



# ELEKTRONIKA

*dla wszystkich*

nr 9/2024 (344) • wrzesień • [www.elportal.pl](http://www.elportal.pl)

**DIY PLUS**  
tylko dla prenumeratorów

## Wzmacniacz mocy 500 W

### PROJEKTY dla elektroników

- ▶ Wzmacniacz zasięgu pilota na podczerwień
- ▶ Wielokanałowy „ochraniacz głośników”
- ▶ Stroboskop RGB z Arduino. Kolorowa adaptacja użytecznego instrumentu
- ▶ Urządzenie do wyznaczania temperatury cewki głośnika

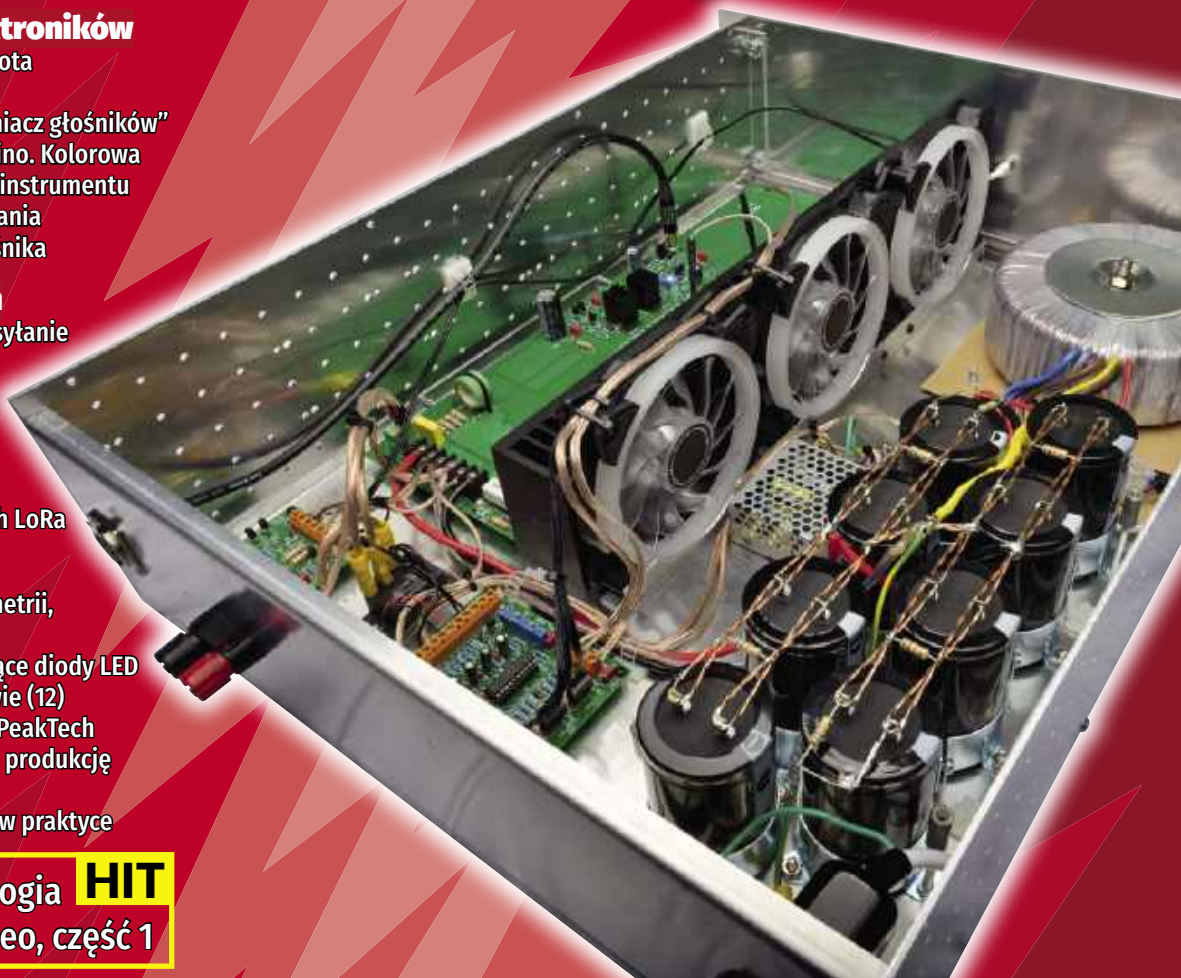
### DIY dla wszystkich

- ▶ Proste i bezpieczne wysyłanie często odświeżanych danych w dowolne miejsce
- ▶ Sieć mesh dalekiego zasięgu oparta na niedrogich modułach LoRa

### TUTORIALE

- ▶ Audio OUT: Kwestia symetrii, część 2
- ▶ Ekscytacje Maxa: Migające diody LED i śliniacy się inżynierowie (12)
- ▶ Miernik LCR 3730 firmy PeakTech
- ▶ Projektowanie PCB pod produkcję seryjną
- ▶ Przetworniki ADC i DAC w praktyce

Historia i technologia **HIT**  
wyświetlaczy wideo, część 1



ISSN 1425-1698 Indeks 33362X

9 771425 169245  
**16,90 zł** (w tym 8% VAT)

Pomocna dłoń



[automatykaB2B.pl](http://automatykaB2B.pl)

**EP.com.pl**

Największy  
portal dla  
elektroników  
konstruktorów

[eprasa.pl/bc7aea75f](http://eprasa.pl/bc7aea75f)

**FIRMA PIEKARZ**  
CZĘŚCI ELEKTRONICZNE

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

[www.piekarz.pl](http://www.piekarz.pl)



# Elektor Bestsellers

SAVE UP TO  
26% NOW!



[www.elektor.com/sale/deals](http://www.elektor.com/sale/deals)



–20%  
NA START  
162,30 zł

–30%  
po pierwszym roku  
prenumeraty  
142,00 zł

–40%  
po drugim roku  
prenumeraty  
121,70 zł

–50%  
po trzecim roku  
nieprzerwanej prenumeraty  
101,40 zł

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

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

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

Im dłużej jesteś z nami, tym więcej oszczędzasz:

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

### Jak otrzymać rabat za lojalność?

Zaloguj się na swoje konto prenumeratora na [www.UlubionyKiosk.pl](http://www.UlubionyKiosk.pl) i zamów prenumeratę, korzystając z przycisku PRZEDŁUŻ w zakładce „Prenumeraty”.

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

Wybierz prenumeratę cyfrową PDF i ciesz się dostępem do czasopisma nawet 7 dni przed oficjalną premierą w kioskach. Oszczędzaj czas i pieniądze – skorzystaj z **rabatu 30%** na roczną e-prenumeratę w cenie 113,40 zł.

Dodatkowa oferta dla prenumeratorów wersji drukowanej: jeśli już subskrybujesz wersję papierową, możesz dokupić równoległe e-wydania w cenie 32,40 zł/rok – z **niesamowitym rabatem 80%**.

## Zyskaj nieograniczony dostęp do zasobów dla pasjonatów elektroniki!

Tylko prenumeratorzy mają pełny dostęp do:

- cyfrowego archiwum *Elektroniki dla Wszystkich* na [www.elportal.pl/archiwum](http://www.elportal.pl/archiwum)
- projektów DIY+ na [www.elportal.pl/diy](http://www.elportal.pl/diy)

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

## ARCHIWUM



**Zacznij korzystać z pełnych zasobów już dziś!**



7



16



26



35



66

## Projekty dla elektroników:

|                                                           |    |
|-----------------------------------------------------------|----|
| Wzmacniacz mocy 500 W, część 1 .....                      | 7  |
| Wzmacniacz zasięgu pilota na podczerwień .....            | 16 |
| Wielokanałowy „ochraniacz głośników” .....                | 26 |
| Stroboskop RGB z Arduino.                                 |    |
| Kolorowa adaptacja użytecznego instrumentu .....          | 32 |
| Urządzenie do wyznaczania temperatury cewki głośnika..... | 35 |

## Tutoriale:

|                                                                          |    |
|--------------------------------------------------------------------------|----|
| Audio OUT: Kwestia symetrii, część 2.....                                | 38 |
| Ekscytacje Maxa:                                                         |    |
| • Migające diody LED i śliniacy się inżynierowie (12).....               | 42 |
| • Sprytne porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania .. | 44 |
| Miernik LCR 3730 firmy PeakTech.....                                     | 46 |
| Projektowanie PCB pod produkcję seryjną.....                             | 50 |
| Edukacja w EdW dla szkół i uczelni:                                      |    |
| Wykład 22 – Konwersja analogowo-cyfrowa.....                             | 54 |
| Przetworniki ADC i DAC w praktyce .....                                  | 58 |
| Historia i technologia wyświetlaczy wideo, część 1 .....                 | 66 |

## DIY dla wszystkich:

|                                                                                  |    |
|----------------------------------------------------------------------------------|----|
| Proste i bezpieczne wysyłanie często odświeżanych danych w dowolne miejsce ..... | 77 |
| Sieć mesh dalekiego zasięgu oparta na niedrogich modułach LoRa .....             | 80 |

## Elektronika dla Wszystkich – Junior:

|                                                                                                           |    |
|-----------------------------------------------------------------------------------------------------------|----|
| Trzecie spotkanie z najmłodszymi pasjonatami elektroniki.....                                             | 84 |
| Na zdjęciu na okładce Łukasz z Wrocławia, uczestnik kółka zainteresowań „Młodych Entuzjastów Elektroniki” |    |

**DIY PLUS** tylko dla prenumeratorów zamawiających prenumeratę na [www.UlubionyKiosk.pl](http://www.UlubionyKiosk.pl)

|                                                                                        |    |
|----------------------------------------------------------------------------------------|----|
| Pojemnościowy czujnik wilgotności do konwertera wyjścia analogowego .....              | 91 |
| Mostek H dla wysokiej mocy szczotkowego silnika prądu stałego z czujnikiem prądu ..... | 91 |

## Rubryki stałe:

|                   |   |
|-------------------|---|
| Prenumerata ..... | 3 |
| Od redakcji.....  | 5 |
| Poczt.....        | 6 |

## A za miesiąc w październikowym EdW



### \* Klasyczne metronomy LED

Dwa projekty metronomów elektronicznych, świetnie imitujących rozwiązania mechaniczne. Oba układy, każdy o nieco odmiennej konstrukcji, zapalają i wygaszają serię diod LED, czemu towarzyszą charakterystyczne dźwięki uderzeń znane z urządzeń klasycznych.

### \* SMD Tweezers

Prosta, ale bardzo pomysłowa konstrukcja pęsety do sprawdzania rezystorów, diod i kondensatorów w obudowach SMD. Przy udziale zasilanego z baterii pastylkowej, ośmionóżkowego mikrokontrolera PIC16F15214 wyniki pomiarów prezentowane są na miniaturowym wyświetlaczu OLED.

### \* 500 W wzmacniacz audio, część 2

Kontynuujemy temat wzmacniacza audio o mocy 500 W RMS. W tej części omówione zostaną szczegóły budowy, w tym obłożenie i montaż PCB, obróbka radiatora oraz nawinięcie cewki indukcyjnej dotyczące najważniejszego modułu wzmacniacza.

### \* Kolejna porcja intrygujących projektów DIY.

### \* Wiele artykułów w Twoich ulubionych cyklach Tutoriali.

### \* Nie zabraknie oczywiście treści dla najmłodszych Czytelników EdW, którzy w ramach kolejnego spotkania grupy Juniorów, pod opieką troskliwego mentora będą mieli okazję zbudować własnymi rękoma kolejną niebanalną zabawkę.

**W kioskach  
od 30 września**

## Z nową mocą (500 W RMS) w pourlopową i powakacyjną codzienność!

Czas urlopów dla większości z nas bezpowrotnie przeminął lub też chyli się ku końcowi.

Wrzesień to czas powrotów do pracowniczych biur i szkolnych ław, ale możliwość bezpiecznego powrotu z wakacyjnych wojaży jest również swego rodzaju błogosławieństwem. Wraz z zespołem redakcyjnym, nie tracąc czasu, przygotowaliśmy dla Was kolejną porcję gorących tematów. Jak co miesiąc, wracamy z nowymi inspiracjami i ciekawymi projektami.

Tematem okładkowym tego wydania jest wzmacniacz audio, który przy zastosowaniu odpowiednich głośników oferuje niebywałą moc aż 500 W RMS (!). Wzmacniacz ten charakteryzuje się ponadto niskim poziomem szumów i zniekształceń, a także wyposażony został w szereg zaawansowanych zabezpieczeń oraz system zautomatyzowanego chłodzenia.

Pozostając przy tematyce audio, prezentujemy też niezwykle przydatny układ do wielokanałowej ochrony głośników przed ich uszkodzeniem na skutek pojawienia się napięcia stałego na wyjściu wzmacniacza. Kolejnym opisanym urządzeniem jest wzmacniacz zasięgu klasycznego pilota na podczerwień, który, dzięki zamianie toru podczerwieni na tor radiowy, umożliwia sterowanie urządzeniami nawet z innego pomieszczenia.

To jednak nie wszystko! W ramach 22. wykładu na warsztat wzięliśmy teorię konwersji analogowo-cyfrowej, po czym głos oddaliśmy znawcy tematu, który opowie o zgodnym ze sztuką wykorzystywaniu tych przetworników w praktyce. Zaprezentujemy także stroboskop RGB z Arduino, ciekawą i niebanalną zabawkę, która łączy zagadnienia elektroniki, fizyki i mechaniki. Projekt tematycznie łączy pokolenia, albowiem będzie można wspominać czasy wykorzystywania stroboskopów do regulacji zapłonu, skorzystać z dobrodziejstw druku 3D i pobawić się pisaniem programów dla Arduino. W rubryce DIY czekają natomiast dwa godne uwagi opracowania projektowe z zastosowaniem modułów radiowych LoRa (Long Range). Jedno z nich opisuje sieć mesh, podobną do ZigBee, ale zrealizowaną w oparciu o moduły LoRa, dzięki czemu zasięg sieci wzrasta z około 100 m (ZigBee) do kilkudziesięciu kilometrów (LoRa). Ponadto w tym numerze zaprosimy Was w fascynującą podróż przez historię i technologię przesyłania na odległość i wyświetlania ruchomych obrazów, a chwilę później, w ramach artykułu obiecanego jednemu z Czytelników w rubryce „Pocztą” z numeru EdW 7/2024, zabierzemy Was na przechadzkę alejkami współczesnej hali produkcyjnej. W przemysłowych klimatach zaskoczmy Was kilkoma wybranymi aspektami specyfiki projektowania płytek drukowanych pod wymagania produkcji seryjnej.

Młodszych Czytelników zapraszam na kolejne spotkanie z cyklu EdW Junior. Tym razem zmontujemy intrygującą zabawkę, która pozwoli pobawić się nieco dźwiękiem.

Życzymy Wam, drodzy Czytelnicy, młodzi ciałem i duchem, przyjemnej lektury i inspirujących chwil z naszym ulubionym czasopiśmie!

**Mariusz Ciszewski**





## „Repolonizacja” EdW

... Lwią część wydań „EdW” stanowią przedruki z artykułów zagranicznych. Trzeba przyznać, że jakość tłumaczeń uległa radykalnej poprawie. Co mnie jednak martwi to fakt, że publikacje autorów polskich stanowią mniej niż 20% objętości pisma (patrz „EdW” 8/2024). Może to utwierdzać Czytelników w równie szkodliwym co nieprawdziwym przekonaniu, że polska myśl techniczna jest gorsza od zagranicznej. Myślę, że trzeba by „zrepolonizować” czasopismo!

Mój pomysł jest taki, by zajrzeć na kilka krajowych forów internetowych poświęconych elektronice i pozyskać do współpracy co aktywniejszych ich członków. Ci ludzie regularnie publikują fachowe posty, redagowane nienaganną polszczyzną, i czynią to za darmo. Dlaczego nie mieliby tego robić na łamach gazety? Skoro już o pieniądzach mowa – nie byłoby źle zwiększyć w „EdW” stawkę honorarium, np. dwukrotnie. To w perspektywie roku przyciągnęłoby do czasopisma grono ok. 10 stałych współpracowników, utrzymujących się w znacznym stopniu z pisania artykułów. Ja sam chętnie poświęciłbym na taką współpracę sporą część czasu zawodowego. Zgodzą się Państwo, że dysponując stabilnym zespołem autorów można naprawdę bardzo wiele dokonać.

Czego Redakcji z całego serca życzę.

**Jarosław Ziembicki**

**Red.** Ten list pobudził mnie do podzielenia się z Czytelnikami kilkoma refleksjami.

1. Nie da się realizować pomysłów, na które nie pozwala budżet EdW. A budżet jest jaki jest, bo... cena EdW jest mniej więcej 3 razy niższa niż cena analogicznych czasopism na Zachodzie („Silicon Chip”, „Elektor”, „Practical Electronics”), a koszty papieru i druku są w Polsce identyczne jak na Zachodzie. Tego parytetu cen nie da się zmienić,

dopóki zarobki w Polsce będą 2–3 razy niższe niż na Zachodzie. Wprawdzie EdW zawsze było i jest wydawane na zasadach non profit (przychody z reklam są do zaniechania), ale Wydawnictwa AVT nie stać na dopłacanie do czasopisma, dlatego nie ma możliwości dwukrotnego zwiększenia honorarium za artykuły. Zresztą, czy wysokość honorarium jest rzeczywiście decydującym motywatorem dla autorów? Jeszcze powrócę do tego tematu. Czytelnicy doradzają nam często, żeby wydawać tylko elektroniczną wersję EdW, unikając horrendalnych wydatków na papier i druk. Nie posłuchamy tych rad. Z doświadczeń innych wydawnictw i trochę własnych wiemy, że taki ruch redukuje wielokrotnie zasięg czytelnictwa. Znamienne, że nawet istotnemu wzrostowi prenumeraty e-wydań EdW w ostatnich 2 latach towarzyszył podobny wzrost prenumeraty wydań papierowych, zamawianej również przez większość prenumeratorów e-wydań. W ogóle na świecie obserwuje się renesans czytelnictwa papierowych wydań książek, a w USA niektóre czasopisma (m.in. „Newsweek”), które przed laty zrezygnowały z wydań papierowych, teraz do nich wracają.

2. Czy Polska myśl techniczna jest gorsza od zagranicznej? Nie jest gorsza, jeśli jednak mamy rozmawiać serio, to udział polskiej myśli technicznej w rozwoju elektroniki światowej jest mniej więcej taki jak udział Polski w sporcie olimpijskim. W roku 80. minionego wieku Amerykanie prognozowali, że w przyszłości na świecie będzie pracowało zaledwie 5000 konstruktorów elektroników. Tak też się stało, jeśli przyjąć dość ortodoksyjny pogląd, że kreatywne rozwiązania projektowe powstają w toku projektowania układów scalonych, a użytkownicy układów scalonych tylko korzystają z rozwiązań „zaszytych” w tych układach, aplikując je w różnych urządzeniach. Zresztą producenci układów scalonych publikują opasłe kilkusetstronicowe podręczniki opisujące dokładnie wszystkie aspekty i przykłady ich aplikacji. Projektantowi urządzenia nie pozostaje zbyt wiele pola do wykazania się kreatywną „myślą techniczną”. Wystarczy, że się nie poleni i dokładnie przestudiuje podręcznik producenta układów scalonych. Wszyscy konstruktorzy urządzeń korzystają z „myśli technicznej” powstającej w biurach konstrukcyjnych producentów układów scalonych. Zatem prawdziwa „myśl techniczna” w elektronice powstaje w kilku zaledwie krajach na świecie. Jeśli spojrzymy na lokalizację trzydziestu najważniejszych

fabryk układów scalonych, to okazuje się, że w USA jest 13, na Dalekim Wschodzie 13 i 4 w UE. Oczywiście, w aplikacjach też czasami udaje się błysnąć kreatywnym projektowaniem, czego przykładami mogą być pomysły na Arduino i Raspberry Pi, tak ważne dla rozwoju ruchu DIY.

A gdzie lokuje się polska myśl techniczna? W mediach od czasu do czasu można przeczytać tytuły:

– Czy Polska może zostać liderem w produkcji mikroprocesorów?

– Za trzy lata może powstać polski mikroprocesor. Koszt to 150 mln zł.

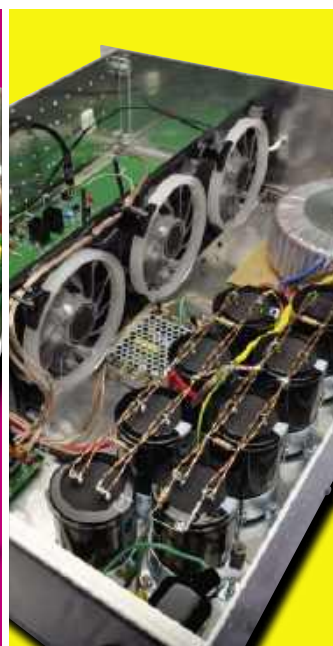
Nie dajmy się ogłupić. Polska nie jest w stanie produkować mikroprocesorów. Musiałaby powstać linia technologiczna, której koszt wynosi ładnie kilka miliardów dolarów. Natomiast Polska może sobie zamówić serię mikroprocesorów według własnej specyfikacji i ewentualnie własnego (do pewnego stopnia) projektu. Ten sposób działania nazywa się fabless, czyli nie mając możliwości produkcji można skorzystać z usługi producenta realizującego produkcję w trybie Custom Design (na zamówienie). Stopień opracowania projektu może być różny od poziomu specyfikacji, poprzez architekturę do topologii, przy czym projektanci korzystają z reguł projektowych producenta. Zatem Polska może sobie zaprojektować mikroprocesor (za 150 mln zł?! – ho, ho), ale nie może go wyprodukować (nie ma czym), a obiecywanie, że „Polska może zostać liderem w produkcji mikroprocesorów” to kpina.

3. Czy misją EdW jest krzewienie polskiej myśli technicznej? Nie. EdW jest tworzone przez pasjonatów dla pasjonatów. Pasjonaci, którzy mają coś do przekazania czytelnikom i potrafią to robić, swoimi projektami i tutorialami dają frajdę innym pasjonatom elektroniki, od początkujących do zaawansowanych. Umyślnie nie piszę o funkcji dydaktycznej EdW, bo najważniejsze jest, by osoba ucząca się elektroniki robiła to z pasją. Elektronika jest za trudna, żeby się jej uczyć „na siłę”. Publikujemy w EdW materiały z czterech najlepszych redakcji na świecie, więc nie widzę powodów merytorycznych dla „repolonizacji” EdW, ale zawsze zapraszamy Autorów, którzy czują wewnętrzną potrzebę podzielenia się swoją pasją z innymi. Ta potrzeba dzielenia się wiedzą z innymi powinna być motywacją silniejszą niż pieniądze. Tacy ludzie są w szkołach i w wolontariatach dla dzieci. Pomagamy im w Patronatach dla szkół i w rubryce EdW Junior.

**Wiesław Marciniak**



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



# Wzmacniacz mocy 500 W, część 1

**Przedstawiamy wzmacniacz audio dostarczający potężnych mocy przy zachowaniu czystego dźwięku o niskim poziomie szumów i zniekształceń. Dostarcza on 500 watów RMS przy zastosowaniu głośników 4  $\Omega$  lub 270 watów w przypadku użycia głośników 8  $\Omega$ . Został zaprojektowany z myślą o dużej wytrzymałości i wyposażony został w ochronę linii obciążenia i tranzystorów wyjściowych oraz wentylator z regulacją prędkości, który włącza się tylko wtedy, gdy jest to konieczne. Dzięki dwóm takim wzmacniaczom można uzyskać 1000 W mocy na pojedynczym głośniku 8  $\Omega$ . Znalazienie głośnika, który poradzi sobie z taką mocą, to prawdziwe wyzwanie!**

**Opisywany wzmacniacz 500 W jest wielki pod każdym względem.** Jest on fizycznie duży, a utrzymanie temperatury pod kontrolą wymaga zastosowania dwóch radiatorów umieszczonych jeden na drugim. Konieczny jest też duży zasilacz z transformatorem o mocy 800 VA. Wzmacniacz i zasilacz mieszczą się w obudowie typu rack o wymiarze 3RU, a więc również dużej.

Wzmacniacz dostarcza ogromną moc. Idealnie nadaje się do systemów nagłośnieniowych, gdzie wysoka moc może być niezbędna do nagłośnienia dużego pomieszczenia. Dobrze sprawdzi się również przy dostarczaniu mocy do mniej efektywnych głośników o niższych impedancjach. Jak wspomniano powyżej, używany w trybie mostka, może dostarczyć nieco ponad 1000 W mocy audio na kanał. Można zbudować dwie pary dla systemu dźwiękowego tak potężnego, że pojedynczy musiałby być podłączony do dwóch różnych punktów zasilania! Dwa takie wzmacniacze mogą również stanowić podstawę niesamowitego systemu stereo do użytku w dużym pomieszczeniu odsłuchowym.

Można by pomyśleć, że system stereo o mocy 500 W na kanał to po prostu za dużo. To, czy jest to prawda, zależy od tego, jakiej muzyki lubimy słuchać i jak wydajne są posiadane głośniki. Jeśli preferowana jest muzyka rockowa z nieco ograniczonym zakresem dynamiki, to dzięki temu wzmacniaczowi będzie możliwe odtwarzanie jej odpowiednio głośno. Wzmacniacz jest więc idealny do muzyki, która

z definicji musi być głośna. Należy jednak pamiętać, aby nie ogłuchnąć z powodu ekstremalnych poziomów dźwięku możliwych przy tak dużym wzmacniaczu. Konieczne może być również zapewnienie ochrony uszu sąsiadom!

Nie dotyczy to tylko rockmanów. Muzyka klasyczna również wymaga dużej mocy. Nie dlatego, żeby była głośno odtwarzana, ale dlatego, że zapas mocy pozwala na odtworzenie

## Cechy i specyfikacja:

- Moc wyjściowa: >500 W przy 4  $\Omega$ , >270 W przy 8  $\Omega$  – rysunek 3
- Pasma przenoszenia: +0,-0,1 dB w zakresie 20 Hz...20 kHz (-3 dB @97 kHz) – rysunek 1
- Stosunek sygnału do szumu: 112 dB w odniesieniu do 500 W przy 4  $\Omega$  lub 250 W przy 8  $\Omega$
- Całkowite zniekształcenia harmoniczne (4  $\Omega$ ): <0,005% @1 kHz dla 1,5 W...350 W – rysunki 2 i 3
- Całkowite zniekształcenia harmoniczne (8  $\Omega$ ): <0,025% @1 kHz dla 2 W...270 W – rysunki 2 i 3
- Impedancja wejściowa: 10 k $\Omega$  || 4,7 nF
- Czuość wejściowa: 1,015 V RMS dla 500 W przy 4  $\Omega$ , 1,055 V RMS dla 270 W przy 8  $\Omega$
- Zasilanie:  $\pm$ 80 V nominalnie z transformatora 800 VA 55-0-55V
- Prąd spoczynkowy/zasilanie: 94 mA, 15 W
- Zabezpieczenia: bezpieczniki DC, śledzenie termiczne o podwójnym nachyleniu, ograniczenie prądu SOA, diody zabezpieczające na wyjściach
- Inne funkcje: zerowanie offsetu wyjściowego, wskaźniki przepalonych bezpieczników, wbudowany wskaźnik zasilania

szerokiego zakresu dynamiki, jaki jest uzyskiwany w salach koncertowych. Duża moc bez zniekształceń jest potrzebna do wytworzenia wysokich szczytowych poziomów głośności, takich dźwięków, jak potężne uderzenia bębna takowego lub dźwięków organów piszczałkowych, przy niskim poziomie szumów ze wzmacniacza, po to, by nie zagłuszyć cichych fragmentów.

Wytworzenie tak dużej mocy nie jest łatwe. W naszym wzmacniaczu zastosowano 12 tranzystorów wyjściowych. Wszystkie są zamontowane na radiatorze o szerokości 400 mm. Główna płyta drukowana ma również spore wymiary: 402 mm × 124 mm. Ostateczna instalacja w obudowie rack 3U mierzy 559 mm × 432 mm × 133,5 mm i waży nieco ponad 12 kg.

W tym odcinku skoncentrujemy się na opisie modułu wzmacniacza. W kolejnych odcinkach zostaną podane również pełne szczegóły montażu tego modułu, a także zostanie opisany odpowiedni zasilacz. Następnie pokażemy, jak zbudować moduł wzmacniacza, zasilacz, chłodzenie wentylatorem z regulacją prędkości (który wyłącza się przy odpowiednim obciążeniu), zabezpieczenie głośników i detektor klipów, wszystko to w jednej aluminiowej obudowie 3RU do montażu w szafie typu rack.

## Wydajność

Główne parametry wydajności zostały podsumowane w ramce specyfikacji i na **rysunkach 1...3**. Wynika z nich, że wzmacniacz charakteryzuje się nie tylko dużą mocą, ale zapewnia również wysoką jakość dźwięku.

Po pierwsze, pasmo przenoszenia jest płaskie jak linijka od 20 Hz do 20 kHz,

z zaledwie 0,1 decybelowym spadkiem przy 20 kHz. Na obciążeniu 4 Ω wydzielana jest moc 500 W. Przy typowych poziomach mocy, od 1,5 W do 350 W, całkowite zniekształcenia harmoniczne plus szum (THD+N) wynoszą poniżej 0,007% przy 1 kHz.

Dla obciążenia 8 Ω, maksymalna moc wynosi około 270 W do początku obcinania, ze zniekształceniami THD+N mniejszymi niż 0,004% przy 1 kHz. W idealnych warunkach jest on zbliżony do tego, co nazwalibyśmy „jakością CD” przy około 0,002% THD+N.

Jak widać na rysunku 2, zniekształcenia nieco rosną wraz z częstotliwością. W rzeczywistości są one znacznie niższe niż podano powyżej przy bardziej typowych zakresach częstotliwości audio dla większości instrumentów, wynoszących około 100...500 Hz. Powyżej 1 kHz, zniekształcenia wzrastają umiarkowanie, choć nadal są stosunkowo niskie nawet przy 10 kHz, powyżej których filtry w naszym sprzęcie testowym zaczynają tłumić harmoniczne.

Wynik THD+N poniżej 0,05% dla 266 W na 3 Ω pokazuje, że wydajność tego wzmacniacza nie pogarsza się znacząco nawet w trudnych warunkach, sterując niższymi impedancjami obciążenia niż można by się spodziewać w przypadku większości głośników 4 Ω o dużej mocy.

Prawdopodobnie najważniejszym atutem tego wzmacniacza dużej mocy jest bardzo dobry stosunek sygnału do szumu wynoszący 112 dB. Oznacza to, że można uzyskać bardzo wysoki poziom wyjściowy, w tym głośne transjenty, bez irytującego syczenia w tle przez resztę czasu.

## Szczegóły układu

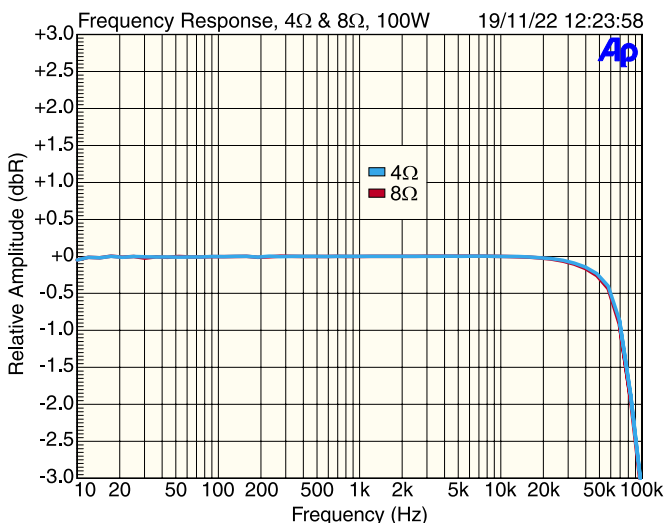
Pełny schemat wzmacniacza pokazano na **rysunku 4**. Pomijając dużą liczbę tranzystorów wyjściowych, układ jest podobny w konfiguracji do wielu poprzednich wzmacniaczy, w tym wzmacniaczy Ultra-LD Mk.2 do Mk. 4 (Silikon Chip sierpień i wrzesień 2008, lipiec-wrzesień 2011 i lipiec-wrzesień 2015).

Jedną z głównych różnic jest dodanie ochrony bezpiecznego obszaru pracy (SOA) dla tranzystorów wyjściowych. Pomaga to zapobiec ich uszkodzeniu w przypadku zwarcia wzmacniacza lub obciążenia przekraczającego ich bezpieczny obszar roboczy. Nie jest to tylko ochrona przed zwarcie, działa w całym zakresie roboczym wzmacniacza.

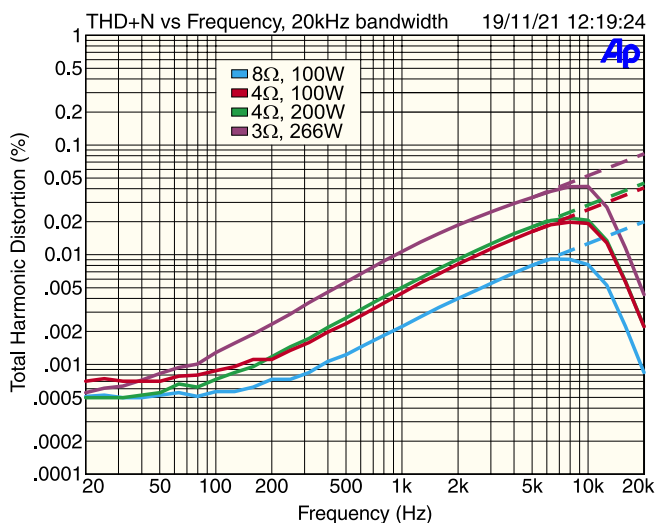
W przeszłości słyszeliśmy, że zabezpieczenie SOA pogarsza wydajność wzmacniacza, ale testowaliśmy ten wzmacniacz z włączonym i odłączonym zabezpieczeniem i nie odnotowaliśmy żadnych różnic. Nie trzeba więc martwić się o jego wpływ na jakość dźwięku.

Linie zasilające mają napięcie ±80 V lub łącznie 160 V. Tak wysokie napięcie narzuca stosowanie wytrzymałych tranzystorów, zwłaszcza tranzystorów wyjściowych i sterujących, które wymagają dużego współczynnika SOA. Można by było użyć tranzystorów NJL3281D/NJL3282D ThermalTrak, które są stosowane we wzmacniaczach Ultra-LD. Do zapewnienia solidności konstrukcji potrzebnych byłoby jednak 12 takich tranzystorów na stronę czyli łącznie 24.

Tranzystory ThermalTrak mają dwie główne zalety: dobrą liniowość i wbudowaną diodę polaryzującą. Dioda w obudowie tranzystora



**Rysunek 1.** Pasma przenoszenia tego wzmacniacza jest wyjątkowo płaskie, różniąc się o mniej niż 1/20 dB między 20 Hz a 20 kHz. Górny punkt -3 dB znajduje się tuż poniżej 100 kHz. Dolny punkt -3 dB nie jest widoczny na tym wykresie, ale prawdopodobnie znajduje się w okolicach 1 Hz. Aktywny wstępny filtr poddźwiękowy byłby niezbędny, aby zapobiec nadmiernemu rozciągnięciu pasma, jeśli wzmacniacz jest używany do bezpośredniego sterowania subwoofera



**Rysunek 2.** Wykresy THD+N dla obciążeń 8 Ω, 4 Ω i 3 Ω (dla 4 Ω pokazano dwa różne poziomy mocy) przy paśmie przenoszenia 20 Hz...22 kHz. Widać, że zniekształcenia bazowe w dużej mierze zależą od impedancji obciążenia i stale rosną wraz z częstotliwością powyżej około 100 Hz. Krzywa 3 Ω jest przedstawiona głównie jako „najgorszy przypadek”. Z wykresu wynika więc, że możliwe jest dostarczanie mocy do nawet bardzo niskich impedancji obciążenia bez większych uciążliwości



We wzmacniaczu 500 W zostały zastosowane dwa nasze poprzednie projekty: wentylator chłodzący i zabezpieczenie głośników (luty 2022; [siliconchip.com.au/Article/15195](https://siliconchip.com.au/Article/15195)) oraz wskaźnik przesterowania wzmacniacza (marzec 2022; [siliconchip.com.au/Article/15240](https://siliconchip.com.au/Article/15240))

umożliwia dokładną kontrolę prądu spoczynkowego (jałowego) przy zmianach temperatury. Niestety, sama liczba wymaganych tranzystorów sprawiłaby, że wzmacniacz byłby niepraktycznie duży i drogi, więc w praktyce byłyby nieodpowiednie. Zamiast nich zostały zastosowane tranzystory MJW21196/MJW21195, co pozwoliło użyć zaledwie sześciu sztuk na stronę. Było to możliwe dzięki odpowiedniej krzywej SOA tych tranzystorów.

Sygnał wejściowy jest sprzężony zmiennoprądowo poprzez niespolaryzowany kondensator elektrolityczny 47  $\mu\text{F}$  i elementy ograniczające wysoką częstotliwość, koralek ferrytowy FB1 i rezystor 22  $\Omega$  do bazy tranzystora Q1. Rezystor wejściowy 22  $\Omega$  i kondensator 4,7 nF stanowią filtr dolnoprzepustowy z tłumieniem  $-6$  dB/oktawę powyżej 1,5 MHz.

Tranzystor Q1 jest częścią wejściowej pary różnicowej Q1 i Q2, które są niskoszumowymi tranzystorami PNP Toshiba 2SA1312. Są one odpowiedzialne za bardzo niski poziom szumów resztkowych wzmacniacza.

Tranzystory 2SA1312 stają się coraz trudniejsze do zdobycia, ale zapewniłyśmy naszym czytelnikom pewne zapasy, ponieważ nie mogliśmy znaleźć żadnych odpowiedników. Można uznać, że praktyka producentów zaprzestających produkcji elementów bez wprowadzenia bezpośredniego zamiennika jest bardzo frustrująca i utrudnia zadania konstruktorom.

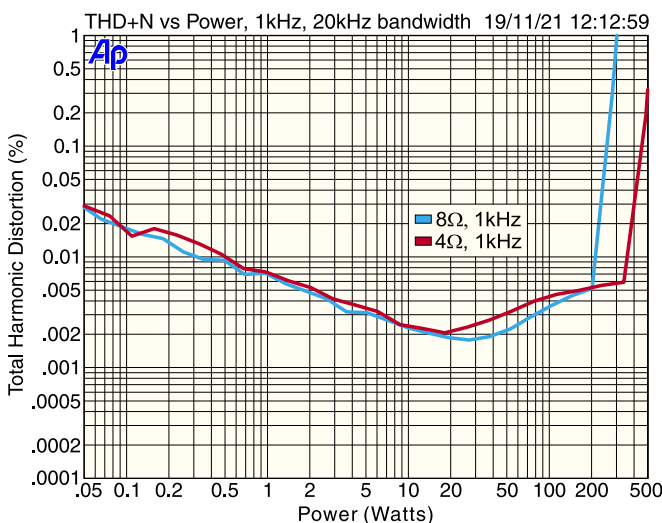
Rezystor polaryzujący dla Q1 i szeregowy rezystor sprzężenia zwrotnego do bazy Q2 są ustawione na stosunkowo niską wartość 10 k $\Omega$ , aby zminimalizować impedancję źródła sygnału, a tym samym zmniejszyć szum

termiczny. Rezystancja wejściowa 10 k $\Omega$  i kondensator wejściowy 47  $\mu\text{F}$  zapewniają obniżenie częstotliwości przy 0,34 Hz.

Wzmocnienie wzmacniacza jest ustawiane przez stosunek rezystorów sprzężenia zwrotnego 10 k $\Omega$  i 220  $\Omega$  na bazie Q2. Wzmocnienie to jest 46-krotne (33 dB), podczas gdy kondensator 2200  $\mu\text{F}$  ustawia tłumienie niskich częstotliwości (punkt  $-3$  dB) w pętli sprzężenia zwrotnego na 0,33 Hz. Stosunkowo wysokie wzmocnienie pomaga utrzymać stabilność wzmacniacza oraz sprawia, że czułość wejściowa jest rozsądna i wynosi około 1 V RMS dla pełnej mocy wyjściowej.

## Kondensatory sprzęgające

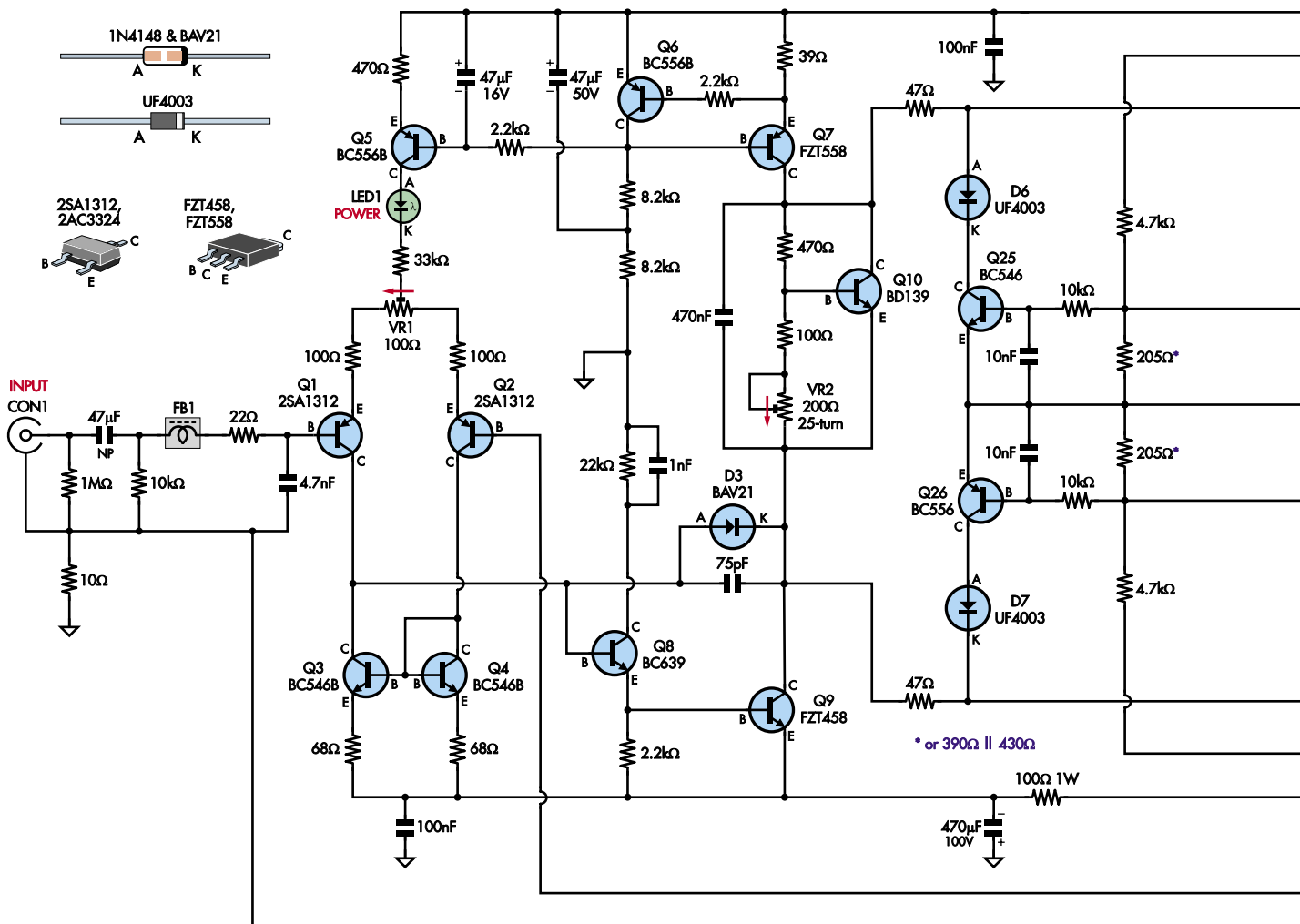
Kondensator elektrolityczny o dużej pojemności dla sprzężenia wejściowego (47  $\mu\text{F}$ )



Rysunek 3. Zniekształcenia THD+N w funkcji mocy przy 1 kHz. Zniekształcenia zaczynają rosnąć powyżej 350 W dla obciążeń 4  $\Omega$ , ale dostarczana jest moc 500 W (a momentami nawet więcej) bez rażących zniekształceń. Wydajność jest całkiem dobra w średnim zakresie mocy, od kilku watów do kilkuset watów. Wzmacniacz zapewni „jakość CD” przy 8  $\Omega$  do około 200 W. Można dwukrotnie zwiększyć liczby na osi poziomej i sprawdzić krzywą 4  $\Omega$  dla wydajności 8  $\Omega$  w mostku!



Gotowy moduł wzmacniacza zamontowany w obudowie 3RU z radiatorami i wentylatorami. Należy zwrócić uwagę na 120-milimetrowe wentylatory PWM przymocowane do radiatora, ponieważ większe nie zmieściłyby się w obudowie z założoną pokrywą



Rysunek 4. Podstawową różnicą między tym wzmacniaczem a kilkoma ostatnimi projektami autora jest sama liczba tranzystorów wyjściowych (sześć par) i dodanie układu zabezpieczającego SOA na każdą linię obciążenia. Obwód zabezpieczający opiera się na układach napięcia referencyjnego REF1 i REF2, tranzystorach Q25 i Q26 oraz powiązanej sieci rezystorów, w tym szeregu rezystorów 3,3 kΩ podłączonych do emitera każdego tranzystora wyjściowego

i pętli sprzężenia zwrotnego (2200 µF) eliminuje wszelkie efekty zniekształceń w paśmie przepustowym audio zależnych od tego kondensatora, a także minimalizuje impedancję źródła.

Spróbujmy to wyjaśnić. Jeśli użyjemy mniejszego kondensatora wejściowego, powiedzmy 2,2 µF, jego impedancja wyniesie 14,47 Ω przy 50 Hz. Będzie to miało niewielki wpływ na pasmo przenoszenia dźwięku, ale oznacza znaczny wzrost impedancji źródła przy niskich częstotliwościach. Dla porównania, użyty przez nas kondensator wejściowy 47 µF ma impedancję tylko 67,7 Ω przy 50 Hz. Oznacza to również, że napięcie na tych kondensatorach jest minimalne w porównaniu do sygnałów audio, więc właściwa nieliniowość kondensatorów elektrolitycznych nie ma znaczenia.

Diody D1 i D2 są umieszczone równolegle do kondensatora sprzężenia zwrotnego 2200 µF jako zabezpieczenie na wypadek uszkodzenia wzmacniacza, w wyniku którego wyjście zostałoby podciągnięte do linii -80 V. W takiej sytuacji kondensator otrzymałby znaczne napięcie wsteczne.

Zamiast jednej zostały zastosowane dwie diody, co zapewnia brak zniekształceń dźwięku spowodowanych nieliniowymi efektami pojedynczego złącza diody przy maksymalnym poziomie sygnału sprzężenia zwrotnego wynoszącym około 1 V. W takim układzie, w normalnych warunkach pracy dioda nie przewodzi.

## Stopień wzmocnienia napięcia

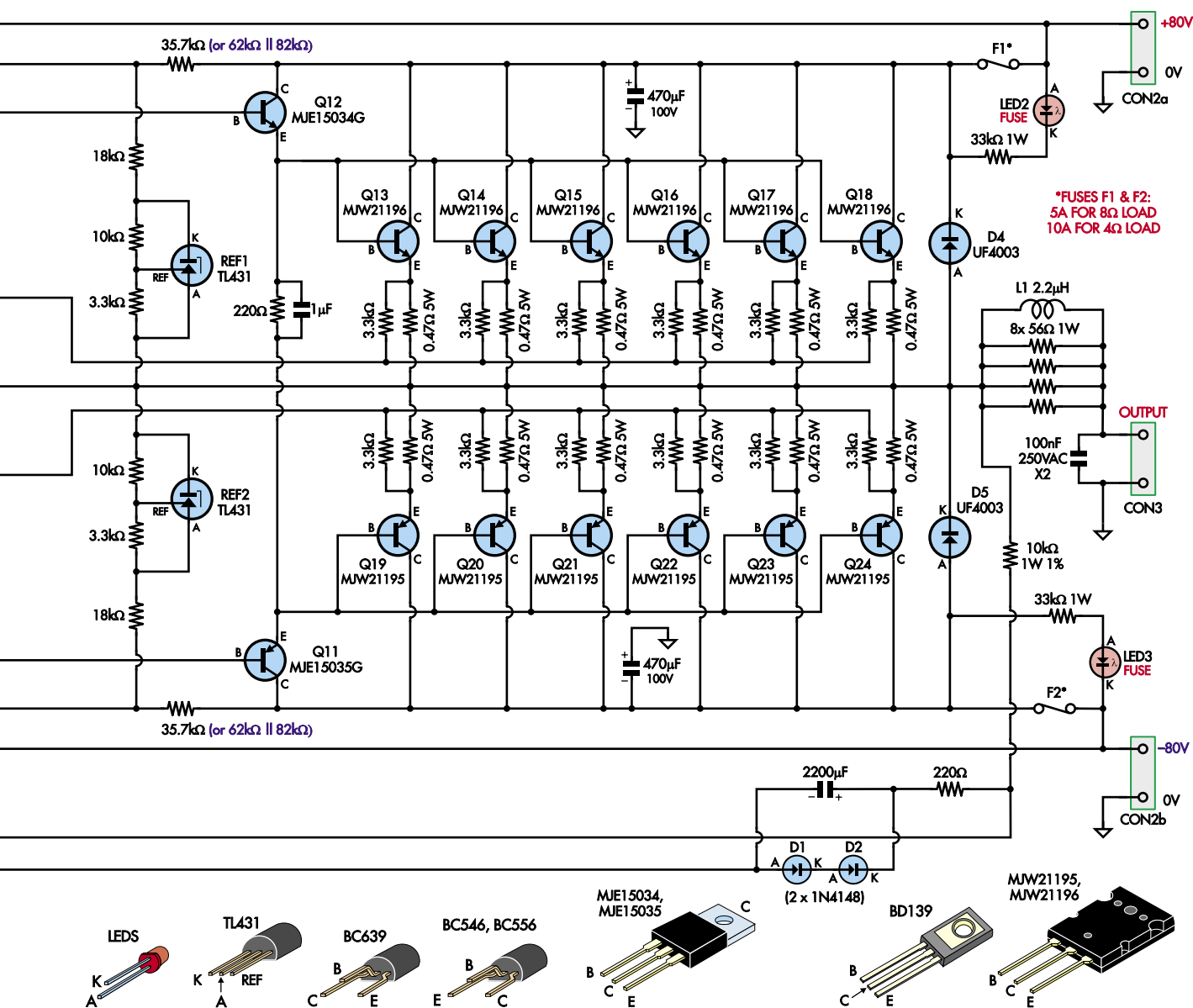
Większą część wzmocnienia napięciowego wzmacniacza zapewnia tranzystor Q9, sterowany przez wtórnik emiterowy Q8 z kolektora

Q1. Razem tranzystory te tworzą stopień wzmocnienia napięcia (VAS). Q8 buforuje kolektor Q1, aby zminimalizować nieliniowość.

Q9 pracuje bez rezystora emiterowego, aby zmaksymalizować wzmocnienie, a także zmaksymalizować zmiany napięcia wyjściowego. Takie maksymalne napięcie uzyskiwanego we wzmacniaczu napięcia jest konieczne, jeśli oczekujemy najwyższej mocy ze stopni wyjściowych.

## Lustro prądowe

Obciążeniami kolektora Q1 i Q2 są tranzystory NPN Q3 i Q4, które działają jako lustro prądowe. Q4 działa jak szybka dioda odcinająca, zapewniając napięcie bazy Q3 równe spadkowi napięcia baza-emiter tranzystora Q4 (około 0,6 V) i powiększone o spadek napięcia na jego rezystorze emiterowym 68 Ω.



Jeśli tranzystor Q2 pobiera więcej niż jego udział prądu emitera z Q5, napięcie na bazie Q3 wzrasta, więc prąd kolektora Q3 również wzrasta. Zmusza to Q1 do pobierania nieco większego prądu i powstrzymuje Q2 przed pobieraniem większej części prądu niż należny mu udział. Ponieważ Q3 odzwierciedla prąd Q4, tranzystor Q1 jest wyposażony w obciążenie kolektora, które ma wyższą impedancję niż miałyby to miejsce w innym przypadku. W rezultacie obserwujemy zwiększone wzmocnienie i lepszą liniowość różnicowego stopnia wejściowego.

Podobnie, obciążenie kolektora tranzystora Q9 jest obciążeniem stałoprądowym składającym się z tranzystorów Q6 i Q7. Co ciekawe, napięcie polaryzacji bazy dla źródła prądu stałego tranzystora Q5 jest również ustawiane przez Q6. Tranzystor Q5 jest

źródłem stałoprądowym dla wejściowej pary różnicowej Q1 i Q2 i ustala prąd płynący przez te tranzystory. Dioda LED1 jest dołączona do tego układu jako niezależny wskaźnik włączenia zasilania.

Powodem nieco skomplikowanego układu polaryzacji tranzystorów Q5, Q6 i Q7 jest znaczna poprawa współczynnika odrzucenia zasilania (PSRR) wzmacniacza. Podobnie, PSRR jest poprawiany przez filtr obejściowy składający się z rezystora 100  $\Omega$  1 W i kondensatora 470  $\mu$ F 100 V w ujemnej linii zasilania.

Dlaczego współczynnik PSRR jest tak ważny? Ponieważ wzmacniacz pracuje w klasie AB, pobiera duże asymetryczne prądy z dodatniej i ujemnej linii zasilającej. Prądy są asymetryczne w tym sensie, że w danym momencie pobierane są z jednej lub drugiej linii. Przebiegi będą miały podobny kształt dla

fali sinusoidalnej, tylko są przesunięte w czasie względem siebie.

Na przykład, gdy dodatnia połowa stopnia wyjściowego (Q13 do Q18) przewodzi, przebieg prądu jest efektywnie dodatnią półfalą sygnału, tj. następuje prostowanie. Podobnie, gdy ujemna połowa stopnia wyjściowego (Q19...Q24) przewodzi, prąd jest ujemną półfalą sygnału.

Mamy więc półfalowe tętnienia prostownicze sygnału nałożone na linie zasilające, a także tętnienia 100 Hz z samego zasilacza. Współczynnik PSRR wzmacniacza może być bardzo wysoki przy niskich częstotliwościach, jednak przy wysokich częstotliwościach zawsze jest gorszy. Jeśli te tętnienia napięcia przedostaną się do wcześniejszych stopni wzmacniacza, spowodują zniekształcenia, należy więc je tam zminimalizować.

## Wykaz elementów:

### Moduł wzmacniacza 500 W (pojedynczy)

- 1 dwustronna płytka drukowana – kod 01107021, 402 mm × 124 mm
- 2 radiatory o szerokości 200 mm [Altronics H0536].
- 2 małe radiatory do montażu na PCB [Jaycar HH8516].
- 12 silikonowych podkładek izolacyjnych TOP-3
- 3 silikonowe podkładki izolacyjne TO-220
- 2 tuleje izolacyjne dla tranzystorów TO-220
- 4 zaciski bezpiecznikowe M205 (dla F1 i F2)
- 2 bezwzględne bezpieczniki ceramiczne M205 (5 A dla obciążenia 8 Ω, 10 A dla obciążenia 4 Ω) (F1, F2)
- 1 koralik ferrytowy (FB1) [Jaycar LF1250, Altronics L5250A]
- 1 6-pinowy zacisk śrubowy do montażu na płytce drukowanej z barierkami (CON2) [Altronics P2106]
- 1 2-pinowe, wtykowe, pionowe gniazdo zaciskowe (CON3) [Altronics P2572, Jaycar HM3112].
- 1 2-pinowy wtykowy zacisk śrubowy (CON3) [Altronics P2512, Jaycar HM3122].
- 1 pionowe gniazdo RCA (phono) montowane na płytce drukowanej (CON1) [Altronics P0131].
- 1 szpulka dla cewki L1 [Altronics L5305, Jaycar LF1062].
- 1 przewód miedziany emaliowany o długości 2 m i średnicy 1,25 mm (dla uzwojenia L1)
- 1 ocynowany drut miedziany o długości 60 mm i średnicy 0,7 mm (do połączeń drutowych)
- 12 śrub z łbem walcowym M3 × 20 mm
- 5 śrub z łbem walcowym M3 × 15 mm
- 6 śrub z łbem walcowym M3 × 6 mm
- 17 nakrętek sześciokątnych M3
- 12 stalowych podkładek M3
- 6 podkładek dystansowych M3 z gwintem 9 mm
- 2 zaciski tranzystorowe [Altronics H7300, Jaycar HH8600]
- 1 rurka termokurczliwa o długości 15 mm i średnicy 25 mm (dla L1)
- 1 rurka termokurczliwa o długości 60 mm i średnicy 1 mm (dla łącz przewodowych)
- 1 mała tubka pasty termoprzewodzącej do radiatorów

### Półprzewodniki:

- 6 tranzystorów NPN MJW21196 250 V 16 A (Q13...Q18) [element14 1700966] \*
- 6 tranzystorów PNP MJW21195 250 V 16 A (Q19...Q24) [RS 790-5410] \*
- 1 tranzystor PNP MJE15035G 350 V 4 A (Q11) [Mouser 863-MJE15035G] \*
- 1 tranzystor NPN MJE15034G 350 V 4 A (Q12) [Mouser 863-MJE15034G] \*
- 1 tranzystor PNP FZT558TA 400 V 300 mA (Q7) [RS 669-7388P] \*
- 1 tranzystor FZT458TA 400 V 300 mA NPN (Q9) [RS 669-7326] \*
- 2 2SA1312 120 V 100 mA niskoszumowe tranzystory PNP (Q1, Q2) \*
- 3 tranzystory BC546 65 V 100 mA NPN (Q3, Q4, Q25)
- 1 BC639 80 V 500 mA tranzystor NPN (Q8)
- 3 tranzystory PNP BC556 65 V 100 mA (Q5, Q6, Q26)
- 1 BD139 80 V 1,5 A tranzystor NPN (Q10)
- 2 diody sygnałowe 1N4148 75 V 200 mA (D1, D2)
- 4 ultraszybkie diody przełączające UF4003 200 V 1 A \* (D4...D7)
- 1 BAV21 250 V 250 mA dioda przełączająca o niskiej pojemności \* (D3) [RS 436-7846]
- 2 programowalne źródła napięcia referencyjnego TL431, TO-92 (REF1, REF2) [element14 3009364] \*
- 1 zielona dioda LED 5 mm (LED1)
- 2 czerwone diody LED 5 mm (LED2, LED3)

### Kondensatory:

- 1 2200 ΩF 16 V lub elektrolityczny o niskim ESR 10 V
- 3 470 μF 100 V elektrolityczny [element14 3464457]
- 1 47 μF elektrolityczny niespolaryzowany (NP/BP)
- 1 47 μF 50 V elektrolityczny
- 1 47 μF 16 V elektrolityczny
- 1 1 μF 100 V poliestrowy MKT
- 1 470 nF 100 V poliestrowy MKT
- 2 100 nF 100 V poliestrowy MKT
- 1 100 nF 250 V AC metalizowany polipropylenowy klasy X2
- 2 10 nF 100 V poliestrowy MKT
- 1 4,7 nF poliestrowy MKT
- 1 1 nF 100 V MKT poliestrowy
- 1 75 pF 200 V COG [Mouser 80-C315C750JCG lub 80-C325C750KAG5TA] \*

### Rezystory: (wszystkie cienkowarstwowe 1/4 W, 1%, o ile nie określono inaczej)

- |                                                                               |                             |                           |                                                        |                             |
|-------------------------------------------------------------------------------|-----------------------------|---------------------------|--------------------------------------------------------|-----------------------------|
| 1 1 MΩ                                                                        | 2 35,7 kΩ                   | 1 (lub 2,82 kΩ i 2,62 kΩ) | 1 33 kΩ                                                | 2 33 kΩ 1 W 5% (węglowy OK) |
| 1 22 kΩ                                                                       | 2 18 kΩ                     | 5 10 kΩ                   | 1 10 kΩ 1 W 1% cienkofoliowy [Yageo MFR1W5FTE52-10K] * |                             |
| 2 8,2 kΩ                                                                      | 2 4,7 kΩ                    | 14 3,3 kΩ                 | 3 2,2 kΩ                                               |                             |
| 2 470 Ω                                                                       | 2 220 Ω                     | 2 205 Ω                   | 1 (lub 2 430 Ω i 2 390 Ω)                              |                             |
| 3 100 Ω                                                                       | 1 100 Ω 1 W 5% (węglowy OK) |                           |                                                        |                             |
| 2 68 Ω                                                                        |                             |                           |                                                        |                             |
| 2 68 Ω 5 W 5% drutowy (do celów testowych)                                    |                             |                           |                                                        |                             |
| 8 56 Ω 1 W 5% (węglowy OK)                                                    |                             |                           |                                                        |                             |
| 2 47 Ω                                                                        | 1 39 Ω                      | 1 22 Ω                    | 1 10 Ω                                                 | 12 0,47 Ω 5 W 5% drutowy    |
| 1 jednoobrotowy potencjometr montażowy poziomy 100 Ω (VR1) [Altronics R2591]  |                             |                           |                                                        |                             |
| 1 wieloobrotowy potencjometr montażowy poziomy 200 Ω (VR2) [Altronics R2372A] |                             |                           |                                                        |                             |

\* Części te są również dostępne w krótkim zestawie Silicon Chip (nr kat. SC6019) do wyczerpania zapasów.

