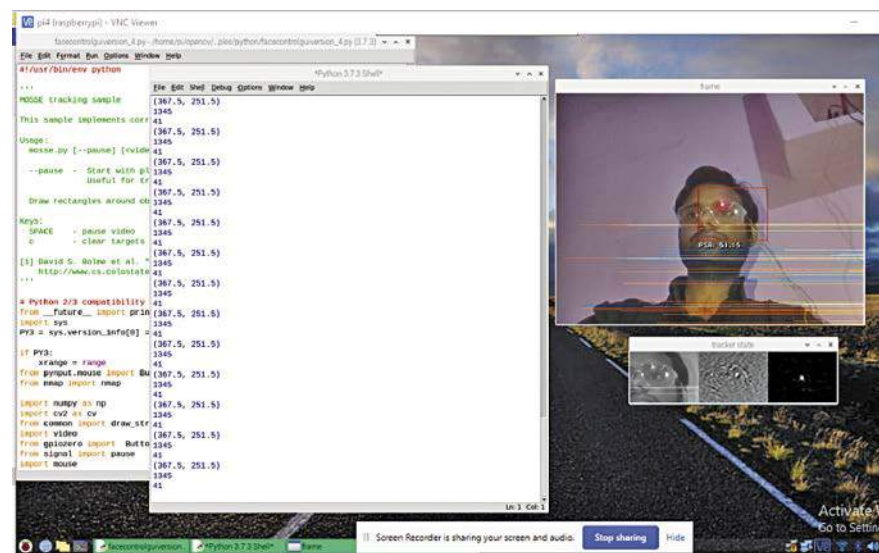
```
cv.imshow('frame', vis)
ch = cv.waitKey(10)
if ch == 27:
    break
if ch == ord(' '):
    self.paused = not self.paused
if ch == ord('c'):
    self.trackers = []
if button.is_pressed:
    press=0
    # print("eyeblink")
    for x in range(0, 5):
        if button.is_pressed:
            press =press+1
            #print(press)
            if press==5:
                mouse.click(Button.left, 2)

    #print("mousepress")
```

Rysunek 10. Fragment kodu zastępujący kliknięcie przyciskiem myszy



Rysunek 11. Autor testuje oprogramowanie zastępujące działanie tradycyjnej myszy



Rysunek 12.

ikonkę na ekranie. Należy zatem powiązać informację z „detektora oka” z oprogramowaniem GUI.

Wymagania w zakresie oprogramowania

Układ wykorzystywał mikrokomputer Raspberry Pi i będzie bazował na oprogramowaniu Pythona. Należy więc przygotować pamięć „SD card” z systemem operacyjnym Raspbian i sprawdzić czy program działa w środowisku Pythona IDLE. Teraz należy zainstalować następujące biblioteki:

- Open CV
- pynput
- numpy
- gpiozero

Instalację tą można wykonać wpisując następujące komendy:

```
sudo pip3 install python-opencv
sudo pip3 install numpy
sudo pip3 install gpiozero
```

```
sudo pip3 install pynput
sudo pip3 install nmap
```

Raspberry Pi będzie też potrzebował składowanie GitHub, którą można łączyć wpisując komendę:

```
git clone https://github.com/
opencv/opencv
```

Teraz można przystąpić do ostatniego kroku „kodowania”.

Dołączenie kodu programu

W bibliotece OpenCV znajdziesz przykładowe oprogramowanie myszy, które należy zmodyfikować. Po wybraniu kolejno OpenCV folder → platforms folder → python znajdziesz „kod” mouse.py. Należy skopiować ten kod i wkleić do pliku tekstowego, któremu należy nadać nazwę headcontrol-GUI.py. W oprogramowaniu GUI pracującym pod systemem operacyjnym Raspberry Pi trzeba utworzyć wirtualny port myszy. Tym zajmie się ściągnięty moduł biblioteki pynput.

Obsługę mrugnięcia okiem „przetworzonym” na kliknięcie przyciskiem myszy realizuje moduł gpiozero. Modyfikacja kodu programu sprowadza się do wstawienia komendy „if”, czy status oka jest = blink (powieka opuszczona). Jeśli tak, należy dodać opóźnienie (powiedzmy) czterech sekund sprawdzając, czy oko jest nadal zamknięte. Jeśli warunek jest spełniony (=tak), należy przejść do komendy pynput, która w wirtualnej myszy ma wykonać reakcję odpowiadającą kliknięciu lewym przyciskiem myszy.