Wykaz elementów zasilacza, obudowy, okablowania itp. zostanie przedstawiona w następnym odcinku.

Dioda D3 jest dołączona w celu poprawienia wydajności odzyskiwania, gdy wzmacniacz jest doprowadzany do twardego obcinania. Sprawia ona, że powrót do stanu wyjściowego po obciążeniu napięcia ujemnego jest tak samo czysty i szybki, jak po obciążeniu napięcia dodatniego, poprawiając symetrię sygnału i redukując dzwonienie w takich warunkach. Do tej roli jest użyta dioda BAV21 o niskiej pojemności 2 pF przy 1 MHz, dzięki czemu nie wpływa na jakość dźwięku.

## Sprężenie zwrotne i kompensacja

Jak wspomniano, elementy sprzężenia zwrotnego na bazie Q2 ustawiają wzmocnienie wzmacniacza w zamkniętej pętli. Dolny koniec pętli sprzężenia zwrotnego jest podłączony do masy przez kondensator elektrolityczny 2200 μF. Ponieważ zmniejsza to wzmocnienie prądu stałego do jedności, napięcie offsetowe na wyjściu wzmacniacza jest znacznie niższe niż w innym wypadku (38-krotnie).

Kondensator kompensacyjny 75 pF podłączony między kolektorem Q9 i bazą Q8 zapobiega oscylacjom poprzez ograniczenie szybkości narastania.

Rezystor 22 kΩ w kolektorze Q8 ogranicza prąd płynący przez Q9 w warunkach błędu. Jeśli wyjście wzmacniacza zostanie zwarte, spróbuje on podciągnąć napięcie wyjściowe w górę lub w dół tak mocno, jak to możliwe, w zależności od polaryzacji wyjściowego napięcia offsetowego.

Jeśli spróbuje go podciągnąć, prąd wyjściowy jest z natury ograniczony przez źródło prądu 15 mA sterujące Q9 z Q7. Jeśli jednak spróbuje podciągnąć w dół, tranzystor Q9 może pobierać znacznie większy prąd. Rezystor 22 kΩ ogranicza prąd bazy tranzystora Q9, a tym samym jego prąd kolektora i rozproszenie. Kondensator równoległy 1 nF jest wymagany do utrzymania niskiej impedancji kolektora AC, co poprawia stabilność.

## Stopień sterujący

Sygnał wyjściowy ze stopnia wzmacniacza napięcia Q9 jest sprzężony z tranzystorami sterującymi Q11 i Q12 poprzez rezystory 47 Ω. Rezystory 47 Ω działają jako ograniczniki, aby zapobiec pasywnym oscylacjom w stopniu wyjściowym. Są one również potrzebne, aby umożliwić obwodowi ochrony linii obciążenia zastąpienie wysterowania z VAS.

Tranzystor Q10 ustawia napięcie stałe między Q7 i Q9, co określa prąd spoczynkowy i moc w stopniach wyjściowych. Zapewnia polaryzację około 2,3 V między bazami Q7 i Q9, dzięki czemu tranzystory te zawsze lekko przewodzą, nawet bez sygnału wejściowego.

Q10 jest „mnożnikiem  $V_{BE}$ ”, mnożącym napięcie między bazą a emiterem przez stosunek rezystancji kolektor-emiter i baza-emiter. Potencjometr montażowy VR2 zmienia wynikowe napięcie kolektor-emiter, w rzeczywistości jest on regulowany w celu ustawienia prądu spoczynkowego przez tranzystory wyjściowe.

Ważne jest, aby napięcie polaryzacji wytwarzane przez Q10 zmieniało się wraz z temperaturą tranzystorów stopnia wyjściowego. Gdy tranzystory wyjściowe stają się bardziej gorące, a ich napięcie baza-emiter spada, napięcie kolektor-emiter Q10 również powinno spaść, tak aby prąd spoczynkowy był taki sam lub mniejszy jak w niższych temperaturach, zapobiegając niebezpieczeństwu niekontrolowanego wzrostu temperatury.

## Stopień wyjściowy

Stopień wyjściowy wzmacniacza jest komplementarnym symetrycznym wtórnikiem emiterowym składającym się z sześciu tranzystorów NPN (Q13...18) i sześciu tranzystorów PNP (Q19...Q24).

Każdy tranzystor mocy w stopniu wyjściowym ma rezystor emiterowy o wartości 0,47  $\Omega$ , co wymusza mniej więcej równy podział prądu obciążenia między tranzystorami. Rezystory emiterowe pomagają również w niewielkim stopniu ustabilizować prąd spoczynkowy i nieznacznie poprawiają charakterystykę częstotliwościową stopnia wyjściowego, zapewniając prądowe sprzężenie zwrotne.

## Regulacja offsetu wyjścia

Regulację offsetu przeprowadzamy potencjometrem montażowym 100  $\Omega$  (VR1) pomiędzy emiterami pary wejściowej Q1 i Q2. Potencjometr VR1 reguluje równowagę prądu między parą wejściową, co powoduje zmianę przesunięcia prądu stałego na wyjściu. Potencjometr jest ustawiony tak, aby przesunięcie DC (offset) było jak najbliżej 0 V. Wartość ta powinna być utrzymywana w granicach  $\pm 5$  mV.

Ogólnie, offset warto utrzymywać na niskim poziomie, ale jest on szczególnie krytyczny w przypadku użycia transformatora podwyższającego napięcie do pracy z linią 100 V. Wynika to z faktu, że rezystancja DC uzwojenia pierwotnego transformatora jest znacznie niższa niż rezystancja uzwojenia głośnika, co mogłoby doprowadzić do przepływu przez uzwojenie znacznego prądu stałego.

## Zabezpieczenie wyjść

Kluczowe znaczenie ma zapobieganie pracy tranzystorów wyjściowych poza ich bezpiecznym obszarem działania. Wzmacniacz o dużej mocy, taki jak ten, może być narażony na nadmierne obciążenie, co czasami prowadzi do przekroczenia jego limitów.

Na **rysunku 5** zostały przedstawione wykresy prądu kolektora w funkcji napięcia kolektor-emiter ( $V_{CE}$ ) dla sześciu równoległych tranzystorów wyjściowych MJW21196 i MJW21195. Spośród tych dwóch typów, MJW21195 (PNP) ma niższą krzywą SOA, z niższym dopuszczalnym prądem powyżej 150 V niż komplementarny MJW21196, więc ta właśnie krzywa została wykreślona (ciągła linia zielona).

Krzywa SOA opiera się na temperaturze złącza tranzystora wynoszącej 150°C i temperaturze obudowy wynoszącej 25°C. Utrzymanie takiej temperatury obudowy jest jednak dość trudne, zwłaszcza gdy tranzystory rozpraszają znaczną moc.

Rzeczywista temperatura obudowy tranzystora zależy od rozpraszania, rezystancji termicznej złącza każdego tranzystora do jego obudowy (0,7°C/W) oraz rezystancji termicznej obudowy do otoczenia, która jest określana przez radiator i wentylatory. Posiadanie dużego radiatora z wymuszonym przez wentylator obiegiem powietrza znacznie pomaga utrzymać niską temperaturę tranzystorów.

Należy uwzględnić temperaturę otoczenia (warunki w których będzie pracował wzmacniacz) i zadbać o to, aby tranzystory nie pracowały z mocą przekraczającą ich maksymalną moc znamionową, która wynosi 200 W przy 25°C i maleje o 1,43 W na każdy stopień Celsjusza. W warunkach nadmiernej temperatury krzywa mocy znamionowej może ulec zmianie, co nie pozostanie

**Pierwszy (niniejszy) odcinek dotyczący wzmacniacza 500 W zawiera opis działania modułu wzmacniacza. Montaż i testowanie zostaną omówione w kolejnych częściach**

bez wpływu na bezpieczny obszar pracy (SOA) tranzystorów.

Została wykreślona zarówno krzywa mocy w temperaturze obudowy 25°C (krzywa zielona), jak i krzywa mocy w temperaturze obudowy 50°C (krzywa fioletowa). Wprawdzie sześć tranzystorów o mocy 200 W pozwala na uzyskanie łącznej mocy 1200 W w temperaturze 25°C, lecz jedynie 985 W jest dopuszczalne przy temperaturze obudowy 50°C.

Krzywe zostały wykreślone przy skądinąd słusznym założeniu, że każdy z sześciu równoległych tranzystorów równo dzieli prąd. Założenie jest słuszne, ponieważ każdy z nich ma rezystor emiterowy o stosunkowo wysokiej wartości. Jeśli jeden z tranzystorów mocy pobiera więcej niż wynosi jego udział w prądzie obciążenia, spadek napięcia na jego rezystorze emiterowym będzie proporcjonalnie wyższy. Spowoduje to zdlawienie tranzystora, aż jego prąd powróci do poziomu pozostałych.

Niebieskie i czerwone krzywe przedstawiają obciążenia rezystancyjne 8  $\Omega$  i 4  $\Omega$  (linie proste), które zakładają, że obciążenie jest czysto rezystancyjne. W praktyce, gdy obciążeniem są głośniki, nie jest to prawda. Występująca w nich znaczna reaktancja bierna powoduje, że impedancja zmienia się wraz z częstotliwością.

Zakrzywione niebieskie i czerwone linie pokazują krzywe impedancji obciążenia przy założeniu, że składowa rezystancyjna i reaktancyjna są równe. Wykresy pokazują najgorszą impedancję występującą w zakresie częstotliwości roboczych.

Na przykład, dla głośnika  $4\ \Omega$ , wykresamy krzywą z rezystancją  $2,83\ \Omega$  i reaktancją  $2,83\ \Omega$ , która jest o  $90^\circ$  poza fazą („j” jest jak „i” w matematyce, jest to jednostka urojona o wartości  $\sqrt{-1}$ , tworząca urojoną składową impedancji).

Obliczenie całkowitej impedancji można przedstawić jako dwie składowe tworzące dwa boki trójkąta prostokątnego o długości przeciwprostokątnej równej sumie, która w tym przypadku wynosi  $4\ \Omega$  lub  $8\ \Omega$ .

Wykresy te dotyczą dość dużego obciążenia wzmacniacza. Zazwyczaj głośnik nie wykazuje takiego obciążenia, ale chcemy mieć pewność, że wzmacniacz nie zostanie uszkodzony, dlatego uwzględniamy najgorszy przypadek.

Należy zauważyć, że zakrzywione wykresy impedancji zbliżają się do krzywej SOA bardziej niż w przypadku obciążeń

czysto rezystancyjnych. Należy również zauważyć, że w podwyższonych temperaturach dopuszczalna krzywa rozpraszania zbliża się do wykresu składowej reaktancyjnej impedancji  $4\ \Omega$ , szczególnie w obszarze  $V_{CE}$  od  $60\ \text{V}$  do  $100\ \text{V}$ . Przy temperaturach obudowy powyżej  $50^\circ\text{C}$ , dopuszczalne rozpraszanie tranzystora może zostać przekroczone.

Zapobiegają temu dwie linie ochronne na wykresie. Przerywana zielona linia dotyczy temperatury obudowy tranzystora wynoszącej  $25^\circ\text{C}$ , natomiast przerywana fioletowa linia dotyczy temperatury obudowy wynoszącej  $50^\circ\text{C}$ . Linie pokazują punkty na wykresie, w których tranzystory wyjściowe są chronione poprzez zmniejszenie ich bazowego wystrojenia, jeśli obciążenie osiągnie linię ochronną.

Linie ochronne przesuwają się bliżej krzywej impedancji  $4\ \Omega$  wraz ze wzrostem temperatury.

Ponadto linie ochronne mają podwójne nachylenie z jedną linią prostą między osią Y a małym okręgiem (kropką) i drugą linią między tą kropką a osią X. Należy pamiętać, że aby zapobiec fałszywemu ograniczeniu w pobliżu zera prądu wyjściowego, w punkcie, w którym linia styka się z osią X, musi występować co najmniej całkowite napięcie zasilania ( $160\ \text{V}$ ).

Wraz ze wzrostem temperatury napięcie na osi zerowego prądu maleje. Jednak nawet krzywa  $50^\circ\text{C}$  spotyka się z osią powyżej  $160\ \text{V}$ , przy  $165\ \text{V}$ . Jeśli wzmacniacz znacznie się nagrzeje, być może powyżej  $60^\circ\text{C}$ , wyjście prawdopodobnie zostanie odcięte, ale może to nie jest złe rozwiązanie, ponieważ jest to znak, że system chłodzenia mógł zawieść.

Chociaż różnica między dwoma nachyleniami krzywej ochrony jest subtelna, jest to konieczne, aby ściślej podążać za krzywą mocy znamionowej, a tym samym zapobiegać wkraczaniu krzywej ochrony w temperaturze  $50^\circ\text{C}$  i wyższej na krzywą impedancji  $4\ \Omega$  przy napięciu około  $70\ \text{V}$ .

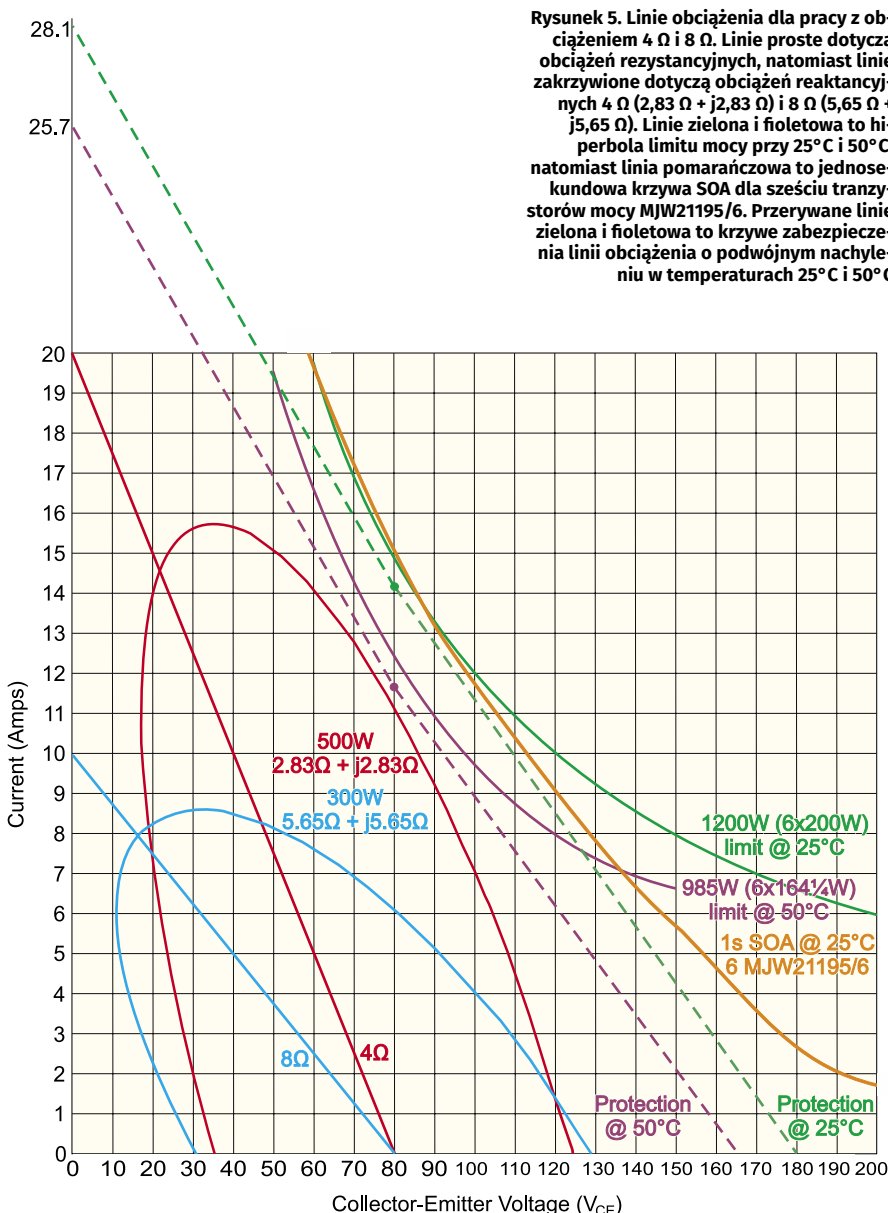
## Układ zabezpieczający SOA

Schemat zabezpieczenia typu foldback o podwójnym nachyleniu jest oparty na artykule „The Safe Operating Area (SOA) Protection of Linear Audio Power Amplifiers” autorstwa Michaela Kiwanuki, B.Sc. (Hons) Electronic Engineering, który jest dostępny na stronie [siliconchip.com.au/link/abc4](http://siliconchip.com.au/link/abc4).

Napięcie zasilania, napięcie wyjściowe i prąd płynący przez tranzystory wyjściowe są monitorowane w celu zapewnienia ochrony linii obciążenia w całym zakresie napięcia i prądu wzmacniacza.

Tranzystory Q25 i Q26 oraz diody D6 i D7 zapewniają funkcję zabezpieczającą. Tranzystor Q25 (NPN) może wyłączać tranzystory MJW21196, natomiast Q26 (PNP) działa na tranzystory MJW21195. Diody są dołączone w celu zapobieżenia bocznikowania sygnału sterującego przez Q25 i Q26, gdy są one odwrotnie polaryzowane. Dzieje się tak przy każdym półcyklu sygnału sterującego tranzystorami.

Obwody w otoczeniu tranzystorów Q25 i Q26 są zasadniczo identyczne. W warunkach normalnych Q25 i Q26 są wyłączone i nie wpływają na pracę wzmacniacza. Jeśli jednak obciążenie przekroczy krzywą ochronną, tranzystory Q25 i/lub Q26 włączają się, aby obniżyć zasilanie tranzystorów wyjściowych, ograniczając tym samym prąd wyjściowy i chroniąc tranzystory. Chroni to również układ przed zwarciami. Tranzystory Q25 i Q26 są zamontowane na radiatorze wzmacniacza, dzięki czemu krzywe obwodu zabezpieczającego zmieniają się w zależności od temperatury.



Mówiąc bardziej szczegółowo, napięcie na każdym rezystorze emiterowym stopnia wyjściowego o wartości  $0,47\ \Omega$  jest monitorowane przez zestaw rezystorów  $3,3\ k\Omega$ . Napięcia te są uśredniane (równoważne sumowaniu) na bazach tranzystorów Q25 lub Q26. Dzielniki rezystancyjne utworzone z par równoległych rezystorów zapewniają monitorowanie napięcia wyjściowego i napięcia zasilania poprzez doprowadzenie dodatkowego prądu do tych punktów sumowania. W efekcie, dzielniki te sprawiają, że wraz ze spadkiem napięcia na zestawie rezystorów wyjściowych (z powodu zmniejszonego napięcia zasilania, albo wychylenia wyjścia bliżej linii zasilającej), obwód zabezpieczający staje się mniej wrażliwy i wymaga wyższego prądu wyjściowego do wyzwolenia. Podobnie, wraz ze wzrostem  $V_{CE}$ , prąd wyzwala się maleje, tworząc krzywe pokazane na rysunku 5.

Podwójne nachylenie w obwodzie zabezpieczającym jest tworzone przez napięcie referencyjne REF1 dla dodatniej połowy układu i REF2 dla połowy ujemnej. Prąd polaryzacji tych tranzystorów jest dostarczany przez rezystory szeregowo  $18\ k\Omega$ . REF1 i REF2 są regulowanymi źródłami napięcia odniesienia. Rezystory  $10\ k\Omega$  i  $3,3\ k\Omega$  ustalają na nich napięcie  $10\ V$ .

Układ zabezpieczający opiera się na napięciu baza-emiter tranzystorów Q25/Q26 wynoszącym około  $0,6\ V$  w temperaturze  $25^\circ C$ . Napięcie to spada do  $0,55\ V$  w temperaturze  $50^\circ C$ , więc tranzystory te włączają się przy mniejszym przyłożonym napięciu w wyższych temperaturach. Powoduje to przesunięcie linii ochronnej w dół wraz ze wzrostem temperatury, zgodnie z ruchem w dół krzywej mocy znamionowej tranzystorów wyjściowych.

Diody D4 i D5 pomiędzy wyjściem wzmacniacza a liniami zasilającymi są również częścią układu zabezpieczającego. Pochłaniają one wszelkie duże skoki generowane przez indukcyjność głośnika, gdy układ zabezpieczający odcina zasilanie tranzystorów wyjściowych. Diody D4 i D5 to diody Fast recovery, dołączone w celu zapewnienia ich działania przy wysokich częstotliwościach i dużej mocy. Są one jeszcze bardziej krytyczne w przypadku zasilania transformatora liniowego, ponieważ jego indukcyjność pierwotna jest prawdopodobnie znacznie wyższa niż w przypadku jakiegokolwiek obciążenia głośnika.

## Wyjściowy filtr RLC

Wyjściowy filtr RLC (rezystor, cewka i kondensator), składający się z dławika powietrznego  $2,2\ \mu H$ , ośmiu równoległych



Zbliżenie obwodów front-end modułu wzmacniacza 500 W

rezystorów  $56\ \Omega$  (tworzących rezystancję  $7\ \Omega$ ) i kondensatora  $100\ nF$ . Ten filtr wyjściowy skutecznie izoluje wzmacniacz od wszelkich dużych reaktancji pojemnościowych w obciążeniu, zapewniając tym samym bezwzględną stabilność.

Pomaga również tłumić wszelkie sygnały radiowe indukowane na przewodach głośnikowych i zapobiega ich przesyłaniu z powrotem do pierwszych stopni wzmacniacza, gdzie mogą przekazywać zakłócenia radiowe.

## Zabezpieczenie bezpiecznikiem

Linie zasilania stopnia wyjściowego są zasilane przez bezpieczniki F1 i F2 z głównych linii zasilania  $+80\ V$  i  $-80\ V$ . Zapewniają one „ostatnią deskę ratunku” dla wzmacniacza, ograniczając uszkodzenia w przypadku poważnej usterki. Zalecane są bezpieczniki ceramiczne. Dioda LED2 jest wskaźnikiem przepalonego bezpiecznika dla F1, a dioda

LED3 dla F2. Zapalają się, gdy bezpiecznik jest przepalony, ponieważ nie zawsze jest to oczywiste, zwłaszcza w przypadku bezpieczników ceramicznych.

## W następnym odcinku

Następny artykuł będzie zawierał pełne szczegóły budowy modułu wzmacniacza, w tym wiercenie radiatora i instrukcje dotyczące nawijania cewki indukcyjnej L1.

Następnie pokażemy, jak zbudować odpowiedni zasilacz, zamontować go wraz z modulem wzmacniacza w obudowie i okablować wszystko wraz ze sterownikiem wentylatora i wentylatorami. Zamieścimy również instrukcję montażu osłony głośnika i wskaźnika przesterowania. ■

John Clarke

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. [www.siliconchip.com.au](http://www.siliconchip.com.au)



Gotowa obudowa jest bardzo prosta, zawiera jedynie przycisk zasilania i wskaźnik przesterowania LED z przodu oraz wejście/wyjście audio i gniazdo zasilania z tyłu

Większość pilotów zdalnego sterowania wykorzystuje impulsy światła podczerwonego do sterowania urządzeniami. Zwykle działa to niezawodnie tylko do kilku metrów i jest łatwo blokowane przez meble, ludzi, rośliny... prawie wszystko. Przekształcenie pilota na podczerwień tak, aby korzystał z częstotliwości UHF sprawi, że radykalnie zostanie zwiększony jego zasięg. Będzie działał nawet wtedy, gdy jakaś przeszkoda znajdzie się między pilotem a urządzeniem, niezależnie od tego, w którą stronę jest skierowany!



## Wzmacniacz zasięgu pilota na podczerwień

W większości przypadków piloty na podczerwień działają bardzo dobrze. Zdarzają się jednak sytuacje, w których są całkowicie bezużyteczne. Może to wynikać z faktu, że między pilotem a sterowanym urządzeniem znajduje się przeszkoda, lub odbiornik w urządzeniu może być nieodpowiednio umieszczony, co utrudnia skierowanie na niego wiązki podczerwieni.

Czasami nawet chcielibyśmy użyć pilota zdalnego sterowania w innym pomieszczeniu niż urządzenie, które ma być sterowane.

Konieczne może być też ustawienie urządzenia w taki sposób, aby odbiornik nie był skierowany w kierunku miejsca, w którym zwykle znajduje się użytkownik. Na przykład projektor, zwykle znajduje się za operatorem. Czasami można odbijać sygnały podczerwieni od ekranu do projektora, ale nie zawsze taka metoda działa niezawodnie.

Niezależnie od tego, dlaczego sygnał podczerwieni nie sprawdza się w danym zastosowaniu, opisane urządzenie jest świetnym rozwiązaniem.

Umożliwia ono konwersję sygnału pilota na podczerwień w taki sposób, aby nadawał sygnały radiowe UHF zamiast światła podczerwonego. Kolejne małe pudełko umieszczone przed odbiornikiem podczerwieni w urządzeniu docelowym odbiera te sygnały radiowe i przesyła podczerwień bezpośrednio do odbiornika.

Jeśli chcemy sterować kilkoma urządzeniami, do przekazywania sygnałów do każdego z nich można użyć pojedynczego konwertera sygnału UHF na podczerwień. Pod warunkiem, że urządzenia znajdują się w pobliżu, więc światło z jednego nadajnika może dotrzeć do wszystkich odbiorników.

### Koncepcja

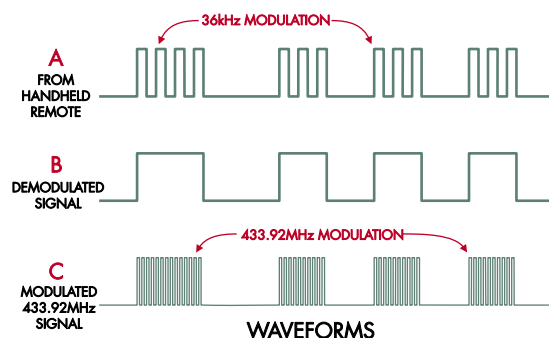
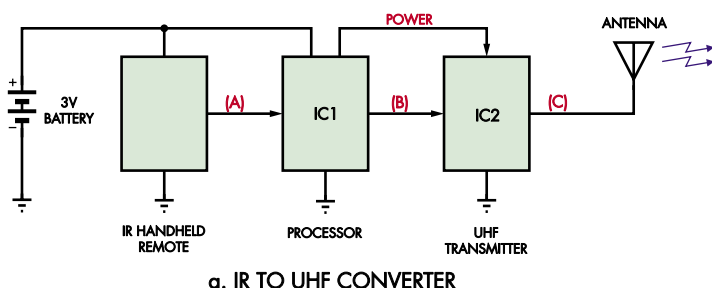
Ogólny układ wzmacniacza zasięgu przedstawiono na **rysunku 1**. Na rysunku 1a została zilustrowana zasada działania konwertera IR na UHF, natomiast rysunek 1b przedstawia zasadę działania konwertera UHF na podczerwień.

Konwerter IR na UHF monitoruje sygnał, który normalnie byłby podawany do diody LED IR w pilocie zdalnego sterowania. Po naciśnięciu dowolnego przycisku na pilocie, wytwarza on modulowany sygnał o częstotliwości ~36 kHz do sterowania tą diodą. Zamiast tego układ IC1 demoduluje ten sygnał, a jego wyjście (kształt przebiegu B) jest pokazane na **oscylogramie 1 i 2**.

„Demodulacja” konwertuje serię krótkich impulsów 36 kHz na sygnał, który jest wysoki, gdy impulsy są obecne i niski w przeciwnym razie.

Gdy IC1 wykryje, że odbiera sygnał, zasila nadajnik UHF (IC2) i wysyła zdemodulowany sygnał do wejścia nadajnika UHF. W rezultacie nadajnik UHF wytwarza modulowany sygnał 433,92 MHz do anteny nadawczej. Jest to przebieg C.

Oryginalny sygnał modulowany 36 kHz jest konwertowany na sygnał modulowany 433,92 MHz do transmisji bezprzewodowej.



**Rysunek 1a.** Przedłużacz zasięgu pilota składa się z dwóch części. Pierwszą z nich jest konwerter IR na UHF, który jest zasilany z baterii pilota i konwertuje sygnał sterujący diodą LED IR na transmisję UHF. Druga to konwerter UHF na IR, który odbiera sygnały UHF i steruje diodą podczerwoną z odpowiednią modulacją w celu sterowania docelowym urządzeniem (urządzeniami)

## Konwerter IR na UHF

- Zasięg transmisji: 25 m przez jedną ścianę Hardiplank i Gyprock
- Opóźnienie sygnału: 56  $\mu$ s
- Okres wyłączenia zasilania nadajnika UHF: 600 ms po ostatnim sygnale
- Prąd czuwania: 80 nA typowo przy zasilaniu 3 V (90 nA zmierzone)
- Prąd roboczy: średnio 8 mA podczas transmisji

## Konwerter UHF na podczerwień

- Wykrywanie poprawnej transmisji: wymaga minimalnego okresu ciszy 3 ms
- Podświetlenie diody LED potwierdzenia: 654 ms limitu czasu po prawidłowym sygnale
- Częstotliwość modulacji: 32,4 kHz do 41,4 kHz w 32 krokach
- Cykl pracy modulacji: 33,3%
- Pobór prądu: blisko 50 mA podczas odbioru sygnału
- Zasięg transmisji w podczerwieni: typowo 2 m do odbiornika urządzenia



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/d2dw8>

Materiały dodatkowe są również dostępne na stronie [elportal.pl/do-pobrania](http://elportal.pl/do-pobrania)

Odpowiedni konwerter UHF na podczerwień ma odbiornik UHF (RX1), który dostarcza zdemodulowany przebieg, pokazany jako przebieg D. Jest on zgodny z przebiegiem B – patrz **oscylogram 3**. Mikrokontroler IC1 na drugiej płytce generuje następnie nową nośną 36 kHz używaną do wytworzenia modulowanego przebiegu E, który odpowiada oryginalnemu kształtowi przebiegu A, jak pokazano w **oscylogramach 4 i 5**.

Modulowany sygnał steruje następnie diodą LED IR, która wysyła sygnał do urządzenia (urządzeń) sterowanego (sterowanych) podczerwienią.

Należy pamiętać, że częstotliwość 36 kHz jest charakterystyczna dla pilotów pracujących w podczerwieni. Częstotliwość modulacji końcowego wyjścia podczerwieni można dostosować do częstotliwości oryginalnego pilota zdalnego sterowania, ponieważ pilot zdalnego sterowania może generować inną częstotliwość w zakresie od około 32 kHz do 41 kHz.

Oryginalny sygnał pilota jest powielany na wyjściu konwertera UHF na podczerwień. Urządzenie odbierające sygnał nie jest „świadome”, że nastąpiło jakieś przetwarzanie.

## Wcześniejsze projekty

Warto zauważyć, że podobny projekt był już opublikowany pod tytułem „Add a UHF link to a universal remote control” (lipiec 2013; [siliconchip.com.au/Article/3846](http://siliconchip.com.au/Article/3846)). Chociaż tamten projekt jest nadal aktualny, ten ma znacznie mniejsze wymiary nadajnika, który w przeciwieństwie do tego z 2013 roku można zamontować wewnątrz nawet bardzo małego pilota na podczerwień.

Próba zainstalowania wcześniejszego projektu w małym pilocie do projektora LCD skończyła się niepowodzeniem, dlatego cały układ konwertera IR-UHF został przeprojektowany i oparty na elementach do montażu powierzchniowego.

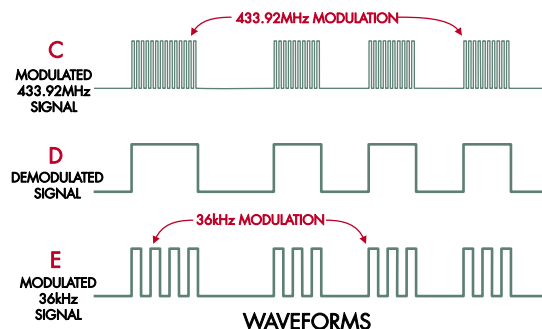
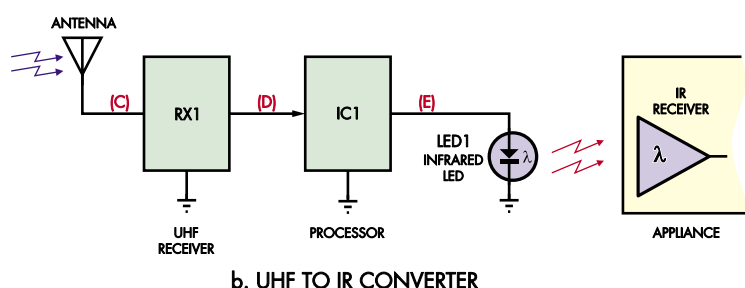
Zamiast korzystać z dużego, gotowego modułu nadajnika UHF, można teraz używać bardzo małego układu scalonego nadajnika UHF z kilkoma elementami dyskretnymi.

## Żywotność baterii pilota

Pojawia się pytanie, co stanie się z żywotnością baterii zmodyfikowanego pilota. Czy bateria ulegnie rozładowaniu w krótkim czasie po dodaniu układu nadajnika UHF?

Jak wykazały testy, po przełączeniu obwodów w tryb uśpienia, gdy pilot nie jest używany, będzie to miało tylko znikomy wpływ. Typowy pilot na podczerwień pobiera około 1...2  $\mu$ A z baterii w sposób ciągły i około 10...20 mA podczas transmisji w podczerwieni. Dodatkowy pobór mocy nadajnika UHF nie ma prawie żadnego wpływu na te wartości.

Po zainstalowaniu konwertera IR na UHF, prąd w trybie czuwania wzrósł o zaledwie 90 nA! Prąd pobierany po naciśnięciu



Rysunek 1b. Przebiegi po prawej stronie, zarówno tu, jak i na rysunku 1a obok, pokazują, jak oryginalny sygnał sterujący diodą IR LED jest demodulowany, następnie remodulowany do 433,92 MHz, następnie demodulowany, i ponownie remodulowany do około 36 kHz, aby sterować diodą IR LED



W celu wytworzenia nośnej UHF jego częstotliwość jest mnożona przez 32 w układzie IC2. Tak więc używając rezonatora 13,56 MHz uzyskujemy częstotliwość nośną 433,92 MHz. Jest ona odpowiednia do częstotliwości nośnej używanej w większości modułów nadajnika/odbiornika UHF ASK, które są dostępne do transmisji danych UHF o niskiej mocy.

Układ MICRF113 i powiązane z nim elementy są niewielkie i mieszczą się w znacznie mniejszej przestrzeni niż większość dostępnych gotowych modułów nadajników UHF.

Prąd zasilający dla stopnia wyjściowego RF IC2 jest dostarczany przez dwie szeregowo połączone cewki 220 nH, działające również jako obciążenie 440 nH sterownika. Kondensator szeregowy 12 pF i cewka 68 nH oraz kondensator 5 pF połączony do masy działają jako filtr, który usuwa drugą i trzecią harmoniczną z sygnału UHF, zanim przejdzie on do anteny.

Używane są dwie cewki 220 nH zamiast jednej cewki 470 nH, ponieważ cewki 220 nH są łatwiejsze do pozyskania. Każda cewka indukcyjna użyta w układzie musi mieć częstotliwość rezonansu własnego (SR) powyżej 433,92 MHz, w przeciwnym razie nie będzie działać jako cewka indukcyjna przy tej częstotliwości.

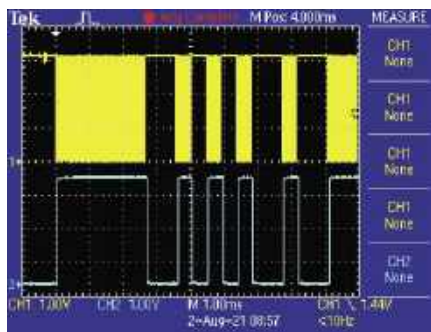
## Zasilanie układu IC2

Linia zasilająca IC2 (nóżka 3) jest filtrowana kondensatorem ceramicznym 1  $\mu$ F, natomiast kondensator 100 nF filtruje zasilanie stopnia wyjściowego. Oba te kondensatory są zasadniczo połączone równolegle, ale znajdują się w różnych miejscach na płycie drukowanej, dzięki czemu zasilanie każdej części jest filtrowane bezpośrednio przy jej połączeniu z linią zasilania.

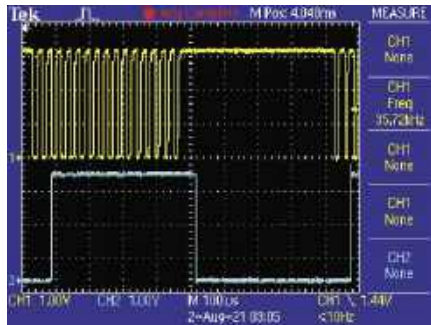
Pomiędzy sygnałem ASK a zasilaniem IC2 została umieszczona dioda Schottky'ego D2, która pomaga utrzymać stabilne napięcie zasilania, gdy układ scalony nadaje. Gdy IC2 pracuje, napięcie na nóżce 3 spada. Prąd płynący z nóżki 4 układu IC1 przez diodę D2 wspomaga utrzymanie stabilnego napięcia zasilania dla IC2.

Chociaż układ IC2 działa poprawnie nawet przy napięciu 1,8 V, dla optymalnej wydajności zaleca się utrzymywanie napięcia zasilania blisko 3 V.

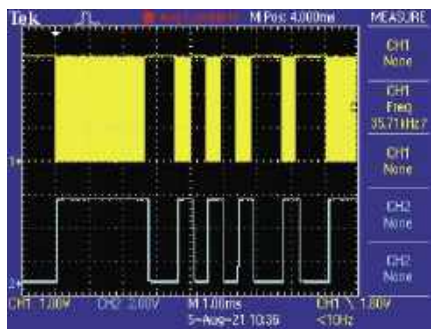
Zasilanie układu IC1 jest filtrowane przez kolejny kondensator ceramiczny 100 nF. Dioda D1 jest dołączona na wypadek, gdyby ogniwa w pilocie zostały włożone odwrotnie, powodując odwrotną polaryzację. W takim przypadku dioda D1 przewodzi i zmniejsza napięcie wsteczne przyłożone do układu IC1, zapobiegając jego uszkodzeniu (przynajmniej w krótkim okresie).



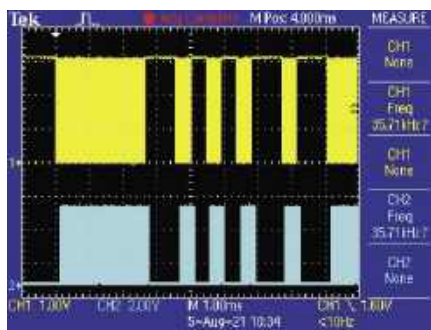
Oscylogram 1. Górny żółty przebieg to sygnał sterujący diodą LED IR pilota zdalnego sterowania, zdjęty z nóżki 1 układu scalonego IC1. Jest to seria impulsów o częstotliwości 36 kHz. Dolny niebieski przebieg jest pobrany z wyjścia IC1 na nóżce 4, które steruje wejściem ASK (nóżka 6) nadajnika MICRF113 434 MHz (IC2). Sygnał ten jest wysoki, gdy na wejściu występuje sygnał 36 kHz, a niski w przeciwnym przypadku



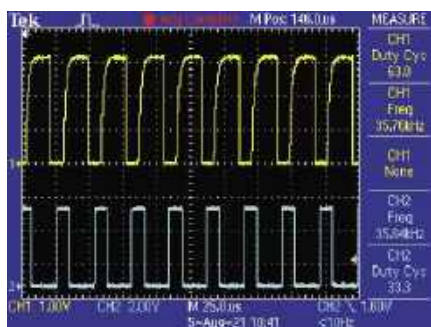
Oscylogram 2. Jest to ten sam zrzut, co oscylogram 1, ale z rozciągniętą podstawą czasu, więc widoczna jest modulacja 36 kHz. Należy zwrócić uwagę na opóźnienie około 56  $\mu$ s między wejściem IC1 odbierającym impulsy 36 kHz i wytwarzającym demodulowane impulsy na swoim wyjściu. Nie powoduje to zniekształcenia sygnału, ponieważ jest on symetryczny



Oscylogram 3. Górny żółty przebieg to sygnał podczerwieni z pilota, jak na oscylogramie 1, ale dolny przebieg jest sygnałem wyjściowym z odbiornika UHF w konwerterze UHF-IR, tj. po przejściu przez łączące bezprzewodowe



Oscylogram 4. Górny żółty przebieg to sygnał sterujący diodą LED IR z oryginalnego pilota na podczerwień, a niebieski przebieg dolny to sygnał sterujący diodą LED IR w konwerterze UHF-IR. Oba przebiegi są zasadniczo takie same, z wyjątkiem niewielkiego opóźnienia drugiego przebiegu różnych poziomów napięcia ze względu na układ UHF-IR zasilany napięciem 5 V zamiast 3 V. Odwrócenie sygnału nie ma znaczenia



Oscylogram 5. Rozciągnięty przebieg z oscylogramu 4 pokazujący modulację obu sygnałów. Czas narastania oryginalnego przebiegu u góry jest długi ze względu na niski prąd podciągania z nóżki 1 układu PIC10LF322. Dolny niebieski przebieg toysterowanie diody LED IR z konwertera UHF-IR. Częstotliwość została ustawiona na około 36 kHz, aby odpowiadała częstotliwości typowej dla pilotów działających w podczerwieni. Górny przebieg jest odwrócony w porównaniu z dolnym, ponieważ oryginalna dioda LED IR w pilocie była włączana niskim stanem na wyjściu, a dioda LED IR w konwerterze UHF-IR jest aktywna w stanie wysokim

## Konwerter UHF na podczerwień

W celu sterowania obsługiwanym urządzeniem sygnał UHF musi być wykryty i przekonwertowany z powrotem na strumień impulsów podczerwieni. Układ konwertera UHF na podczerwień pokazano na **rysunku 3**. Składa się on z odbiornika UHF RX1, mikrokontrolera PIC12F617 (IC1) i podczerwonej diody LED (LED1).

Układ jest zasilany z gniazda DC CON1 lub gniazda micro-B USB CON2. Odbiornik UHF jest zasilany w sposób ciągły, gotowy do odbioru transmisji z konwertera IR na UHF zamontowanego w pilocie.

Przy braku sygnału, dane wyjściowe z odbiornika UHF to tylko losowy szum o amplitudzie 5 V. W tym stanie odbiornik działa z maksymalnym wzmocnieniem dzięki automatycznej regulacji wzmocnienia (AGC). Po odebraniu sygnału UHF, AGC zmniejsza czułość odbiornika tak, że wykryty sygnał jest zasadniczo wolny od szumów. Jest on podawany do wejścia GP5 (nóżka 2) mikrokontrolera PIC IC1.

Aby określić, czy sygnał jest prawidłowy, mikrokontroler IC1 sprawdza okresy, w których linia danych z odbiornika UHF ma wartość 0 V przez co najmniej 3 ms. Wskazuje to, że AGC zmniejszył czułość odbiornika i że zachodzi transmisja.

Dane wyjściowe z odbiornika UHF odpowiadają danym przesyłanym do nadajnika UHF. Sygnał danych, częściowo staje się

przebiegiem potwierdzenia, który steruje diodą LED2 poprzez wyjście cyfrowe GP0. Rezystor 1 k $\Omega$  ogranicza prąd diody LED do około 3 mA.

Układ scalony IC1 zasila diodę IR LED (LED1) z połączonych równolegle wyjść GP1 i GP2. Takie połączenie ma na celu zapewnienie wystarczającego prądu. Rezystor 220  $\Omega$  ogranicza ten prąd do około 18 mA.

Sygnał sterujący podczerwonymi diodami LED musi zawierać taką samą lub podobną modulację, jak ta używana w oryginalnym pilocie. Gdy wyjście danych z odbiornika UHF przechodzi w stan wysoki, wyjścia GP1 i GP2 są sterowane sygnałami modulowanymi szerokością impulsu. Cykl pracy wynosi 33,3%, więc pozostają one w stanie wysokim przez 1/3 czasu i niskim przez 2/3 czasu.

Wejście GP4 układu IC1 monitoruje napięcie ustawione przez potencjometr montażowy VR1, podłączony do linii zasilającej 5 V. Napięcie na jego suwaku jest przetwarzane na wartość cyfrową w układzie IC1, umożliwiając dostosowanie częstotliwości nośnej podczerwieni do oryginalnego nadajnika. Zakres regulacji wynosi od 32,4 kHz do 41,4 kHz w 32 krokach. Ustawienie VR1 w pozycji środkowej odpowiada częstotliwości 37 kHz.

Zazwyczaj ustawienie zbliżone do środkowego jest zadowalające, ale niektóre urządzenia mogą do niezawodnego działania wymagać innej częstotliwości nośnej.

Drugie wyjście jest dostępne za pośrednictwem gniazda jack 3,5 mm CON3 dla

zewnętrznej diody LED IR (w razie potrzeby). Dioda ta może być zamontowana w pobliżu odbiornika podczerwieni obsługiwanego urządzenia.

Zasilanie zasilacza wtyczkowego 9...12 VDC jest doprowadzane przez diodę D1, zapewniając ochronę przed odwrotną polaryzacją. 3-nóżkowy stabilizator 78L05 zapewnia zasilanie 5 V dla RX1 i IC1. Zasilanie przez złącze USB jest doprowadzane do linii zasilającej 5 V przez rezystor 4,7  $\Omega$ . Rezystor ten zapobiega nadmiernemu przepływowi prądu między wyjściem REG1 a linią 5 V z USB, jeśli oba są podłączone.

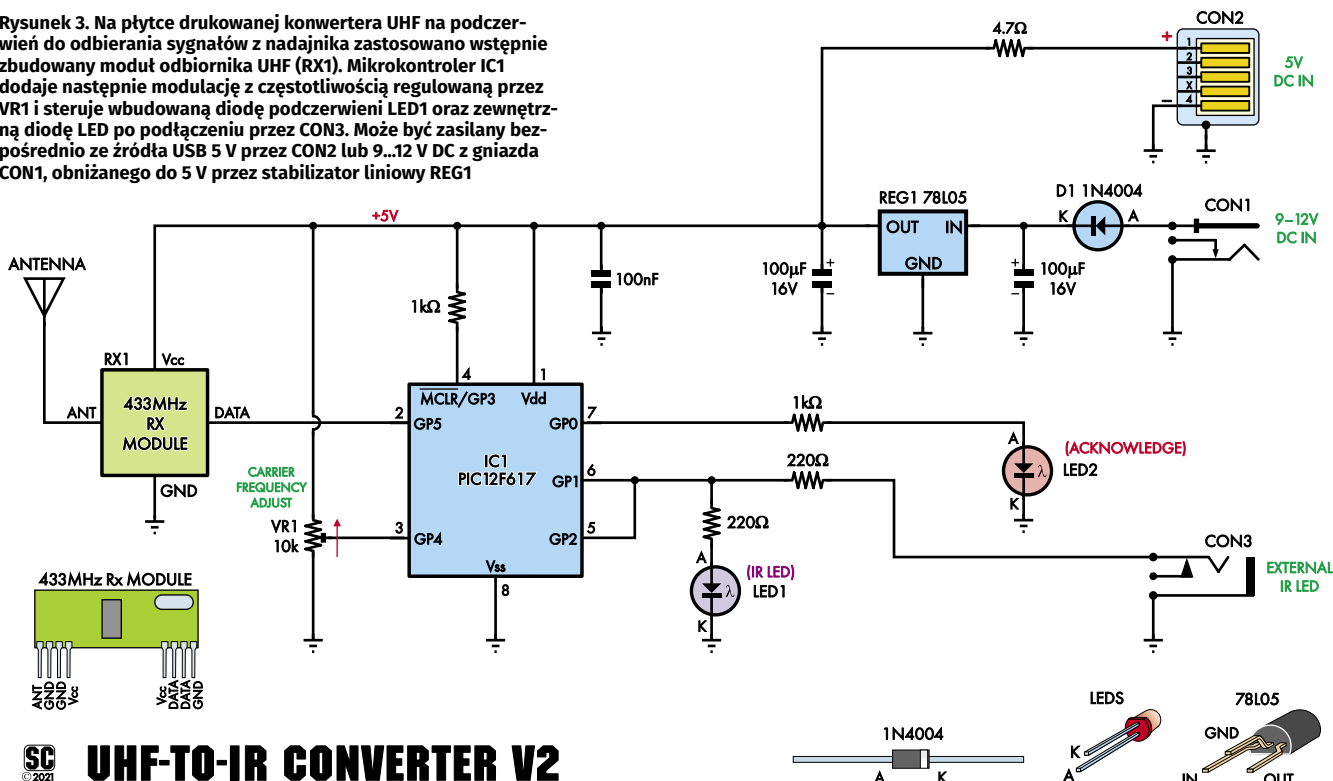
## Budowa

Płyta PCB konwertera IR na UHF ma oznaczenie 15109212. Ma ona wymiary: 15 mm  $\times$  12 mm. Posiada elementy zamontowane po obu stronach. Aby zobaczyć, jak zamontować części należy zapoznać się ze schematami PCB przedstawionymi na **rysunkach 4a i b**.

Jeśli układ IC1 nie został zaprogramowany, należy to zrobić przed jego zamontowaniem. Jeśli czytelnik nie ma odpowiedniego programatora mikrokontrolerów PIC10LF322, może je nabyć w sklepie internetowym Silikon Chip w wersji zaprogramowanej.

Montaż należy rozpocząć od zamontowania elementów do montażu powierzchniowego na górnej stronie płytki drukowanej. Można je przylutować za pomocą lutownicy z ciekłym grotem. Podczas montażu na pewno

**Rysunek 3.** Na płycie drukowanej konwertera UHF na podczerwień do odbierania sygnałów z nadajnika zastosowano wstępnie zbudowany moduł odbiornika UHF (RX1). Mikrokontroler IC1 dodaje następnie modulację z częstotliwością regulowaną przez VR1 i steruje wbudowaną diodą podczerwieni LED1 oraz zewnętrzną diodą LED po podłączeniu przez CON3. Może być zasilany bezpośrednio ze źródła USB 5 V przez CON2 lub 9...12 V DC z gniazda CON1, obniżanego do 5 V przez stabilizator liniowy REG1





przyda się lupa lub okulary. Do przytrzymywania elementów pomocna może być również precyzyjna pęseta z ostrymi końcami.

Łatwiej będzie najpierw zamontować dwie cewki 220 nH. Najpierw lutujemy jedno wyprowadzenie i sprawdzamy ułożenie cewki. W razie potrzeby ponownie podgrzewamy lutowany pad i ewentualnie korygujemy jej położenie. Następnie lutujemy drugie wyprowadzenie.

W kolejnym kroku montujemy dwa układy scalone. Układy IC1 i IC2 są umieszczone w taki sposób, aby mała kropka na nóżce 1 pokrywała się z kropką na płytce drukowanej. Gdy układ scalony jest zestawiony z nóżką 1 w lewym dolnym rogu, napis na górnej powierzchni układu scalonego będzie skierowany w prawo. IC1 będzie oznaczony LF, po którym następują dwa numery kodu identyfikacyjnego. IC2 będzie miał wytrawione „F\_113” na górnej powierzchni.

Układy scalone powinny być rozmieszczone na płytce drukowanej z kropką oznakowania nóżki 1 w lewym górnym rogu. Dla każdego układu scalonego lutujemy najpierw jeden pad, a następnie sprawdzamy wyrównanie pozostałych. W razie potrzeby należy skorygować ułożenie elementów poprzez ponowne podgrzanie punktu lutowniczego przed przylutowaniem pozostałych nóżek. Wszelkie zwarcia między pinami można usunąć za pomocą plecionki lutowniczej. Aby skutecznie usunąć nadmiar cyny można użyć topnika w postaci pasty.

Teraz, przed zamontowaniem rezonatora kwarcowego X1 można wlutować diodę D2. Trzeba upewnić się, że jest ona zorientowana tak, jak pokazano na **rysunku 4a**.

Następnie można zainstalować pozostałe elementy montowane od góry. Wiele kondensatorów i cewek indukcyjnych w obudowach do montażu powierzchniowego nie ma oznakowania, więc trzeba będzie polegać na oznaczeniach podawanych na opakowaniu.



**Przed zamontowaniem konwertera IR na UHF wewnątrz pilota należy sprawdzić, czy potrzebny jest rezystor ściągaający sygnał sterowania diodą IR do masy**



Elementy montujemy pojedynczo, aby uniknąć ich pomieszczenia.

Użyto tu kondensatorów i rezystora w nieco mniejszych obudowach 0805 (M2012) w porównaniu do obudów 1206 (M3216), użytych gdzieś indziej. Zapobiegamy w ten sposób tworzeniu podczas ich montażu przypadkowych mostków lutowniczych do sąsiednich komponentów.

Możliwe jest również zgubienie elementów, należy więc zachować ostrożność i, jeśli to możliwe, zaopatrzyć się w części zamiennie (rezystory SMD i kondensatory są zazwyczaj bardzo tanie i sprzedawane w zestawach).

Zalecane jest zamontowanie cewki 68 nH po wlutowaniu kondensatorów 12 pF i 5 pF. W przeciwnym razie cewka ta może zostać przypadkowo odlutowana.

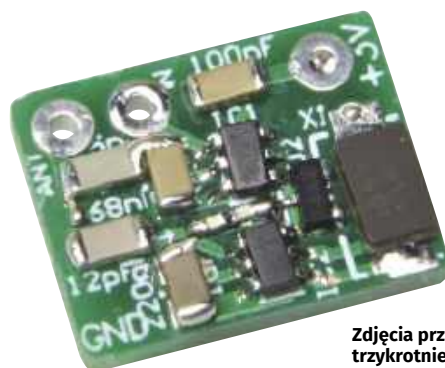
Odwracamy teraz płytkę spodnią stroną do góry. Znajdują się tam dwa kondensatory

18 pF, jeden rezystor 1 kΩ i dioda D1. Pasek katody powinien znajdować się w pozycji pokazanej na **rysunku 4b**. Rezystor może nie być wymagany, więc na razie należy go pominąć.

Aby upewnić się, że elementy zostały przyłutowane prawidłowo, można prześledzić połączenia z innymi sekcjami płytki drukowanej w miejscach gdzie powinna być zachowana ciągłość. Na przykład, nóżka 3 układu IC1 powinna mieć niską rezystancję do nóżki 3 układu IC2. Dodatkowo sprawdzamy, czy nie ma zwarcia między nóżkami elementów na płytce drukowanej, które nie powinny być połączone.

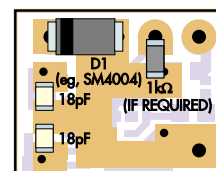
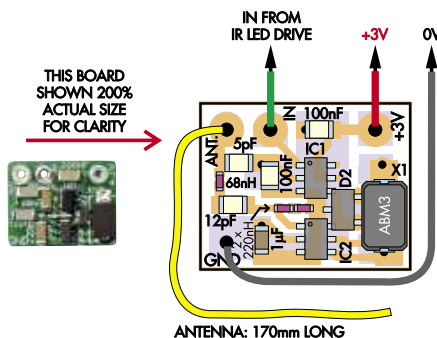
## Podciąganie lub ściągnięcie

Jak wspomniano, pilot zdalnego sterowania może włączać diodę LED IR napięciem niskim lub wysokim. Sposób sterowania diodą LED IR decyduje o tym, czy należy zainstalować rezystor pull-down 1 kΩ. Jeśli rezystor pull-down



**Zdjęcia przedstawiają górną i dolną stronę płytki drukowanej IR-UHF w około trzykrotnie większym rozmiarze**

**Rysunek 4 (po prawej).** Płytkę drukowaną konwertera IR na UHF jest upakowana tak, aby zmieściła się w prawie każdej obudowie pilota zdalnego sterowania. Nie należy zbytnio przejmować się mostkami powstałymi na nóżkach IC1 i IC2 podczas ich lutowania, ponieważ można to dość łatwo naprawić za pomocą plecionki lutowniczej i topnika w postaci pasty, ale trzeba uważać, aby zamontować układy scalone w odpowiednim kierunku i nie pomieszać ich. Cewka indukcyjna 68 nH jest niewielka, więc należy uważać, aby jej nie zgubić. Po przylutowaniu cewki należy sprawdzić, czy między zaciskiem anteny a lewym wyprowadzeniem kondensatora 12 pF występuje niska rezystancja



nie jest zainstalowany, wewnętrzny rezystor pull-up w układzie IC1 zostanie automatycznie aktywowany.

Aby to ustalić, należy najpierw otworzyć obudowę pilota. Niektóre obudowy pilotów są zabezpieczone za pomocą śrub, które są łatwe do znalezienia, ale istnieją również takie, które mają śruby ukryte pod ogniwami. Aby to sprawdzić, należy otworzyć pojemnik baterii i wyjąć ogniwa. Po ich wyjęciu i odkręceniu wszystkich odnalezionych śrub otwieramy obudowę pilota, delikatnie oddzielając część górną od dolnej wzdłuż boków przy użyciu cienkiego narzędzia.

Teraz należy zlokalizować dodatnie i ujemne zaciski baterii. Aby sprawdzić, czy rezystor jest potrzebny, wystarczy wykonać kilka pomiarów za pomocą multimetru.

Po pierwsze, należy sprawdzić rezystancję między biegunem dodatnim baterii a anodą (+) diody IR LED. Jeśli jest ona niska (poniżej 30  $\Omega$ ), można oczekiwać, że rezystor pull-down nie będzie potrzebny. Wynika to z faktu, że dioda LED IR jest sterowana za pomocą katody.

Jeśli rezystancja między katodą (-) diody LED a ujemnym zaciskiem baterii jest niska (mniej niż 30  $\Omega$ ), oznacza to, że dioda IR LED jest załączana stanem wysokim, więc potrzebny będzie rezystor ściągający 1 k $\Omega$ .

Po przylutowaniu rezystora ściągającego (jeśli jest potrzebny), polutowaną płytkę można zamontować w obudowie pilota. Dioda IR LED powinna zostać zdemonstrowana.

Teraz dołączamy zasilania: + do +3 V na pinie, GND do zacisku 0 V i IN do pinu

sterowania IR LED-em w układzie scalonym pilota (np. do punktu, w którym przylutowano diodę IR LED). Do znalezienia odpowiedniego punktu może być konieczne przesłedzenie połączeń na płytce drukowanej.

Umieszczamy płytkę drukowaną w odpowiednim wolnym miejscu w pilocie, lutujemy przewód antenowy i prowadzimy go wokół obudowy w miejscu, w którym nie zostanie przytrzaśnięty podczas ponownego montażu.

Zauważmy, że chociaż określiliśmy długość przewodu antenowego na 170 mm, zasięg transmisji nie ucierpi znacząco, jeśli przewód zostanie skrócony. Okazało się, że przewód antenowy o długości 53 mm zmniejszył zasięg tylko o 5 m w porównaniu do przewodu o długości 170 mm.

Na koniec zamykamy obudowę i ponownie wkręcamy śruby zabezpieczające, jeśli były zastosowane.

## Moduł konwertera UHF na podczerwień

Konwerter UHF na podczerwień jest zbudowany na dwustronnej płytce drukowanej (kod 151009211) o wymiarach 79 mm  $\times$  47 mm. Całość mieści się w plastikowym pudełku UB5 o wymiarach 83 mm  $\times$  54 mm  $\times$  31 mm. Można do niej przymocować etykietę panelu pokrywy o wymiarach 78 mm  $\times$  49 mm.

Jeśli układ IC1 dla tej płytki drukowanej nie został jeszcze zaprogramowany, należy to zrobić teraz przed kontynuowaniem prac. Jeśli czytelnik nie dysponuje odpowiednim programatorem, podobnie jak w przypadku układu SMD, możliwe jest zakupienie układu PIC12F617-I/P w wersji wstępnie zaprogramowanej.

Na **rysunku 5** przedstawiono schemat montażowy. Zaczynamy od montażu gniazda micro USB w wersji SMD. Umieszczamy wyprowadzenia gniazda na odpowiednich polach lutowniczych na płytce, a następnie przylutowujemy jeden z pinów, aby unieruchomić gniazdo na miejscu.

Ponownie sprawdzamy ustawienie nóżek sygnałowych na padach przed ich przylutowaniem, a następnie lutujemy pozostałe wyprowadzenia. Jeśli pozycjonowanie jest nieprawidłowe, lut na nóżce pozycjonującej można ponownie podgrzać, a złącze poprawić.

Sprawdzamy czy na wyprowadzeniach sygnałowych nie ma mostków lutowniczych. Jeśli zostaną zauważone, trzeba je usunąć za pomocą plecionki lutowniczej. Należy po tym ponownie upewnić się, czy nóżki są nadal przylutowane do płytki drukowanej.

Następnie montujemy rezystory. Do rozpoznania ich wartości można sprawdzać umieszczone na obudowie kody, ale dobrym

### Wykaz elementów

#### Konwerter IR na UHF

- 1 dwustronna PCB – kod 15109212, 15 mm  $\times$  12 mm
- 1 rezonator kwarcowy do montażu powierzchniowego 13,56 MHz (X1) [RS Components 171-0468]
- 2 cewki indukcyjne 220 nH 500 MHz, M1005/0402 SMD (L1) [RS 741-3797]
- 1 cewka 68 nH 1,2 GHz, M0602/0201 SMD (L2) [element14 3386563]
- 1 przewód połączeniowy o długości 170 mm (do anteny)
- 1 czerwony przewód połączeniowy o długości 200 mm
- 1 zielony przewód połączeniowy o długości 200 mm
- 1 niebieski przewód połączeniowy o długości 200 mm

#### Półprzewodniki

- 1 PIC10LF322-I/OT 8-bitowy mikrokontroler programowany 1510921M.HEX, SOT-23-6 (IC1) [Silicon Chip Online Shop]
- 1 MICRF113YM6 układ nadajnika ASK UHF, SOT-23-6 (IC2) [RS 177-3314P]
- 1 dioda SMD 1 A, DO-214AC (D1) [SM4004 lub GS1G; Altronics Y0174, Jaycar ZR1003].
- 1 BAT54S\* matosygnatowa dioda Schottky'ego, SOT-23 (D2) [Altronics Y0075].
- \* Odpowiednie są wszystkie modele BAT54, BAT54S, BAT54C, BAT54FILMY i BAT54SFIMLY

#### Kondensatory (wszystkie ceramiczne SMD M2012/0805)

- 1 1  $\mu$ F 16 V X7R (preferowany) lub Y5V [Altronics R8650]
- 2 100 nF 50 V X7R (preferowany) lub Y5V [Altronics R8638]
- 2 18 pF 50 V COG/NPO [Altronics R8533]
- 1 12 pF 50 V COG/NPO [Altronics R8527]
- 1 4,7 pF lub 5 pF 50 V COG/NPO [Altronics R8512]

#### Rezystory

- 1 1 k $\Omega$  SMD M2012/0805 1/8 W (może nie być wymagany; patrz tekst) [Altronics R1220].
- 1 10 k $\Omega$  do 470 k $\Omega$  1/4 W rezystor osiowy (do testowania)

#### Konwerter UHF na podczerwień

- 1 dwustronna PCB – kod PCB 15109211, 79 mm  $\times$  47 mm
- 1 UB5 Jiffy box, 83  $\times$  54 mm  $\times$  31 mm
- 1 etykieta na wieczko, 78 mm  $\times$  49 mm
- 1 moduł odbiornika 433,92 MHz (RX1) [Jaycar ZW3102, Altronics Z6905A].
- 1 gniazdo zasilania do montażu na płytce drukowanej pasujące do zestawu wtyczek (CON1)
- 1 gniazdo USB micro-USB SMD Type-B (CON2) [Jaycar PS0922, Altronics P1309].
- 1 gniazdo jack 3,5 mm z przełącznikiem do montażu na płytce drukowanej (CON3) [Jaycar PS0133, Altronics P0092].
- 1 8-nóżkowa podstawa DIL IC (dla IC1)
- 1 przewód połączeniowy o długości 170 mm
- 1 miniaturowy poziomy potencjometr montażowy 10 k $\Omega$  (VR1)

#### Półprzewodniki

- 1 8-bitowy mikrokontroler PIC12F617-I/P, DIP-8, zaprogramowany programem 1510921A.hex (IC1) [Silicon Chip Online Shop]
- 1 78L05 stabilizator liniowy 5 V 100 mA, TO-92 (REG1)
- 1 dioda LED IR 3 mm (LED1)
- 1 czerwona dioda LED 3 mm (LED2)
- 1 dioda 1N4004 400 V 1 A (D1)

#### Kondensatory

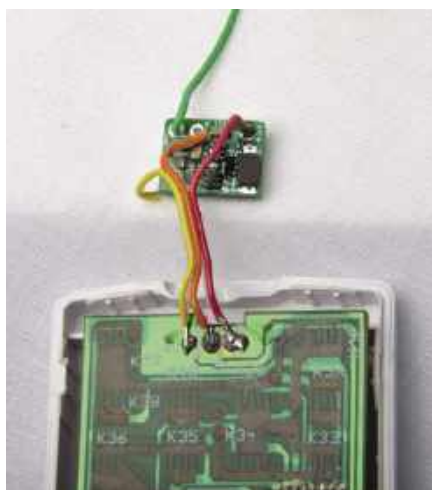
- 2 100  $\mu$ F 16 V PC elektrolityczny
- 1 100 nF 63 V MKT poliestrowy

#### Rezystory (wszystkie 1/4 W 1% cienkowarstwowe osiowe)

- 2 1 k $\Omega$       2 220  $\Omega$       1 4,7  $\Omega$

#### Opcjonalne części do zbudowania diody IR na kablu przedłużającym

- 1 wtyczka mono jack 3,5 mm
- 1 przewód ekranowany jednożyłowy o długości 1 m
- 1 dioda LED IR 3 mm
- 1 rurka termokurczliwa o długości 100 mm i średnicy 3 mm



Dioda IR LED w pilocie jest zastąpiona naszą płytką PCB IR-UHF. Płytką ta może być następnie pokryta folią termokurczliwą i umieszczona w obudowie pilota



Płytkę drukowaną UHF-IR można zamontować wewnątrz obudowy UB5 i umieścić w pobliżu urządzenia odbiorczego. Konieczne będzie wywiercenie otworów w obudowie UB5 na gniazda i diody LED, jak pokazano na rysunku 6

rozwiązaniem jest również sprawdzenie rezystancji za pomocą multimetru.

Następnie montujemy diodę D1, pamiętając o zachowaniu właściwego kierunku montażu, po czym montujemy kondensatory. Tylko dwa kondensatory elektrolityczne 100  $\mu$ F są spolaryzowane. Oprócz upewnienia się, że ich dłuższe wyprowadzenia trafiają do otworów oznaczonych symbolami +, należy je wygiąć, aby umożliwić domknięcie pokrywki, gdy płytka drukowana jest zamontowana w obudowie.

Następnie można zamontować stabilizator REG1 i gniazdo DC (CON1), gniazdo jack 3,5 mm (CON2) i potencjometr montażowy VR1 (jego suwak ustawiamy w środkowym położeniu). Teraz montujemy odbiornik UHF (RX1), upewniając się, że jest włożony we właściwym kierunku.

## Instalacja diod LED

Dioda LED1 musi być zamontowana na całej długości wyprowadzeń (25 mm) bez ich obcinania, aby można ją było później zgiąć, a jej soczewkę wsunąć przez otwór z boku pudełka (nad gniazdem 3,5 mm). Dioda LED2

jest montowana z górną częścią soczewki 20 mm nad powierzchnią płytki drukowanej. Trzeba upewnić się, że diody LED są prawidłowo zorientowane, a ich anodowe (dłuższe) wyprowadzenia podłączone są do punktów lutowniczych oznaczonych „A”.

Kolejny element, który lutujemy to 8-nóżkowa podstawka DIL dla IC1, ale na tym etapie nie wkładamy w nią mikrokontrolera. Ten krok nastąpi później, po przetestowaniu zasilania. Montaż płytki drukowanej kończymy lutując do niej przewód antenowy o długości 170 mm wykonany z kawałka przewodu izolowanego.

## Montaż końcowy

Płytkę PCB montujemy bezpośrednio pomiędzy uźebrowane ścianki obudowy UB5. Przed zamontowaniem płytki, należy wyciąć lub wywiercić otwory na gniazdo USB, gniazdo DC, gniazdo 3,5 mm oraz na dwie diody LED. Schematy wiercenia pokazano na **rysunku 6**.

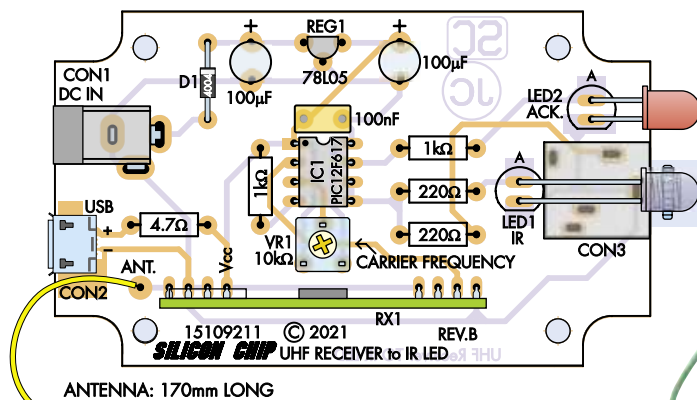
Najpierw można wywiercić otwór na gniazdo DC. Znajduje się on 6,5 mm w dół od górnej krawędzi podstawy na lewym końcu.

Otwór ten należy wywiercić małym wiertłem pilotującym, a następnie ostrożnie powiększyć do 6,5 mm za pomocą rozwiertaka stożkowego. Otwór gniazda 3,5 mm jest wyśrodkowany wzdłuż osi poziomej na drugim końcu obudowy, 10,5 mm poniżej krawędzi. Ponownie używamy wiertła pilotującego i rozwiercamy do 6,5 mm.

Otwór dla diody LED1 można następnie wywiercić 3,5 mm poniżej krawędzi, bezpośrednio nad otworem gniazda. Powinien być wywiercony na głębokość 3 mm, a następnie wiercimy podobny otwór dla diody LED2 około 12 mm w prawo. Prostokątne wycięcie USB można najpierw wywiercić, a następnie nadać mu odpowiedni kształt za pomocą pilnika igłowego.

Teraz wsuwamy płytkę drukowaną w szczeliny w bocznych żebrach pudełka (najpierw gniazdo jack 3,5 mm wkładamy w jego otwór). Po umieszczeniu płytki na miejscu, zginamy dwie diody LED i wpychamy je przez odpowiednie otwory. Zabezpieczamy całość, przykręcając nakrętkę do gniazda jack.

Etykiętę pokrywki można pobrać (w formacie PDF) ze strony [siliconchip.com.au](http://siliconchip.com.au) (zakładka „Shop”, a następnie „Panel artwork”)



Rysunek 5. Montaż tej płytki jest prosty, ponieważ elementy są znacznie większe niż na drugiej płytce. Należy zwrócić uwagę na orientację odbiornika UHF, IC1 i diody D1



◀ Przedłużacz można podłączyć do konwertera UHF-IR za pośrednictwem gniazda jack 3,5 mm (CON3). Na rysunku 7 zawarto szczegóły dotyczące konstrukcji tego kabla

▼ Etykietę panelu przedniego zdalnie sterowanego wzmacniacza zasięgu można pobrać w formacie PDF w skali 1:1 ze strony [siliconchip.com.au](http://siliconchip.com.au)



i wydrukować na odpowiedniej etykiecie (patrz informacje na temat tworzenia etykiet na stronie [siliconchip.com.au/Help/FrontPanels](http://siliconchip.com.au/Help/FrontPanels)) i przymocować do pokrywy. Cztery narożne otwory na śruby obudowy można wyciąć za pomocą skalpela.

## Tworzenie przedłużacza

W zależności od rozmieszczenia urządzeń, warto przygotować kabel z wtyczką jack 3,5 mm na jednym końcu i zewnętrzną diodą IR LED na drugim. Szczegóły przedstawiono na **rysunku 7**. Konieczne będzie użycie jednożyłowego przewodu ekranowanego o odpowiedniej długości, a wyprowadzenia diody powinny być odizolowane od siebie za pomocą rurki termokurczliwej. Rurka powinna mieć większą średnicę, aby zakryć koniec kabla, w tym obie nóżki diody i część soczewki, jak pokazano poniżej.

## Testowanie

Na czas testów należy zdemontować układ IC1. Następnie podłączamy zasilanie i sprawdzamy, czy napięcie między nóżkami 1 i 8 podstawki pod ten układ wynosi 5 V. Jeśli nie, sprawdzamy polaryzację zasilania i upewniamy się, że D1 i REG1 są prawidłowo zamontowane.

Jeśli zmierzone napięcie wynosi 5 V, wyłączamy zasilanie i umieszczamy w podstawce układ IC1 ze znacznikiem pierwszego wyprowadzenia ustawionym w kierunku sąsiedniego kondensatora 100 nF. Ponownie podłączamy zasilanie i sprawdzamy, czy czerwona dioda LED potwierdzenia miga po naciśnięciu dowolnego z przycisków pilota.

Następnie można przetestować urządzenie. Konwerter UHF na podczerwień musi mieć diodę podczerwona skierowaną w stronę urządzenia w zasięgu około 1 m. Jeśli nie działa, potrzebna jest regulacja VR1 podczas obsługi pilota, aż urządzenie zareaguje. Zwykle ustawienie

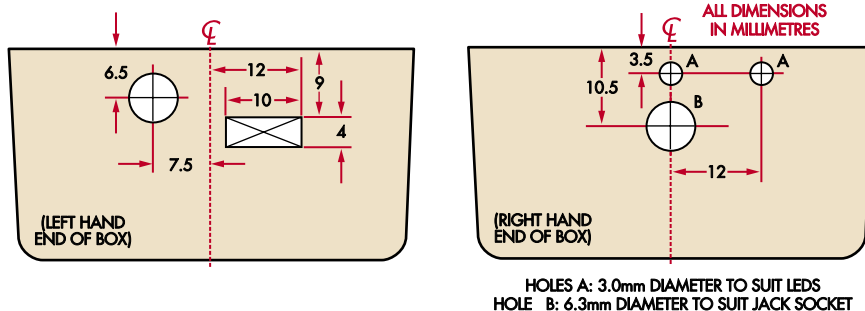
VR1 w połowie (odpowiadające częstotliwości nośnej około 37 kHz) będzie odpowiednie.

Gdy urządzenie działa prawidłowo, można spróbować użyć pilota do sterowania urządzeniem z innego pomieszczenia. Zasięg na wolnym powietrzu powinien wynosić 20...25 m, ale wewnątrz domu będzie

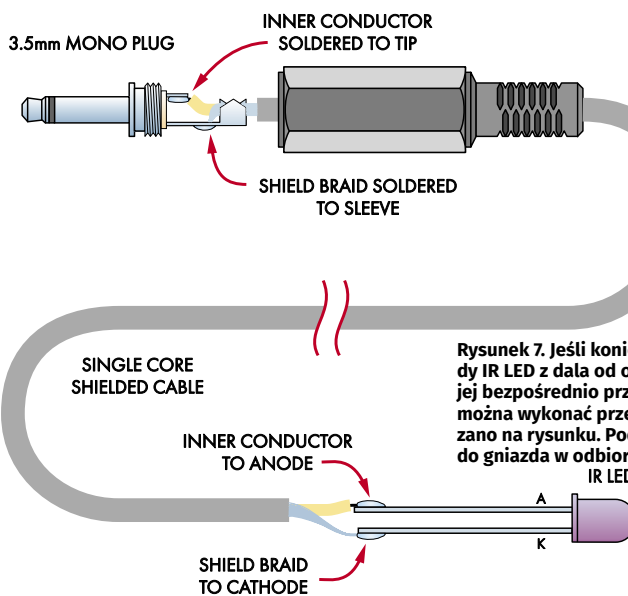
on mniejszy, w zależności od przeszkód (ścian itp.) między pilotem a konwerterem UHF na podczerwień. ■

John Clarke

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. [www.siliconchip.com.au](http://www.siliconchip.com.au)



**Rysunek 6.** Szablon do wiercenia otworów, które należy wywiercić lub wyciąć w obudowie UB5 jiffy. Otwór na gniazdo jack w prawym końcu pudełka można pominąć, jeśli przedłużacz podczerwieni nie jest używany, i podobnie, wystarczy wykonać tylko jeden otwór w lewym końcu, w zależności od tego, czy do zasilania będzie używane gniazdo USB czy gniazdo zasilacza



**Rysunek 7.** Jeśli konieczne jest zamontowanie diody IR LED z dala od odbiornika (np. zamontowanie jej bezpośrednio przed odbiornikiem urządzenia), można wykonać przedłużacz tak, jak to pokazano na rysunku. Podłącza się go bezpośrednio do gniazda w odbiorniku

Subscribe to Elektor's newsletter and get the chance to

# WIN

a Raspberry Pi Pico W board



[www.elektor.com/eda](http://www.elektor.com/eda)



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



Be one of the 10 fortunate winners!



**elektor**  
design > share > earn

Podczas użytkowania systemu audio z wieloma głośnikami może dochodzić do zdarzeń powodujących ich uszkodzenia. Warto wówczas stosować odpowiednie układy zabezpieczające. Opisywany w artykule „ochraniacz głośników”, w połączeniu ze wzmacniaczem Hummingbird (opublikowanym w zeszłym miesiącu), doskonale sprawdza się w systemach głośnikowych stereo z aktywną zwrotnicą lub do systemów dźwięku przestrzennego, w których używanych jest wiele głośników.



Materiały dodatkowe dostępne są na stronie Silicon Chip: <https://tiny.pl/d2dwv>  
Materiały dodatkowe są również dostępne na stronie [elportal.pl/do-pobrania](https://elportal.pl/do-pobrania)



## Wielokanałowy „ochraniacz głośników”

Czy drogie głośniki w używanym systemie audio są chronione na najgorszy przypadek, gdy awaria wzmacniacza spowoduje podanie na ich cewki napięcia stałego? Jeśli nie, prosty i skuteczny układ opisywany w artykule będzie z pewnością bardzo przydatny. Świetnie sprawdzi się również w połączeniu z modułami wzmacniaczy Hummingbird. Posiada układ opóźnionego załączenia i może chronić od jednego do sześciu kanałów. Mieści się na płytce drukowanej o wymiarach zaledwie 67×120 mm dla maksymalnie sześciu kanałów lub 67×91 mm dla wersji czterokanałowej.

Można śmiało powiedzieć, że w ciągu kilku lat budowania sprzętu hi-fi i nagłośnionego, autor nie zniszczył zbyt wielu głośników. Kiedy jednak dochodziło do takich zdarzeń, zawsze było to kosztowne, bolesne i kłopotliwe.

Doświadczenie oglądania wzmacniacza o mocy 60 W dostarczającego napięcie 40 V DC do cewki bardzo drogiego głośnika, którego zakup wymagał miesięcy oszczędności, pozostało w pamięci na długo. Był to głośnik o mocy 250 W, ale nie miał szans przeżyć po podaniu na niego napięcia stałego 40 V! Cewka głośnika spłonęła w ciągu sekundy, znacznie szybciej, niż człowiek byłby w stanie zareagować i odłączyć zasilanie.

W przypadku głośników istnieją dwa zasadnicze czynniki niszczące:

1. Nadmierne wychylenie membrany, szczególnie w obudowach wentylowanych poniżej rezonansu bez

odpowiedniego filtrowania poddźwięków. Jest to niezawodny sposób na zabicie głośnika basowego. Rozwiązaniem tego problemu jest aktywna zwrotnica zaprezentowana w wydaniach z października i listopada 2021 roku (w wydaniu 7/2024 i 8/2024 „Elektroniki dla Wszystkich”), która zawiera filtr poddźwiękowy.

2. Napięcie stałe pojawiające się na wyjściu wzmacniacza z powodu awarii wzmacniacza, albo błędów konstruktora (czy kiedykolwiek zostawiłeś wyjęty bezpiecznik lub zapomniawszy podłączyć jakiegoś kabla?).

Projekt opisany w artykule rozwiązuje kwestię numer 2.

Można zapytać: a co z przeciążeniem głośnika? Czy, jeśli ono wystąpi, głośnik też nie wybuchnie? Z doświadczenia autora wynika, że wymaga to heroicznego wysiłku, jeśli zwrotnica jest prawidłowo skonfigurowana,

pozostawiamy więc kwestię regulacji głośności użytkownikowi.

### Impuls

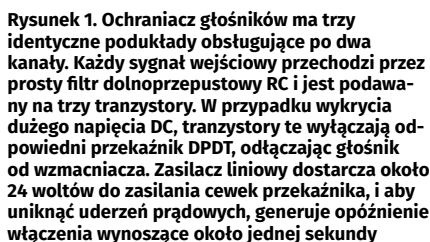
Budując aktywną zwrotnicę w połączeniu z sześcioma modułami wzmacniacza Hummingbird, okazało się, że brakuje w niej trochę miejsca. Próbując zmieścić wszystko wraz z zasilaczami w obudowie 2RU o głębokości 330 mm, i przenosząc projekt ze szkiców na papierze do oprogramowania CAD „komputer powiedział”, że konieczne jest zmniejszenie osłony głośników. Nie było rady, trzeba było to zalecenie zrealizować.

Opisywany układ ochroni głośnik przed większością awarii wzmacniacza. Skromna inwestycja zwróci się przy pierwszej aktywacji, z drugiej strony, jest to jeden z takich projektów, wobec których ma się nadzieję, że nigdy nie będą musiały zadziałać.

Istnieje wiele sposobów podejścia do ochrony głośników. Ten projekt ma na celu

#### Podstawowe parametry i funkcje:

- Chroni do sześciu głośników pojedynczych lub trzy pary głośników stereo
- Bardzo prosty, niewielki rozmiar i niski koszt budowy
- Możliwość zasilania z tego samego źródła co wzmacniacze (do ±40 V DC)
- W przypadku usterki polegającej na pojawieniu się na wyjściu głośnikowym napięcia stałego, odłącza głośnik/głośniki w ciągu 100 ms przy napięciu 30 V
- Zapewnia 1...2-sekundowe opóźnienie włączenia, pozwalając na ustabilizowanie się wyjść wzmacniacza
- Niewrażliwy na sygnały AC o niskiej częstotliwości
- Wykorzystuje przekładniki DPDT ze stykami o obciążalności 10 A przy napięciu 28 V DC



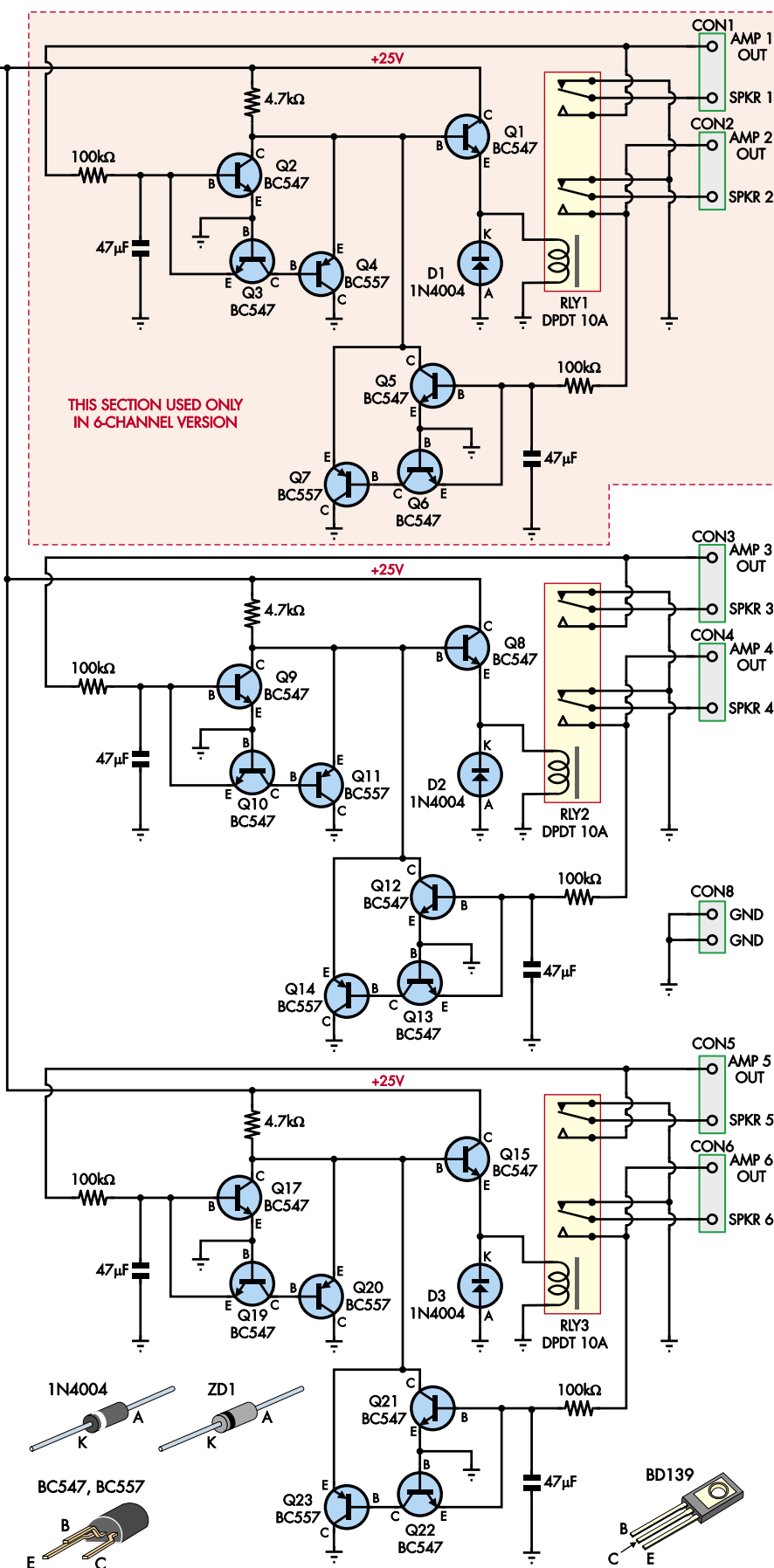
## Szczegóły układu

Główna część układu ochraniacza głośników jest elegancka, ale na pierwszy rzut oka może nie być oczywiste, jak działa.

Jego pierwszym zadaniem jest wykrycie napięcia stałego na wyjściu wzmacniacza, podłączonego do jednego z zacisków AMP×OUT po prawej stronie. Służą do tego: rezystor wejściowy 100 kΩ i kondensator bipolarny 47 μF, które tworzą filtr dolnoprzepustowy RC z punktem -3 dB (częstotliwość graniczna) przy częstotliwości 0,25 Hz. Wyjście tego filtra steruje detektor DC, który wyzwala się przy napięciu  $V_{BE}$  tranzystora, czyli około 0,6 V.

Do regularnej pracy wzmacniacz nie może więc wytwarzać składowej stałej większej niż 0,6 V na wyjściu tego filtra. Wybierając 10 Hz jako „bezpieczną” granicę niskiej częstotliwości i zakładając, że wzmacniacz może dostarczyć 100 W na 4  $\Omega$ , możemy obliczyć, że na wyjściu filtra pojawi się tylko 135 mV. Tak więc układ chroniący nie będzie wyzwalał podczas normalnego użytkowania wzmacniacza.

Ale powiedzmy, że wzmacniacz jest uszkodzony i dostarcza na wyjście napięcie linii 40 V DC (o dowolnej polaryzacji) zamiast przebiegu zmiennego. W takim przypadku po 100 ms (0,1 s) wyjście filtra dolnoprzepustowego osiągnie 0,84 V, co z pewnością wyzwoli następujący po nim obwód wykrywania



napięcia stałego. Filtr ten działa identycznie zarówno dla napięć dodatnich, jak i ujemnych.

Przy napięciu 40 V na obciążeniu 8  $\Omega$  przez 100 ms, do cewki głośnika zostanie dostarczone 20 J energii (impedancja spadnie z czasem, zbliżając się do wartości rezystancji, ale jest to wystarczająco dobre przybliżenie). W efekcie wystąpi solidny trzask, który prawdopodobnie sprawi, że słuchacz podskoczy w tym momencie, ale nie spowoduje, że cokolwiek się spali. Co więcej, jeśli usterka występuje w momencie włączania, głośnik po prostu nigdy nie zostanie dołączony.

Detektor DC składa się w sumie z trzech tranzystorów. W przypadku najwyższej sekcji na rysunku 1 są to tranzystory Q2, Q3 i Q4. Dodatkowo napięcie stałe wykrywane jest przez tranzystor Q2, którego kolektor jest dołączony za pośrednictwem rezystora 4,7 k $\Omega$  do napięcia +25 V. Dodatkowo napięcie z filtra, większe niż około 0,6 V, spowoduje włączenie tego tranzystora i w konsekwencji obniżenie napięcia na bazie Q1, wtórnika emiterowego, a tym samym wyłączenie przełącznika.

Tranzystory Q3 i Q4 wykrywają napięcia ujemne. Tranzystor NPN Q3 jest połączony w konfiguracji wspólnej bazy. Jego baza jest połączona z masą, a emiter jest wejściem. Ujemne napięcie wejściowe pobierze prąd z 0 V przez złącze baza-emiter, powodując spadek prądu na kolektorze. Ponieważ prąd pobierany przez kolektor trafia do emitera, musi być utrzymywany na niskim poziomie. Dlatego ten niewielki prąd jest buforowany przez Q4, tranzystor PNP podłączony jako wtórnik emiterowy. Emiter Q4 łączy się z tym samym rezystorem, co kolektor Q1, więc ujemne napięcie stałe z filtra podobnie obniża napięcie na bazie Q1, wyłączając przełącznik.

W układzie istnieje równowaga między ustawieniem niskiej częstotliwości odciecia a minimalnym napięciem DC, przy którym obwód wyłączy przełącznik. 47  $\mu$ F to rozsądne maksimum dla kondensatora filtrującego, więc wszelkie poprawki najlepiej wykonywać poprzez zmianę wartości rezystorów wejściowych.

Wybrana została rezystancja 100 k $\Omega$ , gwarantująca brak problemów z fałszywym wyzwaniem w przypadku zastosowań o bardzo dużej mocy i bardzo niskiej częstotliwości. Ale jeśli nie jest chroniony subwoofer, każda wartość większa niż 33 k $\Omega$  powinna być poprawna, a jako bonus, niższe wartości zapewnią szybsze wyłączenie w przypadku awarii.

Rezystory 100 k $\Omega$  wpływają również na wartość progową napięcia DC, która spowoduje wyzwolenie detektora. Usterka front-endu wzmacniacza może spowodować, że na wyjściu pojawi się kilka woltów napięcia stałego, a jeśli zostanie ono przyłożone do głośnika

przez wystarczająco długi czas, może dojść do jego przegrzania i uszkodzenia. Idealnie byłoby więc wykrywać również ten stan, a nie tylko usterkę, w której następuje natychmiastowe połączenie z jedną z linii zasilających.

Zakładając, że minimalny współczynnik hFE tranzystora wynosi 120 i że przełącznik wyłączy się przy napięciu 20 V na rezystorze 4,7 k $\Omega$  (pozostawiając zaledwie kilka woltów na cewce przełącznika), przełącznik zostanie wyłączony, jeśli prąd bazy tranzystora wzrośnie co najmniej do 20 V/4,7 k $\Omega$  = 35  $\mu$ A (około). Oznacza to, że aby wyłączyć przełącznik, napięcie stałe ze wzmacniacza musi wynosić co najmniej 3,5 V (35  $\mu$ A  $\times$  100 k $\Omega$ ).

Dotyczy to jednak najgorszej wartości hFE. Zalecamy stosowanie tranzystorów BC547B lub BC547C, które mają wyższe gwarantowane wartości hFE i wyłączą przełącznik przy około 1,5 VDC na wejściu. Obniżenie rezystorów wejściowych jeszcze bardziej zmniejszyłoby to napięcie wyzwajające.

Nasz ochraniacz głośników odłącza głośnik za każdym razem, gdy wykryty zostanie prąd stały. Zastosowane przełączniki są solidne i powinny być w stanie przerwać duży prąd, którego można spodziewać się w przypadku modułu wzmacniacza Hummingbird lub podobnego. Istnieje jednak możliwość, że po odłączeniu napięcia i prąd będą wystarczająco wysokie, aby utworzyć łuk między stykami przełącznika.

Normalnie zamknięty styk przełącznika służy do bocznikowania tego prądu do masy, gdy głośnik jest odłączony. Jeśli więc powstanie łuk elektryczny i prąd będzie nadal płynął, bezpiecznik DC wzmacniacza dla tej linii

zasilającej przepali się, a łuk elektryczny zgaśnie. Prawdopodobnie tranzystor wyjściowy będzie już przepalony, więc bezpiecznik nie będzie dla użytkownika wielkim zmartwieniem.

Na jednej płytce zostały umieszczone trzy zestawy układów, umożliwiając monitorowanie i ochronę sześciu standardowych kanałów wzmacniacza. Wybrany przełącznik ma standardowe wyprowadzenia i jest dostępny u wielu dostawców. Należy jednak upewnić się, że posiadany egzemplarz ma odpowiednią wersję. Zalecamy użycie przełączników z cewką na napięcie 24 V DC, chociaż 12 V też będą działać, pod warunkiem, że zostanie odpowiednio zmieniony układ stabilizacji napięcia stałego, a tranzystor BD139 poradzi sobie z rozproszeniem temperatury (patrz poniżej).

Jako odniesienie masy w omawianym układzie wybieramy pin masy zasilania. Łączy się on z uziemieniem chronionych wzmacniaczy mocy. Ponieważ wejścia są już sparowane, zabezpieczenie to dobrze sprawdzi się do ochrony przed napięciem stałym w zmostkowanym wzmacniaczu.

## Zasilanie

Zasilacz to prosty stabilizator liniowy generujący około 25 V DC. Przełączniki wymagają 24 V na swoich cewkach, co jest odpowiednie dla wzmacniaczy o różnych napięciach linii zasilającej. Przełącznik można dostosować do zasilania poniżej  $\pm$ 25 V lub powyżej  $\pm$ 40 V (patrz poniżej).

Zasilacz zapewnia opóźnienie włączenia wynoszące około jednej sekundy. Dzieje się tak, ponieważ rezystor 47 k $\Omega$  opóźnia ładowanie kondensatora 47  $\mu$ F w bazie Q18. Dotyczy

### Wykaz elementów:

- 1 dwustronna płytka PCB – kod 01101221, 67 mm  $\times$  121,5 mm
- 3 (2) przełączniki do montażu na płytce drukowanej/podstawce, styk 30 V DC 10 A styk DPDT, cewka 24 V DC cewka (RLY1-RLY3) [Altronics S4313, Jaycar SY4007].
- 8 (6) 2-pinowych mini listew zaciskowych o rozstawie 5,08 mm (CON1...CON8)
- 1 żebrowany radiator o wysokości 44 mm i wymiarach 16,5 mm  $\times$  10 mm do montażu na płytce drukowanej (HS1; dla Q16) [Altronics H0645].
- 1 silikonowa podkładka izolacyjna TO-126 lub TO-220 i tuleja izolacyjna [Altronics H7230, Jaycar HP1176].
- 1 śruba z tłem walcowym M3 3 mm  $\times$  10 mm
- 1 podkładka sprężysta M3
- 1 nakrętka sześciokątna M3
- 4 gwintowane podkładki dystansowe i 8 wkrętów (w zależności od instalacji)

### Półprzewodniki:

- 16 (11) BC547B/C 50 V 100 mA tranzystory NPN, TO-92 (Q1...Q3, Q5, Q6, Q8...Q10, Q12, Q13, Q15, Q17...Q19, Q21, Q22)
- 6 (4) tranzystory PNP BC557B/C 50 V 100 mA, TO-92 [BC558-9B/C również będą działać] (Q4, Q7, Q11, Q14, Q20, Q23)
- 1 BD139 80 V 1,5 A tranzystor NPN, TO-126 (Q16)
- 1 dioda Zenera 27 V 1 W (ZD1) [1N4750]
- 3 (2) Diody 1N4004 400 V 1A (D1...D3)

### Kondensatory:

- 6 (4) 47  $\mu$ F bipolarne/niepolarizowane elektrolityczne [Jaycar RY6820]
- 3 47  $\mu$ F 50 V elektrolityczne [Altronics R4807 lub Jaycar RE6344]

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

- 6 (4) 33 k $\Omega$ ...100 k $\Omega$  (patrz tekst, jeśli nie ma pewności, należy użyć 100 k $\Omega$ )
- 1 47 k $\Omega$
- 3 (2) 4,7 k $\Omega$

(n) dla wersji czterokanałowej (kod PCB 01101222, 67 mm  $\times$  91 mm), wymagane ilości są wymienione w nawiasach na czerwono.

to wszystkich kanałów chronionych przez płytkę. Wraz ze wzrostem napięcia zasilania opóźnienie włączenia nieznacznie maleje, ponieważ kondensator ładuje się szybciej. W razie potrzeby można to skompensować, zwiększając wartość rezystora.

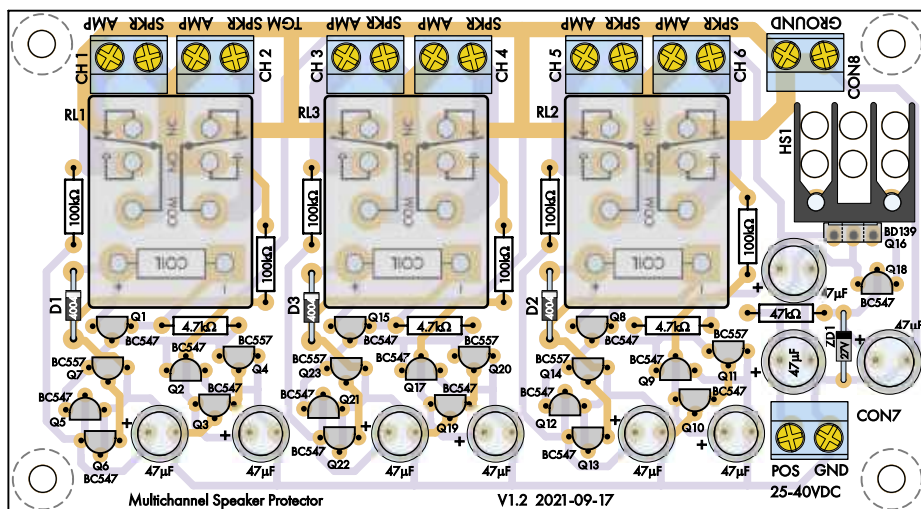
Jeśli używany wzmacniacz ma linie zasilające poniżej  $\pm 25$  V, można użyć przekaźniki w wariacie z cewką 12 V DC i dokonać niezbędnych zmian w układzie stabilizatora (w lokalizacji ZD4 należy wówczas zamontować diodę Zenera o napięciu przebicia równym 15 V). Podobnie, jeśli napięcie linii zasilających jest wyższe, układ powinien działać poprawnie. Trzeba tylko uważać na rozmiar radiatora. Podany radiator Altronics H0655 powinien być odpowiedni dla każdego normalnego napięcia zasilającego.

## Budowa

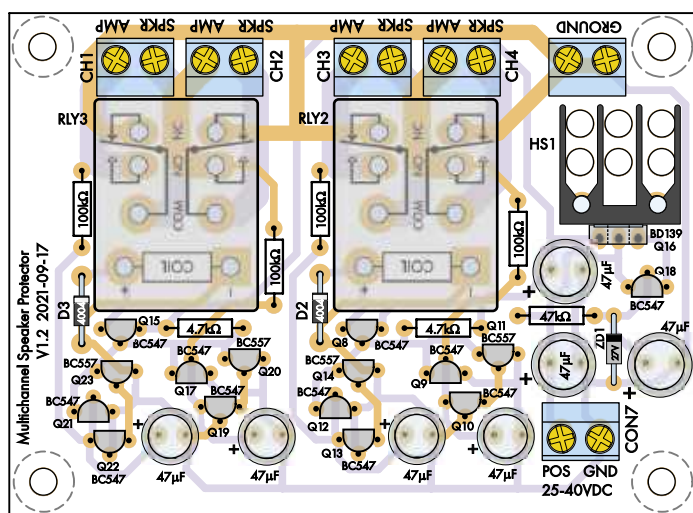
Konstrukcja jest prosta. Dostępne są dwie płytki drukowane: wersja sześciokanałowa (kod 01101221, 67×120 mm) i wersja czterokanałowa (kod 01101222, 67×91 mm).

W artykule została opisana wersja sześciokanałowa. Wersja czterokanałowa jest identyczna, z wyjątkiem pominięcia jednego przekaźnika i powiązanych z nim elementów, dzięki czemu płytkę drukowaną jest mniejsza. Podczas montażu należy zapoznać się z opisem na płytce – rysunek 2 lub rysunek 3.

Montaż zaczynamy od rezystorów i diod. Trzeba upewnić się, że diody są podłączone we właściwy sposób. Następnie montujemy dwupinowe złącza, jedno na zasilanie, kolejne na każdą parę wejść/wyjść i zacisk uziemienia (którego nie widać na wersji płytki widocznej na zdjęciu poniżej, ale jest obecne na rysunku



Rysunki 2 i 3. Mniejsza płytkę do ochrony maksymalnie czterech kanałów lub nieco większą płytkę dla pięciu lub sześciu kanałów. Montaż jest prosty, wszystkie elementy są typu przewlekane i powinny być montowane w kolejności od najniższego do najwyższego (na samym końcu przekaźniki i radiator dla Q16)



montażowym), aby zapobiec wyładowaniom łukowym w przekaźnikach (CON8).

Teraz nadszedł czas na wlutowanie tranzystorów BC5XX. Warto dopilnować, aby

były zamontowane na tej samej wysokości. Zadbamy w ten sposób o schludny wygląd płytki.

Następne są kondensatory. Trzy spolaryzowane kondensatory 47 µF muszą mieć napięcie znamionowe 50 V i wszystkie są montowane w ten sam sposób, z dłuższymi dodatnimi wyprowadzeniami do otworów na płytce oznaczonych jako +. Sześć bipolarnych/niepolaryzowanych kondensatorów elektrolitycznych 47 µF nie wymaga wysokiego napięcia znamionowego, ponieważ nigdy nie będzie na nich więcej niż 0,6 V – są montowane na płytce drukowanej w dowolnej orientacji.

Przyszła kolej na mocowanie tranzystora BD139 do radiatora za pomocą podkładki izolacyjnej, śruby M3 10 mm, podkładki zabezpieczającej i nakrętki. Radiator lutujemy do płytki drukowanej, ale trzeba pamiętać, że podczas lutowania elementów do dużych punktów lutowniczych konieczne jest dostarczenie sporej ilości ciepła.

Nakoniec montujemy przekaźniki. Płytkę PCB ma dla tych komponentów otwory o średnicy



Płytkę ochraniacza głośników będzie wyglądać mniej więcej tak. Należy zwrócić uwagę na otwory wywiercone w płytce pod radiatorem, aby umożliwić konwekcję powietrza od spodu. W prototypowej wersji brakuje bloku zacisków CON8 (GROUND). Można go pominąć, ale jeśli zostanie dotknięte do uziemienia wzmacniacza, zapewni lepszą ochronę głośników



**Wielokanałowy ochraniacz głośników jest dostępny w wersji sześciokanałowej (na zdjęciu) i mniejszej czterokanałowej. Wersja czterokanałowa będzie odpowiednia dla dwudrożnego systemu głośników stereo z aktywną zwrotnicą lub zmostkowanym wzmacniaczem stereo**

1,5 mm, co jest minimum zalecanym dla tej rodziny przekaźników – urządzenia Altronics i Jaycar pozostawiają sporo miejsca w otworach. Należy je dobrze przylutować.

## Testowanie

Po zamontowaniu wszystkich części nadszedł czas na przetestowanie urządzenia. Podczas wstępnych testów należy pozostawić odłączone zaciski wzmacniacza (CON1...CON6).

Najpierw dołączamy zasilanie i sprawdzamy, czy na wyjściu stabilizatora występuje napięcie 25 V. Wygodnym punktem pomiaru jest nóżka najbliższego rezystora 4,7 kΩ tuż przy tranzystorze Q11. Jako punkt masy można wybrać anodę dowolnej z diod D1...D3. Odczyt powinien wynosić od 24 V do 26 V dla napięcia wejściowego powyżej 32 V DC.

Jeśli takiego napięcia nie ma lub jest nieprawidłowe, trzeba sprawdzić, czy na ZD4 jest około 27 V. Jeśli nie, sprawdzamy rezystor 47 kΩ i tranzystory Q16 i Q18. Sprawdzamy, czy te dwa tranzystory są odpowiedniego typu i czy są prawidłowo wlutowane. Trzeba ponadto sprawdzić, czy nie ma zwarcia – czy BD139 się nagrzewa?

Po sekundzie lub dwóch przekaźniki powinny kliknąć. Jeśli tak się nie stanie, należy sprawdzić napięcia na bazach tranzystorów Q1, Q8 i Q15 (tranzystory sterujące przekaźnikami). Czy napięcia te utrzymują się w granicach około jednego lub dwóch woltów poniżej tego, które występuje na linii 25 V? Jeśli nie,

trzeba sprawdzić, czy są to właściwe tranzystory, i czy są prawidłowo wlutowane.

Następnie sprawdzamy napięcie na wyjściu filtrów RC. Napięcie na obu końcach rezystorów 100 kΩ (lub innej wartości, na które można było je zmienić) powinno być bliskie 0 V. Jeśli tak nie jest, sprawdzamy, czy elementy BC547 i BC557 znajdują się we właściwych miejscach i czy są zamontowane w odpowiednim kierunku.

Zakładając, że przekaźniki się włączą, sprawdzimy, czy z kolei wyłączą się prawidłowo. Najłatwiejszym sposobem na wykonanie tego testu jest wykorzystanie baterii 9 V. Ujemny biegun baterii podłączamy do zacisku uziemienia na płycie ochraniacza głośników. Dodatni biegun baterii łączymy do zacisku „AMP” każdego kanału zabezpieczenia DC. Odpowiedni przekaźnik powinien wyłączyć się niemal natychmiast.

Powtarzamy tę czynność dla wszystkich kanałów, sprawdzając, czy przekaźniki przełączają się szybko. Następnie powtarzamy test z baterią podłączoną na odwrót (tj. biegun dodatni do GND i ujemny do zacisków AMP).

Jeśli kanał nie przełącza się zgodnie z oczekiwaniami, należy zmierzyć napięcie na obu końcach rezystora 100 kΩ. Jedno powinno wynosić  $\pm 9$  V, a drugie  $\pm 0,6$  V.

Jeśli koniec  $\pm 9$  V nie jest prawidłowy, gdzieś w tym obszarze występuje zwarcie lub przerwa w obwodzie. Jeśli drugi koniec nie jest bliski  $+0,6$  V, sprawdzamy dwa tranzystory NPN w detektorze DC (np. Q9 i Q10), zwłaszcza ich kierunek montażu. Sprawdzamy powiązany tranzystor PNP (np. Q11) pod kątem usterki, w której nie występuje napięcie  $-0,6$  V.

Wprawdzie nie jest zalecane sprawdzenie, czy układ zabezpieczenia rzeczywiście chroni głośnik, ale podczas uruchamiania prototypu do wejścia „AMP” ochraniacza głośników został dołączony subwoofer 4 W wraz z zasilaczem  $-34$  V z ograniczeniem 6 A. Pojawił się solidny trzask i kliknięcie, gdy przekaźnik uratował subwoofer przed nieuchronnym zniszczeniem. Cały test był monitorowany



**Oscylogram 1. Niebieski przebieg to napięcie na głośniku 4 Ω (zero woltów u góry), żółty przebieg natomiast to napięcie  $-34$  V przyłożone do zabezpieczenia głośnika ze źródła zasilania. Głośnik zostaje odłączony w czasie krótszym niż 80 ms. Wtórne napięcie, którego można by się było spodziewać na skutek powrotu membrany do pozycji początkowej gdy przekaźnik rozłączy głośnik nie wystąpi, jeśli złącze CON8 zostanie podłączone do uziemienia. Dodatkowo membrana głośnika zostanie wyhamowana w pozycji początkowej**



Należy pamiętać o założeniu podkładki izolacyjnej na tranzystor BD139 (widoczny powyżej) zapobiegającej jego zwarcia do radiatora. Należy również zwrócić uwagę na zastosowanie podkładek sprężystych na wszystkich śrubach, aby nie poluzowały się z powodu wibracji lub ruchu.

za pomocą oscyloskopu, a wynik pokazano na **oscylogramie 1**.

W czasie tego testu wystąpił błysk wyładowania łukowego, gdy głośnik został odłączony, co nie jest jednak zaskoczeniem w przypadku przerwania przepływu bardzo wysokiego prądu stałego. Nie należy tego próbować w domu, by bez powodu nie narażać głośnika na zbędne uduy.

## Zastosowanie

Ochroniacz głośników musi być podłączony do uziemienia wzmacniacza mocy poprzez dostępny zacisk (CON8) i zasilany napięciem 30...40 V DC za pomocą złącza zasilania (CON7). Jeśli wzmacniacz ma wyższe napięcie dodatnie, do obniżenia napięcia zasilania ochroniacza głośników można użyć rezystora

drutowego o mocy 5 W. Sześciokanałowy ochroniacz pobiera około 100 mA, więc rezystor 100  $\Omega$  5 W obniży napięcie o 10 V i rozproszy około jednego wata mocy.

Podłączamy zaciski oznaczone „AMP” do wzmacniacza, a odpowiadające im zaciski „SPKR” do zacisków głośników. Nieużywane kanały, na przykład, jeśli wzmacniacz jest 3- lub 5-kanałowy można pozostawić niepodłączone. Po zainstalowaniu układu wewnątrz wzmacniacza można o nim zapomnieć i mieć nadzieję, że nigdy nie usłyszymy niespodziewanego kliknięcia przekazników. Jeśli jednak tak się stanie, na pewno będziemy z tego zadowoleni! ■

**Phil Prosser**

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. [www.siliconchip.com.au](http://www.siliconchip.com.au)



Oscylogram 2. Czas reakcji zabezpieczenia na napięcie 20 V DC. Krok napięcia wejściowego wynosi  $t=0$ , a napięcie wyjściowe zaczyna spadać przed  $t=60$  ms. Osiąga 0 V przed  $t=80$  ms. Opóźnienie 80 ms wynika ze stałej czasowej RC filtra do osiągnięcia napięcia 0,6 V



Ze smutkiem i żalem przyjęliśmy wiadomość o śmierci Pawła Sujko, naszego redakcyjnego Kolegi. Gdy dwa lata temu dokonywała się transformacja „Elektroniki dla Wszystkich”, nie wszyscy się tym zachwycali. Było sporo nedoróbek. Wśród krytycznych listów od Czytelników wyróżniał się jeden – od Pawła Sujko. Nie tylko „wypunktował” błędy, ale też zauważył duży potencjał rozwojowy

EdW po zmianach. Chętnie przystał na propozycję współpracy redakcyjnej. Fantastyczny wzrost czytelnictwa EdW przez te dwa lata to także jego zasługa. Mówił o sobie, że jest „saperem dyżurnym” – znajdował każdą „minę” podłożoną przez autora, tłumacza, czy chochlika. Jestem starym belfrem i nie widziałem w życiu kogoś, kto tak jak On potrafił znaleźć każdy słaby punkt w projekcie i zaproponować lepsze rozwiązanie.

Robił to z lekkością eksperta i erudyty o bezgranicznej wiedzy w elektronice i szerokich horyzontach, wykraczających daleko poza elektronikę. W Jego krytycznych komentarzach dostrzeżliśmy szansę wzbogacania artykułów źródłowych o wartość dodaną. Tak zrodziła się idea rubryki „Pokój Nauczycielski”.

Ze smutkiem patrzymy na puste krzesło w „Pokoju Nauczycielskim”.

Będzie nam Go brakowało.

Wiesław Marciniak



# Stroboskop RGB z Arduino

## Kolorowa adaptacja użytecznego instrumentu

**W dawnych czasach, przed pojawieniem się nowoczesnej elektroniki pod maską samochodu, do regulacji kąta wyprzedzenia zapłonu w silnikach używano stroboskopów. Obecnie rzadko jest to konieczne – w nowoczesnych pojazdach czas zapłonu jest regulowany przez jednostkę sterującą silnika (ECU) wspomaganą przez niezliczone czujniki.**

Kiedyś wszystko było inne – niekoniecznie lepsze, ale inne. Jeszcze nie tak dawno można było po prostu otworzyć maskę pojazdu i samodzielnie dokonać drobnych napraw i regulacji. Dzięki nowoczesnym metodom projektowania komputerowego, możliwe jest teraz zamontowanie wszystkich elementów pod maską tak kompaktowo i blisko siebie, że nawet wymiana prostej lampy jest dużym wyzwaniem. Znalazły się tam również komputery zbudowane z wielu połączonych ze sobą mikrokontrolerów.

Przenieśmy się jednak pamięcią do lat 70. lub 80. ubiegłego wieku. Dla prawidłowego działania silnika spalinowego (benzynowego), świece zapłonowe wszystkich cylindrów muszą generować iskry dokładnie we właściwym czasie. Zajmował się tym rozdzielacz (mechaniczne urządzenie, często bardzo wrażliwe na wilgoć), który musiał być precyzyjnie zsynchronizowany z silnikiem. W tym miejscu do gry wkraczał stroboskop – lub światło rozrządu, jak go wówczas nazywano.

Staromodny rozdzielacz miał taką samą liczbę styków, jak liczba cylindrów w silniku, a za każdym razem, gdy styk się otwierał, energia elektryczna zmagazynowana w cewce zapłonowej wytwarzała iskłę na końcówce świecy zapłonowej. To z kolei powodowało zapłon mieszanki paliwowo-powietrznej w cylindrze. Dzięki takiemu rozwiązaniu silnik pracował prawidłowo.

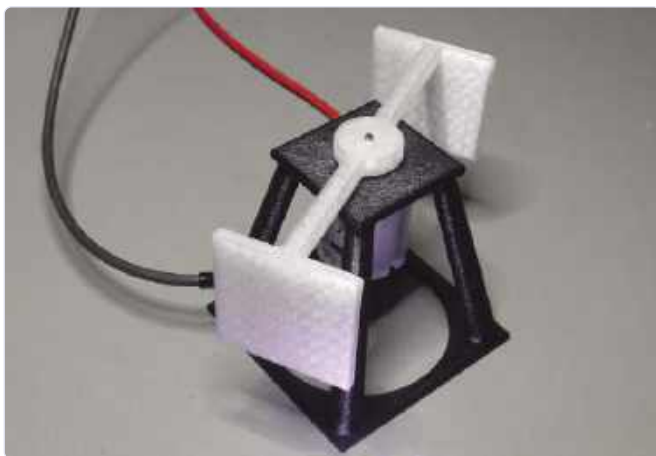
Stroboskop używany do regulacji czasu miał ksenonową lub neonową lampę błyskową i kabel wyzwalający. Kabel wyzwalający (z czujnikiem

indukcyjnym) był podłączony do kabla referencyjnej świecy zapłonowej, dzięki czemu za każdym razem, gdy referencyjna świeca zapłonowa odpalała, lampka błyskowa była wyzwalana i wytwarzała bardzo krótki błysk światła. Na kole pasowym wału korbowego znajdował się biały znak (wąski pasek), a nad kołem pasowym, na czymś w rodzaju kątownika, znajdował się odpowiadający mu stały znak. Te dwa znaki miały być dokładnie wyrównane względem siebie, gdy referencyjna świeca zapłonowa została odpalona, a efekt stroboskopowy światła rozrządu sprawiał, że wyglądało to tak, jakby ten ruchomy w rzeczywistości znak stał nieruchomo. Jeśli ruchomy znak na kole pasowym wału korbowego nie był wyrównany ze stałym znakiem, ale był przesunięty w lewo lub w prawo, oznaczało to, że układ zapłonu wymaga regulacji.

### Wyzwalanie? Kto go potrzebuje?

W opisanym powyżej zastosowaniu stroboskop jest synchronizowany sygnałem wyzwalającym z mierzonym ruchem obrotowym. Jednak sygnał wyzwalający nie zawsze jest konieczny.

Jeśli zmierzona ma być prędkość obrotowa silnika (RPM – *revolutions per minute*) za pomocą niesynchronizowanego stroboskopu, należy dostosować częstotliwość błysków światła tak, aby znak na kole pasowym wydawał się praktycznie nieruchomy i był widoczny tylko w jednym miejscu. Jeśli obserwuje się kilka prawie nieruchomych znaków w tym samym czasie, oznacza to, że częstotliwość błysków jest



Fotografia 1. Rama ze śmigłem służącym jako reflektor

wielokrotnością RPM. Taki efekt będzie obserwowany również wtedy, gdy RPM będzie wielokrotnością częstotliwości błysków.

Jeśli koło pasowe obraca się zgodnie z ruchem wskazówek zegara, a znacznik dryfuje w kierunku zgodnym z ruchem wskazówek zegara, częstotliwość błysków jest nieco zbyt niska. Każdy błysk następuje zbyt późno, przez co znacznik wydaje się poruszać zgodnie z kierunkiem obrotu. Jeśli koło pasowe obraca się zgodnie z ruchem wskazówek zegara, a znacznik dryfuje w kierunku przeciwnym do ruchu wskazówek zegara, częstotliwość błysków jest zbyt wysoka, więc każdy błysk musi być zbyt wcześnie.