Konstrukcja urządzenia i przetestowanie działania

Dla umocowania sensora w pobliżu oka wykorzystano okulary. Najwygodniejszy wydaje się wybór okularów ochronnych bez soczewek, tj. z „naturalnymi szklami” (aczkolwiek okulary tego typu wykonane są w całości z plastiku). Czujnik należy przykleić na „prawo oko”, a przewody podłączyć do pinów portu GPIO płytki Raspberry Pi. Ponadto z Raspberry trzeba połączyć kamerę, którą należy umieścić w takiej odległości od twarzy (centralnie), aby wizjer obejmował całą twarz. Należy założyć okulary i po uruchomieniu kodu programu zgrać czerwony punkt świetlny (jak na rysunku 3) ze wskaźnikiem myszy na ekranie. Testowanie urządzenia sprowadza się do sprawdzenia czy ruchy góra/dół i lewo/prawo gałką ocną przesuną kursor po całym ekranie, a zamknięcie prawego oka na cztery sekundy „zatwierdzi” położenie kursora.

Na rysunku 1 mamy fragment kodu, gdzie czerwony punkt na okularach autora przetworzony jest na współrzędne kursora na ekranie. Na rysunku 2 autor testuje „blink sensor” zamykając prawe oko. Na rysunku 2 i 3 widzimy prototyp urządzenia połączony z Raspberry Pi, a na rysunku 4 „image processing” punktu świetlnego. Na rysunku 5 mamy także odwzorowanie współrzędnych ekranu, a na rysunkach od 6 do 12, fragmenty kodu w programie Pythona. ■

Ashwini Kumar Sinha

Ten artykuł był wcześniej opublikowany na łamach „EFY”, czerwiec 2022 (efymag.com)

REKLAMA

PIPEK DRĘCZYCIEL

AVTEDU625

sklep.avt.pl

Przedstawiamy początkowe fragmenty dwóch projektów ze zbioru kilkudziesięciu projektów dostępnych wyłącznie dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl. W rubryce **DIY PLUS** zamieszczamy aktualnie najciekawsze projekty publikowane w Internecie w formacie open source. Prenumeratorów EdW zapraszamy do zapoznania się na www.elportal.pl z niezwykle inspirującymi zasobami rubryki **DIY PLUS**.

Uniwersalny wskaźnik włączonego biegu manualnej skrzyni biegów w motocyklu

Położenie dźwigni manualnej skrzyni biegów motocykla, w odróżnieniu od samochodowej nie reprezentuje jednoznacznie włączonego biegu. Biegi przełącza się lewą stopą i trudno wykonać inny ruch, jak tylko w górę lub w dół. Prezentowany projekt pozwala na jednoznaczną indykację numeru włączonego biegu na jednocyfrowym wyświetlaczu siedmiosegmentowym. Ponadto jest na tyle uniwersalny, że pozwoli na montaż w niemal każdym typie/modelu motocykla.

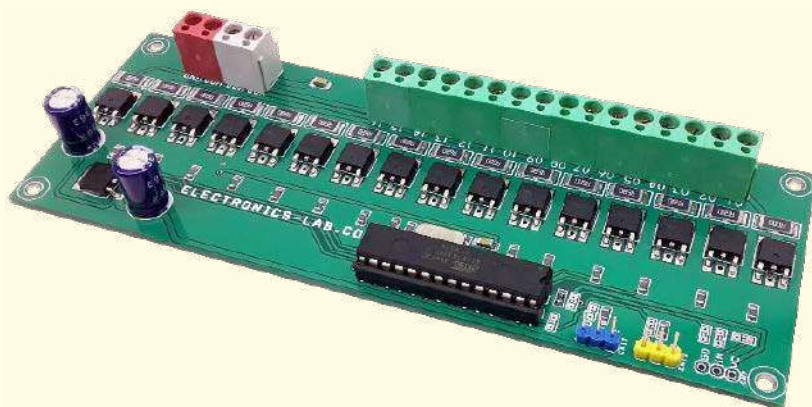
Dokończenie artykułu na stronie:
<https://bit.ly/3Wj7uhE>