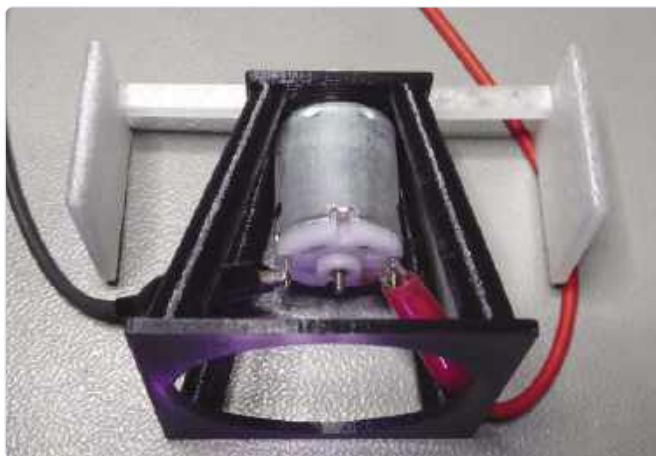
Dostosowując częstotliwość błysków stroboskopu tak, aby był wyraźnie widoczny tylko jeden nieruchomy znak, można zmierzyć prędkość obrotową. Każdy błysk musi być wystarczająco krótki, aby wytworzyć wyraźne odbicie (innymi słowy, wyraźnie widoczny znak). Jeśli czas trwania błysku jest zbyt długi w porównaniu do szybkości błysku, znak będzie wyglądał na „rozmaźany”.

## Zabawmy się trochę

Co się stanie, jeśli wygenerujemy kilka błysków o tej samej częstotliwości, ale w różnych kolorach i z przesunięciem fazowym między poszczególnymi błyskami? A co zobaczymy, jeśli użyjemy również obracającego się białego obiektu i zmienimy częstotliwość, przesunięcie fazowe i czas trwania różnych kolorowych błysków?

Nie jest to trudne do osiągnięcia za pomocą mikrokontrolera, więc nic nie stoi na przeszkodzie, aby spróbować. Możemy użyć diody LED RGB do różnych kolorowych błysków lub, jeszcze lepiej, trzech oddzielnych diod LED, aby uzyskać jaśniejsze błyski. Aby kontrolować częstotliwość błysków, przesunięcie fazowe i czas trwania błysku, możemy użyć Arduino Pro Mini, ponieważ ma ono wystarczającą liczbę wejść/wyjść do pracy i jest więcej niż wystarczająco szybki do takiego zastosowania. Diody LED, jak można się spodziewać, mogą być sterowane sygnałami o modulowanej szerokości impulsu (PWM).

Moglibyśmy użyć trzech modułów PWM Arduino Pro Mini, ale działałyby one w oparciu o trzy różne timery. Utrudnia to ich synchronizację i programowanie przesunięć fazowych. Co więcej, potrzebujemy sygnałów PWM o stosunkowo niskiej częstotliwości, więc rozdzielczość 16 bitów to zdecydowanie za dużo. Sygnały PWM generowane programowo (softPWM) są całkowicie satysfakcjonujące dla naszych celów. W przypadku softPWM sygnały są wyprowadzane na normalne cyfrowe piny I/O. Są one ustawiane i zerowane za pomocą licznika, który generuje przerwanie za każdym razem, gdy osiągnie określony stan. Pozwala to skonfigurować interwał, który jest wystarczająco krótki, aby dostosować szerokość impulsu lub przesunięcie fazowe sygnałów PWM z wystarczającą precyzją.



Fotografia 2. Widok na silnik w zbliżeniu

## Zastosowanie w praktyce

Aby zrealizować opisaną powyżej koncepcję, autor zastosował silnik prądu stałego 12 V (numer części firmy Velleman – MOT3N) i zaprojektował dla niego niewielką ramę. Silnik obraca swego rodzaju śmigło z dwiema pionowymi łopatkami, które działają jak reflektory. Obie części zostały wykonane metodą druku 3D, a odpowiednie pliki znajdują się w zestawie plików do pobrania dla tego projektu. Wygląd śmigiełek przedstawiono na **fotografiach 1 i 2**.

Należy zauważyć, że prędkość bez obciążenia silnika używanego przez autora wynosi około 11500 obrotów na minutę. Jest to prędkość zdecydowanie zbyt duża dla śmigła wydrukowanego na drukarce 3D. Za pierwszym razem, gdy autor testował tę konfigurację, części dosłownie latały mu koło nosa. Jest to oczywiście dość ryzykowne, więc dobrym pomysłem będzie założenie podczas prób okularów ochronnych lub gogli. Należy to potraktować jako ostrzeżenie!

Z tego powodu autor zasiliał silnik napięciem stałym 3 V, a więc niższym od znamionowego. Napięcie to można regulować, w celu zsynchronizowania prędkości silnika z częstotliwością błysków. Przy napięciu zasilania 3 V prędkość obrotowa silnika wynosi około 2100 obr/min, co odpowiada częstotliwości 35 Hz. Śmigło ma symetryczną strukturę z dwiema odbijającymi powierzchniami łopatek, a taka budowa ułatwia wyważenie tego elementu. Z tego powodu częstotliwość błysków musi być dwa razy większa (70 Hz), co odpowiada okresowi około 14 ms.

Najpierw jednak przyjrzyjmy się układowi. Schemat pokazany na **rysunku 3** jest wzorem prostoty. Trzy diody LED RGB są włączane i wyłączane przez niewyszukane tranzystory sterowane za pomocą Arduino Pro Mini. Napięcie zasilania wynosi 5 V DC, a pobór prądu około 500 mA.

Do tego projektu nie zaprojektowano żadnej płytki drukowanej, a wszystko można zbudować w mgnieniu oka na płytce prototypowej. Model zbudowany przez autora został pokazany na **fotografii 4**.

Czas zadziałania stroboskopu jest kontrolowany za pomocą przerwanienia timera, które jest wyzwalane co 100 μs. Interwał 0,1 ms jest zatem rozdzielczością do regulacji szybkości błysku i okazało się zapewniać odpowiednią zdolność sterowania do przesuwania „nieruchomego” obrazu śmigła.

```
...
// Konfiguracja 16-bitowego timera1 do normalnej pracy z
// przerwaniami 100 us
// 16MHz/1 = 6,25 ns. Więc aby uzyskać 100us mu-
// simy pozwolić
// zliczyć 1600 tyknięć.
TCCR1A = 0;
```

```
TCCR1B = _BV(CS10); //preskaler dziel-
ony przez 1
TCNT1 = 0xFFFF - 1600; // przepełnienie
wynosi 65535 = 0xFFFF
TIMSK1 = _BV(TOIE1); // przerwanie
od przepełnienia
TCNT1 = 0;
sei();
}
ISR (TIMER1_OVF_vect)
{
TCNT1 = 0xFFFF - 1600; //100 us
przerwania
...

```

Okres błysku jest zakodowany na stałe w oprogramowaniu i ustawiony na 144, co jest liczbą przerwań timera między dwoma kolejnymi błyskami. Oznacza to, że okres wynosi  $144 \times 100 \mu s = 14,4 \text{ ms}$ . Czas trwania błysku dla każdej diody LED jest ustawiony na stałą wartość 8, co odpowiada  $8 \times 100 \mu s = 800 \mu s$ . Wypełnienie przebiegu (stosunek znak/przerwa) wynosi zatem  $800 \mu s / 14,4 \text{ ms} = 5,5\%$ .

```
#define STROBE_PERIOD 144
```

```
TSoftPwm Pwm[] = {TSoftPwm(ID_RED, 0,
8, STROBE_PERIOD),
TSoftPwm(ID_GREEN, 0, 8, STROBE_PERIOD),
TSoftPwm(ID_BLUE, 0, 8, STROBE_PERIOD)};
```

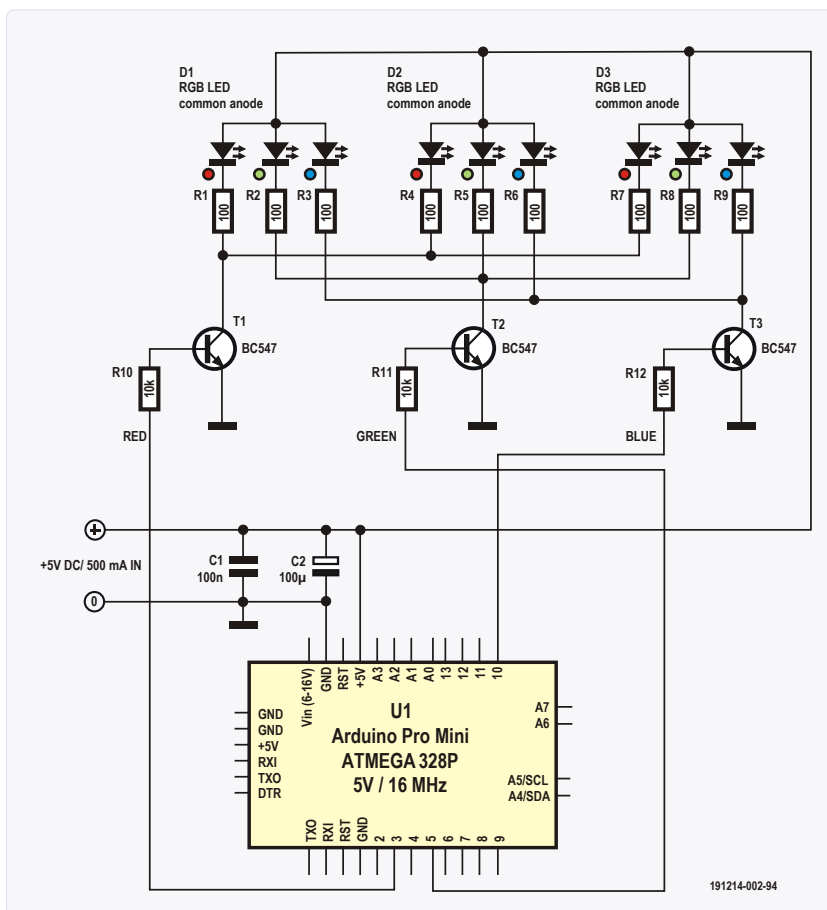
Cykl pracy po odpowiednim dostosowaniu prędkości silnika zapewnia ładny, ostry obraz.

Zarówno program jak i schemat (które można pobrać ze strony projektu [4]) są zaledwie jedną z możliwych koncepcji, dlatego brak tutaj ładnego interfejsu użytkownika i bardziej wymyślnych funkcji. Dostępne oprogramowanie demonstruje, co można zrobić, bawiąc się parametrami szybkości błysku, czasu trwania błysku i przesunięcia fazowego. Oprogramowanie jest obszernie i jasno skomentowane, więc nie ma potrzeby dublowania tych opisów w artykule.

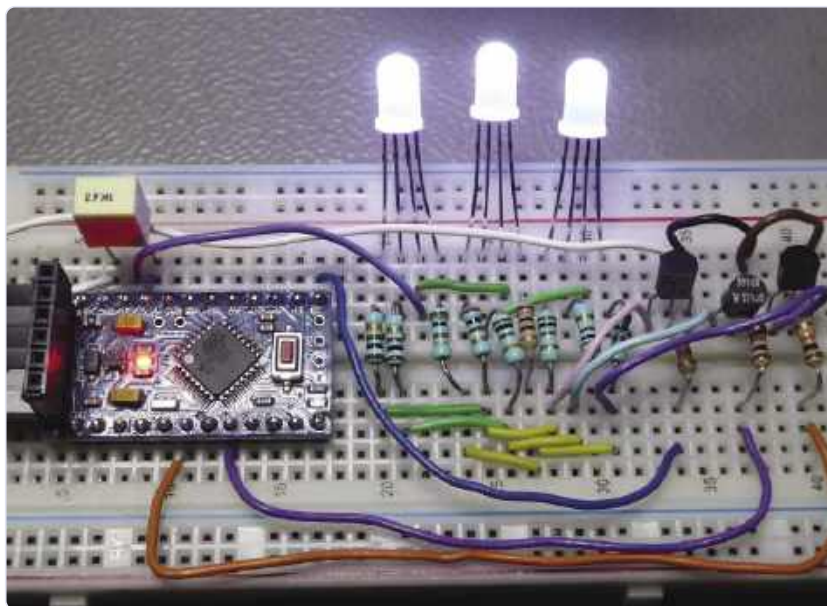
Oczywiście niesynchronizowany stroboskop pozostawia wiele do życzenia. Zasadniczo nie powinno być trudne zamontowanie małego czujnika przerwania wiązki światła obok śmigła i wykorzystanie sygnału z czujnika do wyzwalania stroboskopu. Pozostawimy to jednak jako zadanie domowe dla naszych zainteresowanych czytelników.

Autor zamieścił na YouTube dwa filmy, na których można podziwiać działanie stroboskopu [1] oraz interakcję pomiędzy sygnałami sterującymi diodami LED na ekranie oscyloskopu [2]. Projekt jest również dostępny na stronie Elektor Labs [3]. ■

**Roel Arits** (Holandia)



Rysunek 3. Schemat układu



Fotografia 4. Konstrukcja nie jest krytyczna, a płytka prototypowa w zupełności wystarczy do początkowych eksperymentów

Powiązane adresy internetowe

[1] Wideo demo 1. <https://youtu.be/WHv6DCWM82k>

[2] Wideo demo 2. <https://youtu.be/3tYqUin4-vQ>

[3] Strona projektu na Elektor Labs: <https://tiny.pl/dp333>

[4] Strona projektu dla tego artykułu: <https://tiny.pl/dlhhh>

# Urządzenie do wyznaczania temperatury cewki głośnika

W artykule przedstawiono opis techniczny prostego urządzenia służącego do wyznaczania temperatury cewki głośnika. Urządzenie umożliwia wykonanie pomiarów przeprowadzanych metodą pośrednią wykorzystującą liniowe zmiany rezystancji miedzianego uzwojenia cewki głośnika w funkcji temperatury. Do budowy urządzenia posłużono się popularnym modułem Arduino Nano.

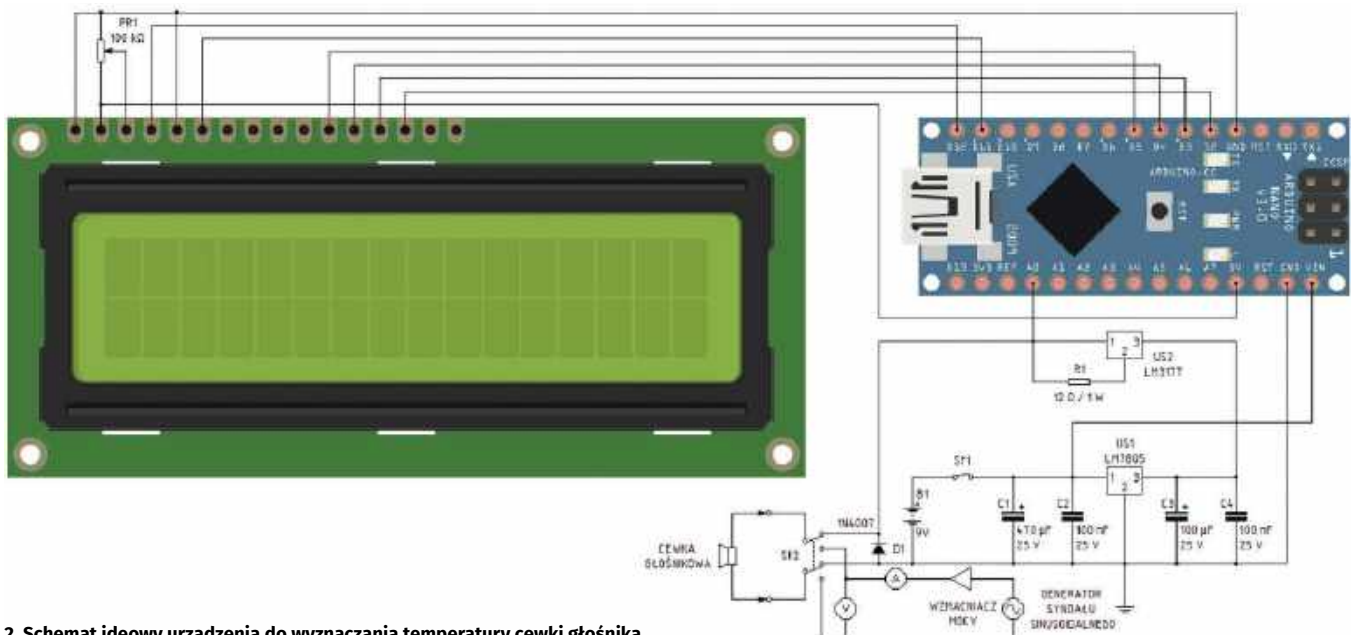
Większość urządzeń służących do wyznaczania temperatury cewki głośnika ma skomplikowaną budowę i bazuje na nadążnym pomiarze składowej rezystancyjnej uzwojenia cewki głośnika podczas pobudzania go do drgań sygnałem sinusoidalnym z generatora wzmacnionym przez wzmacniacz mocy. Zasada działania tego typu urządzeń opiera się na jednoczesnym podaniu na zaciski głośnika sygnału sinusoidalnego oraz składowej stałej o niewielkiej wartości celem umożliwienia przeprowadzenia pomiaru rezystancji metodą techniczną, tzn. na podstawie pomiaru spadku napięcia przy stałym wymuszeniu prądowym. Tego typu rozwiązanie posiada jednak dwie znaczące wady. Pierwszą z nich stanowi konieczność zastosowania filtra dolnoprzepustowego wysokiego rzędu celem odseparowania napięcia stałego od sygnału sinusoidalnego. Drugą wadę stanowią nienaturalne warunki pracy głośnika przy współpracy z zaprojektowanym w ten sposób urządzeniem. Membrana pod wpływem przepływu prądu



1. Wygląd zewnętrzny urządzenia do wyznaczania temperatury cewki głośnika

stałego przez uzwojenie cewki głośnika jest wychylona z położenia równowagi o pewną stałą odległość, względem której oscyluje w wyniku pobudzenia jej do drgań sygnałem sinusoidalnym. Opis techniczny takiego urządzenia

można znaleźć w książce pt. *Current-driving of loudspeakers. Eliminating major distortion and interference effects by the physically correct operation method*, której autorem jest fiński uczonec – mgr inż. Esa Meriläinen.



2. Schemat ideowy urządzenia do wyznaczania temperatury cewki głośnika



3. Kod źródłowy programu komputerowego urządzenia do wyznaczania temperatury cewki głośnika



4. Urządzenie do wyznaczania temperatury cewki głośnika mierzące rezystancję uzwojenia cewki głośnika wysokotonowego typu GDWK 11/100 w temperaturze pokojowej

## Opis układu

Zasada działania urządzenia opisanego w treści artykułu różni się w sposób zasadniczy od opisywanych dotychczas układów elektronicznych. Podstawową różnicę stanowi rozdzielenie sygnału sinusoidalnego od sygnału pomiarowego, dzięki czemu głośnik może pracować w naturalnych dla siebie warunkach i drgania jego membrany odbywają się względem położenia swobodnego. Urządzenie bazuje na projekcie miliomomierza wykonanego w oparciu o platformę Arduino Nano. Oryginalny projekt można znaleźć na stronie internetowej: <https://tiny.pl/dlfv5>.

Projekt ten został jednak zmodyfikowany przez wzgląd na konieczność pracy z cewkami głośnikowymi charakteryzującymi się obecnością składowej reakcyjnej (induktancji) pochodzącej od indukcyjności ich uzwojeń. Dodatkowo zmieniono sposób zasilania na bardziej energooszczędny i zrezygnowano z podświetlenia wyświetlacza LCD 2×16 w celu umożliwienia zasilania bateryjnego. Celem dalszego zmniejszenia poboru prądu przez urządzenie, z modułu Arduino Nano wylutowano nawet oporniki doprowadzające napięcia zasilające do czterech diod LED. Kolejną modyfikację stanowi odseparowanie obwodu pomiarowego od układu stabilizacji znajdującego się na płytce modułu Arduino Nano, zrealizowane za pośrednictwem dodatkowego stabilizatora napięcia. Układ zasilany jest przy pomocy płaskiej baterii dziewięciowoltowej zamontowanej w koszyczku znajdującym się wewnątrz

obudowy. Zwarcie włącznika St1 powoduje uruchomienie urządzenia. Kondensatory elektrolityczne unipolarne C1 oraz C3 filtrują napięcie zasilające, natomiast kondensatory ceramiczne C2 i C4 pełnią rolę przeciwzakłóceń. Scalony stabilizator napięcia US1 obniża napięcie zasilania do poziomu +5 V względem masy. Napięcie to służy do zasilania źródła prądowego zbudowanego w oparciu o układ scalony US2. Opornik R1 ustala wartość prądu na 104 mA. Napięcie referencyjne układu US2 wynosi 1,25 V. Korzystając z prawa Ohma, dzieląc to napięcie przez wartość rezystancji opornika R1 (12 Ω) uzyskamy taką właśnie wartość prądu. Dokładność pomiaru rezystancji zależna jest od tolerancji wykonania opornika R1. Pomiar rezystancji odbywa się metodą techniczną. Spadek napięcia na obciążeniu mierzony jest przez wejście analogowe A0 modułu Arduino Nano i jest przeliczany programowo na wartość rezystancji. Pomiar odbywa się z dokładnością do dwóch miejsc po przecinku natomiast najwyższą możliwą do zmierzenia wartością jest rezystancja równa 50 Ω. Dioda D1 służy zabezpieczeniu układu źródła prądowego przed przepięciem, jakie może wystąpić podczas rozłączania cewek o dużych wartościach indukcyjności. Kontrast wyświetlacza LCD 2×16 można regulować potencjometrem montażowym PR1. Przełącznik St2 służy do wyboru rodzaju pracy. Pierwszym krokiem jest zmierzenie rezystancji uzwojenia cewki głośnika w temperaturze pokojowej (około

20°C). Następnie po przełączeniu przełącznika St2 głośnik jest pobudzany do drgań sygnałem sinusoidalnym o wybranej częstotliwości i mocy określonej przez wzmacniacz. W układzie podłączony jest woltomierz napięcia przemiennego i amperomierz prądu przemiennego służące do wyznaczania wartości mocy zgodnie z prawem Joule'a-Lenza (moc RMS stanowi iloczyn napięcia skutecznego i prądu skutecznego). Po wygrzaniu głośnika przez określony czas (jego wartość można odczytać z norm) możemy szybko przełączyć przełącznik St2 na pozycję służącą do pomiaru rezystancji i sprawdzić o jaką wartość wzrosła rezystancja uzwojenia cewki głośnika. Współczynnik temperaturowy dla miedzi wynosi +0,0039/1°C. Oznacza to, że przykładowo dla głośnika wysokotonowego GDWK 11/100 wzrost temperatury o 40°C powyżej temperatury pokojowej wywoła zmianę rezystancji uzwojenia cewki z 6,43 Ω na 7,43 Ω.

## Program komputerowy

Kod źródłowy programu komputerowego urządzenia do wyznaczania temperatury cewki głośnika został napisany w języku C++ i bazuje na oryginalnym projekcie pochodzącym z Internetu. Na początek zostaje dodana biblioteka obsługi wyświetlacza LCD 2×16. Następnie zainicjowana zostaje komunikacja szeregową o wartości 9600 bitów na sekundę. W górnym wierszu wyświetlacza wyświetlony zostaje napis

## DOME TWEETER LOUDSPEAKER GDWK 11/100

9 5155 132

### TECHNICAL DATA

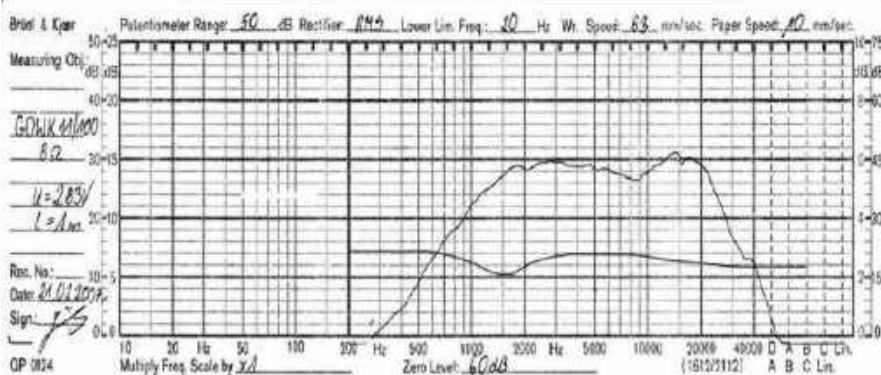
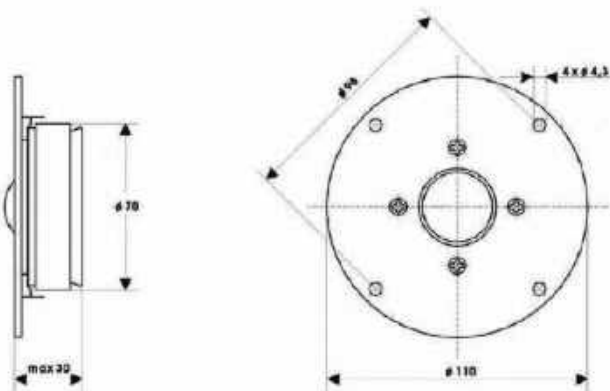
Rated impedance  
Voice coil resistance  
Rated frequency range  
Resonance frequency  
Recommended crossover frequency  
Power handling capacity, P1/P2  
measured with filter  
Max. power  
Sensitivity  
Flux density  
Energy in air gap  
Air - gap  
Voice coil height  
Magnet  
- material  
- dimensions  
- mass  
Mass of loudspeaker

8  $\Omega$   
5,3  $\Omega$   
4 to 20 kHz  
1,6 kHz  
2,8 kHz  
80/3 W  
16/5 W  
89 dB  
1,16 T  
95,6 mJ  
25/30,745 mm  
1,9 mm  
ferite  
70/32/15 mm  
0,23 kg  
0,6 kg

Membrane material - fabric

### AVAILABLE VERSIONS

- GDWK 11/100 - 8  $\Omega$ , catalogue number 9 5155 132 01



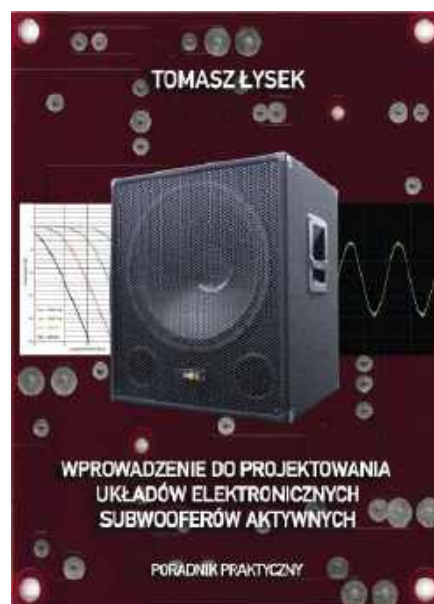
5. Karta katalogowa głośnika wysokotonowego typu GDWK 11/100  
(źródło: <https://www.skleptonsil.pl>)

„miliomierz”. Pętla szczytuje stan wejścia analogowego AO i przelicza go na wartość napięcia. Następnie program wyznacza wartość rezystancji dzieląc zmienną „voltage” przez wartość prądu wynoszącą 104 mA. W następnej kolejności mamy instrukcję warunkową, która w przypadku gdy wynik pomiaru będzie wyższy od 50  $\Omega$  wyświetli w dolnym wierszu napis „brak rezystancji”. W przypadku gdy

wynik pomiaru będzie niższy od 50  $\Omega$  w dolnym wierszu wyświetli się wynik pomiaru wraz z napisem „ohm”. Obydwa warianty instrukcji warunkowej mają zadeklarowane opóźnienia o wartości 1000 ms.

### Podsumowanie i wnioski

Konstrukcja urządzenia została maksymalnie uproszczona aby przyspieszyć



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

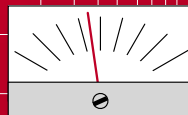
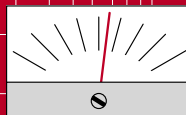
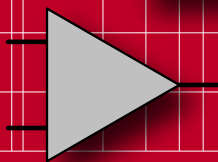
wykonywanie pomiarów. Jego główną zaletą stanowi możliwość bardzo szybkiego przełączania pomiędzy trybem podgrzewania a trybem pomiarowym. Dodatkowo warto zauważyć, że w trybie podgrzewania żaden z zacisków głośnika nie jest podłączony do masy, co z kolei umożliwia stosowanie nowoczesnych wzmacniaczy mostkowych, w których żaden z zacisków wyjściowych nie znajduje się na potencjale masy. Urządzenie może zostać wykorzystane do przeprowadzania badań niszczących mających na celu określenie dopuszczalnej mocy sygnału elektrycznego jaki może zostać doprowadzony do zacisków głośnika. Zarówno kod źródłowy jak i sam schemat ideowy urządzenia może zostać zmodyfikowany w taki sposób, aby wynik pomiaru rezystancji uzwojenia cewki głośnika w temperaturze pokojowej został wpisany do pamięci mikroprocesora a następnie, aby zarejestrowany przyrost rezystancji został przeliczony na temperaturę na podstawie współczynnika temperaturowego miedzi. Urządzenie to z całą pewnością można polecić entuzjastom elektroniki i elektroakustyki.

### Książka o układach elektronicznych do subwooferów aktywnych

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

mgr inż. Tomasz Łysek

# AUDIO OUT



## Kwestia symetrii, część 2

W zeszłym miesiącu przyjrzelśmy się podstawowym ideom przyświecającym przewodom symetrycznym do przedwzmacniaczy mikrofonowych, w tym złączom XLR i odpowiednim odmianom kabli. W tym miesiącu przejdziemy przez procedurę ich faktycznego tworzenia.

### Tworzenie własnych kabli – dlaczego warto to robić?

Można by się zastanawiać, dlaczego warto poświęcać czas na samodzielne wykonanie kabli, skoro można je po prostu kupić. Po pierwsze, zyskujesz satysfakcję, wiedzę i samowystarczalność, wynikające z samodzielnego tworzenia i naprawiania sprzętu. Po drugie, możesz stworzyć kabel naprawdę wysokiej jakości, który będzie trwały. Na przykład, stosując odpowiednie lutowie ołowiowe 60/40, które nie pęka. W Wielkiej Brytanii wszystkie dostępne na rynku kable musiały zostać wykonane przy użyciu bezołowiowego lutu na bazie cyny, który nie jest tak trwały. Wreszcie, chociaż dostępne są gotowe kable mikrofonowe Star Quad, ich cena wynosi około 25 funtów – samodzielnie można je wykonać znacznie taniej.

### Zdejmowanie izolacji

Najlepiej zdejmować izolację z przewodów za pomocą dedykowanego narzędzia, takiego jak AB Mk 1, które umożliwia pracę z krótkimi przewodami bez nadmiernego ciągnięcia. **Rysunek 18** przedstawia to narzędzie w akcji.



**Rysunek 18.** Efektywne zdejmowanie izolacji z przewodów sygnałowych wymaga dedykowanego narzędzia, które poradzi sobie z krótkimi przewodami. Używam tego AB Mk 1 od 1983 roku. Są one nadal dostępne w Grove Sales

Grubą zewnętrzną powłokę lub płaszcz kabla można zwykle „naciąć” za pomocą cążek bocznych, jak pokazano na **rysunku 19**. Sztuka polega na tym, aby nie ciąć zbyt głęboko. Jeśli oplot lub wewnętrzne przewody zostaną nacięte, trzeba je odciąć i zacząć od nowa. Należy naciąć około 80% grubości powłoki a następnie delikatnie pociągnąć kabel, aż pozostała nieodcięta część pęknie. Pomocne może być również odciąganie osłony na zewnątrz podczas nacinania kabla. Te techniki wymagają praktyki i mogą pozostawić nieestetyczną, poszarpaną krawędź. Alternatywnie, można użyć nożyków do tapet lub skalpeli, które – przy odpowiedniej wprawie – mogą dać lepsze rezultaty. Osobiście jednak nie przepadam za tym podejściem, ponieważ elektronika ociekająca krwią jest raczej nieśmaczna. Dostępne są profesjonalne narzędzia do zdejmowania izolacji z przewodów koncentrycznych (**rysunek 20**), które mogą dawać lepszy rezultat, ale nie wszystkie z nich działają tak, jak powinny. Każdy inżynier ma szufladę pełną narzędzi kupionych, a następnie w niej porzuconych. Moja szuflada wydaje się zawierać więcej bezużytecznych cążek, i szczypiec do gięcia przewodów niż jakichkolwiek innych narzędzi. Istnieje niedrogie narzędzie



**Rysunek 19.** Cążki boczne mogą być używane do „nacinania” zewnętrznej powłoki w celu jej usunięcia – dobre do małych partii kabli



**Rysunek 20.** Zestaw narzędzi do ściągania izolacji



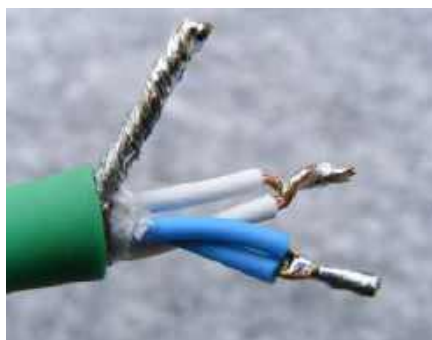
**Rysunek 21.** Zaskakująco eleganckie cięcie izolacji wykonane tanim narzędziem z niebieskim ostrzem z **rysunku 20**



**Rysunek 22.** Narzędzie MK02 z obrotowym ostrzem. Okrągła tuleja umożliwia ustawienie głębokości cięcia, a czarny przycisk służy do zmiany kierunku ostrza



Rysunek 23. Mantra druciarza: nasuń wszystkie wymagane końcówki na kabel przed lutowaniem



Rysunek 24. Skręć i pocynuj przewody, a następnie przytnij je na odpowiednią długość. Prawidłowo przygotowane zakończenie kabla przy użyciu Star Quad

z niebieskim plastikowym ostrzem ściskającym, które może dać dobre rezultaty (pokazane na **rysunku 21**) do okazjonalnego użytku. Osobiście używam narzędzia z obrotowym ostrzem MK02 z Ripley-Tools.com – pokazanego na **rysunku 22**. Należy ustawić głębokość ostrza, wcisnąć je w izolację, a następnie obrócić, aby wykonać cięcie. Potem należy nacisnąć przycisk, aby obrócić ostrze o 90 stopni i wykonać nacięcie wzdłuż osłony podczas jej zdejmowania. Warto się w nie zaopatrzyć, jeśli często pracujesz z przewodami.



Rysunek 25. Przylutuj „gorący” przewód, podgrzewając styk wypełniony lutem. Następnie włóż „zimny” przewód i przylutuj. Na koniec przylutuj krótszy ekran, wciskając go, pozostawiając „luzy” w przewodach sygnałowych



Rysunek 26. Tak powinna wyglądać prawidłowo przylutowana wkładka wtyku standardowego kabla mikrofonowego

## Problemy z zębami

Nie używaj zębów do ściągania izolacji! Musiałem wydać u dentysty 3000 funtów na naprawę mostków, a mój przyjaciel został porażony prądem tak mocno, że uderzył ciałem o szafkę. Szafka do dziś ma ogromne wgniecenie. Żartuję sobie, że jeśli kiedykolwiek będę potrzebował protezy, zafunduję sobie izolowany, wielozakresowy tytanowy zestaw do zdejmowania izolacji, na wzór „Szczęk” z filmów o Jamesie Bondzie.

## Pamiętaj o obudowie i koszulkach

Nie ma nic bardziej irytującego podczas tworzenia kabli niż wykonanie pięknego lutu i zorientowanie się, że zapomniałeś najpierw włożyć na kabel obudowę wtyku lub wsunąć na lutowaną żyłę koszulkę termokurczliwą (**rysunek 23**). Niestety, wciąż zdarza mi się to od czasu do czasu.

## Cynowanie

Zawsze warto jest skręcić i ocynować przewody przed lutowaniem. Wysokiej jakości kable mają cynowane przewody, które mają srebrzysty wygląd. Tańsze kable nie są cynowane i mają charakterystyczny pomarańczowy kolor gołej miedzi. Takie kable należy koniecznie najpierw pocynować. Skręcanie i cynowanie przed lutowaniem pozwala zapobiec wystawianiu pojedynczych drobnych żyłek składających się na cały rdzeń, które mogą powodować losowe trzaski i inne zakłócenia podczas pracy z takim kablem. Po cynowaniu zawsze należy przyciąć przewody na odpowiednią długość, aby umożliwić prawidłowe zaizolowanie wykonanych lutów.

## Długość przewodu

W przypadku złączy XLR przewody powinny mieć około 13 mm długości. W tej długości powinny się również zawrzeć odizolowane na długości 3 mm i pocynowane końcówki. Nieizolowany przewód ekranu musi być krótszy i mieć długość 10 mm. Zawsze lepiej jest skręcić i pocynować dłuższą część przewodu przed jego przycięciem do wymaganych 3 mm. Pozwala to na łatwiejsze trzymanie przewodu palcami i uzyskanie odpowiedniego okrągłego przekroju.



Rysunek 27. Czysty i ocynowany styk lutowniczy. Aby uzyskać najlepsze wyniki, należy użyć 3% aktywowanego lutu cynowo-ołowiowego 60/40 na bazie kalafonii. Zanieczyszczone topniki i europejskie lutowie bezołowiowe na bazie cyny nie są dobre. Najlepszy ze wszystkich lut to Multicore o niskiej temperaturze topnienia (LMP) z 2% zawartością srebra 62/36/2

Skręcone przewody często są nieco rozłożone na końcu skręcenia, co może utrudniać ich wkładanie, na przykład, w otwór złącza do zalutowania (**rysunek 24**).

W przypadku przewodów zasilania sieciowego przewód uziemiający musi być najdłuższy, aby odłączył się jako ostatni w przypadku naciągnięcia się lub szarpnięcia za przewód. Jest to oczywiście podyktowane względami bezpieczeństwa. Jednak w przypadku przewodów audio obowiązuje odwrotna zasada. Jeśli przewód uziemiający zostanie odłączony jako pierwszy, jedyne, co otrzymamy, to dodatkowy szum. Jest to mniej problematyczne niż utrata sygnału z powodu przerwania przewodów sygnałowych w czasie użytkowania przewodu. Łatwiej jest też namierzyć usterkę. Ponadto przewód ekranujący jest grubszy i mocniejszy niż przewody sygnałowe. Przewody sygnałowe są bardziej elastyczne, dlatego lutuje się je jako pierwsze, a ekran jest lutowany jako ostatni, jak pokazano na **rysunku 25**. Rezultat pokazano na **rysunku 26**.

## Styki lutownicze

W niektórych złączach XLR posrebrzane wyprowadzenia styków, do których lutuje się przewody, mogą się utleniać, jeśli złącza są przechowywane w pobliżu zanieczyszczeń, takich jak opary oleju napędowego lub dym siarkowy. Do ich czyszczenia używam małej mosiężnej szczotki lub pilnika igłowego. Następnie napełniam je lutem i wybijam go, uderzając złączem o blat, gdy lut jest jeszcze płynny (**rysunek 27**). Używaj dużej kolby, takiej jak Weller 60 W. Małe 15 W kolby Antex, których używam do lutowania płytek PCB, nie radzą sobie z pinami XLR o dużej pojemności cieplnej. Należy używać okularów ochronnych, ponieważ zwłaszcza podczas odlutowywania drobne kulki cyny potrafią zostać wystrzelone w kierunku twarzy. **Rysunek 28** pokazuje rezultat niewłaściwego cynowania, który może wynikać z użycia kolby o niewystarczającej pojemności



**Rysunek 28.** Jak tego nie robić. Wszystko, co mogło pójść źle, poszło źle: stopiona izolacja, niepoprawnie cynowane styki, szary bezołowiowy lut i zwarcia spowodowane luźnymi żyłami. Stały uczeń – praca na najniższą możliwą ocenę. Nie zaliczono



**Rysunek 29.** Zwykle nie jest to potrzebne, ale czasami obudowa XLR może być uziemiona do pinu 1 za pomocą wystającego pinu (znacznika)

cieplnej. Przewody „trzymają się” w zasadzie tylko na topniku.

Jedną z rzeczy, które lubię w złączach XLR, jest to, że numery pinów są zawsze oznaczone. Jeśli nie widzisz oznaczeń, potrzebujesz szkła powiększającego. Niektóre złącza XLR mają dodatkowy zacisk do uziemienia metalowej obudowy w celu zapewnienia ekranowania RF. Można go podłączyć do pinu 1,



**Rysunek 30.** Białe koszulki silikonowe używane do zapobiegania zwarciom spowodowanym przez niezainstalowane żyły. Widok przed zaciśnięciem kabla. Jeśli dobrze lutujesz, nie jest to konieczne. Mogą przykrywać pęknięcia bądź luty słabej jakości

jak pokazano na **rysunku 29**. W rzadkich przypadkach, gdy nie jest wymagane żadne połączenie między metalową obudową a pinem 1 (który w niektórych systemach może nieprawidłowo pełnić rolę uziemienia sygnału), złącze można pozostawić niepodłączone. AES zaleca jego niepodłączenie, z obawy o to, że metalowa obudowa dotknie i uziemi coś, czego nie powinna. W moich warsztatowych przewodach testowych zapobiegam temu, stosując kondensator 10 nF, który stanowi połączenie dla RF.

## Silikonowe koszulki izolacyjne

Wielu techników audio lubi nasuwać silikonowe gumowe tuleje na styki, aby zapobiec zwarciom, jak pokazano na **rysunku 30**. Ułatwia to odrobina smaru, takiego jak Hellerine. Większe tulejki są często umieszczane na końcu osłony kabla, aby zakryć niezainstalowane przewody i ułatwić zaciskanie kabli. Jeśli przygotowanie zostało wykonane prawidłowo, nie powinno być odsłoniętych przewodów, a tulejki powinny być zbędne.

## Zacisk kabla

Przed włożeniem śruby mocującej wkładkę należy najpierw dokręcić zacisk kabla w miejscu, w którym wchodzi on do obudowy. Wkładka powinna znajdować się w odległości kilku mm od otworu na śrubę, przed jej wciśnięciem (**rysunek 31**). Chodzi o to, aby zapewnić trochę dodatkowego luzu w złączu. Przewody nigdy nie powinny być ciasno ułożone wewnątrz obudowy wtyku, w przeciwnym razie mogą się złamać.

## Blokowanie śrub za pomocą koszulki termokurczliwej

Aby uniknąć zgubienia śrub w złączach XLR Cannon, zazwyczaj stosuje się koszulkę termokurczliwą, jak pokazano na **rysunku 32**. Dodatkowo, aby zabezpieczyć śruby, zaleca się nałożenie na nie odrobiny lakieru do paznokci.

## Zatrzaski

XLR to złącza zatrzaskowe, co oznacza, że „zaklikują” się one w gnieździe. Aby je



**Rysunek 31.** Najpierw należy dokręcić zacisk kabla przed całkowitym wsunięciem wkładki XLR (patrz otwór na śrubę). Gwarantuje to, że przewody sygnałowe nie są naprężone i nie ciągną za lut



**Rysunek 32.** Jeśli używasz złączy ze śrubami, warto użyć koszulki termokurczliwej, aby je unieruchomić. Koszulka może też zapobiec zwarciom obudowy z innymi elementami

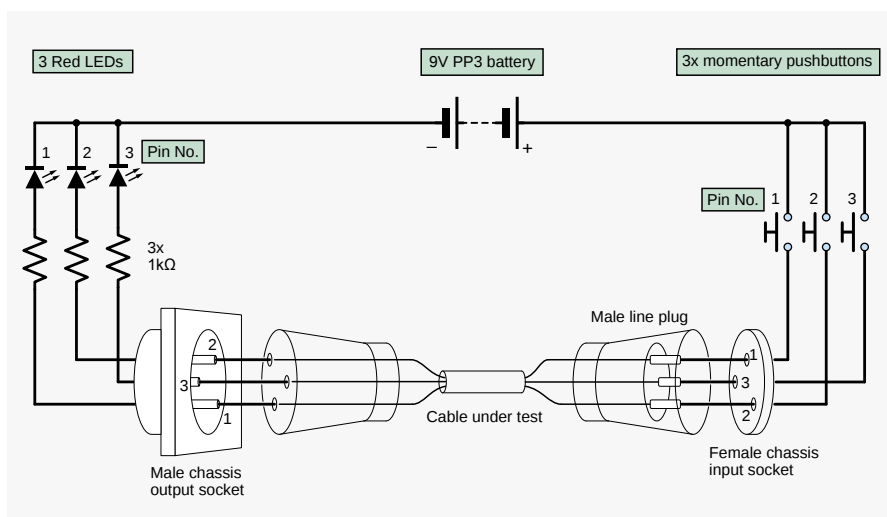
odblokować i wyciągnąć złącze, należy wcisnąć przycisk zatrzasku. Jest on obecny tylko w złączach żeńskich. Zatrzask to mieszane błogosławieństwo: z jednej strony zapobiega przypadkowemu wypięciu się złącza, z drugiej strony, jeśli ktoś potknie się o taki przewód, może to spowodować ściągnięcie sprzętu ze stojaków lub ławek. Te zatrzaski można zdemonstrować – ale nie rób tego na końcu mikrofonu, jeśli masz wokalistę, który lubi bawić się kablem!



**Rysunek 33.** Jeśli wykonujesz wiele przewodów, tester kabli jest rozsądną inwestycją. To urządzenie BSS jest przeznaczone do testowania przewodów XLR i TRS jack



**Rysunek 34.** Stare magnesy głośnikowe są dobrym sposobem trzymania przewodów podczas cynowania i lutowania. Często używam również testera pokazanego na **Rysunek 33** jako sposobu na trzymanie wkładek podczas lutowania



Rysunek 35. Łatwo jest stworzyć własny tester kabli XLR – ten układ jest wystarczający



Rysunek 36. Nie nawijaj kabli w ten sposób, ponieważ powoduje to naprężenie przewodów



Rysunek 37. Związkiwanie kabla na czas przechowywania jeszcze bardziej naraża go na uszkodzenia



Rysunek 38. Technicy sceniczni nawijają kabel skracając nadgarstek podczas zwijania, aby zapobiec zagięciom

## Testowanie

Kable należy zawsze testować pod kątem ciągłości i zwarc. Przykładem prostego testera kabli jest urządzenie firmy Brooke Siren Systems, pokazane na **rysunku 33**. Dwie diody LED świecące jednocześnie wskazują na zwarcie. Jeśli zaświeci się niewłaściwa dioda dla danego przełącznika, może to sugerować, że kilka przewodów zostało skrzyżowanych i doszło do zamiany faz. Podobny tester można również wykonać we własnym zakresie (**rysunek 35**). Testery są również przydatne do trzymania wkładek XLR podczas lutowania. Alternatywnie, czasami używam starych magnesów głośnikowych do trzymania kabli podczas lutowania, jak pokazano na **rysunku 34**. Kable zazwyczaj pękają na końcach złączy – w punktach największego naprężenia mechanicznego. Jeśli posiadasz miernik pojemności, możesz wykryć, który koniec kabla jest uszkodzony, stosując technikę znaną większości operatorów telekomunikacyjnych. Mierzy się pojemność na stykach 2 i 3. Usterka jest zlokalizowana na końcu o najmniejszej pojemności. Jeśli pojemność na obu końcach jest równa, przerwa może znajdować się pośrodku, ale nigdy nie spotkałem się z taką sytuacją.

## Zwijanie kabli

Wiele kabli audio ulega przedwczesnemu zużyciu w wyniku niewłaściwego zwijania. Niewłaściwym sposobem jest ciasne zwijanie go przy łokciu, jak pokazano na **rysunku 36**. Co gorsza, niektórzy ludzie zawiązują kabel na supeł, aby zapobiec jego rozwinięciu (**rysunek 37**). Prawidłowy sposób „roadie” polega na luźnym zawieszeniu zwoju w jednej ręce i delikatnym podawaniu go bez skręcania, jak pokazano na **rysunku 38**. Na koniec należy użyć opaski na rzep, aby kabel się nie rozwijał (**rysunek 39**). W złożonych



Rysunek 39. Opaski na rzep to dobry sposób na zapobieganie rozwijaniu się kabla. Tańszą opcją są nylonowe opaski kablowe



Rysunek 40. Wszystkie duże instalacje wymagają starannego oznakowania. Przydatne są te przypinane etykiety

konfiguracjach konieczne jest oznakowanie kabli. Istnieje wiele systemów. **Rysunek 40** przedstawia rozwiązanie z klipsem. Nie używaj taśmy izolacyjnej PVC, ponieważ klej wyciąga plastifikator z osłony PVC i tworzy lepki bałagan. **Rysunek 41** przedstawia pudełko kabli mikrofonowych typowe dla małego zespołu.

Na koniec, kilka przydatnych linków na temat różnych sposobów wykonywania kabli symetrycznych. Te dwie notatki firmy Rane omawiają wiele sposobów łączenia przewodów i zawierają dobre podsumowanie standardu AES48 (dostępnego bezpłatnie wyłącznie dla członków AES): <https://tiny.pl/d2d91> i <https://tiny.pl/d2d9p>. ■

Jake Rothman

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, październik 2021 ([www.epemag3.com](http://www.epemag3.com))



Rysunek 41. Ten stos kabli mikrofonowych domowej roboty zarabia na swoje utrzymanie od 1984 roku



## Migające diody LED i śliniący się inżynierowie (12)

Jeśli mam jedną wadę (i do niczego się nie przyznaję, rozumiecie), to jest nią to, że łatwo się rozpraszam. W rezultacie wiele moich projektów hobbystycznych zajmuje lata, zanim dojdą do skutku, ponieważ na scenie pojawia się coś innego, co przyciąga moją uwagę. Na przykład, w przypadku mojego Pedagogiczno-Fantasmagorycznego Silnika Progностycznego Inamorata, w chwili pisania tego tekstu zajmowałem się nim prawdopodobnie przez blisko 15 lat.

### Zwleknięcie z prognozowaniem

Niektórzy z moich przyjaciół są skłonni do czynienia niemiłych uwag, posuwając się nawet do rzucania oskarżeń i mówienia mi, że nie powinienem rozpoczynać żadnych nowych projektów, dopóki nie skończę niektórych spośród zaczętych. Czasami mój brak postępów sprawia, że czuję się smutny, ale – z powodów, które wkrótce wyjaśnię – teraz nie czuję się tak źle.

Podczas podróży do Anglii, aby odwiedzić moją matkę w 2018 roku, poszedłem zobaczyć wystawę prac Leonarda da Vinci, która była prezentowana w centrum Sheffield. Mogę tylko powiedzieć, że ten człowiek był geniuszem. Obecnie czytam jego biografię autorstwa Waltera Isaacsona (<https://amzn.to/33D0PGR>), który – jak zawsze – wykonuje świetną robotę. Wcześniej czytałem biografię Waltera o Stevie Jobsie i Einsteinie. Jeśli chodzi o tego pierwszego, myślę, że był genialnym facetem, którego nie polubiłbym zbytnio. Jeśli chodzi o tego drugiego, to chociaż nie zagłębiał się w matematykę, dopiero po przeczytaniu tej książki w pełni zrozumiałem, jak imponującym osiągnięciem była (i nadal jest) ogólna teoria względności Einsteina.

Chodzi o to, że czytając wspomnianą biografię Leonarda, odkryłem, że mamy ze sobą wiele wspólnego. To oczywiste (choć i tak to powiem), że obaj jesteśmy znani

z ciętego dowcipu, eleganckiego ubioru i sportowego wyczucia stylu. Co więcej, dowiedziałem się, że Leonardo zostawił o wiele więcej niedokończonych rzeczy, niż faktycznie ukończył. Ciągłe robił szkice z zamiarem przekształcenia ich w obrazy i nieustannie robił notatki w celu napisania traktatów, ale zanim zdążył coś zrobić, znów był zajęty czymś innym.

Z pewnością nie porównuję się z Leonardem da Vinci (takie porównania z przyjemnością zostawię mojej kochanej mamie, która bez wątpienia stwierdziłaby, że Leonardo jest na dalekim drugim miejscu), ale – mimo wszystko – miło jest wiedzieć, że nie jestem jedyną osobą, która łatwo się rozprasza.

### Końcowe odliczanie

Słowa i melodia The Final Countdown kołatają mi się obecnie po głowie. Piosenka ta, wydana przez szwedzki zespół rockowy Europe w 1986 roku, osiągnęła pierwsze miejsce na listach przebojów w 25 krajach. Nie mogłem się powstrzymać i zrobiłem sobie przerwę, by posłuchać jej na YouTube (<https://youtu.be/NNiTxUEnmKI>).

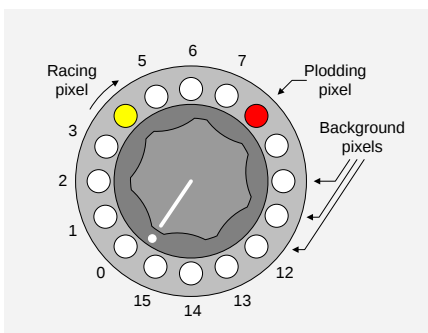
Ta podróż w głąb pamięci została wywołana przez myśli o scenie pod koniec amerykańskiego filmu akcji science fiction Predator z 1987 roku, w którym wystąpił Arnold Schwarzenegger. Myślę o fragmencie, w którym Dutch (Arnold) wyłącza urządzenie maskujące obcego, a następnie mija między krótkowzrocznego stwora pod przeciwną pulapkę, czyniąc go tym samym bardzo niebezpiecznym kosmitą. W odpowiedzi, śmiejąc się maniackalnie, bestia aktywuje urządzenie

autodestrukcyjne, którego sekwencja odliczania jest przedstawiona jako seria obracających się symboli na wyświetlaczu zamontowanym na ramieniu kosmity.

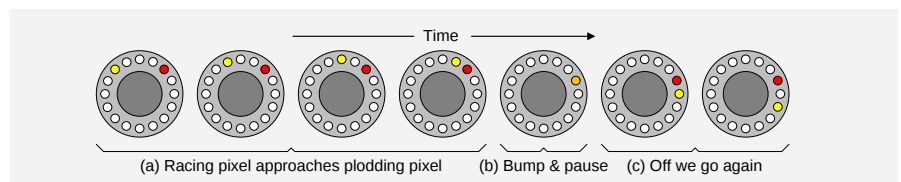
Pomyślałem, że moglibyśmy zaimplementować coś podobnego na pierścieniach NeoPixel, którymi obwieszony jest mój Silnik Progностyczny, jak omówiono w poprzednim odcinku tego cyklu. Ponieważ silnik ma duży czerwony przycisk („Nie naciskać!”), nasza sekwencja odliczania mogłaby być powiązana z aktywacją tego przycisku.

Oczywiście w naszym prototypie mamy tylko jeden pierścień do zabawy, ale będzie to punkt wyjścia. Myślę o zaimplementowaniu czegoś, co nazywam „pikselem powolnym” i „pikselem wyścigowym” (rysunek 1). Pamiętaj, że pokazane tutaj numery pikseli są sposobem, w jaki zdecydowaliśmy, że chcemy nimi manipulować, z pikselem 0 w pozycji „południowo-zachodniej” i numerami pikseli rosnącymi w kierunku zgodnym z ruchem wskazówek zegara. Rzeczywiste numery pikseli na fizycznym pierścieniu są zupełnie inne, więc stosujemy prostą operację konwersji przy użyciu tablicy liczb całkowitych o nazwie `RingXref[]`, aby przetłumaczyć sposób ustawienia pikseli z naszego idealnego sposobu widzenia świata na świat rzeczywisty (więcej szczegółów można znaleźć w artykule i kodzie z zeszłego miesiąca).

Pomysł polega na tym, że piksel wyścigowy porusza się w kierunku zgodnym z ruchem wskazówek zegara (rysunek 2a), aż dotrze do piksela powolnego, w którym to momencie „przeskakuje” piksel powolny do następnego miejsca zgodnego z ruchem



Rysunek 1. Odliczanie pikseli



Rysunek 2. Sekwencja odliczania pierwszego przejścia

wskazówek zegara (rysunek 2b). Następnie zatrzymujemy się na krótki czas, zanim piksel wyścigowy wyruszy ponownie (rysunek 2c).

Wcześniej (EdW, maj 2024) zauważyliśmy, że dobrym pomysłem jest uczynienie naszego kodu i funkcji tak prostymi i uniwersalnymi, jak to tylko możliwe. W związku z tym sposób, w jaki napisałem szkic (program) dla tego konkretnego efektu, polega na tym, że definiujemy cztery kolory – jeden dla pikseli tła (np. biały), jeden dla pikseli wyścigowych (np. żółty), jeden dla pikseli galopujących (np. czerwony) i jeden, którego używamy, gdy piksel wyścigowy wpada na piksel galopujący (np. pomarańczowy). Taki sposób działania pozwala nam szybko i łatwo eksperymentować z różnymi scenariuszami, co ilustruje nagrany właśnie film (<https://bit.ly/3gXpueJ>).

Jeśli chodzi o sam kod, tym razem zrobiliśmy to nieco inaczej (pełny szkic znajduje się w pliku CB-Feb21-01.txt – który jest dostępny na stronie <https://elportal.pl/do-pobrania>).

Na przykład rozważmy nasz kolor uderzeniowy (ten, którego używamy, gdy piksel wyścigowy zderza się z pikselem powolnym). Na podstawie naszych najnowszych programów (EdW, sierpień 2024) wiemy, że użyjemy funkcji `GetGammaCorrectedColor()`, aby skorygować wartość gamma naszego początkowego koloru (w tym przykładzie `COLOR_BUMP_PIXEL`). Następnie użyjemy naszej funkcji `ModifyBrightness()`, aby kontrolować intensywność tego koloru zdefiniowaną przez naszą stałą `BRIGHTNESS`. Wcześniej mogliśmy podzielić tę sekwencję na kilka dyskretnych kroków, jak poniżej:

```
tmpColor = GetGammaCorrectedColor(COLOR_BUMP_PIXEL);
tmpColor = ModifyBrightness(tmpColor, BRIGHTNESS);
ColorBumpPixel = tmpColor;
```

Wdrożyliśmy ten sposób, aby ułatwić zrozumienie tego, co robimy, ale teraz nadszedł czas, aby „zdemontować kółka boczne” (no cóż, tylko jedno kółko na początek) i uczynić rzeczy bardziej zwięzłymi, jak poniżej:

```
ColorBumpPixel =
ModifyBrightness(GetGammaCorrectedColor(COLOR_BUMP_PIXEL),
BRIGHTNESS);
```

Wiem, że na początku takie rzeczy mogą wydawać się nieco przerażające, ale tak naprawdę nie jest tak źle, jak się wydaje. Możemy myśleć o tym jak o warstwach cebuli. Najpierw zaczniemy od zewnątrz i przejdźmy do środka. Jak widzimy, przypisujemy wartość zwróconą z naszej funkcji `ModifyBrightness()` do naszej zmiennej `ColorBumpPixel`. Teraz nasza funkcja `ModifyBrightness()` wymaga dwóch argumentów – koloru, który ma zostać zmodyfikowany, oraz wartości, o jaką ma zostać zmodyfikowany (różnica między argumentami a parametrami została omówiona we wcześniejszym odcinku tego cyklu – EdW, maj 2024. W przypadku pierwszego argumentu, zamiast bezpośredniego określania koloru, używamy wartości koloru zwróconej z naszej funkcji `GetGammaCorrectedColor()`. Nie było tak źle, prawda?

Innym sposobem spojrzenia na to jest myślenie o tym, jak kompilator C/C++ widzi rzeczy, czyli zaczynając od patrzenia na wewnętrzne warstwy cebuli i pracując na zewnątrz. W tym przypadku najpierw określi wartość zwracaną z naszej funkcji `GetGammaCorrectedColor()`, a następnie użyje jej jako pierwszego argumentu naszej funkcji `ModifyBrightness()`.

Kolejną rzeczą, którą zamierzamy zrobić, jest zdefiniowanie dwóch globalnych zmiennych całkowitych o nazwach `PtrPloddingPixel` i `PtrRacingPixel`, których użyjemy do „wskazywania” odpowiednich pikseli (za chwilę zobaczysz, co mam na myśli). W przypadku piksela

wyścigowego chcemy włączyć nowy (następny) piksel i wyłączyć stary (ostatni) piksel, a następnie zaktualizować naszą zmienną `PtrRacingPixel`, aby wskazywała na następny piksel, zatrzymać się na krótkie opóźnienie, a następnie powtórzyć całą czynność, aż nasz piksel wyścigowy wpadnie na nasz piksel powolny.

W poprzednim odcinku przedstawiliśmy sposób określania numeru starego piksela za pomocą operatora `%` modulo, aby obejść warunki brzegowe, gdy przechodzimy z piksela 15 z powrotem do piksela 0 (lub z 0 do 15, jeśli nasz piksel wyścigowy poruszał się w kierunku przeciwnym do ruchu wskazówek zegara). Tym razem moglibyśmy użyć podobnej sztuczki, ale zamiast tego zastosujemy nieco inną technikę.

Być może pamiętasz, że wprowadziliśmy pojęcia `typedef` (defini-cje typów), `enum` (typy wyliczeniowe) i `struct` (struktury) w odcinku z EdW, lipiec 2024. W tym przypadku zaczniemy od zadeklarowania struktury o nazwie `NextLastPair`, a następnie zadeklarujemy tablicę tych elementów o nazwie `PixelPtrs`.

```
typedef struct
{
    int next;
    int last;
} NextLastPair;
NextLastPair PixelPtrs[NUM_NEOS];
```

Pamiętajmy, że liczba pikseli w naszym pierścieniu, reprezentowa-na przez `NUM_NEOS`, wynosi 16. Pomysł polega na tym, że dla dowolnego elementu w tej tablicy, który, jak zakładamy, odpowiada bieżącemu pikselowi, następna wartość będzie numerem sąsiedniego piksela w kierunku zgodnym z ruchem wskazówek zegara, podczas gdy wartość `last` będzie numerem sąsiedniego piksela w kierunku przeciwnym do ruchu wskazówek zegara.

Sprytna część polega na zainicjowaniu tej tablicy w naszej funkcji `setup()`, jak pokazano poniżej. Ta inicjalizacja opiera się na tych samych sztuczkach z operatorem modulo, które wprowadziliśmy w poprzednim odcinku.

```
for (int iNeos = 0; iNeos < NUM_NEOS; iNeos++)
{
    PixelPtrs[iNeos].next = (iNeos + 1) % NUM_NEOS;
    PixelPtrs[iNeos].last = (iNeos + (NUM_NEOS - 1))
        % NUM_NEOS;
}
```

| iNeos | +1 | %16 | iNeos | +15 | %16 |
|-------|----|-----|-------|-----|-----|
| 0     | 1  | 1   | 0     | 15  | 15  |
| 1     | 2  | 2   | 1     | 16  | 0   |
| 2     | 3  | 3   | 2     | 17  | 1   |
| 3     | 4  | 4   | 3     | 18  | 2   |
| 4     | 5  | 5   | 4     | 19  | 3   |
| 5     | 6  | 6   | 5     | 20  | 4   |
| 6     | 7  | 7   | 6     | 21  | 5   |
| 7     | 8  | 8   | 7     | 22  | 6   |
| 8     | 9  | 9   | 8     | 23  | 7   |
| 9     | 10 | 10  | 9     | 24  | 8   |
| 10    | 11 | 11  | 10    | 25  | 9   |
| 11    | 12 | 12  | 11    | 26  | 10  |
| 12    | 13 | 13  | 12    | 27  | 11  |
| 13    | 14 | 14  | 13    | 28  | 12  |
| 14    | 15 | 15  | 14    | 29  | 13  |
| 15    | 16 | 0   | 15    | 30  | 14  |

(a) The next values      (a) The last values

Generowanie następnej i ostatniej wartości

Wiem, że może to być nieco wymagające dla szarych komórek, więc wygenerowałem wyniki, abyś mógł je przejrzeć i podumać (**rysunek 3**). W przypadku wartości `next`, dla każdej wartości iNeos dodajemy 1, a następnie dzielimy wynik przez 16 (`NUM_NEOS`) za pomocą operatora modulo, aby uzyskać resztę, która – dzięki magii matematyki – jest liczbą, którą chcemy. Podobnie w przypadku ostatnich wartości, dla każdej wartości iNeos dodajemy 15 (`NUM_NEOS - 1`) i ponownie dzielimy wynik przez 16 (`NUM_NEOS`) za pomocą operatora modulo, aby uzyskać resztę, która – po raz kolejny – zapewnia nam wymaganą wartość.

Reszta jest bardzo prosta. Wszystko, co musimy zrobić w naszej funkcji `main` `loop()`, to cyklicznie włączać i wyłączać piksele wyścigowe i powolne, w zależności od potrzeb, w tym (ewentualnie) zmieniać kolor piksela powolnego, gdy wpadnie na niego piksel wyścigowy. Cały ten kod można zobaczyć w pliku `CB-Feb21-01.txt`, do którego można uzyskać dostęp na stronie <https://elportal.pl/do-pobrania>.

Oczywiście nasze eksperymenty dotyczyły jednego pierścienia. Silnik prognostyczny ma pięć takich pierścieni, co otwiera drzwi do wielu

innych możliwości. Na przykład za każdym razem, gdy pierwszy (najmniej znaczący) pierścień doświadcza „uderzenia”, może to zwiększyć piksel wyścigowy na następnym pierścieniu i tak dalej w górę łańcucha. Moja głowa aż huczy od pomysłów, ale obawiam się, że będą one musiały poczekać na swoją kolej.

Ale gdzie jest GOL?

Obawiam się, że was zawiodłem i sprowadziłem na manowce. Na koniec poprzedniego odcinka wspomniałem o tym, że w tym miesiącu wykorzystamy naszą tablicę piłeczek pingpongowych 12×12 do zaimplementowania odmiany słynnej „Gry w Życie” (Game of Life, GOL), której autorem jest brytyjski matematyk John Horton Conway (<https://bit.ly/pe-jan21-cgol>). Przykro mi to mówić, ale zostałem zepchnięty na boczny tor przez sekwencję odliczania omówioną powyżej, więc obawiam się, że będziemy musieli zostawić GOL do następnego razu (zwieszam głowę ze wstydu). Do tego czasu, jak zawsze, czekam na wasze komentarze, pytania i sugestie. ■

Clive „Max” Maxfield

## Sprytne porady i sztuczki cyklu Ekscytacje Maxa dotyczące kodowania



Nieżyjący już Stephen Hawkins napisał niesamowitą książkę zatytułowaną *God Created the Integers: The Mathematical Breakthroughs That Changed History* (<https://amzn.to/3lKfV3e>). Leży na półce w moim biurze. Pewnego dnia zamierzam ją przeczytać.

Powodem, dla którego ta książka przyszła mi do głowy, jest to, że wiele rzeczy w programowaniu C/C++ w ogóle, a programowanie dla Arduino w szczególności, to na najbardziej podstawowym poziomie liczby całkowite w takiej czy innej formie.

### Wielka prawda, czy mały fałsz?

Rozważmy na przykład poziomy pinów `LOW` i `HIGH`. Jeśli zdefiniowaliśmy pin jako wyjście cyfrowe, możemy przypisać mu wartość `LOW` lub `HIGH` w następujący sposób:

```
digitalWrite(PinA, LOW);
digitalWrite(PinB, HIGH);
```

W świecie rzeczywistym, zakładając Arduino Uno, pojawią się one na pinie jako wartości odpowiednio 0 V i 5 V. Ale czy wiesz, że możemy osiągnąć dokładnie ten sam efekt używając 0 i 1? Na przykład:

```
digitalWrite(PinA, 0);
digitalWrite(PinB, 1);
```

Jak to możliwe? Cóż, gdyby powiedzieć prawdę, za kulisami („Nie zwracaj uwagi na tego człowieka za kurtyną”), twórcy Arduino, zaimplementowali coś w rodzaju instrukcji `#define`, aby zrównać słowa kluczowe `LOW` z 0 i `HIGH` z 1.

Znajomość tej ciekawostki otwiera drzwi do wielu rzeczy, gdy się nad tym zastanowimy. Na przykład, rozważmy następujący fragment kodu:

```
int tmpValue;

tmpValue = digitalRead(PinA);
```

```
if (tmpValue == HIGH)
{
    // Zrób kilka rzeczy
}
```

Teraz wiemy, że możemy zastąpić `HIGH` liczbą 1 lub wynikiem wyrażenia całkowitoliczbowego, lub wartością całkowitą zwracaną przez funkcję, lub... lista jest długa.

### Wszyscy mamy swoje prawdy

Jak słynnie mówi (lub śpiewa) Poncjusz Piłat w operze rockowej *Jesus Christ Superstar* z 1970 roku, z muzyką Andrew Lloyd Webbera i tekstami Tima Rice’a: „Ale czym jest prawda? Czy prawda jest niezmiennym prawem? Oboje mamy prawdy – czy moja jest taka sama jak twoja?”. Wiem, co ma na myśli i za chwilę ty też będziesz wiedział.

W latach pięćdziesiątych XIX wieku angielski matematyk, filozof i logik George Boole (1815–1864), będący w dużej mierze samoukiem, opracował dział matematyki, który obecnie nazywamy algebrą Boole’a. Zamiarem Boole’a było wykorzystanie technik matematycznych do reprezentowania i testowania argumentów logicznych i filozoficznych. Niestety, znaczenie jego pracy nie zostało w pełni docenione aż do późnych lat trzydziestych XX wieku, kiedy to absolwent MIT, Claude Shannon, zrewolucjonizował elektronikę, przedstawiając pracę magisterską pokazującą, w jaki sposób algebra Boole’a oferuje idealny sposób reprezentowania logicznego działania systemów cyfrowych.

W przypadku Arduino mamy typ danych o nazwie `bool` (lub `boolean`). Zmiennym zadeklarowanym przy użyciu tego typu można przypisać wartości `false` lub `true`. Załóżmy, że zadeklarujemy zmienną typu `boolean` o nazwie `tmpBool` i przypiszemy jej wartość `true` lub `false`. Później możemy wykonać kilka testów, takich jak:

```
if (tmpBool == true)
{
```

```

// Zrób kilka rzeczy
}

if (tmpBool == false)
{
    // Zrób kilka rzeczy
}

```

Zauważ, że sposób, w jaki piszę ten kod – w przeciwieństwie do użycia konstrukcji `if else` – ma na celu zilustrowanie sedna sprawy. Ponadto, jako ciekawostkę, moglibyśmy uczynić powyższe wyrażenia bardziej związanymi i osiągnąć dokładnie ten sam efekt, pisząc je w następujący sposób:

```

if (tmpBool)
{
    // Zrób kilka rzeczy
}

if (!tmpBool)
{
    // Zrób kilka rzeczy
}

```

W tym ostatnim przykładzie rozumiemy, że `!tmpBool` oznacza „Not tmpBool”.

## Chodzi o liczby całkowite

Pod maską kompilator C/C++ – a co za tym idzie Arduino – ma tendencję do traktowania operacji logicznych tak, jakby zmienne były liczbami całkowitymi. Zacznijmy od wartości `false`, która równa się 0. Na tej podstawie wiele osób uważa, że wartość `true` równa się 1. Jest to prawda (bez zamierzonej gry słów), ale w rzeczywistości każda niezerowa liczba całkowita jest uważana za `true`, więc 1, 2, 6, 10, 100... są prawdziwe w sensie boolowskim.

Z kolei zakładając, że zadeklarowaliśmy zmienną całkowitą o nazwie `tmpInt` i zmienną logiczną o nazwie `tmpBool`, pozwala nam to robić takie rzeczy jak:

```

tmpBool = true;
tmpInt = tmpBool;
tmpInt = false;

```

```

tmpInt = 6;

```

```

tmpBool = tmpInt;
tmpBool = 12;

```

Oczywiście jest to bezsensowny kod, ale daje nam wyobrażenie o tym, co możemy zrobić. Wróćmy teraz do rozważenia instrukcji `if()` w następujący sposób:

```

if (condition)
{
    // Zrób kilka rzeczy
}

Instrukcje objęte instrukcją if() zostaną wykonane tylko wtedy, gdy
warunek zwróci wartość true (tj. niezerową). Załóżmy na przykład,
że mamy coś takiego jak poniżej:

if (tmpInt == 6)
{
    // Zrób kilka rzeczy
}

```

Jeśli `tmpInt` zawiera wartość 6, wówczas warunek zwróci wartość `true` (1) i instrukcje powiązane z `if()` zostaną wykonane; w przeciwnym razie warunek zwróci wartość `false` (0), co uniemożliwi wykonanie instrukcji powiązanych z `if()`. Powodem, dla którego o tym wspominam, jest to, że częstym błędem jest używanie pojedynczego „=” (przypisanie) zamiast podwójnego „==” (porównanie); na przykład:

```

if (tmpInt = 6)
{
    // Zrób kilka rzeczy
}

```

W tym przypadku kompilator najpierw oceni wyrażenie i przypisze wartość 6 do `tmpInt` (co zepsuje resztę programu). Ponieważ procesor zobaczy, że końcowym wynikiem „porównania” jest 6, potraktuje to jako prawdę (niezerową), a następnie wykona instrukcje związane z `if()`.

Będziemy nadal rozważać konsekwencje tego wszystkiego w następnym odcinku. ■

Clive „Max” Maxfield

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

REKLAMA

Publikujemy dla projektantów  
i programistów elektroniki

**ELPORTAL.pl**

Znajdziesz nas również na Facebooku: [facebook.com/ElportalPL](https://facebook.com/ElportalPL)



## Miernik LCR 3730 firmy PeakTech

Testujemy teraz typowy europejski produkt: miernik LCR 3730 niemieckiej firmy PeakTech. Nieco droższy niż jego chińskie odpowiedniki, ale jak się okazuje, znacznie bardziej wytrzymały i dokładny.

### Wprowadzenie do miernika 3730

Producentem jest firma PeakTech Prüf- und Messtechnik GmbH z siedzibą w Ahrensburgu w Niemczech. Firma ta dostarcza wiele elektronicznych przyrządów pomiarowych. Jednak najwyraźniej nie są one produkowane w Niemczech, ponieważ z tyłu modelu 3730 znajduje się napis „Made in China”. Może to oznaczać, że PeakTech ma tylko produkty zaprojektowane w Chinach i opatrzone logo PeakTech, a tym samym obsługuje rynek europejski. Ale może to również oznaczać, że firma (podobnie jak większość europejskich firm) sama opracowuje produkty i zleca ich produkcję w Chinach zgodnie z europejskimi standardami. Kto wie?

Urządzenie jest sprzedawane w sklepie AVT pod oznaczeniem UT603 <https://sklep.avt.pl/pl/products/miernik-rlc-mostek-ut603-172941.html>

### W zestawie

Model 3730 (lub UT603) jest dostarczany w pięknie zadrukowanym kartonowym pudełku. W pudełku tym znajdziemy:

- Bateria 9 V
- Dwa krótkie przewody pomiarowe z zaciskami krokodylkowymi
- Instrukcja obsługi

Na stronie sklepu [www.sklep.avt.pl](http://www.sklep.avt.pl) jest dostępny 26-stronnicowy podręcznik w języku angielskim oraz instrukcja obsługi w języku polskim.

### Ogólna charakterystyka miernika

Miernik 3730/UT603 pozwala mierzyć rezystory, kondensatory i cewki. Jako dodatkowe funkcje można mierzyć napięcia przewodzenia diod i wzmocnienie prądowe hFE tranzystorów. Minimalne i maksymalne zakresy pomiarowe to:

- rezystory: 200  $\Omega$  do 20 M $\Omega$
- kondensatory: 2 nF do 600  $\mu$ F
- cewki: 2 mH do 20 H

To, co od razu rzuca się w oczy, to nieco staroświecki wygląd modelu 3730. Zakres pomiarowy trzeba ustawiać ręcznie, większość nowoczesnych mierników robi to w pełni automatycznie. Trzeba również ręcznie przełączać się z L na C, a nawet wymieniać przewody pomiarowe, jeśli trzeba przełączyć się z L/C na R. Wyświetlacz to „staromodny” siedmiosegmentowy monochromatyczny wyświetlacz LCD, który nawet nie wskazuje jednostki miary. Również nieco staroświecki z uwagi na zasilanie: jest on zasilany baterią 9 V,



Tak wygląda 3730 na stole warsztatowym (© PeakTech)

podczas gdy nowoczesne chińskie mierniki są teraz z reguły wyposażone w akumulator, który można ładować za pomocą zasilacza USB 5 V.

Na próżno szukać w nim nowoczesnych bajerów, takich jak funkcja automatycznego wyłączania po określonym czasie czy podświetlanie wyświetlacza w tle. Brakuje również opcji kompensacji rezystancji i pojemności przewodów pomiarowych za pomocą funkcji „REL”.

Wyświetlacz ma zakres do 1999 i wysokość cyfr 21 mm. Urządzenie jest umieszczone w jasnoszarej plastikowej obudowie z niebieską ochronną i elastyczną silikonową osłoną. Duży przełącznik obrotowy jest łatwy w obsłudze i ma wyraźne kliknięcia.

Komora baterii znajduje się z tyłu za pokrywą, którą można otworzyć za pomocą jednej śruby. Rozkładane oparcie jest przydatne, umożliwiając umieszczenie urządzenia na stole warsztatowym w wygodnej do odczytu pozycji.

### Dane techniczne miernika 3730/UT603

Według producenta urządzenie to ma następujące parametry:

- Wyświetlacz: monochromatyczny siedmiosegmentowy wyświetlacz LCD, wysokość cyfr 21 mm, może wskazać od 0 do 1999
- Ustawienie zera: automatyczne
- Prędkość pomiaru: 2,5 pomiaru na sekundę
- Zakresy pomiaru rezystancji:
  - 199,9  $\Omega$ , rozdzielczość: 0,1  $\Omega$ , dokładność:  $\pm[0,8 \% + 3]$
  - 1,999 k $\Omega$ , rozdzielczość: 1  $\Omega$ , dokładność:  $\pm[0,8 \% + 1]$
  - 19,99 k $\Omega$ , rozdzielczość: 10  $\Omega$ , dokładność:  $\pm[0,8 \% + 1]$
  - 199,9 k $\Omega$ , rozdzielczość: 100  $\Omega$ , dokładność:  $\pm[0,8 \% + 1]$
  - 1,999 M $\Omega$ , rozdzielczość: 1 k $\Omega$ , dokładność:  $\pm[0,8 \% + 1]$
  - 19,99 M $\Omega$ , rozdzielczość: 10 k $\Omega$ , dokładność:  $\pm[2,0 \% + 5]$
- Zakresy pomiarowe pojemności:
  - 1,999 nF, rozdzielczość: 1 pF, dokładność:  $\pm[1,0 \% + 5]$
  - 19,99 nF, rozdzielczość: 10 pF, dokładność:  $\pm[1,0 \% + 5]$
  - 199,9 nF, rozdzielczość: 100 pF, dokładność:  $\pm[1,0 \% + 5]$
  - 1,999  $\mu$ F, rozdzielczość: 1 nF, dokładność:  $\pm[4,0 \% + 5]$
  - 19,99  $\mu$ F, rozdzielczość: 10 nF, dokładność:  $\pm[4,0 \% + 5]$
  - 199,9  $\mu$ F, rozdzielczość: 100 nF, dokładność:  $\pm[4,0 \% + 5]$
  - 600  $\mu$ F, rozdzielczość: 1  $\mu$ F, dokładność: nieokreślona
- Sygnał testujący dla pomiarów pojemności: 1 kHz ~ 150 mV do 100 Hz ~ 1,5 mV

- Zakresy pomiarowe indukcyjności:
  - 1,999 mH, rozdzielczość: 0,001 mH, dokładność:  $\pm[2,0 \% + 8]$
  - 19,99 mH, rozdzielczość: 0,01 mH, dokładność:  $\pm[2,0 \% + 8]$
  - 199,9 mH, rozdzielczość: 0,1 mH, dokładność:  $\pm[2,0 \% + 8]$
  - 1,999 H, rozdzielczość: 0,001 H, dokładność:  $\pm[5,0 \% + 5]$
  - 19,99 H, rozdzielczość: 0,01 H, dokładność:  $\pm[5,0 \% + 15]$
- Sygnał testujący dla pomiarów indukcyjności: 1 kHz ~ 150  $\mu$ A do 100 Hz ~ 15  $\mu$ A
- Zabezpieczenie przepięciowe: do 250 V skuteczne
- Bezpiecznik przeciążeniowy: 0,315 A ~ 250 V ~ 20 mm  $\times$  5 mm
- Zakres pomiarowy wzmocnienia  $\beta$  tranzystora: 1000
- Napięcie pomiarowe podczas pomiaru diod: 5,8 V maks.
- Zakres pomiarowy napięcia przewodzenia diod: maks. 1,999 V.
- Prąd pomiaru diod: 1 mA
- Rozdzielczość napięcia przewodzenia diody: 1 mV
- Zasilanie: bateria 9 V
- Wskaźnik niskiego poziomu baterii: tak
- Funkcja REL: nie
- Automatyczne wyłączanie: nie
- Test ciągłości: sygnał dźwiękowy przy  $<10 \Omega$
- Wymiary: 172 mm  $\times$  83 mm  $\times$  38 mm
- Waga: 310 g

### Praca z miernikiem 3730/UT603

Na zdjęciu widać wszystkie elementy sterujące na płycie czołowej. Praca z urządzeniem jest niezwykle prosta. Niebieski przycisk służy do włączania i wyłączania urządzenia. Żółty przycisk umożliwia wybór pomiędzy pomiarem kondensatorów (wciśnięty) lub cewek (nie wciśnięty). Za pomocą dużego pokrętki można wybrać funkcję pomiarową i zakres. Dwa lewe gniazda bananowe 4 mm są przeznaczone do pomiaru L i C, a dwa prawe do wszystkich innych pomiarów.

Aby wyeliminować wpływ pasożytniczej pojemności przewodów pomiarowych, małe kondensatory można wkładać bezpośrednio do lewego gniazda pod przełącznikiem obrotowym. Prawe gniazdo służy do podłączania tranzystorów NPN lub PNP.

### Elektronika w mierniku 3730/UT603

**Otwieranie obudowy.** Aby otworzyć obudowę urządzenia, należy najpierw zdjąć niebieską silikonową osłonę oraz wyjąć baterię 9 V. Następnie, po odkręceniu trzech śrubek, można zdemonstrować obudowę.

Jak pokazuje poniższe zdjęcie, wewnątrz miernika robi bardzo solidne wrażenie. Tylna część obudowy jest pokryta warstwą folii aluminiowej, która pełni rolę ekranu elektromagnetycznego. Ta warstwa jest połączona z górną płytką PCB za pomocą sprężyny. Warto dodać, że elektronika jest rozmieszczona na dwóch płytkach PCB: dolna płytka wypełnia całą przestrzeń obudowy, natomiast górna płytka jest połączona z dolną za pomocą dwóch złączy.

W tylnej części obudowy znajdują się cztery solidne rozpórki, które utrzymują styki w czterech gniazdach bananowych na miejscu i odciążają je mechanicznie podczas podłączania wtyczek. Dobry projekt rozpoznaje się po takich szczegółach! Bateria jest połączona z płytką drukowaną za pomocą dwóch styków sprężynowych.

**Płytką główną.** Dwie strony podstawowej płytki głównej przedstawiono na poniższym rysunku. Prawy obrazek pokazuje stronę, która przylega do frontu przyrządu. Na górze widać pasek stykowy, za pomocą którego elektronika styka się z wyświetlaczem. Poniżej tego paska, oprócz czerwonego kondensatora, znajduje się czarny kleks, pod którym niewątpliwie znajduje się układ sterujący wyświetlaczem. Gdy spojrzysz w dół, natkniesz się na dwa przyciski. Pod tymi



Otwarta obudowa miernika 3730 od PeakTech (© 2024 Jos Verstraten)

przyciskami wytrawiony jest jeden z dwóch bardzo skomplikowanych wzorów, stanowiących element przełącznika obrotowego. Segmenty stykają się z szeregiem styków sprężynowych znajdujących się z tyłu pokrętkła zamontowanego na froncie. To nie jedyna część przełącznika funkcji i zakresu. Na lewym obrazku widać, że po drugiej stronie głównej płytki drukowanej znajduje się druga część przełącznika obrotowego, która obraca się razem z częścią przełącznika na płycie czołowej. Styki przełącznika podążają za wzorami (ścieżkami) na dużej i małej płytce drukowanej.

Po lewej i prawej stronie tego obrotowego przełącznika znajdują się dwa złącza, pozwalające połączyć obie płytki drukowane.

Na samym dole głównej płytki drukowanej znajduje się bezpiecznik, który chroni urządzenie przed nadmiernymi napięciami i prądami podawanymi na wejścia. Zielony element będzie prawdopodobnie współpracować z tym bezpiecznikiem i bez wątpienia jest to warystor typu MOV lub element PTC.

Na głównej płytce drukowanej znajdziemy dziesięć układów scalonych. Producent nie stara się ukryć typów zastosowanych układów. Krótka eksploracja za pomocą naszego mikroskopu LCD daje wyniki:

- 1× CE7660, konwerter napięcia DC na DC,
- 1× HEF4066, poczwórny przełącznik analogowy,



Dwie strony głównej płytki drukowanej (© 2024 Jos Verstraten)



Dwie strony drugiej płytki drukowanej (© 2024 Jos Verstraten)

- 1× 74HC4053, potrójny przełącznik analogowy,
- 1× LM385, źródło napięcia odniesienia,
- 6× MZ145, nieznan.

Dziwne jest to, że nie możemy znaleźć żadnych informacji o MZ145 w Internecie. Prawdopodobnie jest to op-amp.

**Druga płytka drukowana w stylu pin-up.** Płytkę PCB, której obie strony pokazano na poniższym rysunku, jest również dość skomplikowana. Znajdziemy tu również siedem układów scalonych:

- 1× ICM7555, zegar analogowy,
- 1× AX16P36, nieznan,
- 1× CD4024, siedmiostopniowy licznik/rozdzielacz binarny,
- 1× CD4070, poczwórny port EXOR,
- 1× HEF4066, poczwórny przełącznik analogowy,
- 1× TL062, podwójny wzmacniacz operacyjny,
- 1× L062, wzmacniacz operacyjny JFET.

Druga płytka drukowana jest przymocowana do podstawowej płytki drukowanej za pomocą czterech śrub.

## Testy miernika 3730/UT603

**Testowany zakres rezystancji.** Na potrzeby tego testu mamy dostęp do zestawu rezystorów referencyjnych o tolerancji  $\pm 0,01\%$  oraz kilku mniej dokładnych. Jako miernika referencyjnego używamy oczywiście modelu 8842A firmy Fluke. Miernik ten wykorzystuje czteroprzewodową sondę Kelvina, co nie jest możliwe w przypadku 3730. Dla najniższych zakresów od wskazania odejmujemy rezystancję przewodu pomiarowego 0,2  $\Omega$ . Wyniki są ponownie podsumowane w tabeli 1. Zakładamy, że zgadzasz się z nami, że dokładność 3730 jest doskonała!

**Testowane zakresy kondensatorów.** Dzięki zestawowi pięciu precyzyjnych kondensatorów z tolerancją  $\pm 1\%$  i trzech standardowych kondensatorów z tolerancją tylko  $\pm 0,05\%$ , możemy dokładnie odwzorować wydajność 3730 podczas pomiaru takich komponentów. Powyżej 1  $\mu\text{F}$  mierzymy normalne kondensatory elektrolityczne z naszego magazynu. Jako miernika referencyjnego używamy znacznie

| Tabela 1. Dokładność pomiaru rezystancji |                  |                    |
|------------------------------------------|------------------|--------------------|
| Testowany rezystor                       | Wskazania 3730   | Wskazania 8842A    |
| 10 $\Omega$ – $\pm 1\%$                  | 10,0 $\Omega$    | 10,009 $\Omega$    |
| 100 $\Omega$ – $\pm 0,01\%$              | 99,7 $\Omega$    | 100,008 $\Omega$   |
| 1 k $\Omega$ – $\pm 0,01\%$              | 1,001 k $\Omega$ | 1,00011 k $\Omega$ |
| 10 k $\Omega$ – $\pm 0,01\%$             | 9,97 k $\Omega$  | 10,006 k $\Omega$  |
| 100 k $\Omega$ – $\pm 0,01\%$            | 100,0 k $\Omega$ | 100,020 k $\Omega$ |
| 1 M $\Omega$ – $\pm 0,01\%$              | 997 k $\Omega$   | 999,86 k $\Omega$  |
| 10 M $\Omega$ – $\pm 5\%$                | 10,29 M $\Omega$ | 10,1852 M $\Omega$ |

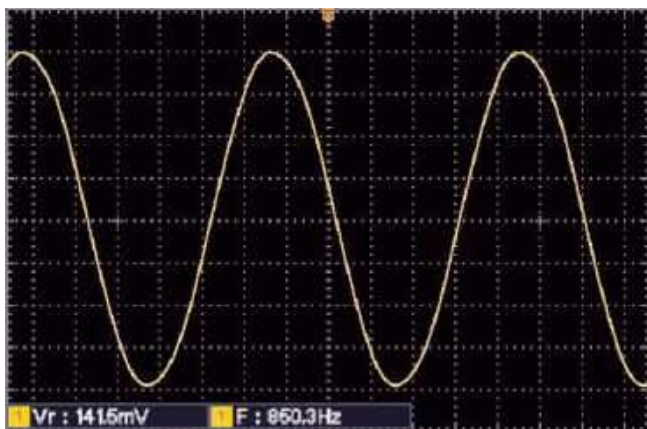
| Tabela 2. Dokładność pomiaru kondensatorów |                |                  |
|--------------------------------------------|----------------|------------------|
| Testowany kondensator                      | Wskazania 3730 | Wskazania ET4401 |
| 100 pF – $\pm 1\%$                         | 98 pF          | 99,083 pF        |
| 1 nF – $\pm 1\%$                           | 1,009 nF       | 1,0056 nF        |
| 10 nF – $\pm 0,05\%$                       | 9,99 nF        | 10,002 nF        |
| 100 nF – $\pm 0,05\%$                      | 100,0 nF       | 100,03 nF        |
| 1 $\mu$ F – $\pm 0,05\%$                   | 1,009 $\mu$ F  | 1,0015 $\mu$ F   |
| 10 $\mu$ F – $\pm 20\%$                    | 9,95 $\mu$ F   | 9,155 $\mu$ F    |
| 100 $\mu$ F – $\pm 20\%$                   | 87,9 $\mu$ F   | 81,52 $\mu$ F    |
| 1 mF – $\pm 20\%$                          | 714 $\mu$ F    | 787,6 $\mu$ F    |

droższego ET4401 firmy East Tester o dokładności  $\pm 0,2\%$  dla kondensatorów nieelektrolitycznych.

ET4401 mierzy za pomocą sond Kelvina i dlatego nie przekłamuje z powodu pasożytniczych pojemności. Model 3739 oczywiście tego nie potrafi. Aby zmniejszyć wpływ pasożytniczych pojemności okablowania, mierzymy kondensatory 100 pF i 1 nF bezpośrednio w gnieździe na płycie czołowej miernika.

Wyniki pomiarów kondensatorów nieelektrolitycznych są doskonałe. Jednak podczas pomiaru kondensatorów elektrolitycznych, nasz egzemplarz 3730 ma większe odchylenia niż określone  $\pm 4,0\%$ .

**Sygnał pomiarowy podczas pomiaru kondensatorów.** Podczas pomiaru kondensatora 100 nF, zgodnie ze specyfikacją, pomiary są wykonywane za pomocą sygnału sinusoidalnego 150 mV o częstotliwości 1 kHz. Oscylogram przedstawia sygnał pomiarowy dla naszej próbki.



Sygnał pomiarowy podczas pomiaru kondensatora 100 nF  
(© 2024 Jos Verstraten)

**Testowane zakresy cewek indukcyjnych.** W tym celu używamy szeregu cewek w zakresie  $\mu$ H i mH z serii L-07HCP firmy Fastron. Te cewki mają tolerancję  $\pm 10\%$ . Wyniki pomiarów ET4401 są wykorzystywane jako odniesienie. Dokonujemy pomiarów z częstotliwością 1 kHz.

Należy pamiętać, że miernik 3730 ma minimalny zakres pomiarowy 2 mH. Cewki w zakresie  $\mu$ H są oczywiście mierzone bardzo niedokładnie.

Nie mamy żadnych cewek o wartości indukcyjności w zakresie henrów. Nasz ET4401 wskazuje takie wartości, gdy mierzymy uzwojenie pierwotne transformatora sieciowego. Ale wyraźnie widać, że zmierzona wartość jest bardzo zależna od częstotliwości, przy której wykonywane są pomiary. Dlatego ustawiliśmy częstotliwość sygnału pomiarowego w ET4401 na 100 Hz, identyczną jak ta, przy której mierzy również 3730.

| Tabela 3. Dokładność pomiaru cewek |                |                  |
|------------------------------------|----------------|------------------|
| Testowana cewka                    | Wskazania 3730 | Wskazania ET4401 |
| 10 $\mu$ H – $\pm 10\%$            | 0,018 mH       | 10,345 $\mu$ H   |
| 100 $\mu$ H – $\pm 10\%$           | 0,109 mH       | 99,669 $\mu$ H   |
| 1 mH – $\pm 10\%$                  | 1,002 mH       | 0,9876 mH        |
| 10 mH – $\pm 10\%$                 | 10,45 mH       | 10,270 mH        |

- Transformator 1
  - Zmierzone za pomocą 3730: 4,23 H
  - Zmierzone za pomocą ET4401: 4,418 H
- Transformator 2
  - Zmierzone za pomocą 3730: 8,68 H
  - Zmierzone za pomocą ET4401: 9,518 H

**Testowanie napięcia przewodzenia diod.** Dzięki tej dodatkowej funkcji 3730, prąd o natężeniu 1 mA płynie przez podłączony komponent, a urządzenie mierzy napięcie, które powstaje na tym komponencie. Zakres pomiarowy sięga jednak tylko do 1,999 V, więc nie można za jego pomocą testować zielonych i niebieskich diod LED. Aby zbadać dokładność tego pomiaru napięcia, podłączamy do miernika dekadę rezystancji wraz z naszym miernikiem 8842A firmy Fluke. Obracamy pokrętkę na dekadzie rezystancji, aż Fluke odczyta około 0,5 V, 1,0 V, 1,5 V i 2,0 V i rejestrujemy wskazanie na 3730. Wyniki są podsumowane w tabeli 4.

| Tabela 4. Dokładność pomiaru napięcia diody |                  |
|---------------------------------------------|------------------|
| Wskazania 3730                              | Wskazania ET4401 |
| 0,490 V                                     | 0,4983 V         |
| 0,990 V                                     | 1,0077 V         |
| 1,478 V                                     | 1,5056 V         |
| 1,968 V                                     | 2,0056 V         |

## Nasza opinia o mierniku 3730/UT603

Strukturalnie rzecz biorąc, 3730, z dwoma płytkami drukowanymi zawierającymi bardzo skomplikowaną strukturę przełączników, jest urządzeniem, które jest bardzo pomysłowo skonstruowane. Chwała projektantom, którzy na to wpadli! Pojawia się jednak pytanie, jak niezawodne będą te wszystkie styki przełączników w dłuższej perspektywie.

Z pewnością nie możemy nazwać modelu 3730 podążającym za trendami. W tym urządzeniu brakuje wszystkich funkcji, które posiadają nowoczesne mierniki LCR. Są to:

- brak zasilania z akumulatora,
- brak automatycznego przełączania zakresów,
- brak wyświetlacza graficznego,
- brak automatycznego wyłączania,
- brak funkcji „REL”. A jednak możemy żyć z tymi ograniczeniami bardzo dobrze, ponieważ...

Najważniejszymi właściwościami każdego miernika LCR są oczywiście zakres i dokładność. Jeśli chodzi o zakres, niestety brakuje nam kilku zakresów  $\mu$ H. Jeśli chodzi o dokładność, mamy tylko pochwały dla tego urządzenia. Wartości naszych komponentów testowych mierzone przez nasz 3730 pozostają w granicach specyfikacji podanych przez producenta w prawie wszystkich przypadkach (z wyjątkiem kondensatorów elektrolitycznych).

Nasz werdykt na temat 3730: doskonale działający, ale nieco staromodny miernik LCR za rozsądną cenę. Możemy polecić to urządzenie bez zastrzeżeń. ■

Jos Verstraten



## Projektowanie PCB pod produkcję seryjną

Jest całkiem prawdopodobne, że jako Czytelnik pisma poświęconego elektronice myślałeś o seryjnym montażu i sprzedaży jakiegoś urządzenia własnego autorstwa. Jednak z pozycji osoby, która pewnego dnia również dopiero wchodziła w tę branżę, idę o zakład, że w niejednym miejscu tekstu poniżej uda mi się Ciebie, drogi Czytelniku, całkiem mocno zaskoczyć. Na wstępie zaznaczę, że „musnę” tu zaledwie kilka wybranych tematów z zakresu produkcji elektroniki. Zagadnienie jest niezwykle szerokie, więc opowiadać o nim, można by godzinami. Zamiast tego, postaram się zaciekawić Cię kilkoma wcale nieoczywistymi faktami.

### Specyfika projektowania pod linię produkcyjną

Niemal świeżo po studiach podjąłem pracę w firmie projektującej elektronikę pod produkcję seryjną, jak również przyjmującą gotowe projekty firm zewnętrznych celem przygotowania ich do produkcji masowej. Byłem zaskoczony, jak wiele procesów musi się zadziać po drodze. Mając za sobą wiele własnych projektów autorskich, opracowywanych na potrzeby publikacji na łamach „Elektroniki Praktycznej”, jak również będąc osobą mającą kilka miesięcy doświadczenia w branży serwisu telefonów komórkowych, komputerowych płyt głównych i nie tylko, wydawało mi się, że o projektowaniu elektroniki wiem całkiem sporo. I owszem, wiedziałem. Tyle, że wiedza dotyczyła projektowania pod montaż własny, hobbystyczny, stricte manualny, który, jak się okazało, z produkcją seryjną niewiele ma wspólnego.

### Nie tylko hobbysta bywa zaskoczony

Lata pracy w branży pozwalają mi twierdzić, że nie tylko hobbysta bywa zaskoczony. Równie zdziwiony potrafi być klient, zwłaszcza start-upowy, który dopiero stawia pierwsze kroki w branży. Ale nie

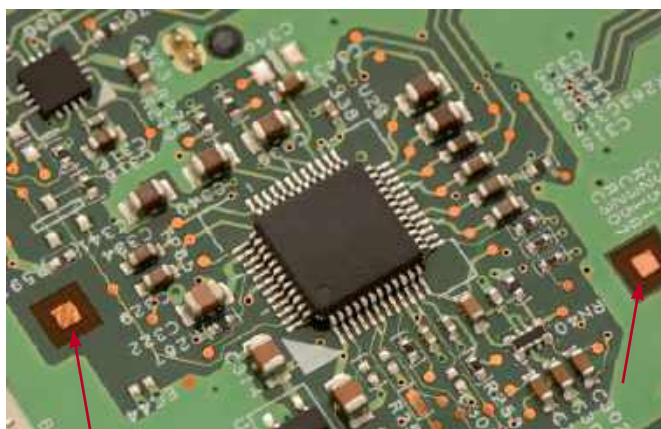
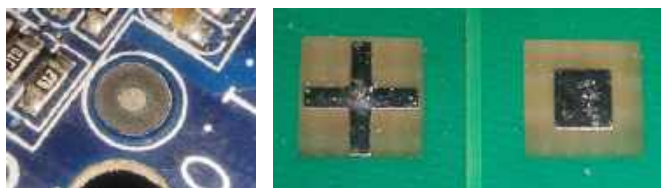
mniej zaskoczeni bywają menedżerowie projektów, święcie przekonani o tym, że skoro klient przychodzi z „gotowym” projektem, to przecież wystarczy wprowadzić dane do maszyn i zacząć go produkować. A tu kłops, tudzież psikus, bo okazuje się, że „gotowego” projektu produkować nie sposób. Albo się da, ale tylko metodami czasochłonnymi, topornymi, manualnymi, z bardzo słabą powtarzalnością i tylko tam, gdzie ludzka siła robocza jest bardzo tania. W innym przypadku, gotowego produktu, z uwagi na koszt produkcji i olbrzymią konkurencję (która robi to lepiej, szybciej, sprawniej i taniej), nie będzie się dało sprzedać.

Poniżej wskażę kilka elementów, których na płytkach PCB projektowanych pod proces automatyczny absolutnie nie może zabraknąć.

**Fiduciale** (fiducial marks) to punkty referencyjne (punkty ustawcze) w postaci znaczników na miedzi, położone najczęściej w dwóch skrajnych, przeciwległych narożnikach laminatu. Maszyny używają fiduciali jako punktów odniesienia. Bez nich nie byłoby możliwe obłożenie płytki komponentami SMD z wystarczającą precyzją, jak również poprawne pozycjonowanie szablonu pasty lutowniczej z płytką PCB, co uniemożliwia precyzyjną aplikację pasty lutowniczej na pola

lutownicze dla komponentów SMD na płytce drukowanej. Fiduciali prawie nigdy nie spotkamy na płytkach PCB zaprojektowanych przez hobbystów. Ci często nawet nie zdają sobie sprawy z ich istnienia, ponieważ w przypadku manualnego montażu punkty ustawcze nie mają najmniejszego znaczenia. Tymczasem, dla montażu automatycznego na linii produkcyjnej są kluczowe.

Znaczniki na warstwach miedzi mogą być realizowane na wiele sposobów.



Punkty ustawcze (fiducial marks)

Standard IPC podpowiada kształt koła o średnicy 1 mm. Niemniej dostawcy usług PCBA (firmy podejmujące się montażu komponentów na PCB) mogą mieć preferowane przez siebie znaczniki, które zagwarantują klientowi najbardziej wydajny i bezawaryjny proces montażu. Aby maszyny mogły z nich skorzystać, obszary te muszą być odsłonięte w warstwie solder maski, wolne od napisów w warstwie opisowej oraz otworów, w tym przelotek. W dodatku ich lokalizacja musi gwarantować ich czytelność dla kamer z nich korzystających, co sąsiedztwo komponentów, zwłaszcza wysokich, lub umiejscowienie

fiduciali na krawędziach transportowych, może skutecznie utrudnić, lub wręcz uniemożliwić. Idealnie byłoby skonsultować ze swoim dostawcą usług PCBA preferowany kształt, lokalizację, zalecane odstępów od krawędzi, otworów, komponentów, solder maski oraz silk screenu na wczesnym etapie projektu płytki PCB, by uniknąć kłopotliwych (a koniecznych) zmian, gdy projekt zostanie uznany (często wyłącznie w przekonaniu klienta) za zakończony. Elementy fiduciala znajdziemy na warstwach miedzi i solder maski. Warto również odpowiednio zaznaczyć ów specyficzny pad na warstwie pasty lutowniczej, gdyż dostawca szablonu musi przygotować ten pad inaczej niż pozostałe (na fiduciale nie nakłada się pasty lutowniczej, jednak są one wykorzystywane do pozycjonowania szablonu względem płytki drukowanej).

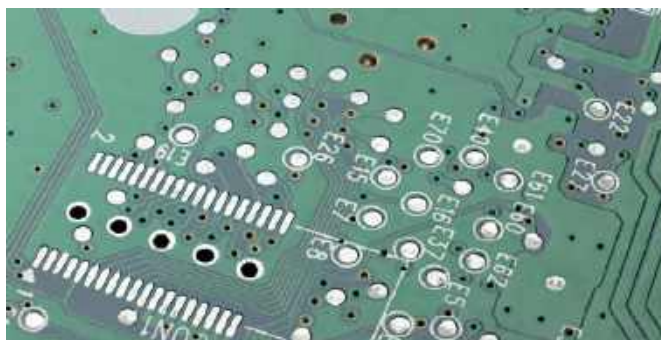
**Otwory pozycjonujące** – kolejny element procesowo niezbędny, którym hobbysta zwyczajowo nie zaprzęta sobie głowy. Otóż płytkę trzeba jakoś unieruchomić w maszynie na czas trwania każdego z procesów montażu. Wiele procesów wykorzystuje sposób polegający na „nasadzeniu” płytki na dwa kołki pozycjonujące. Aby to umożliwić, płytka musi posiadać, dedykowane pod ten proces, niemetalizowane otwory. Muszą być one umiejscowione wewnątrz „ciała” pojedynczej płytki. Ratowanie się próbą implementacji tych otworów, na przykład w panelu łączącym kilka płytek (poza obszarem płytki), jest bezcelowe, ponieważ wiele procesów będzie wymagało zamocowania w maszynie pojedynczej płytki, a więc płytki już po procesie depanelizacji. Dla uzyskania największej stabilności płytki w maszynie, otwory te powinny znajdować się jak najdalej od siebie, podobnie jak w przypadku wspomnianych wcześniej fiduciali, mogą to być skrajne, przeciwległe narożniki płytki. I znów, podobnie jak w przypadku fiduciali, umieszczenie otworów pozycjonujących, będzie ograniczone wieloma wytycznymi wynikającymi z wymagań procesów po stronie dostawcy PCBA. Dlatego, już na wczesnym etapie projektu płytki PCB warto skonsultować ze swoim dostawcą usług PCBA preferowane średnice i kształty otworów pozycjonujących, ich liczbę i lokalizację, zalecane odstępów od krawędzi płytki, od komponentów, od punktów testowych, by uniknąć kłopotliwych zmian, bez których wprowadzenia produkcja urządzenia, albo nie będzie możliwa, albo będzie na tyle droga, że stanie się ona nieopłacalna.

**Poka-yoke** – termin prawdopodobnie nieznany hobbystom i ponownie kluczowy dla produkcji seryjnej. Idea mechanizmu poka-yoke polega na stosowaniu komponentów, których zamontowanie w sposób niewłaściwy (np. niezgodny z polaryzacją) nie będzie fizycznie możliwe. Krótko mówiąc, by pasował do płytki tylko w poprawny sposób. Dotyczy to w zasadzie komponentów montowanych we wszystkich możliwych procesach, ale dla procesu manualnego ma on znaczenie szczególne. Jeśli komponent montowany jest przez człowieka w sposób manualny, nieważne, jak bardzo wyspecjalizowanym będzie on fachowcem. Choćby był arcymistrzem w swojej dziedzinie, z uwagi na żmudny i powtarzalny proces, wcześniej czy później popełni błąd. Nie ma innej możliwości. Chyba, że komponentu nie będzie się dało zamontować w niewłaściwy sposób. I temu właśnie służy idea poka-yoke.

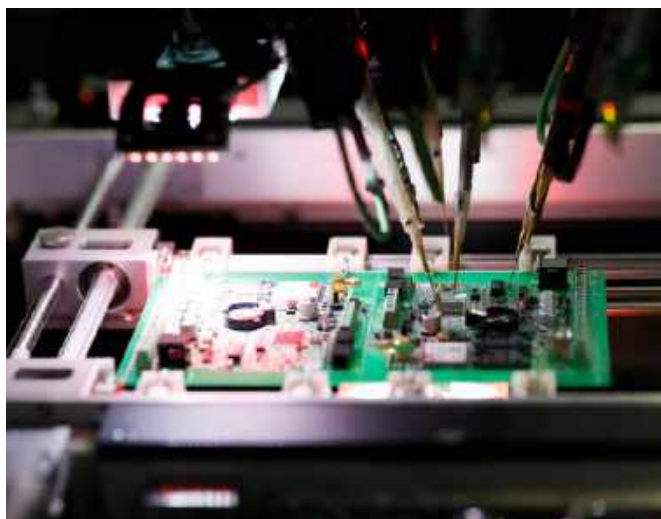
Spora liczba komponentów dostępnych na rynku i przeznaczonych do montażu na PCB jest w taki mechanizm wyposażona, i dotyczy to zarówno komponentów THT (przewlekanych), jak i SMD (montowanych powierzchniowo). Dla wielu komponentów THT, niewyposażonych w ten mechanizm, istnieje zazwyczaj możliwość zaaplikowania go po stronie dostawcy PCBA (np. poprzez odpowiednie uformowanie nóżek). Powoduje to dodatkowy koszt (często wiążący się ze zbudowaniem dedykowanego narzędzia), ale z drugiej strony patrząc, jest dla procesu niezbędne, ponieważ koszt związany z odpadem źle zmontowanych sztuk na linii produkcyjnej byłby, z pewnością, o wiele większy. Świadomy tych tematów klient zastosuje komponent zawierający mechanizm poka-yoke i uniknie kosztu budowy dedykowanego narzędzia po stronie dostawcy PCBA.

**Punkty testowe** – podobnie jak w przypadku wyżej wspomnianych elementów charakterystycznych dla montażu automatycznego, szanse na spotkanie punktów testowych w projektach amatorskich są bardzo niske. Tymczasem na projektach przeznaczonych do produkcji seryjnej, której jednym z etapów są zautomatyzowane testy montażu, mające na celu wyłapać każdą wadliwie zmontowaną sztukę, znajdziemy ich często setki, a nawet tysiące.

Zasada jest tutaj bardzo prosta, każdy net, czyli połączenie pomiędzy dowolnymi dwoma lub więcej komponentami, musi posiadać przynajmniej jeden punkt testowy (net wymagający testów przy większym prądzie, może wymagać kilku punktów testowych na tym samym potencjale). Dzięki punktom testowym, dedykowana aparatura testowa, zbudowana pod dany projekt, jest w stanie wykryć zwarcia, braki połączeń, jak również obecność lub brak, a także wartości



Punkty testowe (test points)

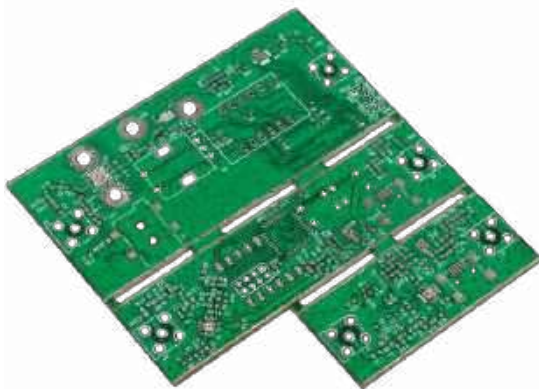


Płytki PCB podczas testów na flying probe stosowanym do partii samplowych (docelowe testery do produkcji wielkoseryjnej wyglądają i działają inaczej)

i właściwości niemal każdego z elementów, które powinny być zostać zamontowane na płytce PCB. Niemal, ponieważ trudno byłoby wykryć, np. dwa połączone ze sobą równolegle kondensatory: 4700  $\mu$ F oraz 100 nF (tego typu potencjalne nieprawidłowości daje się często wykrywać innymi metodami, na przykład za pomocą kamer). Przed przystąpieniem do projektowania PCB, lub na bardzo wczesnym etapie, konieczne skonsultuj ze swoim dostawcą PCBA optymalną (a jeśli Twój projekt jest naprawdę miniaturowy, to również minimalną akceptowalną dla możliwości jego procesów) średnicę punktów testowych oraz dozwolone odstęp między środkami punktów testowych. Zapytaj, jakie muszą być odstęp między krawędzią punktu testowego i otaczających go komponentów, miej też na uwadze, że ta wartość ma prawo być różna, w zależności od wysokości tych komponentów. Jeśli myślisz o użyciu przelotek w roli punktów testowych (otwór w punkcie testowym), zapytaj, czy igły trafiające (lub nie) w środek otworów nie skomplikują, albo też nie uniemożliwią optymalnego testowania produktu. Pamiętaj, że punkty testowe muszą być wolne od solder maski oraz napisów czy innych elementów graficznych na warstwie opisowej (silk screen), które uniemożliwią skuteczny kontakt elektryczny punktu z igłą testera. Jeśli projekt zawiera zwory drutowe, punkt testowy powinien znajdować się na obu jej krańcach, by umożliwić stwierdzenie jej obecności (lub braku). Punkty testowe muszą być zawarte na schemacie, i podobnie jak każdy inny komponent, muszą posiadać indywidualne desygatory (komponenty z jednym wyprowadzeniem, podłączone z odpowiednimi netami, czyli wpięte między odpowiednie komponenty). Muszą być również zawarte, podobnie jak wszystkie inne komponenty, w netliście, by dostawca PCBA był w stanie ich użyć i na ich podstawie wygenerować odpowiednie programy testowe. W miarę możliwości punkty testowe warto zaimplementować po jednej stronie płytki PCB, na przykład po stronie lutowania wyprowadzeń komponentów przewlekanych, często wolnej od innych komponentów potencjalnie narażonych na nieplanowany kontakt z igłami testowymi, bądź innymi elementami testera. Zauważ, że brak punktów testowych uniemożliwi testowanie wyprodukowanych jednostek, a tym samym w zasadzie uniemożliwi montaż produktu. Tymczasem wprowadzenie ich do projektu na zbyt późnym etapie może okazać się kłopotliwe albo wręcz niemożliwe i wiązać się z koniecznością ponownego zaprojektowania płytki drukowanej niemal od zera.

**Obszary do bezpiecznej depanelizacji.** Projektując płytkę drukowaną, mamy na ekranie komputera z uruchomionym dowolnym środowiskiem EDA (Electronic Design Automation), pojedynczą sztukę takowej i z reguły mało prawdopodobne, że na optymalnym ku temu etapie pomyślimy też o tzw. obszarach zabronionych dla komponentów, ścieżek i solder maski. Tymczasem, do maszyn obkładających i lutujących komponenty, najpewniej wkładane będą, nie pojedyncze sztuki tych płytek a formatki zawierające większą ich liczbę.

Te płytki muszą być ze sobą w jakiś sposób wzajemnie połączone, ale też przygotowane do bezpiecznej depanelizacji. Obszary, które będą rozfrezowywane (milling) lub wręcz całe krawędzie w przypadku procesów rozłamywania (V-Cut lub Score Line) muszą zapewniać bezpieczne odległości dla komponentów i miedzi (w tym ścieżek), ale też solder maski, by na skutek naprężeń i niezerowych tolerancji charakteryzujących te procesy nie zostały one uszkodzone podczas rozłamywania lub rozfrezowywania. Potrafi o tym zapominać nawet najbardziej doświadczony projektant, o hobbystach nie wspominając. Ten doświadczony przypomina sobie o tym często pod koniec tworzenia projektu, lub też dopiero gdy na fakt braku możliwości zapewnienia bezpiecznej depanelizacji zwróci mu uwagę dostawca usług PCBA. Wtedy też niejednokrotnie na ostatni moment, zaczyna się gorączkowe poszukiwanie obszarów, w których ścieżki i komponenty udałoby się odsunąć



Panel zawierający kilka płytek, wymagających obszarów wolnych od komponentów, ścieżek, solder maski, co jest niezbędne do bezpiecznego ich rozfrezowania

kilka milimetrów od krawędzi płytki, co w końcowej fazie gęsto upakowanego komponentami projektu potrafi być nie lada wyzwaniem. Nawet jeśli płytka jest na tyle duża, że trafi na maszyny pojedynczo, często okazuje się, że osoba daną płytkę projektująca nie przewidziała, że płytka potrzebuje wolnych od komponentów obszarów (wzdłuż dwóch dłuższych krawędzi) na tzw. transport, które w takim przypadku trzeba będzie dodać do projektu płytki po stronie PCBA (zaprojektować panel dodający do płytki krawędzie transportowe) a następnie (w bezpieczny sposób) je odfrezować (lub odłamać – w zależności od preferencji i sugestii jakościowych dostawcy). Pamiętając o zapewnieniu dwóch dłuższych krawędzi płytki wolnych na szerokości kilku milimetrów od komponentów, unikniemy kosztów konieczności powiększenia jej o dodatkowe obszary, które i tak będą musiały (na końcu procesu wymagających transportu) zostać odfrezowane (bądź odłamane).

**Projekt ramy lutowniczej.** Temat związany z lutowaniem komponentów THT na fali (również na fali selektywnej), który potrafi klienta bardzo zaskoczyć, jeśli po stronie lutowania wyprowadzeń THT znajdują się również komponenty SMD. Oczywiście sytuacja nie jest rzadka i o ile komponenty SMD są zorganizowane w grupy, które na czas lutowania na fali daje się zakryć specjalnie zaprojektowaną ramą lutowniczą problemu nie ma. Problem występuje wtedy, gdy komponenty SMD znajdują się zbyt blisko padów THT. Wówczas, z uwagi na grubość ścianki ramy lutowniczej, może się okazać, że nie da się tą ramą przykryć komponentu SMD, jednocześnie zostawiając odsłonięty do polutowania pad komponentu THT. Innymi słowy, ścianka ramy lutowniczej ma niezerową grubość, i może zwyczajnie nie zmieścić się pomiędzy pad komponentu THT oraz pad komponentu SMD. To uniemożliwi montaż jakiegokolwiek automatyczną metodą, znacznie podrażając proces, potencjalnie windując koszt procesu do granic opłacalności produkcji. Warto o tym pamiętać i odpowiednio wcześniej



Płytkę PCB z dodatkowymi marginesami, dodanymi do projektu przez dostawcę usług PCBA z uwagi na komponenty SMD ustawione zbyt blisko krawędzi transportowych. Klient (z różnych względów) mógł działać świadomie, albo też nie wziął pod uwagę wymagań transportowych dostawcy. Niezależnie od przyczyny, użyta powierzchnia laminatu zwiększyła się, a koszt PCB w przeliczeniu na jedną sztukę płytki wzrósł

skonsultować temat trzymania koniecznych odległości komponentów SMD od padów THT bo zmiany na tym polu zwłaszcza w ciasno upakowanym komponentami projekcie, mogą być bardzo kłopotliwe, bez przeprojektowania połowy płytki.

**Dziesiątki innych tematów.** Iluż z nas na etapie projektowania PCB zastanawiało się do tej pory nad tym, od której strony w daną lokalizację podjedzie głowica maszyny obkładającej komponenty radialne trzymająca w danej chwili, dajmy na to, średniej wielkości kondensator elektrolityczny? I czy nie będzie miała po drodze kolizji z innymi komponentami? A może z powodu niezbyt szczęśliwego ułożenia sąsiadujących ze sobą komponentów, w daną lokalizację w ogóle nie da się dotrzeć? Oglądnięcie kilku filmów pokazujących automatyczny montaż komponentów radialnych ma prawo dać sporo do myślenia (<https://youtu.be/1YGdj1jnDOY>, <https://youtu.be/Z314HXYjiBo>, <https://youtu.be/Dmry2cFUClk>, <https://youtu.be/aeMh00BjX7k>).

A może parametry przygodnie zastosowanego komponentu wykluczą jego maszynowy montaż? Może komponent będzie miał zbyt grube albo zbyt twarde wyprowadzenia, z którymi maszyna sobie nie poradzi? A może będzie miał nietypowy rozstaw nóżek lub też będzie zwyczajnie za wysoki względem tego, co dopuszcza specyfikacja takiej czy innej maszyny? A może wręcz przeciwnie, komponent jest idealny, ale Ty zastosowałeś na płytce PCB średnice otworów, które idealnie sprawdzą się przy montażu manualnym, ale automatycznym już niekoniecznie, i trzeba je zwyczajnie powiększyć zgodnie z wymaganiami (specyfikacją) maszyny?

Tematy, jak widać, nie robią się ani prostsze, ani mniej liczne. Ich liczba urasta w tysiące możliwych a każdy drobiazg może mieć olbrzymie przełożenie na koszt finalnego procesu montażu.

Co gorsza, czyniąc własny design, nawet posiadając niemałą wiedzę w temacie, nie wszystko jesteś w stanie przewidzieć. Nie znasz parku maszynowego przyszłego dostawcy PCBA. Nie bardzo masz wgląd w to, jakiego typu maszynami czy modelami maszyn dysponuje. Nie znasz jego procesów. I choćbyś pękł, nie przewidzisz wszystkich jego możliwości i ograniczeń, tym bardziej że wiele z nich, zupełnie się tego nie spodziewając i nie planując, wygenerujesz własnym designem.