Automatyczne Led-owe oświetlenie klatki schodowej z wykorzystaniem Arduino

Pokazana tu pomysłowa idea oświetlenia klatki schodowej jest przewidziana dla 16 schodów, gdzie na każdym zainstalowano dwunastowoltowy zestaw diod LED. Przy wejściu i zejściu ze schodów zainstalowano czujniki optyczne informujące system, że oświetlenie należy włączyć. Programowe rozwiązanie pozwala na „szeroki koncert życzeń”. W pokazanym systemie przewidziano progresywne rozświetlanie kolejnych stopni w zależności od tego w którą stronę podążasz. Reakcja czujnika wejścia/zejścia gasi kolejne LED-y w odwrotnej kolejności i w narzuconych przez program interwałach czasowych.

Dokończenie artykułu na stronie:
<https://bit.ly/3QK8gTR>



Niektóre projekty aktualnie dostępne tylko dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl:

- | | | |
|--|--|---|
| <ol style="list-style-type: none"> 1. RPi – stacja pogodowa IoT 2. Niskobudżetowy monitor jakości powietrza IoT oparty o RaspberryPi 4 3. Automatyczny system ogrodniczy z NodeMCU i Blynk, ArduFarmBot 2 4. TinyML – Rozpoznawanie ruchu przy pomocy Raspberry Pi Pico 5. Wzmacniacz piezoelektryczny do gitary i skrzypiec 6. Wysokowydajny i niezawodny sterownik bipolarnego silnika krokowego | <ol style="list-style-type: none"> 7. Sterownik silnika prądu stałego z wykorzystaniem przełącznika i mosfetu – interfejs Arduino 8. Przedwzmacniacz do mikrofonu MEMS 9. Super prosty czuły wykrywacz metali 10. Stymulator czaszkowy Arduino (Bio-BrainTuner) 11. Izolowany obwód wykrywania napięcia 250 V AC z pojedynczym wyjściem (wejście 250 V prądu przemiennego, wyjście 5 V) 12. Generator sygnałów AD9833 13. Obserwacja charakterystyk tranzystora | <ol style="list-style-type: none"> 14. Wyświetlacz EKG z użyciem Arduino 15. Łatwy do zbudowania robot kroczący 16. Sonarowy theremin MIDI 17. Zamek elektroniczny na kod 18. Prosty tester tranzystorów 19. Zegar binarny z użyciem Microbit 20. Przetwornik częstotliwości na napięcie (tachometr) – przetwornik częstotliwości na napięcie z czujnikiem magnetycznym o zmiennej reluktancji |
|--|--|---|

Miesięcznik „Elektronika dla Wszystkich” (12 numerów w roku) jest wydawany we współpracy z kilkoma redakcjami zagranicznymi



Wydawnictwo:
AVT-Korporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Adres redakcji:
03-197 Warszawa, ul. Leszczyńska 11
e-mail: edw@elportal.pl, www.elportal.pl

Redaktor merytoryczny
Paweł Sujko

Dział Reklamy:
Katarzyna Gugała
katarzyna.gugala@elportal.pl, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobański
jakub.sobanski@elportal.pl

Copyright AVT-Korporacja Sp. z o.o., Warszawa, ul. Leszczyńska 11. Projekty publikowane w „Elektronice dla Wszystkich” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki dla Wszystkich”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice dla Wszystkich” jest dozwolone wyłącznie po uzyskaniu pisemnej zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice dla Wszystkich”.

DTP, okładka, redakcja strony internetowej www.elportal.pl:
MAD Sp. z o.o.

Prenumerata:
W Wydawnictwie AVT, e-mail: prenumerata@avt.pl
tel. 22 257 84 22, (godz. 10:00–14:00)

W RUCH S.A., e-mail: prenumerata@ruch.com.pl
tel. 801 800 803, 22 717 59 59, www.prenumerata.ruch.com.pl



Ulubiony Kiosk - Twoje internetowe centrum prasy specjalistycznej i hobbystycznej!

Poznaj najważniejsze na rynku tytuły z segmentów

ZDROWIE I RODZINA
DOM, OGRÓD I WNĘTRZA
FOTOGRAFIA, EDUKACJA I HI-TECH
ELEKTRONIKA I AUTOMATYKA
MUZYKA I DŹWIĘK

Sprawdź na **UlubionyKiosk.pl**