**DFM (Design For Manufacturing)** – wybrzmiewa jako projektowanie pod proces produkcyjny. W przypadku klienta usług PCBA/EMS zaprojektowanie urządzenia i odpowiedzialność za design, leży oczywiście po stronie klienta. Tylko jak tu (będąc klientem) projektować pod procesy produkcyjne dziejące się po stronie konkretnego dostawcy, do których bezpośredni dostęp, jak i wiedzę na ich temat, mamy mocno ograniczone?

Z uwagi na powyższe, w praktyce DFM to płaszczyzna spotkania klienta i dostawcy przy wspólnym stole i we wspólnym biznesie. Obojgu bowiem zależy, by produkcja szła szybko i bezproblemowo a jakość montażu była możliwie najwyższa. Klient chce dobrej ceny, a dostawca chce mu taką cenę zapewnić. Dostawca walczy o zadowolonego klienta, a klient walczy o dobrego dostawcę.

Kompetentny dostawca wnikliwie przeanalizuje projekt klienta pod kątem doboru optymalnych procesów, z uwzględnieniem wszystkich możliwości (ale też ograniczeń) parku maszynowego, którym dysponuje. W odpowiedzi poinstruuje klienta wyczerpującym raportem, w jaki sposób ten powinien zmodyfikować projekt, by zapewnić, dla swego wyrobu, efektywny, zarówno cenowo i jakościowo proces montażu.

Finalnie, klient wprowadza modyfikacje w sposób mniej lub bardziej efektywny, tym samym (mniej lub bardziej) pomagając dostawcy zrealizować ich wspólny cel. Dobrze układająca się współpraca przekłada się finalnie na fantastycznie, szybko i bezawaryjnie produkujący się wyrób, ku wzajemnej satysfakcji obojga. ■

Mariusz Ciszewski



**Przetwornik analogowo-cyfrowy (ADC) służy do konwersji sygnału analogowego na kod binarny, którego wartość chwilowa odpowiada chwilowej wartości przetwarzanego sygnału analogowego.**

## Ogólna zasada działania ADC

Digitalizacja lub kwantyzacja to konwersja napięcia analogowego na kod cyfrowy. Wydawać by się mogło, że proces ten jest bezproblemowy: podajemy napięcie analogowe na wejście przetwornika ADC, a na wyjściach tworzone są kody cyfrowe, które reprezentują chwilową wielkość napięcia analogowego. W rzeczywistości jest to dużo bardziej skomplikowane! Podczas kwantyzacji sygnałów analogowych należy wziąć pod uwagę pewne prawa, które omówiono w kolejnych sekcjach.

**Próbkowanie.** Po pierwsze: większość przetworników ADC nie działa w sposób ciągły. Rzeczywiście, gdyby kod cyfrowy na wyjściu stale dostosowywał się do zmian amplitudy analogowego sygnału wejściowego, na wyjściu konwertera panowałby prawdziwy chaos kodowy. W rzeczywistości konwersja napięcia analogowego na kod cyfrowy nie jest procesem ciągłym. Czas konwersji jest rzędu ns nawet dla najszybszych konwerterów. Nowoczesny szybki przetwornik ADC ma częstotliwość próbkowania 50 MSa/s (megapróbek na sekundę). Oznacza to, że taki układ pobiera 50 000 000 próbek sygnału wejściowego na sekundę. Jeden bit wyjściowy będzie reagował nieco szybciej niż drugi. W przypadku ciągłej konwersji, kod cyfrowy zostałby nieumyślnie odczytany w momencie, gdy jeden lub więcej bitów jest w trakcie dostosowywania się do nowej sytuacji na wejściu analogowym. W tym momencie kod cyfrowy może zawierać całkowicie błędne informacje.

Przetworniki ADC bez problemu osiągają częstotliwość próbkowania liczoną w Gsa/s (miliardów próbek na sekundę), przykładem może być układ HMCAD1511 stosowany w wielu oscyloskopach cyfrowych, na przykład w SDS1104X-U firmy Siglent. W przeszłości szybkie przetworniki ADC miały zwyczajnie bardzo wysokie ceny, a ich stosowanie wymagało równie szybkich (i drogich) układów cyfrowych i mikrokontrolerów. Przep. tłum.

**Sample and hold.** Aby uniknąć tego problemu, należy próbować sygnał analogowy w stałych odstępach czasu. Oznacza to, że należy pobierać próbkę chwilowej wartości sygnału wejściowego w regularnych odstępach czasu, przechowywać tę próbkę przez krótki czas w pamięci analogowej, a następnie konwertować tę próbkę na odpowiedni kod cyfrowy. Po upływie czasu konwersji przetwornika ADC i upewnieniu się, że kod cyfrowy dostosował się do nowej wartości próbki, można odczytać kod cyfrowy i dalej go przetwarzać. Następnie można pobrać kolejną próbkę i powtórzyć proces.

Pamięć analogowa, w której próbka analogowego sygnału wejściowego jest krótko przechowywana, jest zawsze wykonywana w formie „sample and hold”, w skrócie S&H. Zasadniczo, chwilowa wielkość sygnału wejściowego jest przechowywana w bardzo małym kondensatorze zintegrowanym z układem ADC w większości przypadków.

Na poniższym rysunku przedstawiono schemat blokowy takiego układu próbkująco-pamiętającego.

**Częstotliwość próbkowania.** Tak więc digitalizacja napięcia analogowego nie jest procesem ciągłym, ale procesem kontrolowanym przez zewnętrzny impuls zegarowy. Częstotliwość tego sygnału zegarowego nazywana jest częstotliwością próbkowania. Częstotliwość ta określa zatem liczbę próbek pobieranych z analogowego napięcia wejściowego na sekundę. Wielkość ta w języku polskim również jest nazywana częstotliwością próbkowania i wyrażana jest w Sa/s, próbkach na sekundę. Gdy mowa o przetwarzaniu dźwięku, zwykle podaje się tę częstotliwość próbkowania, wyrażaną w kHz. W innych zastosowaniach podaje się częstotliwość w Sa/s. Obecnie nie ma problemów z dostępem do przetworników ADC próbkujących z częstotliwością liczoną w MSa/s (milionach próbek na sekundę), a i układy osiągające wartości w GSa/s są dość łatwo dostępne.

**Minimalna wymagana częstotliwość próbkowania.** Ważną kwestią jest to, ile próbek należy pobierać na sekundę, aby prawidłowo próbować analogowy sygnał wejściowy. Należy pamiętać, że te cyfrowe próbki prawdopodobnie zostaną pewnego dnia przekonwertowane z powrotem na sygnał analogowy. W końcu, w wyniku powolnej ewolucji, my, niedoskonalni ludzie, wciąż mamy zmysły działające wyłącznie analogowo. Należy zatem upewnić się, że odzyskany sygnał analogowy jest jak najbardziej zbliżony do oryginalnego sygnału.

Oczywiście im więcej próbek, tym lepsza jakość reprodukcji. W związku z tym interesuje nas tylko pytanie, jak mało próbek należy pobrać, aby zapewnić, że kody cyfrowe mogą być później przekształcone w przetworniku cyfrowo-analogowym na sygnał analogowy podobny do sygnału analogowego na wejściu przetwornika ADC. Poniższy rysunek ponownie graficznie podsumowuje proces ADC+DAC.

Dziesięć próbek jest pobieranych z jednego okresu sygnału sinusoidalnego.

W tym przykładzie można więc powiedzieć, że częstotliwość próbkowania jest dziesięć razy większa niż częstotliwość sygnału. Na dolnym wykresie te dziesięć próbek cyfrowych jest reprezentowanych przez ich odzyskane wartości analogowe na wyjściu przetwornika cyfrowo-analogowego. Oryginalny okres fali sinusoidalnej jest przybliżony przez dziesięć kolejnych napięć krokowych.

W narysowanym przykładzie można rozpoznać kształt oryginalnego sygnału z tej krokowej aproksymacji bez wysilania wyobraźni. Nie można jednak bezgranicznie ograniczać częstotliwości próbkowania.

**Twierdzenie o próbkowaniu.** Matematycznie można wykazać, że częstotliwość próbkowania musi być co najmniej dwa razy większa niż najwyższa częstotliwość sygnału analogowego. To ogólne prawo znane jest jako „twierdzenie o próbkowaniu”. W przypadku próbkowania z niższą częstotliwością, absolutnie niemożliwe jest odzyskanie kształtu oryginalnego sygnału analogowego z kolejnych próbek cyfrowych.

Efekt ten można bardzo ładnie zademonstrować graficznie. Na poniższym rysunku sinusoidalne napięcie analogowe o częstotliwości 15 kHz jest próbkowane z coraz niższymi częstotliwościami. Jeśli następnie przekonwertujesz kolejne kody cyfrowe na napięcie analogowe za pomocą przetwornika cyfrowo-analogowego, otrzymasz sygnał, który powinien być krokowym przybliżeniem sygnału sinusoidalnego. Jest to nadal całkiem poprawne w przypadku próbkowania z częstotliwością 7,5 kHz. Jednak gdy częstotliwość próbkowania staje się coraz niższa, odzyskany sygnał analogowy staje się coraz mniej podobny do oryginalnej fali sinusoidalnej. Przy bardzo niskich częstotliwościach odzyskany sygnał nie jest już nawet rozpoznawalny jako sinusoida, ale staje się trójkątem!

Ta znacznie niższa częstotliwość nazywana jest „częstotliwością aliasu”, a duże zniekształcenie powstałe w tym przypadku „zniekształceniem aliasu”.

**Filtr antyaliasingowy jest absolutnie niezbędny.** Zniekształcenia aliasowe są zawsze obecne w przypadku próbkowania z częstotliwością niższą niż dwukrotność najwyższej częstotliwości sygnału analogowego. Jeśli konieczne jest próbkowanie sygnału wejściowego o znanej stałej częstotliwości, można uniknąć tego zniekształcenia aliasowego, wybierając częstotliwość próbkowania co najmniej dwa razy wyższą.

Jeśli sygnał wejściowy ma nieznaną szerokość pasma, należy sztucznie zapewnić, że pasmo częstotliwości analogowego sygnału wejściowego nigdy nie przekroczy znanej wartości. Następnie można ustawić częstotliwość próbkowania na dwukrotność tej znanej wartości.

Ograniczenie szerokości pasma sygnału analogowego można wykonać za pomocą filtra dolnoprzepustowego o bardzo stromym zboczach. Filtr ten jest zdefiniowany przez określoną częstotliwość odcięcia  $f_0$ . Filtr przepuszcza wszystkie sygnały o częstotliwości mniejszej niż częstotliwość odcięcia i tłumi wszystkie sygnały o częstotliwości wyższej niż częstotliwość odcięcia. Ten filtr dolnoprzepustowy nazywany jest w terminologii ADC/DAC „filtrem antyaliasingowym”. Jest to logiczna nazwa, ponieważ dzięki temu filtrowi nie występują już zniekształcenia aliasowe. Wystarczy, aby częstotliwość próbkowania była równa  $2 \times f_0$  filtra.

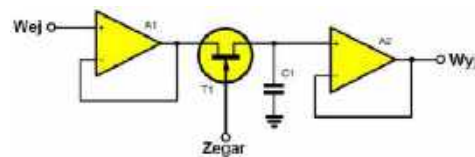
**Schemat blokowy systemu ADC.** Schemat blokowy użytecznego przetwornika analogowo-cyfrowego przedstawiono na poniższym rysunku. Sygnał analogowy, który ma zostać zdigitalizowany, jest najpierw przepuszczany przez filtr antyaliasingowy, a następnie trafia do układu S&H. Wyjście tego obwodu trafia do wejścia analogowego przetwornika ADC.

## Specyfikacja i właściwości przetworników ADC

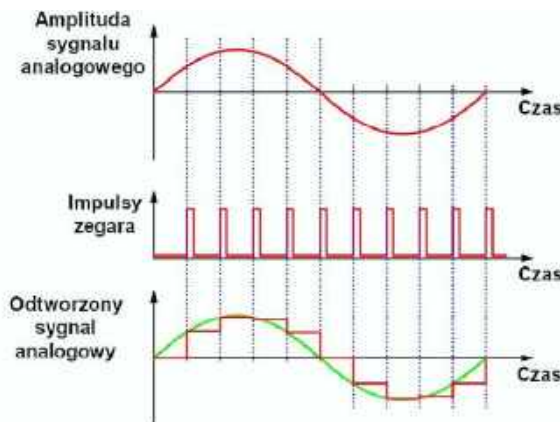
**Długość słowa (liczba bitów).** Długość słowa określa liczbę bitów używanych do konwersji sygnału analogowego na kod cyfrowy. Jeśli pracujesz z systemem ośmiobitowym (powszechnie używany standard), mówisz o 8-bitowej długości słowa.

Praktycznie wszystkie mikrokontrolery używane przez hobbystów posiadają przetworniki ADC 10-bitowe, niektóre zaś nawet 12-bitowe. Na rynku nie brakuje też przetworników ADC 16-bitowych i większych. Układy do konwersji dźwięku występujące w komputerach, smartfonach czy słuchawkach bezprzewodowych operują na słowach 24-bitowych. Warto pamiętać, iż nie trzeba ograniczać wyboru przetwornika ADC do ośmiu bitów, bo tyle ma mikrokontroler. Przyp. tłum.

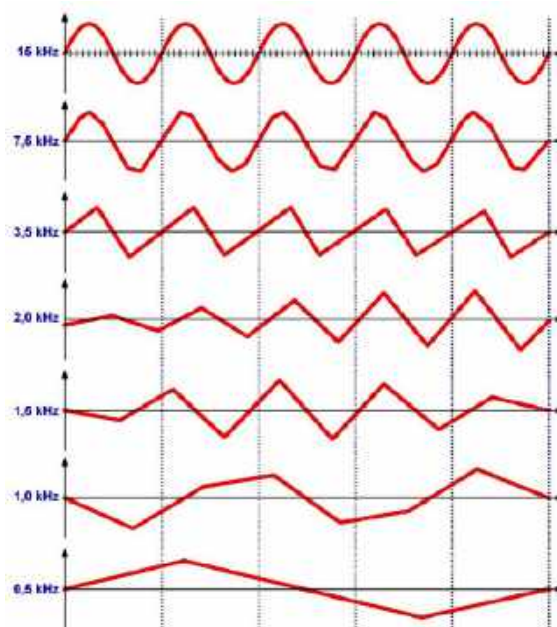
**Waga bitów oraz MSB i LSB.** W poniższym przykładzie rysowane jest rosnące napięcie analogowe piłokształtne o maksymalnej wartości 15 V. Jest ono kwantowane przy użyciu czterobitowego przetwornika ADC. Obwód ten zapewnia standardowy kod binarny. Częstotliwość taktowania jest równa szesnastokrotności częstotliwości sygnału.



1. Zasada działania obwodu sample and hold (© 2021 Jos Verstraten)



2. Graficzne podsumowanie procesu konwersji ADC+DAC (© 2021 Jos Verstraten)



3. Wygląd odtworzonego sygnału analogowego w funkcji czasu dla różnych częstotliwości próbkowania (© 2021 Jos Verstraten)

Tabela pokazuje stan czterech bitów dla każdej próbki. Bit Qa zmienia stan dla każdego wzrostu napięcia o 1 V, bit Qb zmienia stan dla każdego wzrostu o 2 V, Qc dla 4 V, a Qd dla 8 V. Możliwe jest zatem przypisanie określonej „wagi” do każdego bitu. Bit Qa jest „najlżejszy”, a bit Qd „najcięższy”. W końcu zmiana wartości Qa odpowiada jedynie zmianie napięcia analogowego o 1 V. Natomiast zmiana wartości Qd odpowiada analogowej zmianie wartości o 8 V.

Stąd bit Qa nazywany jest „najmniej znaczącym bitem”, w skrócie „LSB”, a bit Qd „najbardziej znaczącym bitem”, w skrócie „MSB”.

**Rozdzielczość.** Rozdzielczość przetwornika ADC wskazuje, w ilu obszarach kwantyzacji można umieścić sygnał analogowy po digitalizacji. Przetwornik jednobitowy może przetworzyć sygnał wejściowy na dwie wartości. Rozdzielczość można obliczyć podnosząc liczbę 2 do potęgi liczby bitów. Tak więc system o długości słowa 16 bitów ma rozdzielczość  $2^{16}=65\,536$ . W takim systemie można podzielić chwilową wartość napięcia analogowego na 65 536 obszarów kwantyzacji.

**Zakres dynamiczny.** Zakres dynamiki reprezentuje logarytmiczny stosunek między minimalną i maksymalną zmianą sygnału, jaką system ADC/DAC może spowodować w sygnale analogowym. Maksymalna zmiana sygnału występuje, gdy wszystkie bity nagle przełączają się ze stanu niskiego na wysoki.

Minimalna zmiana sygnału występuje, gdy tylko najmniej znaczący bit przechodzi z jednego stanu logicznego do drugiego.

Zakres dynamiki wzrasta wraz z liczbą używanych bitów. Im większa długość słowa, tym większy zakres dynamiki. Zakres dynamiki jest wyrażany w dB i jest czasami określany jako stosunek sygnału do szumu przetwornika ADC.

**Rozmiar bitu.** Rozmiar bitu wskazuje, o ile musi spaść lub wzrosnąć analogowy sygnał wejściowy, aby najmniej znaczący bit przetwornika ADC zmienił swój poziom logiczny. Tak więc wielkość bitowa i zakres dynamiczny to w rzeczywistości dwie wielkości opisujące to samo zjawisko fizyczne. Podczas gdy zakres dynamiki reprezentuje stosunek i dlatego jest niezależny od maksymalnej wielkości napięcia, określenie rozmiaru bitu ma sens tylko wtedy, gdy znana jest maksymalna amplituda, którą system musi przetworzyć.

Zależność między długością słowa, rozdzielczością, zakresem dynamicznym i rozmiarem bitów

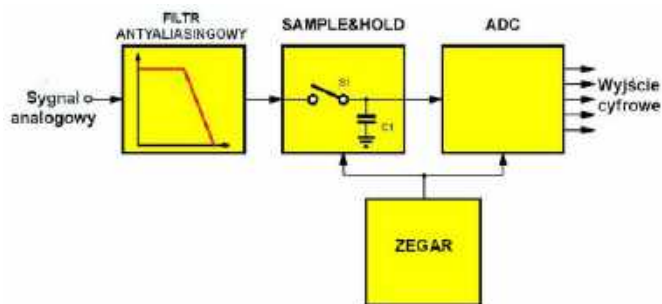
Istnieje matematyczna zależność między czterema omawianymi wielkościami przetwornika ADC. Zależność ta została podsumowana dla niektórych rozdzielczości w tabeli 1.

**Błąd przesunięcia (offsetu).** Jeśli do przetwornika ADC zostanie przyłożone napięcie wynoszące dokładnie 0 V, wszystkie bity powinny mieć stan niski. Czasami jednak kilka najmniej znaczących bitów może przybrać stan wysoki. Nazywa się to błędem przesunięcia (offsetu) przetwornika ADC.

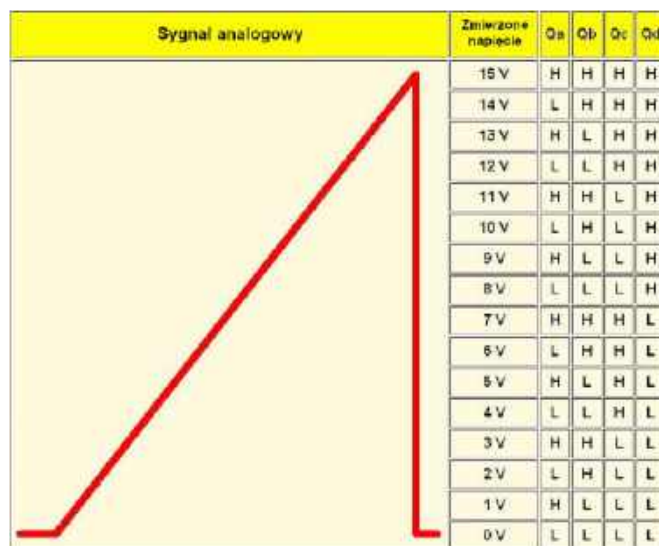
Błąd offsetu może być dodatni lub ujemny, tj. dla wartości napięcia odpowiadającego 1 LSB kod binarny może przyjąć wartości zarówno 000, jak i 010 zamiast poprawnej wartości 001. Przyp. tłum.

**Nieliniowość różnicowa, DNL, błąd histerezy.** Błąd histerezy (właśc. Nieliniowość różnicowa, DNL) jest typowym błędem przetworników ADC. Dla idealnego ADC zmiana wartości napięcia wejściowego o wartość LSB powinna prowadzić do takiej samej zmiany kodu cyfrowego, dla dowolnej wartości napięcia w zakresie pracy przetwornika. W rzeczywistości dla niektórych wartości napięcia do zmiany kodu potrzebna jest zmiana napięcia o wartość większą lub mniejszą od wartości napięcia dla 1 LSB. Błąd DNL określa, jak duża może być różnica między idealnym ADC, a rzeczywistym układem. Co więcej, w niektórych przypadkach błąd ten może zależeć też od tego, czy w kolejnych pomiarach napięcie wejściowe wzrastało, czy opadało. Błąd nieliniowości różnicowej może przybierać różne wartości dla poszczególnych kodów.

**Nieliniowość całkowita, INL.** Jak sama nazwa wskazuje, nieliniowość całkowita przetwornika ADC określa zakres, w jakim charakterystyka wyjściowa odbiega od linii prostej. Zjawisko to zostało wyjaśnione graficznie na poniższym rysunku. Jeśli do przetwornika ADC zostanie przyłożone napięcie piłokształtne, układ scalony musi zwiększyć lub zmniejszyć wartość swojego kodu wyjściowego o jeden bit dla każdej zmiany napięcia o jeden bit wielkości  $\Delta U$ . Graficznie przedstawia to lewy wykres na poniższym rysunku. Zwiększenie wartości kodu o jeden bit jest reprezentowane przez skokową aproksymację napięcia piłokształtnego. Na osi pionowej



4. Schemat blokowy przetwornika ADC (© 2021 Jos Verstraten)



5. Wyjaśnienie terminów „LSB” (Qa) i „MSB” (Qd) (© 2021 Jos Verstraten)

**Tabela 1. Zależność między długością słowa, rozdzielczością, zakresem dynamicznym i rozmiarem bitu**

| Liczba bitów | Rozdzielczość | Zakres dynamiki | Napięcie LSB dla maksymalnej amplitudy 10 V |
|--------------|---------------|-----------------|---------------------------------------------|
| 1            | 2             | 6 dB            | 5 V                                         |
| 2            | 4             | 12 dB           | 2,5 V                                       |
| 4            | 16            | 24 dB           | 625 mV                                      |
| 8            | 256           | 48 dB           | 39,1 mV                                     |
| 10           | 1024          | 60 dB           | 7,65 mV                                     |
| 12           | 4096          | 72 dB           | 2,4 mV                                      |
| 16           | 65536         | 96 dB           | 0,152 mV                                    |

przedstawione są kolejne wielkości bitów  $\Delta U$ , przy których wartość wzrasta lub maleje o jeden krok. W tym idealnym przetworniku ADC każde przejście kodu wynika z identycznego  $\Delta U$ .

Wykres po prawej stronie przedstawia (wyolbrzymione) działanie nieliniowego przetwornika ADC. Teraz do wygenerowania jednego przejścia kodu potrzebne są różne wartości  $\Delta U$ .

Błąd INL można rozpatrywać jako sumę wszystkich błędów DNL. Przyp. tłum.

**Błąd niemonotoniczny.** Przetwornik ADC z błędem niemonotonicznym, jak pokazano na poniższym rysunku, będzie miał krok opadający zamiast kroku rosnącego w jednym lub kilku miejscach swojej charakterystyki wyjściowej. To niemonotoniczne zachowanie jest bezpośrednią konsekwencją systemu o dużej nieliniowości. Dopóki nieliniowość jest mniejsza niż  $\frac{1}{2}$  LSB, obwód ma zagwarantowane monotoniczne działanie. Jeśli nieliniowość jest większa niż  $\frac{1}{2}$  LSB, układ może być niemonotoniczny, ale nie jest to pewne. Niemonotoniczny przetwornik ADC może prowadzić do poważnych błędów systemu w niektórych zastosowaniach.

Problem ten występuje rzadko, i zazwyczaj dotyczy starszych układów ADC lub specyficznych sytuacji. Przyp. tłum.

**Błąd wzmocnienia.** Jeśli na wejście rzeczywistego przetwornika ADC podamy takie napięcie, by na wyjściu wszystkie bity miały stan wysoki, uwzględniając wcześniej błąd offsetu, to różnica między tym napięciem, a napięciem potrzebnym w przypadku idealnego przetwornika ADC jest właśnie błędem wzmocnienia. Wartość tego błędu podawana jest w liczbie bitów LSB. Dla sygnału o wartości połowy maksymalnego napięcia sygnału błąd wzmocnienia też przybiera połowę wartości.

**Błąd pełnoskalowy.** Jest to suma błędów offsetu i wzmocnienia. Wartość tego błędu podawana jest w voltach lub w liczbie bitów LSB.

**Źródło napięcia odniesienia.** Niektóre przetworniki ADC mają zintegrowane źródło napięcia odniesienia, inne zaś wymagają zewnętrznego źródła, a jeszcze inne używają napięcia zasilania jako źródła odniesienia. Wartość napięcia odniesienia określa maksymalną amplitudę mierzonego sygnału oraz wartość napięcia dla najmniej znaczącego bitu. Typowe wartości napięć odniesienia to 1,024 V, 2,048 V i 4,096 V. Dla przykładu 12-bitowy przetwornik ADC ze źródłem odniesienia 4,096 V będzie mierzył napięcie z krokiem 1 mV.

Niezależnie od tego, czy źródło napięcia odniesienia jest częścią przetwornika, czy też jest oddzielnym komponentem, jego parametry mają wpływ na dokładność pomiaru. Parametrami tymi są:

- wartość napięcia odniesienia;
- wstępna dokładność, określająca o ile rzeczywista wartość napięcia odniesienia może różnić się od wartości idealnej w przypadku konkretnego układu;
- współczynnik temperaturowy, określający o ile napięcie źródła odniesienia będzie się zmieniało wraz ze zmianą temperatury;
- poziom szumów własnych, określa amplitudę szumów na wyjściu układu, im niższy tym lepiej;
- regulacja linii, określa o ile napięcie wyjściowe zmieni się wraz ze zmianą napięcia zasilania;
- regulacja obciążenia, określa o ile napięcie wyjściowe zmieni się wraz ze wzrostem poboru prądu ze źródła napięcia odniesienia;
- maksymalny prąd obciążenia;
- PSRR, współczynnik odrzucenia napięcia zasilania, określa stosunek poziomu szumów na zasilaniu źródła napięcia odniesienia względem poziomu szumów na wyjściu, wyrażany w decybelach;
- stabilność długoterminowa, określa o ile napięcie wyjściowe może się zmienić po upływie określonego czasu pracy, zazwyczaj 1000 godzin lub jednego roku.

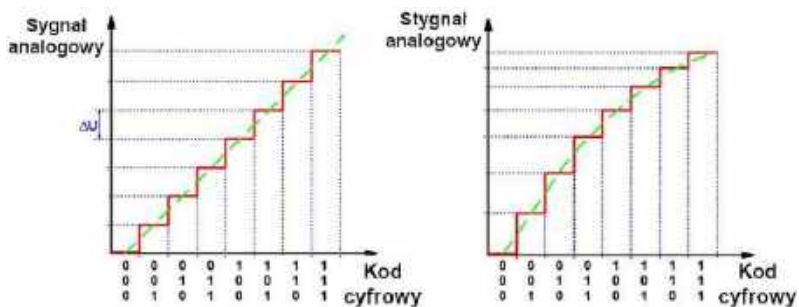
Jakość źródła napięcia odniesienia ma szczególne znaczenie, gdy chcemy mierzyć bezwzględne wartości mierzonego sygnału z dużą precyzją. Jednakże bardzo dobre źródło napięcia odniesienia nie pomoże, jeśli sam przetwornik ADC ma kiepskie parametry.

## Dodatkowe możliwości i parametry przetworników ADC

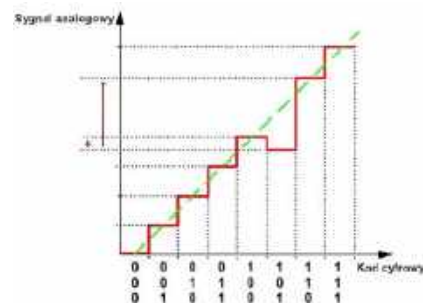
Poza parametrami dotyczącymi samej konwersji każdy przetwornik ADC ma też typowe dla układów scalonych parametry, jak napięcie zasilania i pobór prądu. Większość przetworników ADC ma interfejs szeregowy, zwykle I<sup>2</sup>C lub SPI. Przetwornik audio zazwyczaj posiadają interfejs I<sup>2</sup>S. Niektóre z nich mogą posiadać dodatkowe wejścia i wyjścia: SHUTDOWN czy INT. To pierwsze pozwala mikrokontrolerowi uśpić przetwornik, ograniczając pobór prądu do bardzo niskich wartości, przyjaznych pracy bateryjnej. Drugi sygnał pozwala przetwornikowi zasygnalizować zakończenie konwersji.

Prawie wszystkie obecnie produkowane mikrokontrolery mają wbudowane przetworniki ADC. Niektóre z tych przetworników pozwalają na bardziej zaawansowane tryby pracy, występujące też w niektórych scalonych przetwornikach różnych producentów. Funkcje te, to ciągła konwersja, generowanie przerwania gdy zostanie przekroczona zadana wartość minimalna bądź maksymalna i automatyczne uśrednianie wartości. Przetworniki ADC mogą też posiadać wbudowane wzmacniacze operacyjne o programowalnym wzmocnieniu – dzięki nim układ może konwertować sygnały o małych amplitudach z większą dokładnością. ■

Jos Verstraten



6. Zjawisko nieliniowości wyjaśnione graficznie (© 2021 Jos Verstraten)



7. Zjawisko niemonotoniczności wyjaśnione graficznie (© 2021 Jos Verstraten)

# Przetworniki ADC i DAC w praktyce

**Przetworniki analogowo-cyfrowe (ADC) i cyfrowo-analogowe (DAC) pozwalają w bardzo łatwy sposób integrować mikrokontrolery i układy analogowe. O ile większość hobbystów korzysta z modułów ADC wbudowanych w niemal każdy mikrokontroler dostępny na rynku, o tyle zewnętrzne układy ADC oraz DAC są amatorom praktycznie nieznane. W tym artykule przybliżę praktyczne aspekty stosowania tych ciekawych komponentów.**

Na rynku dostępnych jest bardzo wiele najróżniejszych układów przetworników analogowo-cyfrowych oraz cyfrowo-analogowych. Jeszcze więcej układów integruje te przetworniki celem realizowania swoich funkcji – znanym wielu amatorom przykładem może być termometr cyfrowy DS18B20. Biorąc pod uwagę, w jak wielu układach potrzebna jest konwersja między sygnałami analogowymi, a danymi cyfrowymi, aż dziw bierze, jak rzadko hobbysci korzystają z dedykowanych układów ADC i DAC, i jak jeszcze rzadziej robią to poprawnie.

## Moduł ADC w mikrokontrolerze, a dedykowany układ ADC

Jednym z pytań, które hobbysta może zadać, jest to dotyczące stosowania zewnętrznego układu ADC, gdy mikrokontroler ma już taki moduł w swojej strukturze. To prawda, i w większości zastosowań przetworniki te mogą być wystarczające. Jednakże mają one szereg ograniczeń, z których większość hobbystów nie zdaje sobie sprawy. Wystarczy jednak spojrzeć na specyfikację typowego modułu ADC, i porównać ją ze specyfikacją dedykowanego układu, by zobaczyć pierwszy problem. Dla przykładu niech posłuży mikrokontroler AVR ATmega328P, znany każdemu użytkownikowi Arduino. Nota podaje szereg istotnych parametrów, ale nas interesują cztery:

- nieliniowość całkowita INL: 0,6 LSB typowo;
- nieliniowość różnicowa DNL: 0,3 LSB typowo;
- błąd wzmocnienia:  $\pm 3,5$  LSB;
- błąd przesunięcia (offsetu):  $\pm 3,5$  LSB.

W efekcie bez dokładnej kalibracji 10-bitowy przetwornik ADC w mikrokontrolerze ma w praktyce tylko 8 bitów, bo ostatnim dwóm najmniej znaczącym bitom nie można ufać. Porównajmy to z dedykowanym układem ADC, a dokładniej z ADC101S021 firmy Texas Instruments. Układ ten również oferuje 10-bitową rozdzielczość, ale istotne dla porównania parametry wyglądają tak:

- nieliniowość całkowita INL: od  $-0,13$  LSB do  $+0,14$  LSB;
- nieliniowość różnicowa DNL: od  $-0,09$  LSB do  $+0,12$  LSB;
- błąd wzmocnienia:  $\pm 0,7$  LSB maksymalnie;
- błąd przesunięcia (offsetu):  $\pm 1,0$  LSB maksymalnie.

Już na starcie zyskujemy większą pewność co do wartości przedostatniego bitu, czyli bez kalibracji mamy 9 bitów rozdzielczości. Po kalibracji błędu przesunięcia i wzmocnienia uzyskamy prawdziwie dziesięciobitową dokładność. Ale możemy pójść o krok dalej i zakupić dokładniejszy układ, na przykład 12-bitowy, na przykład MCP3221 firmy Microchip:

- nieliniowość całkowita INL: 0,75 LSB typowo;
- nieliniowość różnicowa DNL: 0,5 LSB typowo;
- błąd wzmocnienia:  $\pm 2$  LSB maksymalnie;
- błąd przesunięcia (offsetu):  $\pm 3$  LSB maksymalnie.

Wygląda to nieco gorzej, niż dla modułu wewnątrz mikrokontrolera, ale pamiętajmy o dwóch dodatkowych bitach – po eliminacji błędów wzmocnienia i offsetu efektywna rozdzielczość wyniesie przynajmniej 10 bitów. Czy można lepiej? Jasne! Układ ADC121S021 firmy Texas Instruments oferuje pod pewnymi względami lepsze parametry:

- nieliniowość całkowita INL: od  $-0,4$  LSB do  $+0,55$  LSB;
- nieliniowość różnicowa DNL: od  $-0,3$  LSB do  $+0,6$  LSB;
- błąd wzmocnienia:  $-1,6$  LSB maksymalnie;
- błąd przesunięcia (offsetu):  $-0,25$  LSB maksymalnie.

Mógłbym wymieniać kolejne, dostępne układy, ale myślę, że nie jest to konieczne, bo łatwo zauważyć, iż praktycznie każdy dedykowany przetwornik ADC sprawdzi się lepiej, niż wewnętrzny moduł ADC w mikrokontrolerze, zwłaszcza gdy zależy nam na ich precyzji.

Warto jednak wspomnieć o tym, iż niektóre mikrokontrolery posiadają dość zaawansowane moduły ADC, potrafiące wykonywać automatyczne pomiary i generować przerwanie tylko wtedy, gdy zostanie przekroczona określona wartość minimalna lub maksymalna. Moduły te potrafią też wykonywać automatycznie operację uśredniania wyników i oversamplingu dla większej rozdzielczości i mniejszych szumów kosztem częstotliwości próbkowania. Oczywiście są i układy potrafiące to samo, a nie będące częścią mikrokontrolera, jak na przykład dwukanałowy, 12-bitowy przetwornik ADC z rozbudowanym systemem kontrolnym, ADS7142 firmy Texas Instruments. Układ ten potrafi prowadzić automatyczną akwizycję danych i zasygnalizować przekroczenie progów minimalnych, maksymalnych oraz bycia między nimi dla obu kanałów niezależnie. Automatycznie koryguje błąd offsetu, choć błąd wzmocnienia jest dość spory i stanowi typowo  $\pm 0,03\%$  pełnego zakresu mierzonych napięć.

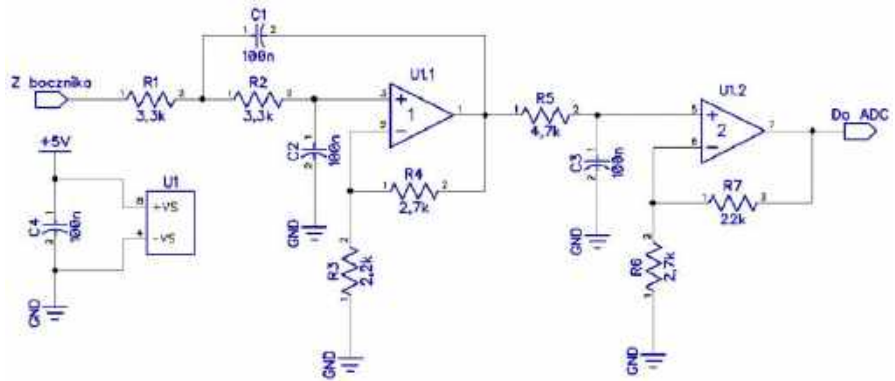
## Rzecz o filtrach i buforach

Wymienione do tej pory układy, w tym także moduł ADC w ATmega328P, były dość szybkie, z typową prędkością próbkowania od 22 kps do 200 kps. Ze względu na sposób działania tych układów, jeśli mierzony sygnał ma wyższą częstotliwość niż połowa częstotliwości próbkowania, uzyskane rezultaty mogą się mocno różnić od rzeczywistości. Nawet przy pomiarze wartości napięcia, które zmieniają się powoli, wszelkie szumy wyższej częstotliwości, niż połowa częstotliwości próbkowania mogą zniweczyć rezultaty. Osobną sprawą jest też sam poziom próbkowanego sygnału. Jeśli na przykład chcemy mierzyć prąd do 1 A na boczniku 0,1  $\Omega$ , i użyjemy modułu ADC w mikrokontrolerze z wewnętrznym napięciem odniesienia wynoszącym 2,048 V, to w zasadzie nasza rozdzielczość pomiaru wyniesie 2 mV, czyli 50 mA! Uważny Czytelnik w tym momencie zawołać może: „Hola, hola, Autorze! Pisałeś wcześniej, iż ostatnie dwa najmniej znaczące bity są bezużyteczne!”, i miałbyś rację. Nasz hipotetyczny amperomierz

bez tych dwóch bitów ma efektywną dokładność do 8 mV albo 200 mA! Toż to multimetr z dyskontu mierzy dokładniej! Dodajmy do tego jeszcze drobny fakt, iż wielu hobbystów nie zdaje sobie z tego sprawy i ślepo wierzy odczytom ADC. Co zatem robić? Ano trzeba zastosować wzmacniacz operacyjny, najlepiej z filtrem dolnoprzepustowym. Dla ułatwienia zadania założmy, iż boczniak pomiarowy jest włączony między obciążenie, a masę układu pomiarowego. Dzięki temu wystarczy nam prosty wzmacniacz nieodwracający z opcjonalnym filtrem. Założmy do tego, iż naszym celem jest dokonywanie pomiarów z częstotliwością 1 ksp/s, i chcemy maksymalizować zakres i dokładność pomiaru naszego modułu ADC.

**Rysunek 1** prezentuje schemat naszego filtra ze wzmacnieniem opartego o układ scalony LMV358 będący unowocześnioną, niskonapięciową wersją LM358. Rezystory R1, R2, R3 i R4, kondensatory C1 i C2, oraz wzmacniacz U1.1 tworzą aktywny filtr drugiego rzędu ze wzmacnieniem wynoszącym z grubsza 2,2 razy. R5 i C3 tworzą dodatkowy filtr pasywny poprawiający nachylenie zbocza wzmacniacza. U1.2 wraz z rezystorami R6 i R7 tworzy drugi stopień wzmacniający o wzmacnieniu około 9,14 razy. **Rysunek 2** przedstawia charakterystykę przenoszenia oraz przesunięcie fazowe dla całego obwodu. Na górnym wykresie dodane są linie 26 dB (wzmocnienie ~20 razy), 23 dB. Jak widać, pasmo jest bliskie teoretycznym założeniom, podobnie wzmacnienie. Dzięki temu układowi możemy mierzyć prąd z rozdzielczością do 10 mA, nawet bez dwóch najmniej znaczących bitów. Z nimi osiągalna rozdzielczość wyniosłaby 2,5 mA. Pokazany układ pozwala mierzyć też prąd DC, stąd ważna uwaga: LMV358 ma napięcie niezrównoważenia  $\pm 1,7$  mV typowo,  $\pm 7$  mV maksymalnie. Drugi stopień wzmacniacza swoje napięcie niezrównoważenia oraz napięcie niezrównoważenia pierwszego stopnia, który pracuje jako bufor. Na szczęście wystarczy wykonać konwersję ADC bez obciążenia boczniaka by zmierzyć sumę błędów niezrównoważenia obu wzmacniaczy, oczywiście ignorując dwa najmniej znaczące bity. Można też te wartości zmierzyć zwykłym multimetrem bezpośrednio w układzie.

Czytelnik może zmienić wartości kluczowych rezystorów i kondensatorów w tym układzie tak, by uzyskać pożądane efekty. Warto pamiętać

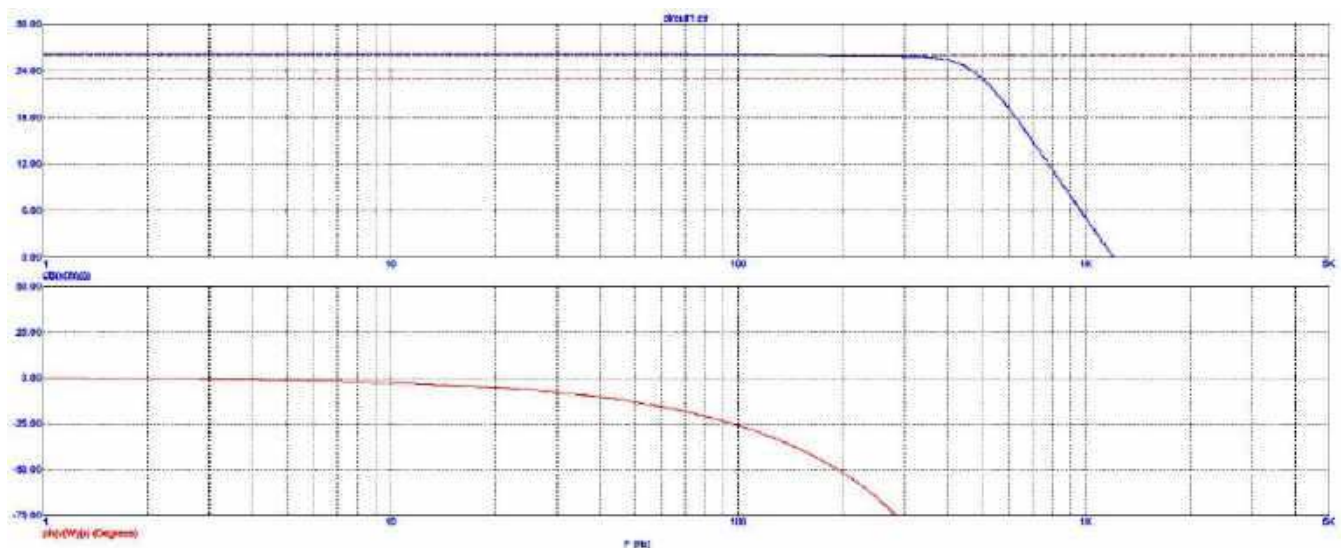


Rysunek 1. Filtr aktywny ze wzmacnieniem do przetwornika ADC

iż sam filtr nie powinien mieć zbyt dużego wzmacnienia, gdyż łatwo może przekształcić się w oscylator. Do symulacji pracy układu i wygenerowania wykresów z rysunku 2 użyty został dostępny za darmo program MicroCap 12.

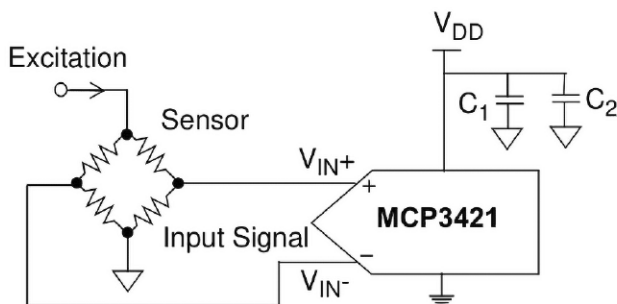
## Więcej! Lepiej! Dokładniej!

Wymienione powyżej układy scalone są relatywnie tanie, cena żadnego z nich nie powinna przekraczać 12 złotych brutto, a większość plasuje się poniżej 9 złotych. To ważne dla amatora, który zazwyczaj ma ograniczone fundusze na swoje hobby. Dodatkowo omawiane układy używały napięcia zasilania jako napięcia odniesienia do konwersji, co dodatkowo pogarsza parametry. Ponadto układy te były w większości przypadków jednokanałowe, co zazwyczaj powinno wystarczyć. A co, jeśli te parametry jednak nie są wystarczające? Co, jeśli serce krzyczy „Więcej bitów!”? Czy hobbysta z głodującym węzłem w kieszeni jest skazany na te żałosne 12 bitów? Nie! Na szczęście w kwocie do dwudziestu złotych można nabyć naprawdę ciekawe układy zaspokajające nasze pomiarowe potrzeby. Do tego ze względu na sposób działania tych układów nie potrzebujemy aktywnych filtrów o stromych zboczach (czasami stosuje się proste filtry RC), a w niektórych przypadkach również wzmacnienia. Te zalety posiadają przetworniki typu Delta-Sigma, oznaczane też greckimi literami  $\Delta\Sigma$ . Przetworniki te oferują znacznie większą rozdzielczość i przyzwolitą dokładność kosztem szybkości działania. Do niedawna problemem była też cena, ale obecnie nie brakuje tanich układów o dobrych parametrach.

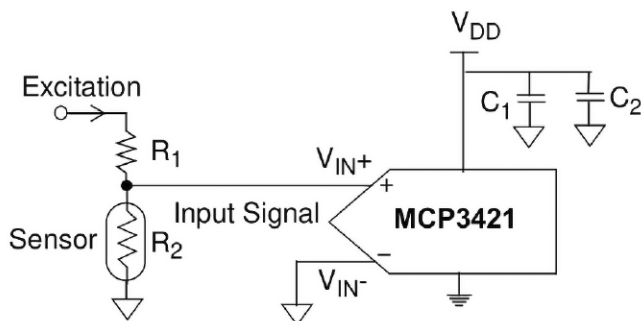


Rysunek 2. Charakterystyka pasma przenoszenia i fazy układu z rysunku 1

## (a) Differential Input Signal Connection:



## (b) Single-ended Input Signal Connection:



$C_1$  : 0.1  $\mu$ F, Ceramic Capacitor

$C_2$  : 10  $\mu$ F, Tantalum Capacitor

**Rysunek 3. Przykłady łączenia przetwornika ADC typu  $\Delta\Sigma$  z wejściem różnicowym do różnych sensorów**

Przyjrzyjmy się kilku z nich: MCP3461/2/4R to rodzina 16-bitowych przetworników  $\Delta\Sigma$  firmy Microchip, poszczególne modele mają 1, 2 lub cztery kanały różnicowe. Maksymalna prędkość próbkowania wynosi 153,6 ksp/s, minimalna zaś 12,5 sp/s. Częstotliwość ta zależy od ustawień wewnętrznego zegara, którego częstotliwość taktowania jest 128 razy wyższa od częstotliwości próbkowania. Układ posiada wewnętrzny wzmacniacz operacyjny o przełączanym wzmacnieniu 0,33x, 1x, 2x, 4x, 8x i 16x. Dodatkowo można dodać „wzmocnienie cyfrowe” 2x lub 4x kosztem nieznacznego pogorszenia parametrów. Układ automatycznie kompensuje błąd offsetu, błąd nieliniowości całkowitej wynosi od  $\pm 7$  ppm FSR do nawet  $\pm 32$  ppm FSR. FSR to pełnoskalowy zakres napięć od 0 V do  $V_{ref} - 1\text{LSB}$ . Błąd wzmacnienia wynosi maksymalnie  $\pm 3\%$ . Warto też wspomnieć o możliwości pracy ciągłej ze skanowaniem kolejnych wejść (z programowalnym opóźnieniem między kolejnymi akwizycjami) i o generowaniu przerwań, gdy dane są dostępne lub gdy rozpocznie się konwersja. Jak widać, układ oferuje bardzo dużo możliwości przy relatywnie niskiej cenie. Jednym z ograniczeń jest niedokładne źródło napięcia odniesienia wewnątrz układu:  $2,4\text{ V} \pm 2\%$ . Na szczęście układy te mają wejścia dla zewnętrznych źródeł napięcia odniesienia. Podobnymi pod względem specyfikacji są układy MCP3561/2/4R – oferują 24 bity rozdzielczości i lepsze parametry odnośnie poziomu szumów, ale reszta specyfikacji nie różni się od ich 16-bitowych kuzynów.

Jeśli jest nam potrzebne coś dokładniejszego, to Microchip oferuje układ MCP3421. Układ ma 18-bitową rozdzielczość, źródło napięcia odniesienia  $2,048\text{ V} \pm 0,05\%$ , programowalny wzmacniacz o wzmacnieniu 1x, 2x, 4x i 8x, oraz błąd INL na poziomie 10 ppm/FSR, gdzie

FSR wynosi 4,096 V/wzmocnienie. Układ ma dwie drobne wady: częstotliwość próbkowania zależy od rozdzielczości: 240 sp/s/12 bitów, 3,75 sp/s/18 bitów. Drugą wadą jest obudowa – SOT23-6. Czytelnicy o słabszym wzroku mogą potrzebować mikroskopu. **Rysunek 3** przedstawia dwa układy pracy tego przetwornika: pomiar różnicowy sensora mostkowego oraz pomiar single-ended. Przykład pochodzi z noty katalogowej układu. Warto tu wspomnieć iż wszystkie przetworniki ADC z wejściami różnicowymi mogą podawać wartości dodatnie i ujemne dla różnicy napięć wejściowych. Można wykorzystać ten fakt do pomiaru prądu przepływającego przez rezystor bocznikowy, determinując przy okazji jego kierunek, pod warunkiem iż napięcie na wejściach nie przekroczy dopuszczalnej wartości. Podobnym układem jest MCP3425, jedyna różnica to liczba bitów, tylko 16, i częstotliwość próbkowania: 15 sp/s. Pozostałe parametry są takie same.

## Kalibracja błędów wzmacnienia i przesunięcia

Nie każdy układ ADC ma możliwość automatycznej kalibracji błędów offsetu i wzmacnienia, ale wiele lepszych układów ma rejestry, do których można wpisać stosowne wartości kalibracyjne. W przypadku prostszych układów i modułów ADC będących częścią mikrokontrolera kalibrację można przeprowadzić w programie. Sam proces jest relatywnie prosty i wymaga jedynie dobrej klasy multimetru i stabilnego, regulowanego źródła napięcia. O tym drugim będzie szerzej w dalszej części artykułu.

Na początek nasz układ z przetwornikiem ADC powinien mieć tryb kalibracji, w którym podaje nam odczytane napięcie w formie surowej wartości, dziesiętnej, szesnastkowej lub binarnej – co kto woli. Do wejścia ADC podłączamy nasze źródło napięcia i multimetr, ustawiając wartość na około 10% maksymalnego napięcia dla wejścia ADC. Załóżmy że mamy do czynienia z 12-bitowym ADC z napięciem odniesienia wynoszącym 4,096 V. Ustawmy więc napięcie na równe 400 mV. Wartość LSB dla tego układu wynosi 1 mV, więc wartość odczytana powinna wynosić dokładnie 400. Ale wartość odczytana wynosi 408, zanotujmy to. Następnie ustawmy napięcie źródła na około 90% napięcia maksymalnego, niech to będzie 3,6 V. Wartość idealna powinna wynosić 3600, ale odczytana wynosi 3664, to też zanotujmy. Teraz musimy policzyć nachylenie charakterystyki przetwornika ADC według następującego wzoru:

$$S = \frac{Kb - Ka}{Vb - Va}$$

Liczymy te wartości dla idealnego ADC i dla zmierzonego ADC:

$$Si = \frac{3600 - 400}{3,6V - 0,4V} = 1000$$

$$Sm = \frac{3664 - 408}{3,6V - 0,4V} = 1017,5$$

Teraz obliczamy błąd przesunięcia:

$$Eo = Sm \cdot (0 - Va) + Ka = 1017,5 \cdot (0 - 0,4V) + 408 = 1\text{LSB}$$

Na koniec obliczamy błąd wzmacnienia, wynosi on:

$$Eg = \frac{Sm - Si}{Si} = \frac{1017,5 - 1000}{1000} = 0,0175 = 1,75\%$$

Aby skorygować błędy przesunięcia i wzmacnienia ADC należy wynik pomiaru przeliczyć według następującego wzoru:

$$K = (K_{adc} - Eo) \cdot \frac{Si}{Sm}$$

gdzie  $K$  to wartość kodu po kalibracji, a  $K_{adc}$  to wartość odczytana z przetwornika. Załóżmy, że  $K_{adc}=1898$  i podstawmy dane do wzoru:

$$K = (1898 - 1) \cdot \frac{1000}{1017,5} = 1864,37$$

Zmierzone napięcie wynosi 1,864 V.

Praktyczna implementacja w programie może wymagać pewnej gimnastyki ze strony programisty, gdyż mikrokontrolery nie za bardzo lubią operacje zmiennoprzecinkowe. Implementacja stałoprzecinkowych operacji wykracza jednak poza zakres tematyczny tego artykułu.

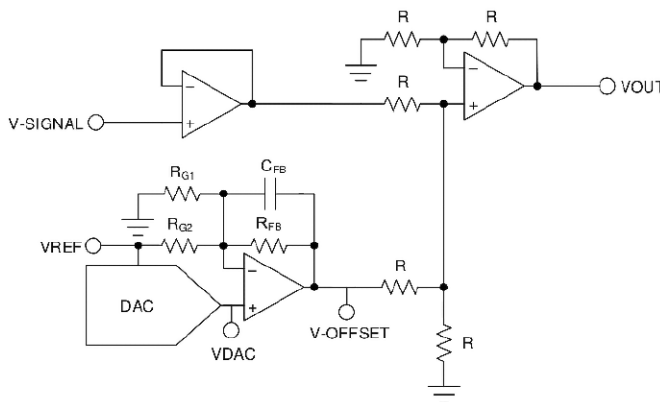
## Przetworniki DAC

Przetworniki DAC robią dokładnie to samo, co przetworniki ADC, tylko na odwrót. Zamieniają kod cyfrowy na napięcie. Czytelnik może się zastanawiać, po co amatorowi przetwornik DAC, skoro zazwyczaj do takich zadań stosuje się moduł PWM mikrokontrolera z prostym filtrem RC na wyjściu. Jasne, w prostych zastosowaniach jest to rozwiązanie wystarczające, nawet w formie programowego PWM. Ale jeśli chcemy mieć dokładną kontrolę nad wartością napięcia wyjściowego, na przykład w regulowanym zasilaczu sterowanym cyfrowo, to rozwiązanie może nie być adekwatne. Dodatkowo przetwornik DAC ze źródłem napięcia odniesienia będzie zapewniał stabilne napięcie niezależnie od wahań i zakłóceń napięcia zasilania całego projektu. Jednym z zastosowań dobrej klasy przetwornika DAC w połączeniu z równie dobrym źródłem napięcia odniesienia może być stworzenie prostego zasilacza do kalibrowania przetworników ADC i innych układów.

Tak samo, jak przetworniki ADC, układy DAC posiadają podobne ograniczenia: błąd całkowitej nieliniowości, różnicowej nieliniowości, przesunięcia i wzmocnienia. Podobnie też potrzebują filtrów RC na wyjściu, jeśli chcemy generować szybkozmiennie przebiegi analogowe. Przetworniki mogą też posiadać wbudowane źródło napięcia odniesienia lub używać zewnętrznego źródła. W tym drugim przypadku niekiedy można zrobić trik polegający na połączeniu wyjścia jednego przetwornika DAC do wejścia referencyjnego drugiego przetwornika, dzięki czemu teoretycznie można podwoić rozdzielczość całego układu. W praktyce trzeba pamiętać o wymienionych już błędach nieliniowości, przesunięcia i wzmocnienia. Do czego to się może przydać? Na przykład do regulacji amplitudy generowanego sygnału analogowego. W takiej sytuacji drugi układ DAC otrzymuje i aktualizuje wartości cały czas, podczas gdy pierwszy jest aktualizowany tylko gdy trzeba zmienić amplitudę. Rozwiązanie takie czasem stosuje się do kontrolowania amplitudy sygnału wyjściowego w scalonych generatorach DDS – jednym z przykładowych układów na to pozwalających jest AD9850 firmy Analog Devices.

Jednym z zastosowań przetworników DAC, zwłaszcza o większej rozdzielczości może być generowanie napięć kontrolnych (CV) dla analogowych syntezy modułowych. Standardem jest napięcie do 10 V ze skalą 1 V na oktawę. Oczywiście będzie to wymagać dodania wzmacniacza operacyjnego do przeskalowania wyjścia DAC do akceptowalnych wartości. Z pomocą przetwornika cyfrowo-analogowego można też generować napięcia i prądy do polaryzacji sensorów wymagających precyzyjnych źródeł napięcia lub prądu. Kolejnym, typowym zastosowaniem jest generowanie napięcia przesunięcia sygnału w obwodach wejściowych oscyloskopów cyfrowych. **Rysunek 4** przedstawia schemat poglądowy takiego obwodu, pochodzi on z noty projektowej SBAA343 firmy Texas Instruments.

Łącząc układ DAC ze wzmacniaczem operacyjnym, tranzystorem mocy i rezystorem pomiarowym można stworzyć sterowane cyfrowo sztuczne obciążenie. Dodatkowy przetwornik ADC może mierzyć prąd obciążenia na boczniku (celem korekcji dla napięcia niezrównoważenia wzmacniacza i błędów układu DAC) i napięcie występujące na całym



**Rysunek 4.** Dodawanie regulowanego przesunięcia do sygnału za pomocą przetwornika DAC i wzmacniaczy operacyjnych za notą aplikacyjną SBAA343

sztucznym obciążeniu. Mikrokontroler nie tylko będzie mógł symulować różne rodzaje zmiennych obciążeń, ale też podawać na bieżąco i logować moc strat. Zastosowaniem takiego układu może być na przykład testowanie pojemności akumulatorów czy testy obciążeniowe dla różnego rodzaju zasilaczy.

## Vref+DAC+ADC= precyzyjne źródło napięcia odniesienia

Na **rysunku 5** przedstawiono częściowy schemat rozwiązania układowego pozwalającego uzyskać precyzyjne źródło napięcia odniesienia z funkcją automatycznej kalibracji. Na schemacie nie uwzględniono kondensatorów filtrujących zasilanie, ani połączeń do mikrokontrolera. Dobrej klasy źródło napięcia odniesienia ADR4525 firmy Analog Devices (U1) dostarcza 2,5 V do podwójnego, 12-bitowego przetwornika DAC MCP4922 (U2). Zamiast ADR4525 można zastosować ADR4540, by mieć napięcie 4,096 V. Dwukanałowy, 16-bitowy przetwornik ADC MCP3422 (U3) mierzy napięcia wyjściowe z przetwornika DAC, dzięki czemu na bieżąco możemy kontrolować napięcie wyjściowe przetwornika. U4 to dowolny, precyzyjny wzmacniacz operacyjny rail-to-rail, pracujący w roli wzmacniacza sumującego oba wyjścia z przetwornika DAC. Rezystory R1 i R2 ustalają wagę tych wyjść. Napięcie wyjściowe tego układu można obliczyć według wzoru:

$$U_{wyj} = \frac{V_{outA} \cdot R2 + V_{outB} \cdot R1}{R1 + R2}$$

W ten sposób możemy sumę wszystkich błędów wyjścia VOUT\_A zniwelować za pomocą napięcia z wyjścia VOUT\_B, jednocześnie

REKLAMA

**ELMAX 1988**

Certyfikat  
Udowodnienia  
Laboratorium

Wzrost: 1,80 m  
Ciężar: 75 kg  
Typ: 1

Zakład produkcyjny:  
85-650 Warka  
ul. M. Rejzendera 17  
tel. 22 761 62 93  
22 761 65 00  
fax 22 761 65 00 w 23  
www.elmax.com.pl  
elmax@elmax.com.pl

**OBWODY DRUKOWANE**

Produkcja, Projektowanie, Montaż

|                                                                            |                                              |                                                                      |                                                         |
|----------------------------------------------------------------------------|----------------------------------------------|----------------------------------------------------------------------|---------------------------------------------------------|
| <p>Platy jednoczynne</p> <p>Platy dwuczynne</p> <p>Platy na zamówienie</p> | <p>Wzrosty</p> <p>Wzrosty</p> <p>Wzrosty</p> | <p>Dokumentacja technologiczna</p> <p>Dokumentacja konstrukcyjna</p> | <p>Montaż elektroniczny</p> <p>Montaż elektroniczny</p> |
| <p>Aktywne kalkulator</p> <p>Pokrycie</p> <p>Platy ciekawe</p>             | <p>Wzrosty</p> <p>Wzrosty</p> <p>Wzrosty</p> | <p>Wzrosty</p> <p>Wzrosty</p> <p>Wzrosty</p>                         | <p>Wzrosty</p> <p>Wzrosty</p> <p>Wzrosty</p>            |

błędy tego drugiego wyjścia będą zredukowane mniej-więcej tysiąckrotnie. Można też rozważyć dodanie kolejnego przetwornika ADC, na przykład MCP3551, który oferuje 22-bitową rozdzielczość i może wykorzystać to samo źródło napięcia odniesienia, co MCP4922. Układ ten monitorowałby wyjście wzmacniacza operacyjnego dla korekcji napięć niezerównoważenia i tolerancji rezystorów. Ostatecznym krokiem, tylko dla pekuniarnych konstruktorów byłoby skorzystanie z usług firmy specjalizującej się w kalibracji różnych instrumentów.

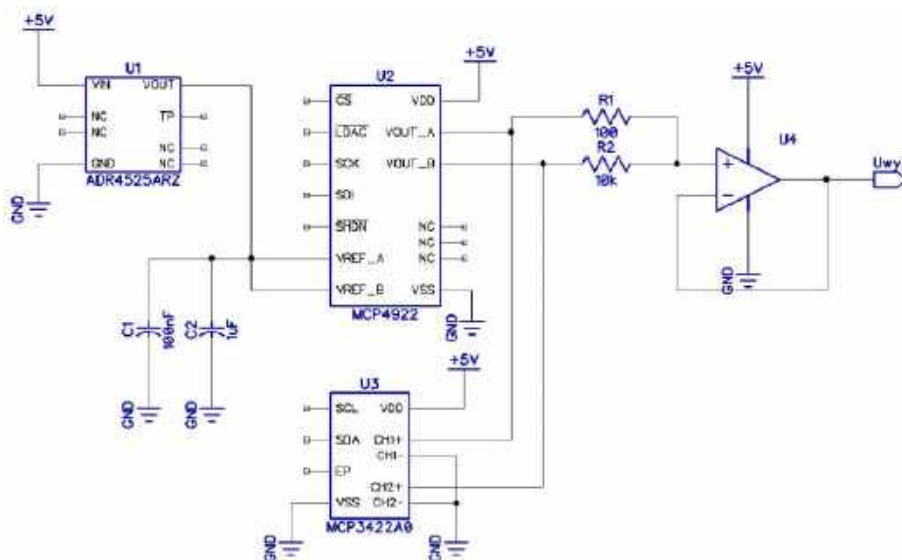
## „Cyfryzacja” zasilacza warsztatowego

Do niedawna niemal wszystkie tanie zasilacze warsztatowe w sprzedaży opierały się o regulację za pomocą potencjometrów. Oczywiście można już kupić zasilacz z cyfrową kontrolą, ale zazwyczaj są to moduły zasilaczy impulsowych, adekwatnych w większości zastosowań. Zasilacz liniowy jednak ma tę zaletę, iż bywa bardziej niskoszumny niż przetwornica impulsowa – istotna sprawa przy pracy z układami o dużym wzmacnieniu i z układami radiowymi. Po co więc dodawać cyfrową regulację do liniowego zasilacza? Choćby dla precyzyjniejszej kontroli i możliwości łatwego ustawiania parametrów. Jednak to zadanie może wymagać zaprojektowania całego zasilacza od zera. Większość zasilaczy z Chin realizuje regulację napięcia za pomocą dzielnika napięcia między wyjściem zasilacza, a wzmacniaczem sterującym tranzystorami mocy. Przeróbka układu wymagałaby za dużo zmian, dlatego lepiej zaprojektować całość od zera.

Załóżmy, że celem jest uzyskanie zasilacza 30 V/5 A, czyli typowych parametrów zasilaczy dostępnych na rynku dzięki naszym przyjaciółom z Chin. Z 12-bitowymi przetwornikami DAC zostawiając sobie odrobinę marginesu możemy uzyskać rozdzielczość regulacji napięcia 7,5 mV oraz prądu 1,25 mA. Z kolei przetworniki ADC mogą pozwolić nie tylko na pomiar napięcia i prądu (najlepiej z nieco wyższą rozdzielczością), ale też na logowanie zmierzonych wartości na bieżąco. Bardziej rozbudowane firmware zasilacza mogłoby też pozwolić na pracę jako ładowarka do akumulatorów w cyklu CC/CV. Takie funkcjonalności oferują zasilacze laboratoryjne sterowane cyfrowo kosztujące sporo powyżej tysiąca złotych. Sprytny hobbysta mógłby się zmieścić w cenie 100...200 PLN, pomijając koszt transformatora.

## Słowo o źródłach napięcia odniesienia do przetworników ADC i DAC

Źródła napięcia odniesienia były w tym artykule wspomniane wielokrotnie, warto zatem przyjrzeć im się bliżej. Każdy oczywiście zwróci wprawę uwagę na wartość napięcia wyjściowego źródła. Nie jest to jednak tak istotny parametr, jak to by się mogło wydawać. Bardziej istotnym, zwłaszcza dla hobbysty, parametrem jest dokładność tego napięcia, tj. o ile napięcie rzeczywiste może się różnić od idealnej wartości po włączeniu układu. Profesjonalista stwierdzi, iż rzeczywistą wartość można zmierzyć i wykorzystać tę informację do kalibracji układu. Oczywiście to prawda, i tak są kalibrowane multimetry stołowe mające 4½, 6½ i więcej cyfr. Dla hobbysty koszt kalibracji może być za duży, a to oznacza, iż musi się opierać na danych producenta układu. Dwa inne, istotne parametry to dryft temperaturowy oraz dryft po upływie tysiąca godzin lub jednego roku pracy. Z powodu wpływu temperatury



Rysunek 5. Precyzyjne, regulowane źródło napięcia odniesienia z możliwością automatycznej kalibracji

na wartość napięcia w niektórych multimetrach źródło napięcia odniesienia było zamknięte w termicznie izolowanym pojemniku z grzałką i termostatem, analogicznie do oscylatorów OCXO. Napięcie z takiego źródła nie będzie takie, jak w temperaturze pokojowej, ale nie będzie też zależne od temperatury otoczenia.

Porównajmy zatem dwa źródła napięcia odniesienia, dla uczciwości oba będą miały to samo napięcie wyjściowe: 2,5 V. Jednym z nich jest układ ADR4525B od Analog Devices kosztujący u jednego z dystrybutorów 46,43 złotych. Drugi układ, to MCP1501-25 firmy Microchip kosztujący u tego samego dystrybutora 5,67 PLN. MCP1501-25 ma zagwarantowaną wartość napięcia wyjściowego w zakresie od 2,975 V do 2,525 V, czyli 0,1%. ADR4525B ma zagwarantowane napięcie wyjściowe 2,5 V  $\pm$  500  $\mu$ V, 0,02%. Współczynnik temperaturowy dla MCP1501-25 wynosi 10 ppm/°C typowo, 50 ppm/°C maksymalnie. Analog Devices zapewnia, iż maksymalny współczynnik temperaturowy dla układu klasy B wynosi 2...4 ppm/°C zależnie od metody pomiaru. Po pierwszych 250 godzinach napięcie może się zmienić o 19 ppm, a po tysiącu o 25 ppm. Microchip nie podaje danych dla długoterminowej stabilności. Jak widać, płacimy osiem razy więcej za znacząco lepszy układ.

Ktoś może zadać pytanie „A co z TL431 albo z diodami Zenera?”, a ja odpowiem, że jako źródła napięcia odniesienia do zadań wymagających precyzji są one bezużyteczne. Zarówno diody Zenera, jak i układy w rodzaju TL431 czy LM4040 są mało dokładne, termicznie niestabilne, a na domiar złego nie do końca niskoszumne. Są dobre do budowy prostych zasilaczy z regulacją potencjometrami, ale do pracy z układami ADC i DAC się nie nadają. Czytelnik dobrze zrobi, jeśli zapomni o tych zdrożnych myślach.

## Komunikacja z układami ADC i DAC

Wszystkie omawiane do tej pory układy używają interfejsów szeregowych do komunikacji z mikrokontrolerem, zazwyczaj SPI lub I²C. Zależnie od częstotliwości pracy układu spotyka się też bardziej rozbudowane protokoły, jak QSPI. Interfejs ten używa czterech linii danych, by w każdym cyklu zegarowym przesyłać po cztery bity. Rozwiązanie to zazwyczaj stosowane jest do łączności między mikrokontrolerem, a szybką pamięcią Flash, jak w słynnych modułach z ESP8266, ale bywa też stosowane właśnie w układach ADC i czasami DAC. Oczywiście mikrokontroler musi wspierać ten interfejs, oraz mieć wystarczającą przepustowość by wykorzystać możliwości, jakie on daje. Bardzo szybkie

przetworniki ADC, jak te stosowane w oscyloskopach cyfrowych, do konfiguracji używają wolniejszych interfejsów szeregowych, ale dane są już przesyłane liniami LVDS. Przykładem takiego układu może być HMCAD1511 firmy Analog Devices, który znajdziemy w oscyloskopach cyfrowych, na przykład w SDS1104X-E i X-U firmy Siglent.

Oczywiście stosuje się też całkowicie równoległe interfejsy komunikacyjne, zwłaszcza gdy układ ADC lub DAC współpracuje z układem FPGA lub ASIC. Dobrym przykładem są szybkie przetworniki DAC AD9776/8/9. Kolejne modele oferują rozdzielczość 12, 14 i 16 bitów, dwa kanały i maksymalną przepustowość 1 Gsps. Do konfiguracji służy interfejs SPI, ale próbki są przekazywane przez dwa 16-bitowe interfejsy równoległe. Same przetworniki DAC oferują naprawdę niezłe parametry biorąc pod uwagę częstotliwość działania. Oczywiście cena też jest stosowna – ponad 250 złotych za układ AD9776, a do tego trzeba doliczyć jeszcze koszt odpowiednio szybkiego układu FPGA i pamięci RAM. Na wyjściach konieczne też będą odpowiednie filtry anty-aliasingowe – zwykły układ RC się nie sprawdzi. Ale właśnie tak buduje się arbitralne generatory funkcyjne osiągające częstotliwości 100 MHz i więcej.

W projektach, w których sygnał audio jest przetwarzany na postać cyfrową, a potem przetwarzany za pomocą funkcji DSP, by potem ponownie zostać przetworzonym na postać analogową, często spotyka się komunikację za pomocą protokołu I<sup>2</sup>S. Ten potomek

formatu I<sup>2</sup>C został stworzony właśnie z myślą o szeregowej transmisji cyfrowego, wielokanałowego dźwięku, i na rynku nie brakuje dedykowanych przetworników ADC i DAC. Pewnym zaskoczeniem dla Czytelników może być fakt, iż spotykane na płytach głównych komputerów stacjonarnych i laptopów układy standardu AC'97 i HD Audio właśnie oferują taki interfejs, i nic nie stoi na przeszkodzie, by taki układ ze starej płyty głównej wylutować, a następnie połączyć z odpowiednio wydajnym mikrokontrolerem posiadającym zarówno odpowiedni interfejs, jak i funkcje DSP i odpowiedni zasób RAMu dla próbek. Zagadnienia przetwarzania sygnałów są skomplikowane, ale zawzięty konstruktor może posiadać odpowiednią wiedzę, by stworzyć na przykład własny multieffekt do gitary. Producenci mikrokontrolerów oferują też gotowe biblioteki DSP oraz różne przykłady ułatwiające naukę.

## Podsumowanie

Ten skromny poradnik nie omawia wszystkich możliwych zastosowań dla przetworników analogowo-cyfrowych i cyfrowo-analogowych. Mam jednak nadzieję, iż zachęci on Czytelników do częstszego korzystania z tych układów w swoich projektach. Być może w przyszłości pojawią się na łamach „Elektroniki dla Wszystkich” projekty zbudowane z użyciem układów ADC i DAC znajdujące zastosowanie w warsztacie elektronika-amatora. Wraz z Redakcją czekam na odzew naszych Czytelników. ■

Paweł Kowalczyk



## Przetworniki ADC/DAC

Rozwiązanie znajdziesz na [www.elportal.pl/quizy](http://www.elportal.pl/quizy)

### 1. Głównym problemem wbudowanego w mikrokontroler przetwornika ADC jest:

- ☐ niska rozdzielczość;
- ☐ niska prędkość przetwarzania, zależna też od zegara mikrokontrolera;
- ☐ duże błędy własne, wymagające kalibracji.

### 2. Jaka jest efektywna rozdzielczość przetwornika ADC w mikrokontrolerze AVR ATmega328P:

- ☐ 8 bitów bez kalibracji;
- ☐ 9 bitów bez kalibracji;
- ☐ 10 bitów, jak podaje nota katalogowa.

### 3. Warto jednak stosować moduły ADC w mikrokontrolerach, ponieważ:

- ☐ są zawsze dostępne;
- ☐ w najnowszych układach nie ustępują zewnętrznym przetwornikom;
- ☐ posiadają szereg zaawansowanych funkcji automatyzujących pomiary i warunkowo generujących przerwania.

### 4. Przetwornik ADC potrzebuje filtra na wejściu gdyż:

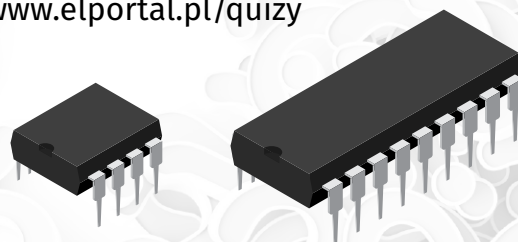
- ☐ szumy w sygnale powodują błędy;
- ☐ pasmo sygnału musi być ograniczone do połowy częstotliwości próbkowania;
- ☐ filtr chroni wejście układu przed uszkodzeniem.

### 5. Używając filtra aktywnego na wzmacniaczach operacyjnych należy uwzględnić i skompensować:

- ☐ błędy nieliniowości pasma przenoszenia;
- ☐ błędy niezrównoważenia wzmacniaczy;
- ☐ przesunięcie fazowe wprowadzane przez wzmacniacze.

### 6. Lepsze parametry oferują nieco droższe przetworniki ADC typu:

- ☐ SAR;
- ☐  $\Sigma\Delta$ ;
- ☐  $\Delta\Sigma$ .



### 7. Przetwornik DAC z własnym źródłem napięcia odniesienia w przeciwieństwie do modułu PWM oferuje:

- ☐ stabilne i niezależne od zasilania napięcie;
- ☐ mniejsze szumy;
- ☐ łatwiejszą implementację.

### 8. Przetworniki DAC można czasem łączyć szeregowo (wyjście jednego do wejścia Vref drugiego) by:

- ☐ mieć regulowane napięcie odniesienia;
- ☐ regulować amplitudę sygnału drugiego przetwornika;
- ☐ zsumować efektywną rozdzielczość obu przetworników.

### 9. Układ połączeń przetwornika DAC i wzmacniacza operacyjnego przedstawiony na rysunku 5 w artykule:

- ☐ pozwala regulować błąd niezrównoważenia wzmacniacza;
- ☐ zwiększa stabilność napięcia na wyjściu przez mieszanie kanałów;
- ☐ pozwala kompensować błędy jednego kanału drugim kanałem zwiększając efektywną rozdzielczość.

### 10. Przy wyborze źródła napięcia odniesienia hobbysta nie mający dostępu do precyzyjnych instrumentów pomiarowych musi pamiętać o:

- ☐ dokładności napięcia odniesienia, stabilności temperaturowej i starzeniu się źródła;
- ☐ tylko o dokładności napięcia odniesienia – układy źródeł są bardziej stabilne termicznie, niż pozostałe układy w typowym projekcie;
- ☐ tylko napięciem wyjściowym źródła – hobbysty nie stać na precyzyjne źródła napięcia odniesienia, więc pozostałe parametry i tak będą kiepskie.

# IDEALNE DO TWOJEGO



## ▲ Dwukanałowy oscyloskop 100 MHz, 1 GSa/s z funkcją generatora sygnału DDS, FNIRSI 1014D

Idealny dla hobbystów, studentów i profesjonalistów. Oferuje pasmo przepustowości 100 MHz, dwa kanały pomiarowe, częstotliwość próbkowania 1 GSa/s oraz duży, czytelny wyświetlacz LCD



## ◀ Dwukanałowy oscyloskop z generatorem 180 MHz FNIRSI DPOX180H

FNIRSI DPOX180H to przenośny, dwukanałowy oscyloskop cyfrowy o szerokości pasma 180 MHz i szybkości próbkowania 500 MS/s. DPOX180H jest idealnym narzędziem dla hobbystów, studentów i techników elektroniki, którzy potrzebują niedrogiego i wszechstronnego oscyloskopu

## ► Cyfrowy oscyloskop 200 kHz LCD 2,4 cala, DSO138 FNIRSI

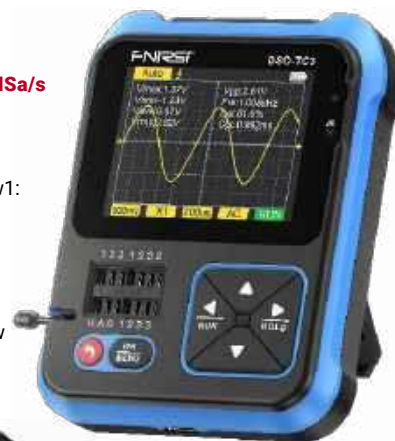
FNIRSI-138 to kompaktowy i niedrogi oscyloskop cyfrowy przeznaczony dla hobbystów, studentów i elektroników. Ma szereg funkcji, które czynią go przydatnym narzędziem do różnych zastosowań

## ► Mikroskop cyfrowy G1200 z wyświetlaczem 7 cali, powiększenie ×1200, tryb foto/wideo

Dzięki zastosowaniu regulowanego kąta oświetlenia, mikroskop G1200 skutecznie redukuje odbłaski, zapewniając czysty i szczegółowy obraz obserwowanego elementu. Jest to szczególnie istotne dla branży serwisu urządzeń elektronicznych. Większość mikroskopów ma tylko dwustopniową regulację powiększenia, która często okazuje się zbyt mała lub zbyt duża do optymalnej obserwacji. Mikroskop G1200 oferuje system płynnej regulacji powiększenia w zakresie 1...1200x. Dzięki temu użytkownik może zawsze uzyskać idealny poziom powiększenia, dostosowany do konkretnego elementu i czynności serwisowej

## ► Oscyloskop 500 kHz 10 MSa/s 3w1: tester elementów, generator sygnału DDS, FNIRSI DSO-TC3

Wielofunkcyjne urządzenie 3w1: oscyloskop/tester elementów LCR/generator sygnałów DDS Oscyloskop FNIRSI DSO-TC3 to przenośne urządzenie służące do wizualizacji i analizy przebiegów sygnałów elektrycznych



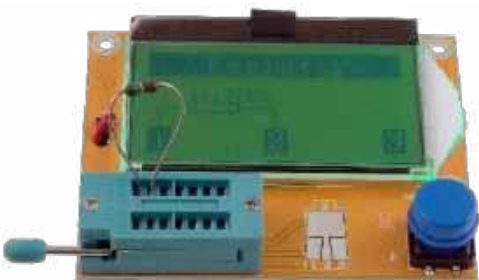
## ◀ Inteligentna lutownica FNIRSI HS-01 z regulacją napięcia pracy grzał 9...20 V

Kompaktowa i stylowa. HS-01 to inteligentna lutownica, która zachwyca swoim kształtem przypominającym ołówek i ergonomicznym designem. Dzięki niewielkiej wadze, niskiemu zużyciu energii i kompaktowej budowie, HS-01 idealnie nadaje się do przenoszenia i pracy w terenie. Antyparzeniowa nasadka zapewniająca bezpieczeństwo i wygodę użytkownika. Wszechstronność i wydajność. Szeroki zakres regulacji temperatury (80...420°C) pozwala na lutowanie różnorodnych elementów i przewodów



# JEGO WARSZTATU

<https://sklep.avt.pl/pl/producers/fnirsi-1713942890.html>



## ◀ Miernik LCR-T4 FNIRSI – tester tranzystorów, rezystorów, cewek, kondensatorów, diod i innych

Miernik elementów elektronicznych LCR-T4 bez obudowy – tester tranzystorów, rezystorów, cewek, diod i innych. Tester elementów elektronicznych z dużym, czytelnym wyświetlaczem z powodzeniem przetestuje diody, rezystory, tranzystory, cewki itd.



## ◀ Dwukanałowy oscyloskop z multimetrem i generatorem 10 MHz FNIRSI 2C23T

FNIRSI 2C23T to w pełni funkcjonalny, praktyczny cyfrowy oscyloskop dwukanałowy 3 w 1, opracowany przez firmę FNIRSI z myślą o zastosowaniach w serwisie i przemyśle. Urządzenie łączy w sobie trzy główne funkcje: oscyloskopu, multimetru i generatora sygnałów. Umożliwia jednocześnie korzystanie z generatora i oscyloskopu



## ◀ Wykrywacz ukrytych przewodów, metali i drewna w ścianach, FNIRSI WD-02

Wykrywacz FNIRSI WD02 to urządzenie zaprojektowane z myślą o rozwiązaniu problemów użytkowników podczas prac wykończeniowych i wiercenia. Produkt pozwala na wykrywanie metali (prętów stalowych, rur miedzianych) ukrytych w różnych miejscach, takich jak ściany, sufity i płyty gipsowo-kartonowe, a także kabli



## ▲ Zgrzewarka ogniw – spawarka punktowa FNIRSI SWM10 z kolorowym wyświetlaczem i funkcją powerbank

FNIRSI SWM10 to przenośna, inteligentna spawarka punktowa (zgrzewarka ogniw), która wykorzystuje technologię spawania punkowego z podwójnym impulsem, aby zapewnić stabilne i trwałe spoiny. Urządzenie ma 1,8-calowy kolorowy wyświetlacz LCD i może spawać blachy niklowe, żelazne i ze stali nierdzewnej. Maksymalny prąd spawania wynosi 1200 A, a spawarka posiada również wbudowany powerbank



## ◀ Inteligentny miernik cyfrowy z dużym wyświetlaczem FNIRSI SMART S1

FNIRSI S1 to inteligentny multimetr cyfrowy o dużym, czytelnym wyświetlaczu LCD. Charakteryzuje się szybkimi pomiarami, podświetleniem ekranu oraz funkcjami takimi jak zabezpieczenie przeciwprzepięciowe i sygnalizacja niskiego napięcia baterii



## ◀ Generator funkcyjny sygnału PWM, kalibrator FNIRSI SG-004A

FNIRSI SG-004A to przenośny, wielofunkcyjny generator sygnałów przeznaczony do kalibracji i testowania urządzeń elektronicznych. Urządzenie to oferuje prostą obsługę i kompaktową konstrukcję, co czyni je idealnym narzędziem do zastosowań terenowych



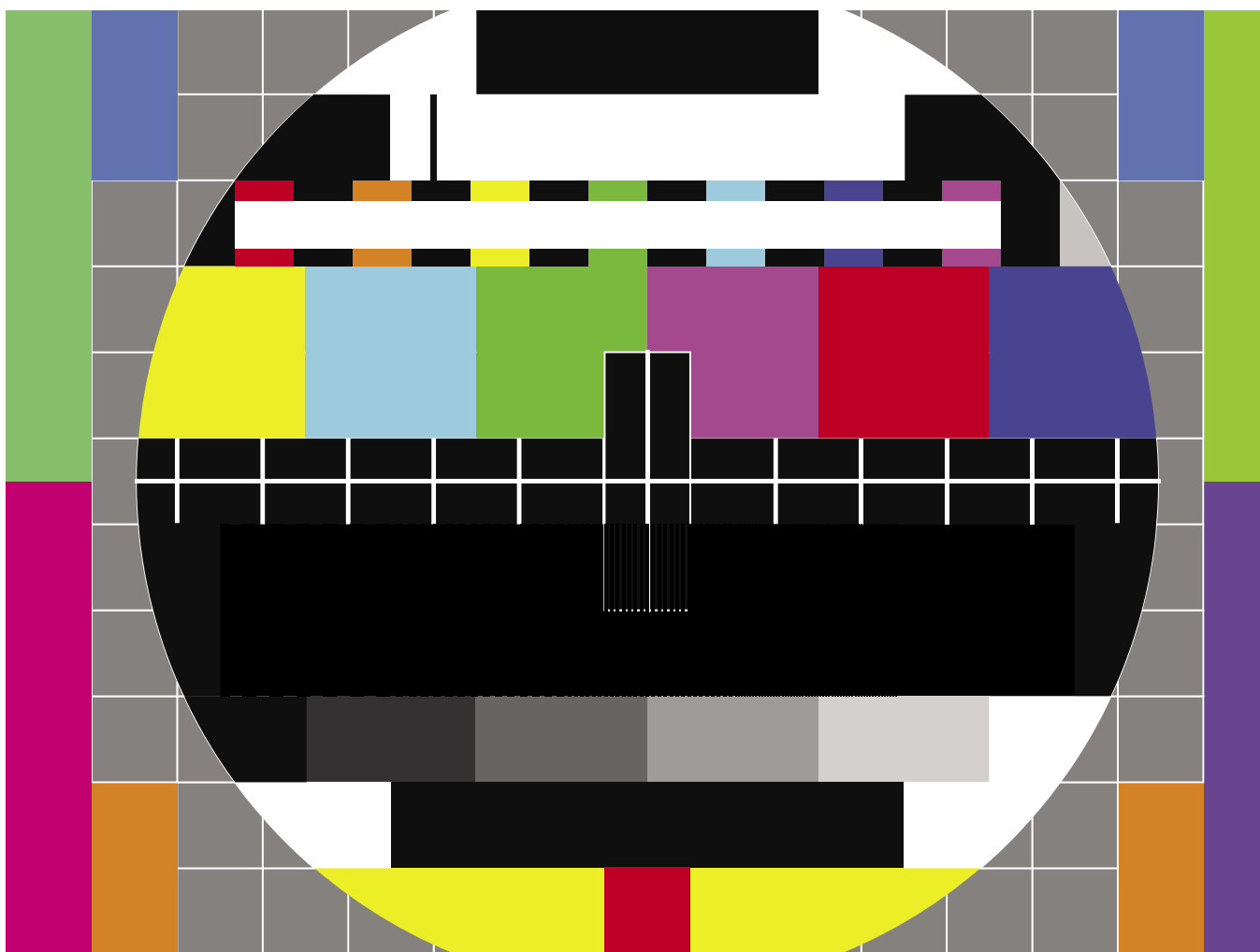
## ▲ Kompaktowy oscyloskop 200kHz FNIRSI DSO152

FNIRSI DSO152 to niezwykle praktyczny i ekonomiczny oscyloskop kieszonkowy, dedykowany branży serwisowej oraz edukacji badawczej



## ▲ Tester rezystancji wewnętrznej akumulatorów FNIRSI HRM-10

HRM-10 FNIRSI to wielofunkcyjne urządzenie przeznaczone do pomiaru i analizy stanu baterii. HRM10 to przenośny, wysoko precyzyjny tester rezystancji wewnętrznej baterii, który dostarcza Ci niezbędnych informacji



# Historia i technologia wyświetlaczy wideo, część 1

**W dwuczęściowym artykule omówimy historię i technologię wyświetlaczy wideo, od dysku Nipkowa i najwcześniejszych ekranów CRT (kineskopowych), po najnowsze wyświetlacze z pikselami kwantowymi. Skupimy się na dwuwymiarowych wyświetlaczach zdolnych do wyświetlania wideo, a nie na prostych wyświetlaczach alfanumerycznych ani technologii obrazowania 3D.**

W pierwszej części zostanie omówiona historia technologii wyświetlaczy do momentu wprowadzenia wyświetlaczy LCD (wyświetlaczy ciekłokrystalicznych) w latach 80. XX. wieku, które obecnie dominują na rynku (choć pojawiają się nowi gracze, tacy jak OLED). Podobnie jak w innych obszarach technologii, w ciągu ostatnich 150 lat obserwowaliśmy wiele rozwiązań innowacyjnych, a postęp trwa nadal.

W drugiej części zostaną omówione wszystkie najnowsze technologie, od wyświetlaczy

LCD po OLED, wyświetlacze z pikselami kwantowymi, wyświetlacze microLED, wyświetlacze EL, DLP, E Ink i inne.

## Dysk Nipkowa

W 1843 roku szkocki wynalazca Alexander Bain opracował pierwsze urządzenie, które umożliwiało zdalne przysyłanie obrazów, wysyłając je telegraficznie za pomocą swojego „elektrycznego telegrafu drukarskiego”. Jednak urządzenie to, podobnie jak „telegraf obrazowy” autorstwa Fredericka

Bakewella z 1848 roku, nie były opłacalne ze względu na bardzo niską jakość obrazu.

Pierwszą komercyjną maszyną faksującą był pantelegraf, wynaleziony przez włoskiego fizyka Giovanniego Caselliego w 1861 roku. Umożliwiał on przysyłanie nieruchomych obrazów, ale nie potrafił przysyłać obrazów ruchomych.

Prawdopodobnie pierwszym urządzeniem wyświetlającym obraz, zdolnym, przynajmniej teoretycznie, do wyświetlania ruchomych obrazów był dysk Nipkowa (**rysunek 1**), wynaleziony w 1883 roku i opatentowany

w 1884 roku. Składał się z obracającego się dysku ze wzorem spiralnych otworów, który mógł być używany zarówno do generowania obrazu przystosowanego do transmisji radiowej lub przewodowej, jak i do odtwarzania obrazu za pomocą innego zsynchronizowanego dysku na końcu odbiorczym.

Zaletą tego pomysłu była podobna konstrukcja urządzenia obrazującego i odbierającego. Wykorzystywany był prosty system obrazowania, wymagający jedynie czujnika światła i modulacji źródła światła. Poza tym miał wysoką rozdzielczość dla każdej linii skanowania.

Wady obejmowały konieczność utrzymywania synchronizacji dysków i praktyczny limit liczby otworów, które mogły być na dysku umieszczone, ograniczając liczbę linii rozdzielczości, zwykle w zakresie 30...100. Eksperymentalnie stosowano jednak do 200 linii.

Ponadto linie skanowania obrazów były zakrzywione ze względu na praktyczne ograniczenia rozmiaru dysku, a tworzone obrazy były małe. Na przykład dysk o średnicy 30...50 cm dawał obraz wielkości znaczka pocztowego.

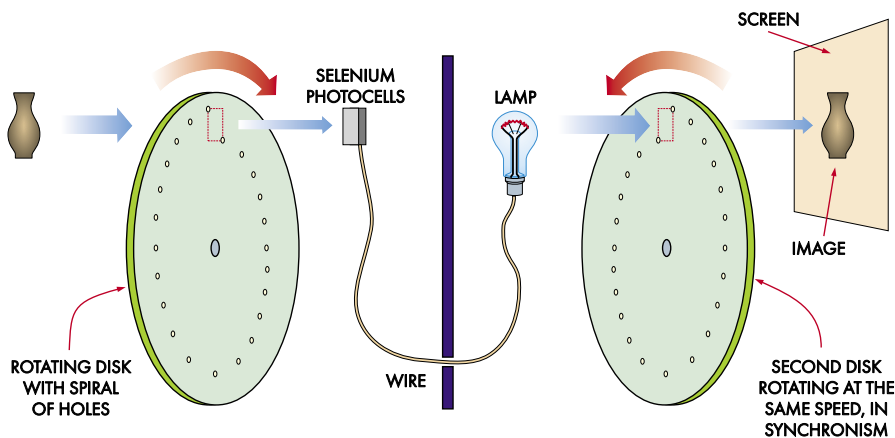
W 1885 roku Henry Sutton z Ballarat Victoria zaprojektował mechaniczny aparat telewizyjny do oglądania na żywo Melbourne Cup w Ballarat. Niestety, nigdy nie zbudował tego urządzenia, ponieważ linie telegraficzne, które proponował użyć nie miały wystarczającej przepustowości do transmisji sygnału. Radio, które miało wymaganą przepustowość, nie było jeszcze wystarczająco praktyczne. Autor nazwał urządzenie „Telephone” i opublikował jego plany w 1890 roku (**rysunek 2**). Wykorzystywało ono dysk Nipkova, fotokomórkę selenową i efekt Kerra (zmiana współczynnika załamania światła materiału w odpowiedzi na przyłożone pole elektryczne).

Dysk Nipkova był istotnym krokiem w kierunku wynalezienia praktycznego telewizora mechanicznego, którego jeden z pierwszych egzemplarzy został zademonstrowany przez Johna Logie Bairda w październiku 1925 roku.

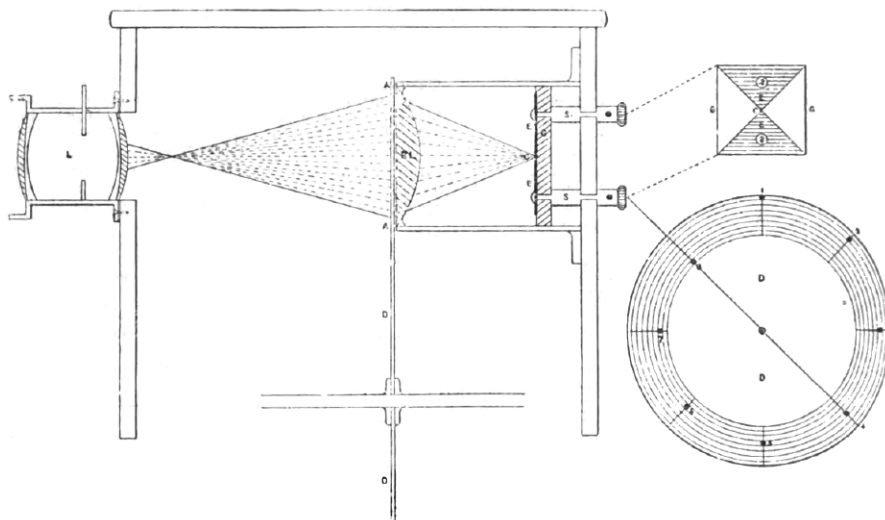
Co ciekawe, koncepcja dysku Nipkova jest do dziś wykorzystywana w jednej z odmian potężnego urządzenia do obrazowania optycznego zwanego mikroskopem konfokalnym.

## Natychmiastowa transmisja ruchomego obrazu

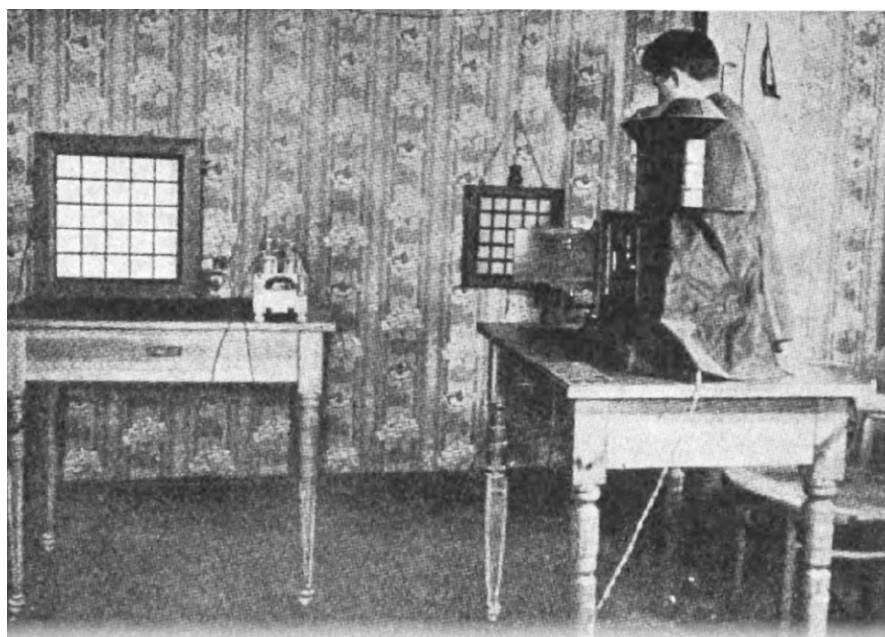
W 1909 roku Niemiec Ernst Ruhmer wynalazł wczesny system telewizyjny (**rysunek 3**). Do wykrywania obrazu został zastosowany układ komórek selenowych. Za pomocą nie do końca ujawnionej metody, modulował on intensywność światła odpowiednich części układu urządzenia wyświetlającego.



Rysunek 1. Reprodukacja obrazów za pomocą dysków Nipkova



Rysunek 2. Nigdy nie skonstruowany aparat „Telephone” australijczyka Henry’ego Suttona z 1885 r.; powieliłmy tylko nadajnik. Z *Telegraphic Journal and Electrical Review*, 7 listopada 1890, str. 550 (<https://hdl.handle.net/2027/mdp.39015012327071>)



Rysunek 3. Wczesny system telewizyjny Ernsta Ruhmera z 1909 roku z matrycą obrazującą 5×5 ogniw selenowych i matrycą odbierającą modulowane światło 5×5. Źródło: *Literary Digest*, 11 września 1909, str. 385 (<https://hdl.handle.net/2027/mdp.39015031441952>)



Rysunek 4. Oryginalny kineskop Brauna z zimną katodą z 1897 r. Z Eugen Nesper, 1921, Handbuch der Drahtlosen Telegraphie und Telephonie, Julius Springer, Berlin, str. 78

Urządzenie demonstracyjne miało matrycę 5×5 zdolną jedynie do wyświetlania prostych kształtów i było niewiarygodnie drogie ze względu na wysoki koszt ogniw selenowych. Każde praktyczne urządzenie zawierające na przykład 4000 ogniw byłoby nieracjonalnie drogie.

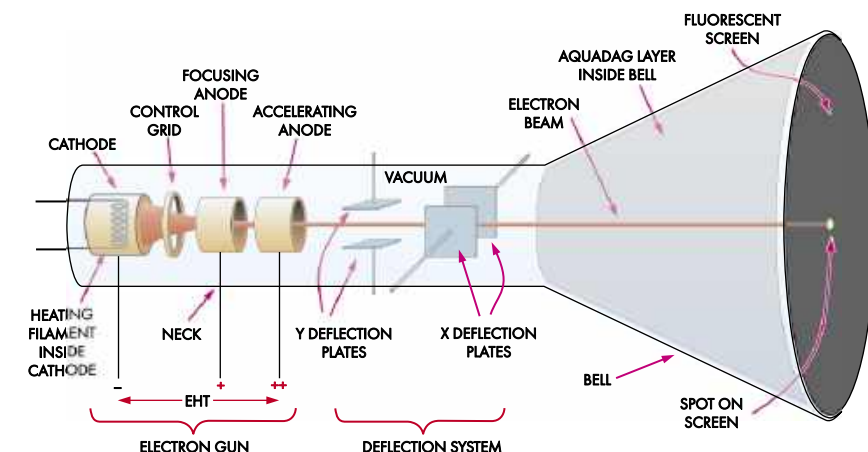
Następnie Francuzi Georges Rignoux i A. Fournier opracowali system zdolny do wyświetlania matrycy 8×8, wystarczającej do wyświetlania liter alfabetu. Mógł on przesyłać kilka pełnych obrazów na sekundę.

Były to niezwykle nowoczesne koncepcje, porównywalne z dzisiejszymi urządzeniami do obrazowania, choć w znacznie niższych rozdzielczościach.

## Kineskop (CRT)

Zdecydowanie najbardziej znanym urządzeniem wyświetlającym XX. wieku była kineskopowa lampa elektronopromieniowa, szeroko stosowana do wyświetlania obrazów telewizyjnych.

Promienie katodowe i niektóre z ich właściwości zostały odkryte wcześniej, ale niemiecki



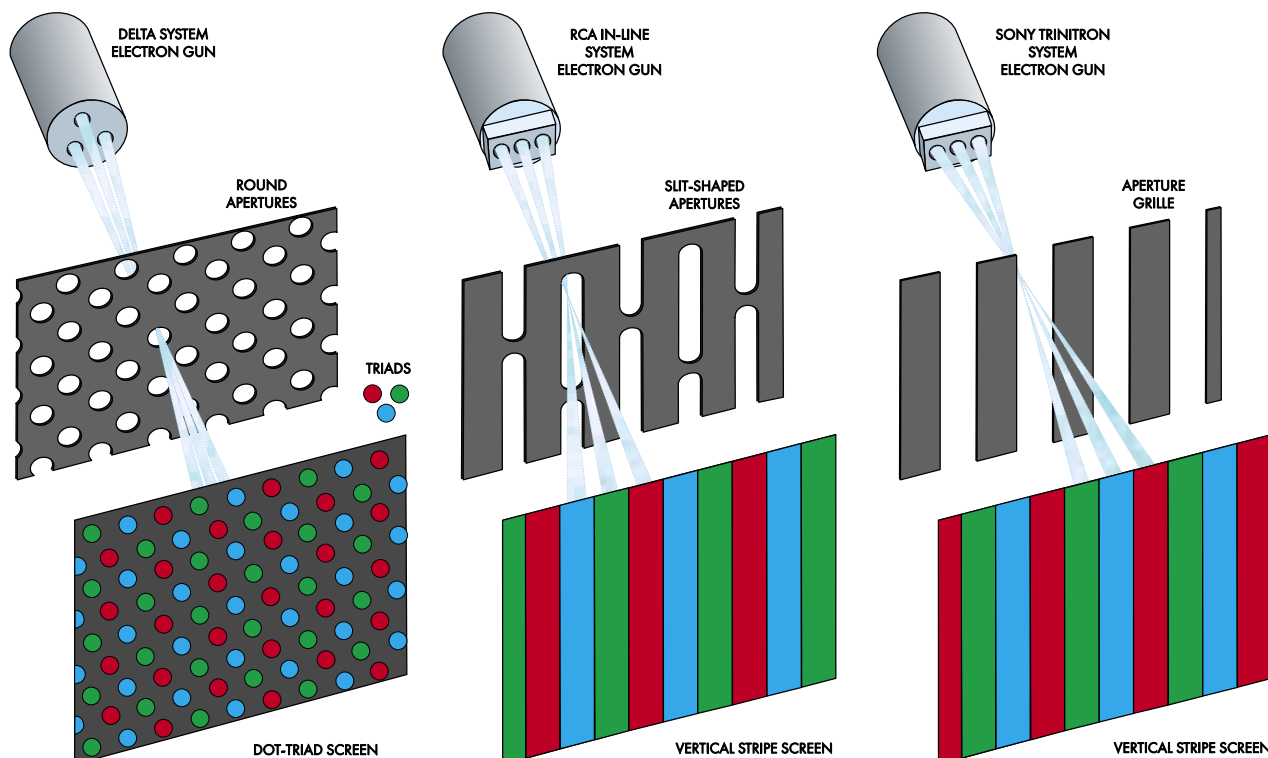
Rysunek 5. Typowy monochromatyczny monitor kineskopowy z elektrostatycznymi płytkami odchyłającymi, stosowanymi standardowo w oscyloskopach. Większość telewizorów wykorzystywała magnetyczne cewki odchyłające na zewnątrz szyjki kineskopu zamiast wewnętrznych płytek odchyłających. EHT oznacza ekstremalnie wysokie napięcie. W kolorowym wyświetlaczu znajdują się trzy działa elektronowe i maska cienia

fizyk Karl Ferdinand Braun wynalazł kineskop w 1897 roku (rysunek 4) i jako pierwszy pomyślał, że można go wykorzystać jako wyświetlacz. W przeciwieństwie do podgrzewanej katody spotykanej w nowocześniejszych urządzeniach, w tym urządzeniu zastosowano katodę zimną.

## Krótką oś czasu głównych zmian w technologii CRT:

- 1876: Eugen Goldstein upowszechnił termin „promienie katodowe”.

- 1897: Lampa Brauna, pierwszy kineskop, została opracowana jako zmodyfikowana lampa Crookesa z ekranem pokrytym luminoforem.
- 1908, 1911: Alan Archibald Campbell-Swinton pisze o „dalekim widzeniu elektrycznym” przy użyciu kineskopu Brauna.
- 1922: John Bertrand Johnson i Harry Weiner Weinart opracowują komercyjny kineskop z gorącą katodą.
- 1926: Kenjiro Takayanagi demonstruje telewizor CRT z 40 liniami.



Rysunek 6. Geometryczne rozmieszczenie dział elektronowych i masek w celu zapewnienia, że każda wiązka koloru trafi we właściwy luminofor. Były na to trzy rozwiązania, z których każde było lepsze od poprzedniego

- **1927:** Takayanagi zwiększa rozdzielczość do 100 linii.
- **1929:** Władimir K. Zworykin wprowadza termin „kineskop”.
- **1932:** Radio Corporation of America (RCA) oznacza lampę elektronopromieniową terminem Cathode Ray Tube – CRT.
- **lata 30. XX wieku:** Allen B. DuMont stworzył pierwsze kineskopy, które mogły działać przez tysiące godzin.
- **1934:** Pierwsze telewizory kineskopowe zostają wyprodukowane przez niemiecką firmę Telefunken.
- **1950:** RCA wprowadza termin „kineskop” do domeny publicznej.
- **1954:** RCA produkuje pierwsze kineskopy kolorowe.
- **1957:** William Ross Aiken otrzymuje amerykański patent 2,795,731 na płaskie kineskopy.
- **1958:** Aiken otrzymuje kolejny amerykański patent (2,837,691) na płaski kineskop.
- **1968:** Wprowadzenie płaskiego kineskopu Sony Trinitron.
- **1987:** Opracowano kineskopy z płaskimi ekranami do monitorów komputerowych.
- **lata 90.:** Sony wypuszcza na rynek monitory CRT o wysokiej rozdzielczości.

Schemat typowego kineskopu pokazano na **rysunku 5**. Jest to lampa próżniowa zawierająca działło elektronowe (katodę czyli elektrodę ujemną) generującą wiązkę elektronów, którą można kierować zarówno w kierunku X (poziowym), jak i Y (pionowym).

**Start-to-finish production operations for making RCA tricolor kinescopes in plant at Lancaster, Penna. Three-gun type 15GP22, now available to manufacturers of home receivers, employs electrostatic focusing and magnetic deflection**



**Rysunek 7. Stworzenie pierwszego komercyjnego kolorowego kineskopu w 1954 r. Źródło: Early Television Museum and Foundation ([www.earlytelevision.org](http://www.earlytelevision.org))**

Działło elektronowe zawiera żarnik, który podgrzewa katodę emitującą elektrony. Siatka steruje przepływem elektronów między katodą a anodą przyspieszającą, regulując tym samym jasność/intensywność. Do anody przyspieszającej przykładane jest napięcie do 20 kV względem katody, co powoduje, że elektrony tworzą wąską wiązkę poruszającą się w kierunku ekranu. Druga anoda utrzymuje skupienie wiązki.

Po tym, jak wiązka opuści zespół działła elektronowego (grzałka, katoda, siatka sterująca, anoda przyspieszająca i anoda ogniskująca), jest ona odchylana i w ten sposób kierowana w celu utworzenia obrazu. Jest to możliwe dzięki cewkom wytwarzającym pole magnetyczne lub elektrostatycznym płytom odchylającym generującym pole elektryczne. Tak czy inaczej, istnieją dwie pary cewek lub płytek do odchylania poziomego i pionowego.

Wiązka elektronów uderza w ekran pokryty luminoforem, emitując światło. W przypadku kolorowego ekranu istnieją trzy wiązki elektronów i trzy różne kolory luminoforu (ułożone jako kropki lub paski), a wiązka elektronów dla każdego koloru uderza tylko w odpowiedni kolor luminoforu.

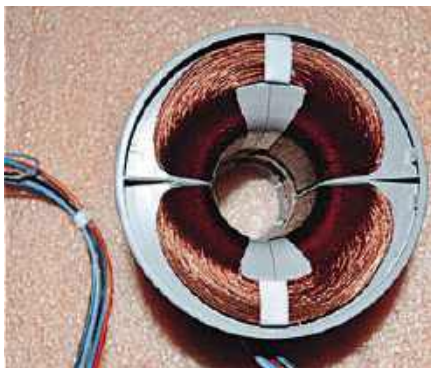
Aby upewnić się, że każda wiązka trafia we właściwy luminofor, stosowana jest specjalna maska (tzw. maska cienia), a każda kolorowa wiązka elektronów jest nieznacznie przesunięta względem pozostałych (**rysunek 6**).

W kolorowych kineskopach CRT wypróbowano wiele rozwiązań mających na celu doprowadzenie do tego, by wiązka elektronów uderzała we właściwy kolor luminoforu. Jednak koncepcja maski cienia firmy RCA, wprowadzona w 1950 roku, doprowadziła do porzucenia wszystkich innych kierunków badań nad kineskopami kolorowymi. Pomysł RCA okazał się najlepszy. Firma RCA wprowadziła na rynek pierwszą kolorową lampę (15GP22) w roku 1954 (**rysunek 7**).

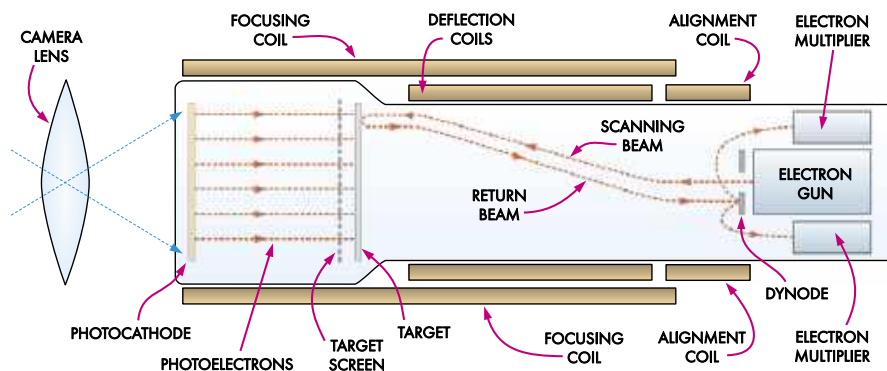
Maski cieni są wykonywane w procesie litograficznym zwanym obróbką fotochemiczną. Koncepcja maski cienia RCA była główną stosowaną do czasu wprowadzenia przez Sony maskownicy przysłaniającej w 1968 roku, która służy temu samemu celowi co maska cienia, ale której zastosowano długie szczeliny zamiast otworów lub małych szczelin.

Począwszy od późnych lat 60., zestawy inne niż Trinitron używały prostokątnych luminoforów i prostokątnych otworów w masce cienia, zamiast triady kropek luminoforu i okrągłych otworów.

Można zastanawiać się, gdzie trafiają wszystkie elektrony po uderzeniu w luminofor. Wnętrze kineskopu (część między



Rysunek 8. Zespół odchyłania magnetycznego (jarzmo) z telewizora CRT. Źródło: JHCOILS



Rysunek 9. Kineskop kamery wideo zwany orticonem obrazu, powszechnie stosowany w amerykańskiej telewizji w latach 1946–1968

szyjką a ekranem) pokryte jest grafitową warstwą przewodzącą o nazwie Aquadag. Zbiera ona elektrony i stanowi część anody. Pomaga również utrzymać jednolite pole elektryczne wewnątrz lampy.

Połączeniem elektrycznym z tą częścią lampy jest duży, wystający przewód przymocowany z boku kineskopu.

## Odchylenie elektryczne vs magnetyczne

Układ pokazany na rysunku 5 ma elektrostatische płytki odchyłające, takie jak stosowane w oscyloskopie.

W większości telewizorów (z wyjątkiem kilku wczesnych typów z małymi lampami) zamiast płytek instalowane są cewki wytwarzające pole magnetyczne. Magnetyczne cewki odchyłające umożliwiają większy kąt odchylenia, a tym samym płyszają lampę, co jest stosowane w telewizorach i monitorach komputerowych

(rysunek 8). Umożliwiają one również przepływ większego prądu wiązki zapewniającego jaśniejszy obraz.

W tradycyjnych oscyloskopach CRT (CRO) krótka lampa nie była uważana za konieczną, ponieważ obraz był mały, więc lampa również była mała. Co jednak ważniejsze, układy były prostsze, ponieważ płytki odchyłania pionowego mogły być sterowane bezpośrednio przez wzmacnione przebiegi sygnału.

Ponadto systemy odchyłania mogą szybciej reagować na sygnały o wysokiej częstotliwości (wielu megaherców), ponieważ elektrostatische płytki odchyłające stanowią jedynie niewielkie obciążenie pojemnościowe w porównaniu z wysoce indukcyjnym obciążeniem magnetycznych cewek odchyłających.

W telewizorze lub monitorze komputerowym obraz jest tworzony poprzez skanowanie linii po linii, od góry do dołu, w tak zwanym wzorze rastrowym. Dzieje się to tak szybko, że dla człowieka nie jest widoczne. Na stronie <https://youtu.be/3BJU2drdtCM> jest umieszczony doskonały film, który wykorzystuje szybkie zdjęcia do zademonstrowania sposobu skanowania rastra.

Z kolei w oscyloskopie wiązka jest cyklicznie przesuwana od lewej do prawej,

natomiast przesuwanie w górę i w dół jest zgodne z przyłożonym napięciem sygnału.

Oprócz wyświetlania obrazu wideo i oscyloskopów, ekrany CRT były stosowane w radarach, monitorach pracy serca, a w niektórych przypadkach stanowiły formę pamięci komputerowej.

Od ich powstania do połowy lat 90. były jedyną praktyczną i powszechną formą wyświetlacza wideo. Ekrany LCD były komercyjnie dostępne w laptopach od początku lat 90., ale wypadały bardzo słabo w porównaniu z kineskopami, nadrabiając zaległości dopiero w końcowych latach XX wieku i pierwszych latach XXI. wieku.

Płaskie telewizory LCD wyprzedziły telewizory CRT po raz pierwszy w roku 2007, a w tym samym roku Sony zaprzestało produkcji słynnej marki CRT Trinitron.

Na przestrzeni lat wyprodukowano wiele odmian kineskopów:

- Niektóre z nich mogą zachowywać obraz do momentu jego skasowania, np. w niektórych oscyloskopach.
- Istniały wyświetlacze wektorowe, które tworzyły obrazy przy użyciu linii narysowanych od punktu do punktu, a nie we wzorze rastrowym. Były one używane we wczesnych monitorach komputerowych

| Tabela 1. Największe komercyjne CRT |                          |
|-------------------------------------|--------------------------|
| 1938                                | przekątna 51 cm/20 cali  |
| 1955                                | przekątna 53 cm/21 cali  |
| 1985                                | przekątna 89 cm/35 cali  |
| 1989                                | przekątna 110 cm/43 cali |



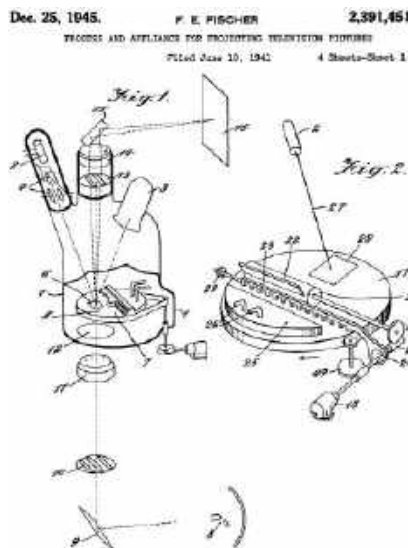
Rysunek 10(a) i (b). Kamera kineskopowa Pye Mk III, po raz pierwszy sprzedana w 1952 roku i używana do telewizyjnych transmisji testowych w Australii oraz do relacjonowania Igrzysk Olimpijskich w Melbourne w 1956 roku. Była wyposażona w silniki i mogła być zdalnie sterowana, w tym ustawiać ostrość, zmieniać obiektyw, a także przechylać i obracać za pomocą odpowiednich przystawek. Źródło: Australian Centre for the Moving Image, [siliconchip.au/link/abf9](http://siliconchip.au/link/abf9)

do projektowania wspomagane komputerowo (CAD), w niektórych grach zręcznościowych i w domowym systemie gier Vectrex.

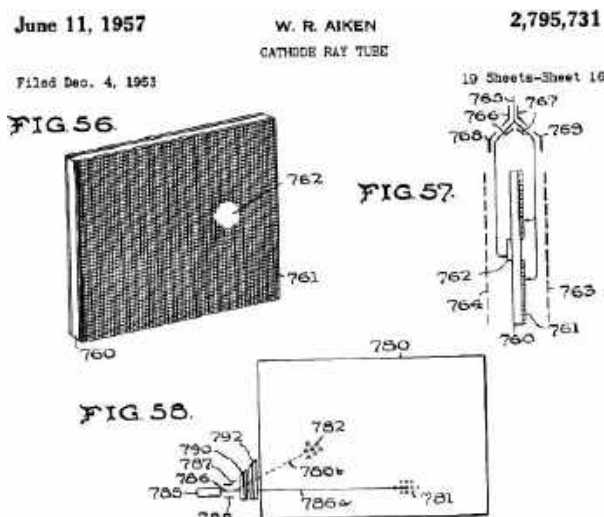
- Kineskopy projekcyjne tworzyły obraz na odległym ekranie pasywnym.
- Lampa do przechowywania danych z końca lat 40. znana jako lampa Williamsa przechowywała dane binarne, zwykle 256...2560 bitów.
- Ulubiony wskaźnik służący do strojenia, magiczne oko, było używane w niektórych radiach lampowych od 1935 roku do lat 60. XX wieku.

Pod koniec ery telewizorów kineskopowych, przez pewien czas konkurowały one jeszcze z telewizorami LCD i plazmowymi. Telewizory kineskopowe z płaskim ekranem zdążyły zyskać popularność, ponieważ początkowo były znacznie tańsze w produkcji. Ostatecznie cena alternatywnych wyświetlaczy spadła, a nieporęczne i ciężkie kineskopy odeszły w zapomnienie.

Obecnie Thomas Electronics ([www.thomaselectronics.com](http://www.thomaselectronics.com)) nadal produkuje i naprawia kineskopy CRT jako zamienniki dla specjalistycznego sprzętu wojskowego i lotniczego. Na tych rynkach często bardziej opłacalne jest utrzymanie starej technologii niż modernizacja platform za pomocą nowych ekranów LCD itp. W takich przypadkach



Rysunek 12. Schemat urządzenia Eidophor z oryginalnego patentu USA (<https://patents.google.com/patent/US2391451A/en>)



Rysunek 11. Proponowany kolorowy płaski ekran radarowy CRT autorstwa Williama Aikena z 1957 r. (<https://patents.google.com/patent/US2795731A/en>)

koszt produkcji nie ma znaczenia, ponieważ koszty badań i rozwoju dla zamienników byłyby ogromne.

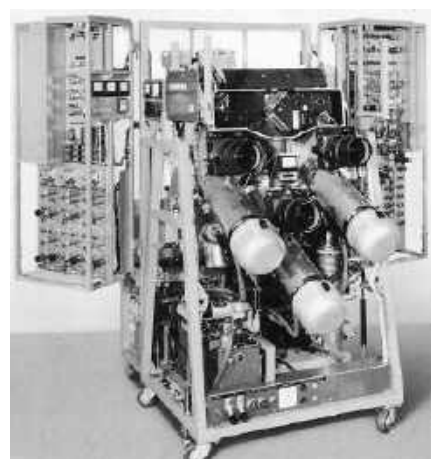
Niektórzy pasjonaci gier komputerowych nadal używają telewizorów CRT, ponieważ mogą one mieć szybszy czas reakcji niż wiele wyświetlaczy LCD, a niektórzy wolą wygląd linii skanowania. Niektóre stare gry wideo (takie jak klasyczne gry zręcznościowe) zostały zaprojektowane specjalnie do oglądania na kineskopach, a można tylko życzyć powodzenia w znalezieniu najnowszego telewizora LCD ze złączem S-Video lub SCART, co jest niezbędne do natywnego wyświetlania wyższych rozdzielczości. Kineskopy poprawnie wyświetlają również nietypowe, przestarzałe rozdzielczości, takie jak 256×224, używane przez stare systemy Nintendo.

## Aiken CRT

William Ross Aiken podjął wczesną próbę zaprojektowania płaskiego kineskopu z działem elektronowym umieszczonym z boku, a nie z tyłu (rysunek 11). W latach 1957 i 1958 uzyskał on amerykańskie patenty na te projekty. Niestety, doszło do sporów patentowych i rozwój został zatrzymany. Po wygaśnięciu patentów pomysł był dalej rozwijany przez Sinclair Electronics i RCA.

## CRT jako kamera

Inny typ kineskopu nie wyświetlał obrazu, ale był używany we wczesnych kamerach telewizyjnych od lat 30. do 80. ubiegłego wieku. Następnie kineskopy kamer wideo zostały zastąpione przez czujniki obrazu CCD (Charge-Coupled Device), wprowadzone do technologii nadawczej w 1984 roku, a następnie przez czujniki CMOS (rozwiązanie CCD).



Rysunek 13. Eidophor model EP 6 bez ostoi, typ używany przez NASA w sali kontroli operacji misji w erze Apollo i później. Źródło: Szwajcarskie Muzeum Narodowe

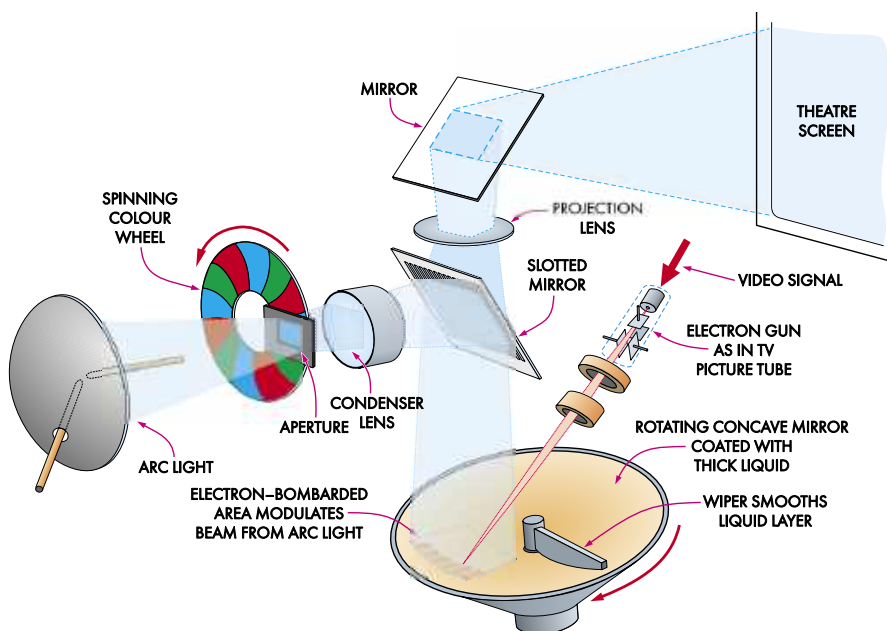
Zasada działania kineskopu CRT w roli kamery wideo polega na tym, że promień katody jest skanowany w poprzek obrazu tworzonego przez fotokatodę. Powracający promień katodowy jest modulowany zgodnie z intensywnością obrazu tworzonego przez fotokatodę – (rysunki 9 i 10). Fotokatoda jest związkem światłoczułym, który emituje elektrony po uderzeniu fotonów ze źródła światła w wyniku efektu fotoelektrycznego.

## Eidophor

Niewiele osób słyszało o projektorze wideo Eidophor. Został on wynaleziony przez szwajcarskiego naukowca Fritza Fischera w 1939 roku, a amerykański patent na niego został przyznany w roku 1945 (rysunek 12 i [siliconchip.au/link/abf7](http://siliconchip.au/link/abf7)).



Rysunek 14. Obraz Eidophor (w środku) w Mission Operations Control Room w Houston, podczas misji Apollo 11 22 lipca 1969 roku. Źródło: NASA, Image id=S69-39815



Rysunek 15. Ścieżka optyczna Eidophor. Oryginalne źródło: [www.ngzh.ch/media/njb/Neujahrsblatt\\_NGZH\\_1961.pdf](http://www.ngzh.ch/media/njb/Neujahrsblatt_NGZH_1961.pdf)

Eidofory były używane podczas wydarzeń publicznych na dużą skalę, w projektorach filmowych i wideo, a przede wszystkim przez NASA w sali kontroli misji podczas misji Apollo.

NASA korzystała z 34 projektorów Eidophor w latach 1965...1969 (rysunki 13 i 14). Mimo swojej złożoności miały one wskaźnik gotowości na poziomie 99,9%. Kosztowały łącznie około 85...90 milionów dzisiejszych dolarów.

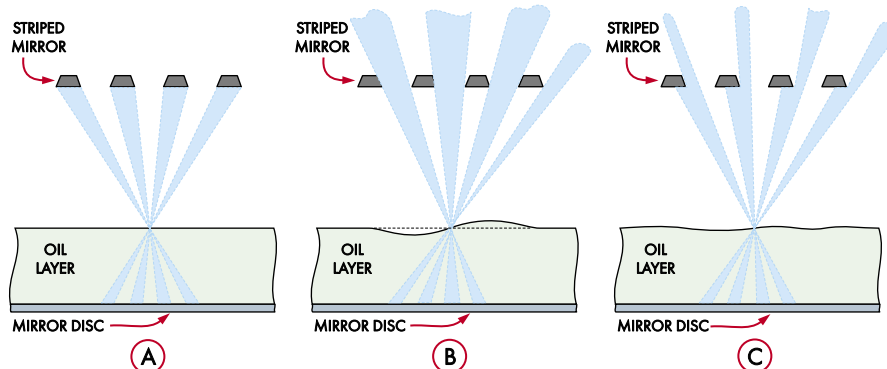
Eidophor był dużym, złożonym, drogim w zakupie i eksploatacji urządzeniem, ale był niezawodny i zapewniał najlepsze wyświetlane obrazy wideo w tamtym czasie. Zasadę ich działania opisujemy niżej.

Dysk lustrzany w komorze próżniowej jest pokryty warstwą oleju o grubości około 14  $\mu\text{m}$ . Wiązka elektronów skanuje powierzchnię oleju w podobny sposób, jak wiązka elektronów na ekranie kineskopu skanuje luminofor.

Ładunek wprowadzony do warstwy oleju powoduje jej deformację w wyniku działania sił elektrostatycznych. Wiązka światła z potężnej lampy łukowej świeci na lustro pokryte olejem, a odbite światło jest rzutowane przez układ optyczny na płaszczyznę obrazu za pomocą lustra paskowego lub podobnego układu (rysunek 15).

Odchylenie wiązki światła w celu utworzenia obrazu jest generowane przez dyfrakcję optyczną lub załamanie światła, gdy przechodzi ono przez cienką warstwę oleju o różnej grubości.

Światło rzucane na lustro pokryte olejem przechodziło przez lustro szczelinowe z naprzemiennymi przezroczystymi i lustrzanymi paskami. W rezultacie światło odbijające się od lustra głównego w obszarach, na które nie oddziałuje wiązka elektronów, odbija się z powrotem do szczelin i jest blokowane, natomiast obszary, w których olej jest zakłócany, powodują, że odbite światło omija szczeliny



Rysunek 16. Działanie zwierciadła paskowego Eidophora. (A) Światło jest odbijane z powrotem bez obrazu i żadne światło nie dociera do płaszczyzny obrazu. (B) Przy silnym obrazie całe światło przechodzi przez przezroczyste paski i jest rzutowane na płaszczyznę obrazu. (C) Przy słabym obrazie część światła jest blokowana, ale nie całe. Oryginalne źródło: [www.ngzh.ch/media/njb/Neujahrsblatt\\_NGZH\\_1961.pdf](http://www.ngzh.ch/media/njb/Neujahrsblatt_NGZH_1961.pdf)



i przechodzi na ekran projekcyjny. Ekran projekcyjny pozostaje więc ciemny w obszarach, w których wiązka elektronów jest odcięta i jest tym jaśniejszy, im wyższa jest intensywność wiązki elektronów w tym obszarze. W przypadku części ekranu, które nie są w pełni jasne lub ciemne, część światła jest odbijana i blokowana, natomiast pozostała część światła dociera do ekranu projekcyjnego. Umożliwia to stopniowanie poziomów intensywności w celu wygenerowania obrazu, jak pokazano na rysunku 16.

Aby usunąć już wyświetlony obraz z oleju w celu przygotowania kolejnego, lustrzana tarcza jest obracana do elektrody, która neutralizuje ładunek cząsteczek oleju, wygładza powierzchnię i zeruje ją w celu przygotowania kolejnego obrazu.

Wczesne Eidophory były monochromatyczne, natomiast późniejsze wersje mogły wyświetlać kolorowe obrazy za pomocą koła kolorów lub trzech projektorów z kolorowymi filtrami.

Fascynujący film o Eidoforach z 1944 roku z angielskimi napisami zatytułowany „Eidophor: Die bildspendende Flüssigkeit (1944)” można obejrzeć na stronie [https://www.youtube.be/w\\_9NhiGekII](https://www.youtube.be/w_9NhiGekII).

## Ekran wyświetlacz NASA misji Apollo

Wiele osób zastanawiało się, w jaki sposób NASA skonfigurowała gigantyczne ekrany w Mission Operations Control Room („Mission Control”) w Johnson Space Center w Houston w Teksasie podczas lądowania



◀ Rysunek 17. Zdjęcie z epoki Apollo przedstawiające salę kontroli misji NASA („Mission Control”), pokazujące duże ekrany. Po lewej i prawej stronie znajdowały się dwa ekrany o wymiarach 10×10 stóp (3 m × 3 m), a pośrodku ekran o wymiarach 20×10 stóp (6 m × 3 m). Obraz wideo Eidophor można zobaczyć po prawej stronie, a obrazy graficzne pośrodku i po prawej stronie. Nie ma pewności co do dwóch obrazów po lewej stronie



Rysunek 18. Duże ekrany przednie kontroli misji NASA z czasów Apollo w późnych latach 60. i wczesnych 70.. Widok z pomieszczenia widokowego dla odwiedzających z tyłu sali kontroli operacji misji. Źródło: NASA (www.nasa.gov/sites/default/files/atoms/files/apollo\_mcc\_press\_release.pdf)

Apollo na Księżycu (rysunki 17 i 18). Niestety, dokumentacja na temat zastosowanej technologii jest bardzo skromna.

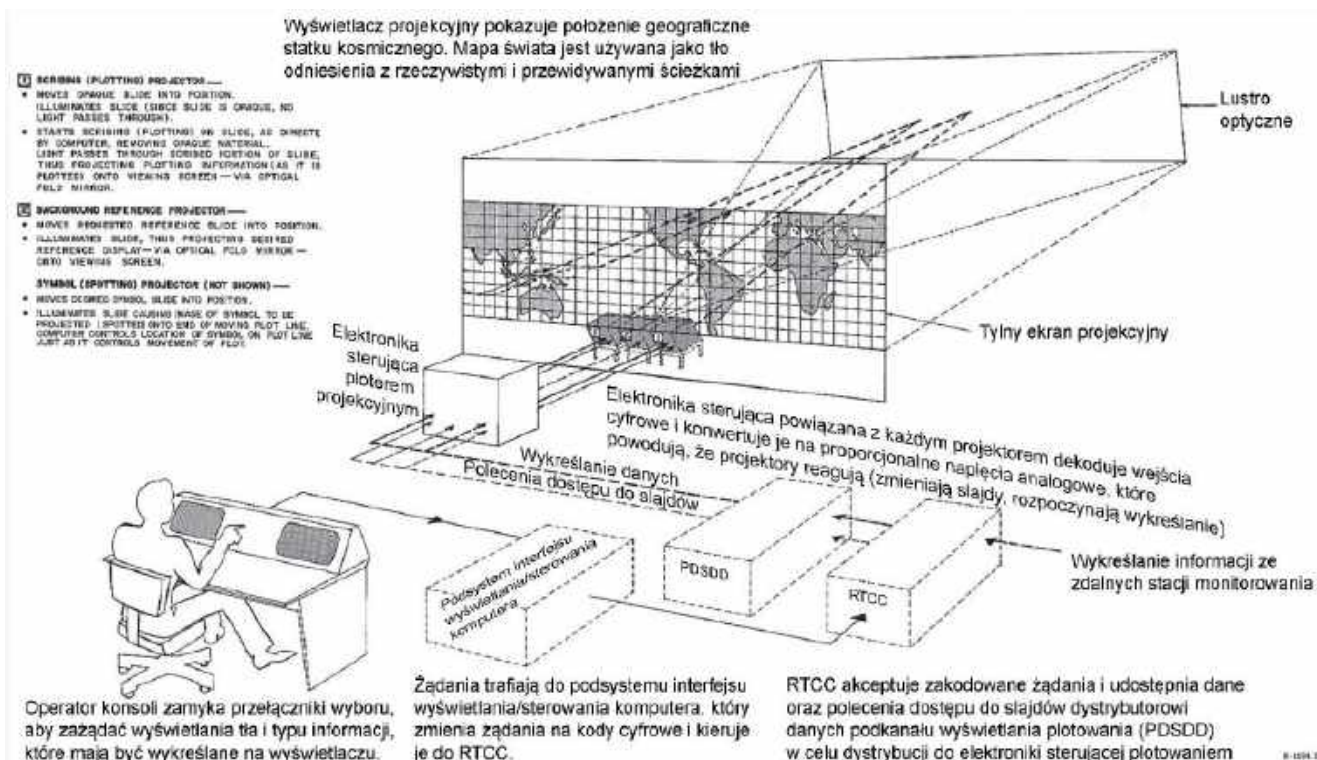
Były to prawdopodobnie pierwsze duże wyświetlacze wideo, które mogło w tamtym czasie zobaczyć tak wiele osób, i jedno z pierwszych, jeśli nie pierwsze, tak wielkie wyświetlacze wideo. Jak więc działały?

NASA używała zarówno projektorów graficznych, jak i projektorów wideo Eidophor. Opisaliśmy już, jak działały projektory Eidophor, więc teraz zajmiemy się specjalnym projektorom graficznym.

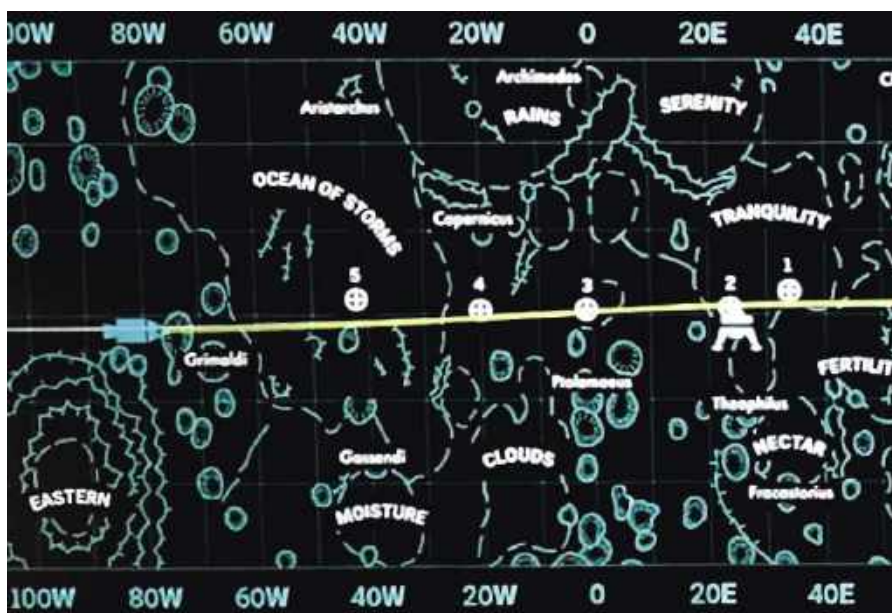
Projektory te zbadała dokładnie youtuberka Fran Blanche. Jej znakomity film zatytułowany „How Mission Control’s Big Displays Worked”

jest dostępny na stronie [https://youtu.be/N2v4kH\\_PsN8](https://youtu.be/N2v4kH_PsN8) i na pewno warto go obejrzeć.

Według tego filmu, system był używany do 1989 roku. Graficzne projektory slajdów wyświetlały mapy Ziemi i Księżyc, strony z podręczników i wszelkie inne materiały, które można było przechowywać na slajdach projektora. Odpowiednie slajdy mogły być



Rysunek 19. Jak działał system projekcji graficznej z ery Apollo w Kontroli Misji NASA. Sprzęt znajdował się za salą Kontroli Misji (zwaną The Pit) oraz w pokoju projekcji podsumowań lub „Jaskini Nietoperza”. Projektory wideo Eidophor nie są pokazane na tym schemacie



Rysunek 20. Obraz przedstawiający mapę Księżyca w tle, linię trajektorii, ikony pojazdów orbitalnych (moduł dowodzenia) i lądujących (LEM), a także inne ikony oznaczone od 1 do 5, przypuszczalnie odpowiadające różnym wydarzeniom związanym z lądowaniem. Został on wykonany z wielu slajdów na wielu projektorach, a kolory zostały wygenerowane przez filtry kolorów

wybijane za pośrednictwem komputera, spośród tych przechowywanych na podajniku (rysunek 19).

Projektory musiały wyraźnie wyświetlać obrazy w jasnym oświetleniu Mission Control. Oznaczało to, że wymagane było niezwykle silne oświetlenie. Ciepło niszczyłoby tradycyjne slajdy wykonane z poliestru. Zastosowano więc szklane slajdy z obrazami w metalowych powłokach. Obszary, w których metalu na slajdzie nie było przepuszczają całe światło, natomiast tam gdzie został naniesiony tak, że był nieprzezroczysty bez gradacji, podobnie jak miedź na płytce drukowanej, światło nie przechodziło.

Kolory były generowane za pomocą filtrów kolorów, a wiele slajdów można było nakładać na siebie z wielu projektorów. Możliwość nakładania slajdów była bardzo ważna. Było to dobre dla statycznych obrazów, ale jak dodano wykresy w czasie rzeczywistym lub dane orbitalne i trajektorie?

Dane orbitalne i trajektorie zostały wygenerowane przez komputery mainframe IBM 360 System 75 (<https://w.wiki/59xB>). Odbierały one dane telemetryczne i przekładały je na wykresy, które można było wyświetlać w czasie rzeczywistym.

Specjalne projektory wykresów pobierały dane z komputera. Wykreślały je za pomocą diamentowego lub podobnego rysika na płótnie X-Y, wpisując je w „pusty” (w pełni metalizowany) slajd, zarysowując linię w metalu. Wcześniej naniesione dane pozostawały do momentu włożenia nowego pustego szkła (rysunek 20). Ikony, takie jak statek kosmiczny, były również przesuwane pod kontrolą komputera, aby pokazać rzeczywistą pozycję statku kosmicznego.

NASA przywróciła do pierwotnego stanu oryginalną salę kontroli misji Apollo, która była używana do 1992 roku ([siliconchip.au/link/abf8](http://siliconchip.au/link/abf8)).



## Sinclair TV80/FTV1 Pocket TV

W 1983 roku Sinclair wypuścił telewizor kieszonkowy TV80 (znany również jako FTV1). Zastosowano w nim kineskop z odchylaniem elektrostatycznym z zamontowanym z boku działem elektronowym, podobnym do powyższego kineskopu Aiken (rysunki 21 i 22).

Był on komercyjną porażką, częściowo ze względu na podobne produkty wypuszczone przez Sony („Watchman”) z innymi producentami używającymi kineskopów, a później wyświetlaczy LCD. Zegarek telewizyjny Seiko LCD T001 został wprowadzony do sprzedaży w 1982 roku, a Casio LCD Pocket Television TV-10 (rysunek 23) w roku 1983.

Więcej informacji na temat TV80 można znaleźć w filmie „Doom on 1983 Sinclair FTV1 TV80 Mini Flat CRT & Teardown” na stronie <https://youtu.be/fEcs52IAI3E> i na blogu Australijczyka Davida L. Jonesa „EEVblog #554 – Sinclair FTV1 TV80 Flat Screen Pocket TV Teardown” na stronie <https://youtu.be/qCJPF6Ei3Vw>.

## Wyświetlacze plazmowe

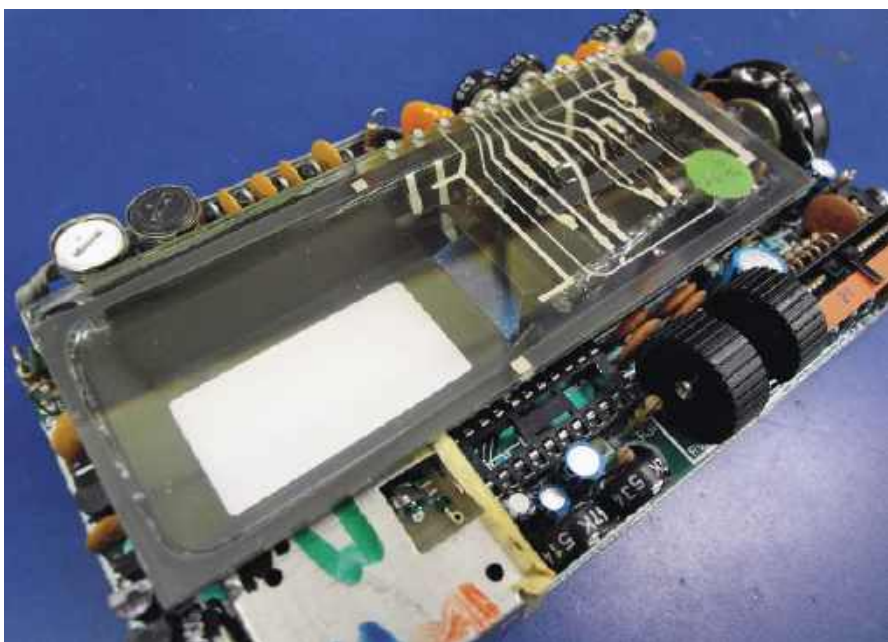
Wyświetlacze plazmowe były pierwszymi wyświetlaczami z płaskim ekranem o przekątnej ponad 80 cm (32 cale) i jako pierwsze zastąpiły wyświetlacze CRT, przynajmniej w przypadku większych rozmiarów. Po roku 2013 ich popularność przebiły wyświetlacze LCD. Wyświetlacze plazmowe są obecnie uważane za przestarzałe i zostały w większości zastąpione na rynku przez wyświetlacze OLED.

Węgierski inżynier Kálmán Tihanyi po raz pierwszy zaproponował wyświetlacz

### Powiązane linki:

- Strona Cathode Ray Tube: [www.crtsite.com](http://www.crtsite.com)
- Kineskopy były kiedyś przebudowywane. Film „The Craft of Picture Tube Rebuilding” jest spojrzeniem na ostatniego producenta kineskopów w USA – <https://youtu.be/W3G7b-Dc004>
- The 8-bit Guy opowiada o modyfikowaniu konsumenckiego telewizora CRT tak, by posiadał wejścia RGB do gier vintage i korzystania z komputerów vintage w filmie „Modding a consumer TV to use RGB input” – [https://youtu.be/DLz6pgvsZ\\_I](https://youtu.be/DLz6pgvsZ_I)
- THE LAST SCAN – Historia desperackiej walki o utrzymanie starych telewizorów przy życiu – [www.theverge.com/2018/2/6/16973914/tvs-crt-restoration-led-gaming-vintage](http://www.theverge.com/2018/2/6/16973914/tvs-crt-restoration-led-gaming-vintage)
- Fascynujący eksperyment z monochromatycznym panelem plazmowym – <https://youtu.be/Oj4tRnLKN6U>
- „vintageTek Demo of the 1930's 905 CRT” – <https://youtu.be/NBeOMSDPuT8>
- „Lampa elektronopromieniowa – jak to działa”, wojskowy film szkoleniowy 16 mm z 1943 r. – <https://youtu.be/GnZSopHjmYQ>
- „Mullard Made for Life Vintage Documentary” – <https://youtu.be/32yYfTVIzBE>
- „Budowa ceramicznego kineskopu Tektronix 1967” – <https://youtu.be/G0DCi5RPe94>

Rysunek 21. Płaski telewizor Sinclair TV80 CRT.  
Źródło: Wikimedia user Binarysequence, CC BY-SA 3.0

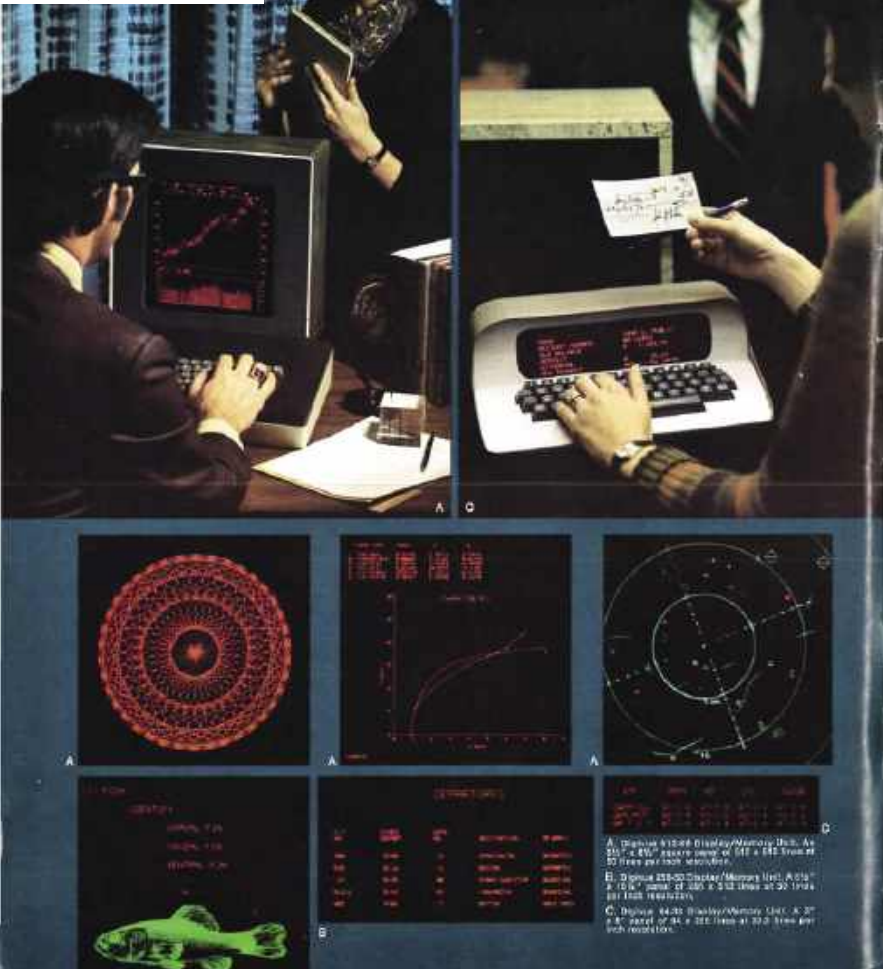


Rysunek 22. Płytką drukowaną Sinclair TV80, z obszarem podglądu kineskopu po lewej i działem elektrycznym po prawej. Źródło: Wikimedia user Binarysequence, CC BY-SA 3.0

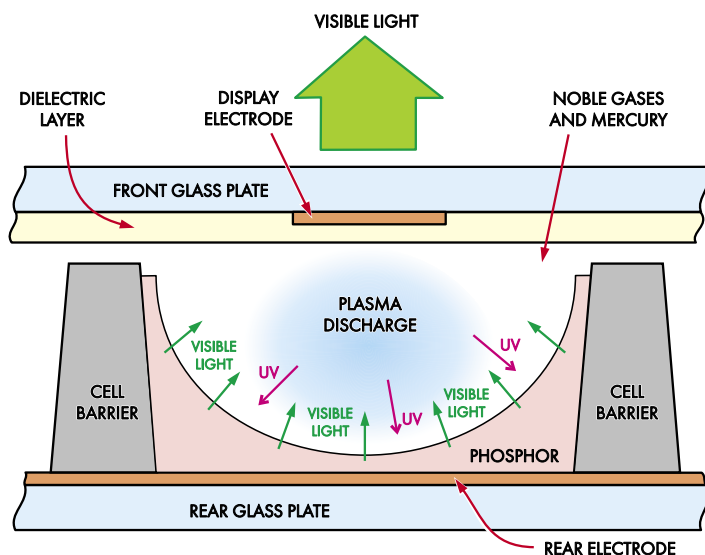
plazmowy w 1936 roku. Pierwszy prototyp wyświetlacza plazmowego został opracowany na Uniwersytecie Illinois w 1964 roku przez Donalda Bitzera, Gene'a Slottowa i Roberta Willsona. Składał się jednak tylko z jednego piksela, więc był całkowicie nieprzydatny z praktycznego punktu widzenia. Minęło jeszcze wiele lat, zanim opracowano pierwszy użyteczny wyświetlacz plazmowy.

Do 1972 roku firma Owens-Illinois Inc. sprzedawała linię monochromatycznych plazmowych monitorów komputerowych lub zespołów wyświetlaczy o rozdzielczości do 512x512 pikseli (rysunek 24). Ich zaletą było to, że były płaskie, nie migotały i nie dryfowały, były w pełni cyfrowe i miały minimalne wymagania dotyczące pamięci, ponieważ wyświetlacz nie wymagał ciągłego odświeżania, jak kineskopy CRT.

Rysunek 24. Reklama pierwszych komercyjnych wyświetlaczy plazmowych z 1972 roku. Źródło: siliconchip.au/link/abfa



Rysunek 23. Kieszonkowy telewizor LCD Casio TV-10, wprowadzony do sprzedaży w tym samym roku co oparty na kineskopie Sinclair TV80.



Rysunek 25. Jedna komórka wyświetlacza plazmowego. Każdy piksel ma trzy komórki, każda z jednym podstawowym kolorem luminoforu, wypełnione gazami szlachetnymi i niewielką ilością rtęci. Wyładowanie plazmowe powoduje emisję światła UV z rtęci, przekształcanego w światło widzialne przez luminofor. Przednie elektrody to przezroczyste przewodniki, takie jak tlenek indowo-cynowy

Niskie wykorzystanie pamięci było znaczące w czasach, gdy pamięć była niezwykle droga i liczył się każdy zaoszczędzony bajt.

Wyświetlacze te były kosztowne, do 2500 USD za jednostkę 512×512. Straciły popularność pod koniec lat 70., gdy pamięć stała się tańsza, dzięki czemu monitory CRT stały się bardziej atrakcyjne, nawet jeśli nie były płaskie.

W 1983 roku IBM wyprodukował 19-calowy (48 cm) monochromatyczny wyświetlacz plazmowy, „panel informacyjny” model 3290, który mógł jednocześnie wyświetlać cztery sesje terminala IBM 3270.

IBM planował zamknąć swoją fabrykę w 1987 roku, ale została ona kupiona przez

Larry'ego Webera, Stephena Globusa i Jamesa Kehoe, którzy założyli nową firmę, Plasmaco. Plasmaco zostało następnie przejęte przez Matsushita (Panasonic) w 1996 roku i nie prowadzi już badań ani nie produkuje wyświetlaczy plazmowych.

W 1992 roku firma Fujitsu wprowadziła na rynek pierwszy pełnokolorowy wyświetlacz plazmowy o przekątnej 21 cali (53 cm). Fujitsu sprzedało pierwszy komercyjny kolorowy telewizor plazmowy w USA w 1997 roku. Miał on przekątną 42 cale (107 cm), rozdzielczość 852×480 pikseli i kosztował 14 999 USD.

W roku 2000 ceny podobnych wyświetlaczy spadły do około 10000 USD. Panasonic

zademonstrował największy wyświetlacz plazmowy o przekątnej 150 cali (3,8 m) w 2008 roku. Miał on 1,8 m wysokości i 3,3 m szerokości.

Po 2006 roku telewizory LCD wyprzedziły plazmy. W 2013 r. Panasonic zaprzestał produkcji wyświetlaczy plazmowych, a w 2014 r. dołączyły do niego firmy LG i Samsung.

Wyświetlacze plazmowe działają podobnie jak żarówki fluorescencyjne. W obojętnym gazie pod niskim ciśnieniem zawierającym niewielką ilość rtęci następuje wyładowanie elektryczne. Rtęć uwalnia światło ultrafioletowe, które następnie uderza w luminofor emitujący światło widzialne odpowiadające kolorowi luminoforu. Każdy piksel wyświetlacza plazmowego składa się z trzech komórek, po jednej dla każdego koloru podstawowego.

Jedno ogniwo wyświetlacza plazmowego pokazano na **rysunku 25**, a zespół wyświetlacza na **rysunku 26**. Ciśnienie gazu wewnątrz każdego ogniwa wynosi około 0,66 bara (2/3 ciśnienia atmosferycznego) z niewielką ilością rtęci w środku. Typowe napięcie zasilające wynosi około 300 V.

Napięcie nie zmienia się, więc aby zmienić intensywność piksela napięcie jest włączane i wyłączane wiele razy na sekundę przy zastosowaniu modulacji szerokości impulsu (PWM).

## ALiS

ALiS (alternate lighting of surfaces) to technologia wyświetlaczy plazmowych opracowana przez Fujitsu i Hitachi w 1999 roku w celu umożliwienia wyświetlaczom o niższej rozdzielczości zapewnienia wyższej rozdzielczości pozornej.

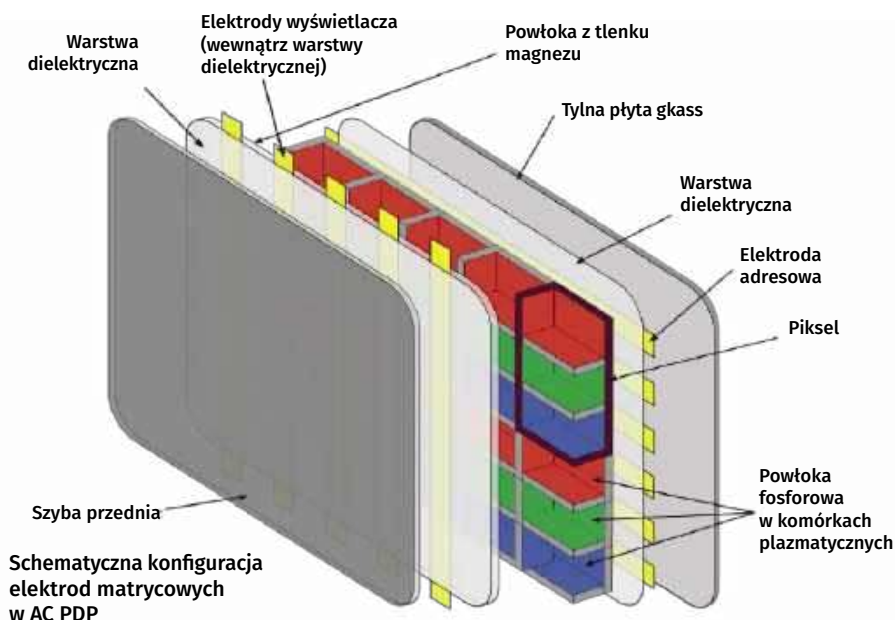
Zamiast skanowania progresywnego, jak w przypadku zwykłego wyświetlacza plazmowego, w którym wszystkie piksele są podświetlane co klatkę, w przypadku skanowania z przeplotem zastosowane jest naprzemienne podświetlanie linii. W ten sposób panel 720-liniowy mógł wyświetlać pozorną rozdzielczość 1080i. Obraz z takiego ekranu miał być również jaśniejszy przy niższym zużyciu energii.

## W następnym odcinku

Część druga będzie dotyczyć ekranów LCD, które przejęły rynek wyświetlaczy w połowie roku 2000. Omówimy również najnowsze i nadchodzące technologie wyświetlania, takie jak OLED i ekrany o wysokim zakresie dynamiki (HDR). ■

dr David Maddison

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. [www.siliconchip.com.au](http://www.siliconchip.com.au)



Schematyczna konfiguracja elektrod matrycowych w AC PDP

Rysunek 26. Panel wyświetlacza plazmowego pokazujący piksel (element obrazu) składający się z trzech komórek oraz pionowych i poziomych elektrod adresujących każdą komórkę. Źródło: Jari Laamanen, Free Art License 1.3

# Proste i bezpieczne wysyłanie często odświeżanych danych w dowolne miejsce

**Czasami potrzebujemy przysyłać paczki często zmieniających się danych, uzyskiwanych na przykład z czujników, z jednego końca świata na drugi. Używane są do tego serwery i usługi online, na przykład takie jak ThingSpeak – usługa platformy analitycznej IoT, która pozwala agregować, wizualizować i analizować strumienie danych na żywo w chmurze. Nasza maszyna może być wystarczająco potężna, aby wysyłać dane co sekundę, ale o ostatecznej przepustowości decyduje wolna ścieżka danych dopuszczona przez serwery online.**

Opisany w artykule projekt może pomóc czytelnikom przysyłać dane, np. z czujników, z dużą częstością powtarzania, a także bezpłatnie i niezależnie od przepustowości serwerów dostarczać je przez WebSocket (WebSocket to protokół komunikacji komputerowej, zapewniający kanały komunikacyjne full-duplex za pośrednictwem pojedynczego połączenia TCP). Osiągamy to za pomocą lepszej, bezpiecznej, ale dostępnej na całym świecie, wieloplatformowej, opartej na chmurze usługi przesyłania wiadomości freemium – komunikatora Telegram. Możemy wysyłać dane co kilka sekund za pomocą taniej płytki rozwojowej ESP32 do kanału bota Telegram, a następnie odbierać je na drugim końcu świata za pomocą prostego kodu napisanego w Pythonie.

W ramach artykułu omówimy zdalną stację telemetryczną, która wysyła dane co cztery sekundy. Dane przychodzą do stacji bazowej zbudowanej w oparciu o płytkę ESP32. ESP32 jest podłączony do Internetu za pomocą sieci Wi-Fi. Następnie wysyła dane do kanału bota Telegramu, a po stronie odbiorcy zwykły komputer podłączony do Internetu uruchamia skrypt napisany w języku Python3, aby w sposób ciągły odbierać dane i używać je do określonych celów. Działanie mogłoby wyglądać następująco:

czujniki w monitorowanym terenie → ESP32 → kanał Telegramu → cyberprzestrzeń → odległy odbiornik → Internet Wi-Fi → Python → odwołanie danych.

Autorski prototyp pokazano na **rysunku 1**. W artykule zamieszczono również wykaz materiałów do jego zbudowania.

Zainstaluj aplikację Telegram na swoim telefonie komórkowym, tablecie lub laptopie. Aplikacja jest dostępna bezpłatnie w Google Play Store (dla telefonów z systemem Android) i App Store (dla iPhone'ów) itp. Do komunikowania się z płytką ESP32 utwórz kanał w aplikacji Telegram.

Po zainstalowaniu aplikacji i skonfigurowaniu konta wyszukaj Botfather w aplikacji. Gdy

tylko otworzysz Botfather, zobaczysz przycisk *Start* lub *Restart*. Otworzy się lista poleceń i ich aplikacji. Kliknij polecenie */newbot*.

Po tym poleceniu należy nadać nazwę botowi. Ja nazwałem go *bera\_arduino*. Następnie ustaw nazwę użytkownika. Podczas ustawiania nazwy użytkownika należy pamiętać, że powinna być unikatowa i kończyć się słowem *bot*, na przykład *bera1bot*. Jak tylko ustawisz nazwę użytkownika, twój bot zostanie utworzony i zobaczysz token API. Zapisz go gdzieś, ponieważ będzie potrzebny później.

Na **rysunku 2** przedstawiono bota Telegram AI, a na **rysunku 3** ustawienie bitu Telegram.

Teraz, aby przygotować oprogramowanie, należy przygotować dwa parametry dla tego kanału (*bera\_arduino*) – *chat\_id* i *bot\_token*. Po utworzeniu nazwy czatu i identyfikatora użytkownika, wydaj polecenie */mybots*, co spowoduje otwarcie twoich botów. Po wskazaniu wybranego bota zostanie wyświetlone okno graficzne, a po naciśnięciu *API\_token* otrzymasz *bot\_token* do zapisu dla swojego bota. Wygląd tokena bota komunikatora Telegram pokazano na **rysunku 4**.

Jedynym potrzebnym parametrem jest Twój *chat\_id*. Aby go uzyskać, znajdź bota o nazwie *@GetIDsBot*. Zostanie wyświetlone okno z opisem Twojego *chat\_id*. **Rysunek 5** przedstawia ID bota w aplikacji Telegram.

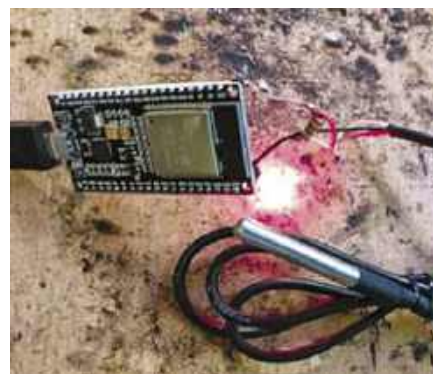
Cóż, teraz masz już dane uwierzytelniające do uruchomienia swojego chat-bota w Telegramie. Zannotuj te dwa upiornie wyglądające numery, aby użyć ich w szkicu.

## Wymagania wstępne dotyczące Pythona

Aby zapewnić gotowość Pythona do pracy z Telegramem, muszą być zainstalowane dwa moduły Pythona za pomocą pip: *python-telegram-bot* i *telethon*.

```
$ pip3 install python-telegram-bot
$ pip3 install telethon
```

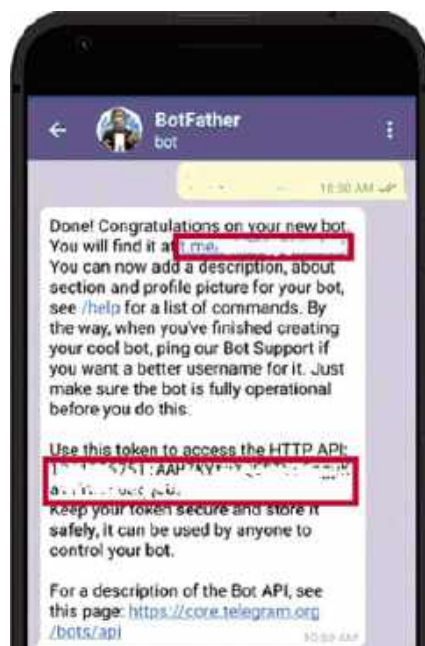
W systemie Linux do instalacji można użyć polecenia *sudo*, natomiast w systemie Windows



Rysunek 1. Prototyp autorski

wystarczy te operacje wyklikać. Skrypt jest stworzony dla wersji Python3. Po uruchomieniu pliku skryptu w tym samym katalogu, w którym znajduje się plik *@bera1bot.session*, dane zaczną pojawiać się bez końca.

```
somnath@somnath-
Satellite-L855:~$ python3
get.py
```



Rysunek 2. API bota Telegram

## Telegram z Pythonem

Powyższa konfiguracja jest wystarczająca do odbierania wiadomości w aplikacji Telegram na telefonie komórkowym. Jednak żeby połączyć Telegram z Pythonem, musisz wykonać poniższe czynności.

W przeglądarce internetowej otwórz stronę <https://my.telegram.org> i wprowadź ten sam numer telefonu komórkowego, który skonfigurowałeś w aplikacji mobilnej (np. +79942345xxxx). Spowoduje to wysłanie wiadomości SMS do bota Telegramu w aplikacji Telegram telefonu (kanał w aplikacji Telegram). Wprowadź ten kod w następnym oknie. Zostaniesz przeniesiony do kolejnego okna, w którym musisz nadać tytuł aplikacji i krótką nazwę (np. bera2 bera2).

Naciśnięcie *Next*, spowoduje wyświetlenie dwóch ważnych identyfikatorów: `api_id`, `api_hash`. Nazwa użytkownika, którą skonfigurowałeś podczas tworzenia bota @bera1bot już tam jest. Teraz, do uzyskania dostępu do wiadomości przychodzących w kanale Telegram w skrypcie Pythona niezbędne są tylko te trzy parametry. Zobaczmy, jak to działa.

Z `TelegramClient(name, api_id, api_hash)` jako klientem, uruchom kod:

```
@client.on(events.NewMessage())
async def handler(event):
```

### Wykaz materiałów

ESP32 – 1  
DS18B20 – 1  
LED 3,5 mm – 1  
Rezystor 5,6 kΩ – 1  
Zasilacz USB – 1



Rysunek 3. Ustawienie bitu telegram

```
async for message in
client.iter_messages(''):
print(message.text() )
client.
run_until_disconnected()
```

Jest to główna pętla do ciągłego odbierania wiadomości, gdy docierają one do kanału. Powyższy kod można przepisać z artykułu, nazwać `receive.py` i uruchomić w Pythonie.

Przy pierwszym uruchomieniu kodu Python na komputerze, poprosi on o parametry sesji. W tym celu zapyta o numer telefonu, którego użyłeś wcześniej. Po jego wprowadzeniu zostanie wysłany kod SMS-em do bota Telegramu w aplikacji Telegram na telefonie. Wprowadź ten kod w następnym oknie. Spowoduje to utworzenie pliku sesji i kod Pythona zacznie działać.

Aby twój kod był naprawdę przenośny, musisz przekazać oba pliki jednocześnie – kod Pythona i plik @bera1bot.session.

## Nadajnik z Arduino

Szkic Arduino jest dosyć prosty. ESP32 jest wyposażony w szereg nazw użytkowników Wi-Fi (ID) i hasła. Najpierw łączy się z odpowiednim połączeniem Wi-Fi. Jeśli się nie powiedzie, uruchomi się ponownie i wybierze następny identyfikator/hasło. Po nawiązaniu połączenia pobiera ciąg danych, a następnie za pomocą `CHAT_ID` i `BOTtoken` tworzy ciąg danych `https`, który jest przekazywany do Telegramu za pomocą `http.GET()`. Dioda LED podłączona do pinu 23 będzie migać za każdym razem, gdy przesyłanie danych zakończy się powodzeniem.

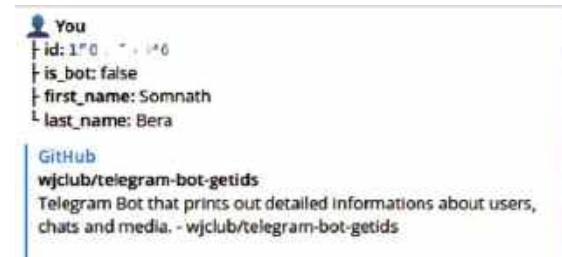
Po wgraniu kodu Arduino, podłącz ESP32 jak pokazano na rysunku 6.

## Przykładowy projekt

Przykładowy projekt jest podzielony na dwie części: nadajnik i odbiornik. Nadajnik jest oparty na ESP32 z podłączonym do niego czujnikiem temperatury DS18B20 i generuje prosty strumień danych CSV zawierający datę i godzinę pobrane w locie z serwera NTP. Dodaje kilka dowolnych strumieni



Rysunek 4. Token bota Telegram



Rysunek 5. Identyfikator bota w aplikacji Telegram

danych z temperaturą dodaną na końcu, takich jak 08/09/2021, 06:35:42, 0746, 0876, 3278, 3563, 2795, 3012.

Ostatnie dane (3012) to temperatura (30,12°C). Trzymając DS18B20 w dłoni, można zauważyć, że wartość zmienia się pod wpływem temperatury dłoni. Cały ciąg jest przesyłany do kanału Telegram za pomocą prostej instrukcji `GET ESP32`:

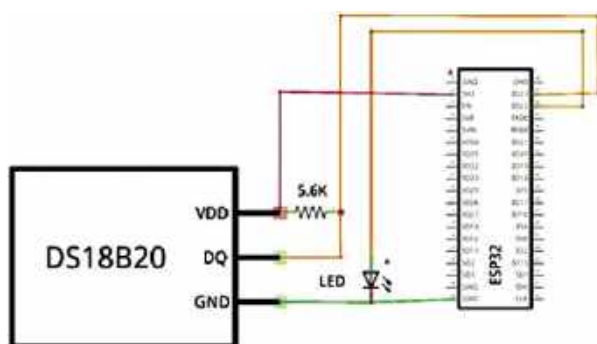
```
https://api.telegram.org/bot\$token/sendMessage?chat\_id=\$chat&text=Hello+World+how+are+you
```

To samo można osiągnąć za pomocą komputera, ale używając ESP32, zyskujemy kompaktowe i niezależne urządzenie, które w sposób ciągły dostarcza dane do Telegramu bez potrzeby korzystania z nieporęcznego komputera. ESP32 zużywa tylko ułamek energii z zasilania 3,3 V, co czyni je bardziej efektywnym rozwiązaniem.

Autor zasugerował tę prostą wersję projektu, którą można łatwo wdrożyć, ale rzeczywisty projekt wykonany przez niego jest znacznie bardziej wyrafinowany. Ma on zdalną jednostkę nadawczą, która wysyła dane z czujników w terenie przez sieć LoRa do innego ESP32, który z kolei przekazuje dane do kanału Telegram za pomocą LoRa z jednej strony i Wi-Fi z drugiej strony – działającą podobnie jak router).

## Odbiornik

To jest właśnie sedno projektu. Możesz otrzymywać wysoce zabezpieczone i zaszyfrowane dane na swój telefon. Pierwsze dane testowe docierają tylko do telefonu komórkowego. Trudno jednak z poziomu telefonu przetworzyć te dane w sposób umożliwiający realizację bardziej zaawansowanych działań, takich jak umieszczenie ich w internetowej



Rysunek 6. Schemat połączeń pomiędzy termometrem oraz modułem ESP32

bazie danych do monitorowania trendów, uruchomienie przekaźnika czy kontrolowanie systemu w oparciu o otrzymywane informacje.

Dlatego należy zastąpić telefon komórkowy komputerem podłączonym do Internetu z zainstalowanym Pythonem3. Skrypt Python zbierze dane, a następnie wyświetli je ponownie na komputerze z taką samą prędkością, z jaką są wysyłane po stronie nadawczej. Ponieważ wszystkie dane są dostępne jako zmienne w Pythonie, można je przetwarzać w dowolny sposób.

Platformy analityczne IoT, takie jak ThingSpeak, mogą zapewnić maksymalnie cztery ciągi danych na minutę, w Telegramie natomiast można przesyłać nawet trzydzieści ciągów danych na minutę, a mimo to wszystkie są dobrze zabezpieczone i zaszyfrowane. Poniżej znajduje się ciąg danych, który dociera co dwie sekundy i jest odbierany w terminalu Python Telegram na drugim końcu świata bez utraty jakiegokolwiek strumienia:

```
08/09/2021 08:22:15 0746 3563 0876 2795 3278 2962
08/09/2021 08:22:16 0746 3563 0876 2795 3278 2962
08/09/2021 08:22:18 0746 3563 0876 2795 3278 2962
08/09/2021 08:22:20 0746 3563 0876 2795 3278 2962
08/09/2021 08:22:21 0746 3563 0876 2795 3278 2962
```

```
08/09/2021 08:22:23 0746 3563 0876 2795 3278 2968
08/09/2021 08:22:25 0746 3563 0876 2795 3278 2968
08/09/2021 08:22:26 0746 3563 0876 2795 3278 2968
08/09/2021 08:22:28 0746 3563 0876 2795 3278 2962
08/09/2021 08:22:30 0746 3563 0876 2795 3278 2962
```

Oprogramowanie arduino\_sketch.ino i python\_script.py można pobrać pod adresem:

<https://www.electronicsforu.com/electronics-projects/hardware-diy/published-telegram-received-python>.

## Testowanie

Po podłączeniu ESP32 do czujnika i zasileniu płytki oraz upewnieniu się, że ESP32 i system odbiornika mają połączenie z siecią, uruchom plik Python po stronie odbiornika. Komunikat z czujnika będzie widoczny w terminalu Pythona. Możesz również uzyskać dane z czujnika na żywo w Telegramie jako czat.

Po pomyślnym nauczeniu się uruchamiania przykładowego projektu, będziesz mógł wdrożyć Telegram do rzeczywistego użytku. Z drugiej strony, po otrzymaniu go w Pythonie3, można go dalej przetwarzać w celu uruchomienia urządzeń lub zapisywania w bazie danych. Dzięki światłowodowemu połączeniu z domowym Internetem autor mógł uzyskać trzydzieści odczytów na minutę (jeden odczyt danych o długości 49 znaków co dwie sekundy), co jest bardzo dobrym wynikiem. Jeśli sieć zwolni, liczba przesyłanych danych zostanie oczywiście zmniejszona.

Uwaga EFY. Można zwiększyć przepustowość nawet „raz na sekundę”, ale oznacza to ogromny ruch w kanale telegramu – co jest deprecjonowane przez Telegram i sprawi, że zostaniesz natychmiast zablokowany. Dlatego nigdy nie próbuj zwiększać przesyłu danych i stosuj przynajmniej dwu sekundy interwał. ■

Somnath Bera

Ten artykuł był wcześniej opublikowany na łamach „EFY”, wrzesień 2022 (efymag.com)

REKLAMA

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

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

Ponadto każdy prenumerator ww. czasopism korzysta z rabatów od 30% do 50% na zakup czasopism z oferty [www.UlubionyKiosk.pl](http://www.UlubionyKiosk.pl)

**K L U B**  
**AVT**  
**ELEKTRONIKA**

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

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

# Sieć mesh dalekiego zasięgu oparta na niedrogich modułach LoRa

**Podobną do opisanej sieć można zbudować na modułach ZigBee. Jednak podczas gdy zasięg ZigBee zazwyczaj nie przekracza 100 m, zasięg LoRa w sprzyjających warunkach może wynosić nawet kilkadziesiąt kilometrów.**

Chociaż moduły radiowe Zigbee są stosunkowo drogie, znajdują liczne zastosowania ze względu na wyjątkową łatwość tworzenia sieci. Radia Zigbee można umieszczać w konfiguracjach mesh (siatka), w których każde radio Zigbee jest połączone jako węzeł z innymi. Jeśli zdarzy się, że jeden lub wiele węzłów ulegnie awarii, sieć będzie nadal działać!

Meshtastic to projekt open source dostępny w serwisie GitHub. Jego sieć LoRa ma bardzo duży zasięg i niskie zużycie energii, jest dobrze szyfrowana i odporna na zakłócenia elektromagnetyczne. Autorski prototyp pokazano na **fotografii 1**.

## Sieć Meshtastic

Projekt Meshtastic.org obejmuje oprogramowanie open source przeznaczone do tworzenia sieci LoRa mesh. LoRa działa najlepiej, gdy moduły mogą się widzieć wzajemnie, dlatego jeśli jeden lub dwa węzły można umieścić w strategicznie wysokim miejscu, węzły będą mogły się bezproblemowo ze sobą komunikować.

Jest to próba stworzenia praktycznej sieci mesh dla czytelników EFY, w której dane z sensorów, pozycja GPS oraz alarmy mogą być łatwo obsługiwane. Cała konfiguracja może być rozbudowana do większych zastosowań.



Fotografia 1. Prototyp autorski

W tym przypadku nie mamy do czynienia z programowaniem Arduino, w którym połączenia można ustawiać programowo w celu uzyskania pożądanego rezultatu, należy więc dokonać poprawnych połączeń pomiędzy modułami: Heltec LoRa oraz ESP32. Choć nie jest to bardzo trudne, jeśli przestrzega się zaleceń instrukcji, można również zakupić gotowy moduł ESP32 Heltec LoRa. Elementy użyte do wykonania tego urządzenia są wymienione w poniższym wykazie materiałów.

Moduł Heltec LoRa jest również dostępny w sklepach internetowych za nieco wyższą cenę. Decydując się na samodzielną budowę, należy zaopatrzyć się w kilka elementów: moduł uruchomieniowy ESP32, wyświetlacz OLED o przekątnej 2,4 cm (0,96 cala), radio SX1278 868 MHz SPI LoRa. Moduł Heltec LoRa został pokazany na **fotografii 2**.

## Budowa

Żeby móc zacząć cokolwiek testować należy utworzyć działającą sieć mesh, składającą się z trzech lub więcej węzłów LoRa. Można też kupić gotowe moduły ESP32 Heltec LoRa dostępne w sklepach. Dodatkowo należy zakupić co najmniej jeden czujnik I<sup>2</sup>C z poniższej listy: BME280, BME680, MCP9808, INA260, INA219, SHTC3, SHT31, PMSA0031.

Do załadowania oprogramowania Meshtastic do każdej płytki ESP32 należy połączyć płytkę do portu USB komputera i przejść na stronę <https://flasher.Meshtastic.org/>. Podany link działa tylko w przeglądarce Chrome lub Edge.

Schemat ideowy projektu LoRa został pokazany na **rysunku 3**. Zastosowano w nim układ SX1278 SPI LoRa (MOD1), moduł ESP32 (MOD2), OLED I<sup>2</sup>C (MOD3), BME280 (MOD4), stabilizator 3,3 V HT7333A (IC1) i kilku innych elementów.

## Wgrywanie firmware'u

Na pierwszym ekranie należy wybrać wersję firmware'u, która będzie zapisywana w pamięci Flash układu (Heltec V1). Na następnym ekranie wybieramy mostek CP2102 USB do UART-a, który jest sprzętowym interfejsem umieszczonym między MCU a komputerem. Następnie należy postępować zgodnie z instrukcjami pojawiającymi się na ekranie

– np. ostrzeżenie o wymazaniu całej zawartości pamięci w urządzeniu itp.

Konieczne trzeba upewnić się, że połączenie internetowe będzie stabilne podczas kluczowych 4 do 5 minut ładowania programu. W przypadku wbudowanej płytki Heltec konieczne może być przytrzymanie przycisku Boot przez całe 4 do 5 minut podczas procedury ładowania firmware'u.

## Aktualizacja oprogramowania

Okazuje się, że aktualizacja oprogramowania firmowego oparta na przeglądarce, w której używany jest wiersz poleceń, nie działa niezawodnie. Wielokrotnie niemal natychmiast pojawia się komunikat informujący o zakończeniu zadania, jednak w rzeczywistości programowanie nie zostało przeprowadzone. Instalacja/aktualizacja firmware'u zajmuje około 4 do 5 minut.

### Polecenia Mac/polecenia Linux

```
pip3 install --upgrade esptool
[upgrade esptool first]

esptool.py chip_id [you will get
a long output with chip id and chip
mac id]

cd ~/Downloads/firmware/ [firmware-
heltec-v1-2.2.12.092e6f2.bin]
./device-install.sh -f firmware-
BOARD-
VERSION.bin [for installing]
./device-update.sh -f firmware-BOARD-
VERSION.bin [for updating]
```

### Polecenia systemu Windows

```
device-install.bat -f firmware-BOARD
VERSION.bin (for installing)
device-update.bat -f firmware-BOARD-
VERSION.bin (for updating)
```



Fotografia 2. Moduł Heltec LoRa

Z tego powodu udostępniono tutaj zarówno firmware, jak i instrukcje wiersza poleceń.

## Testowanie

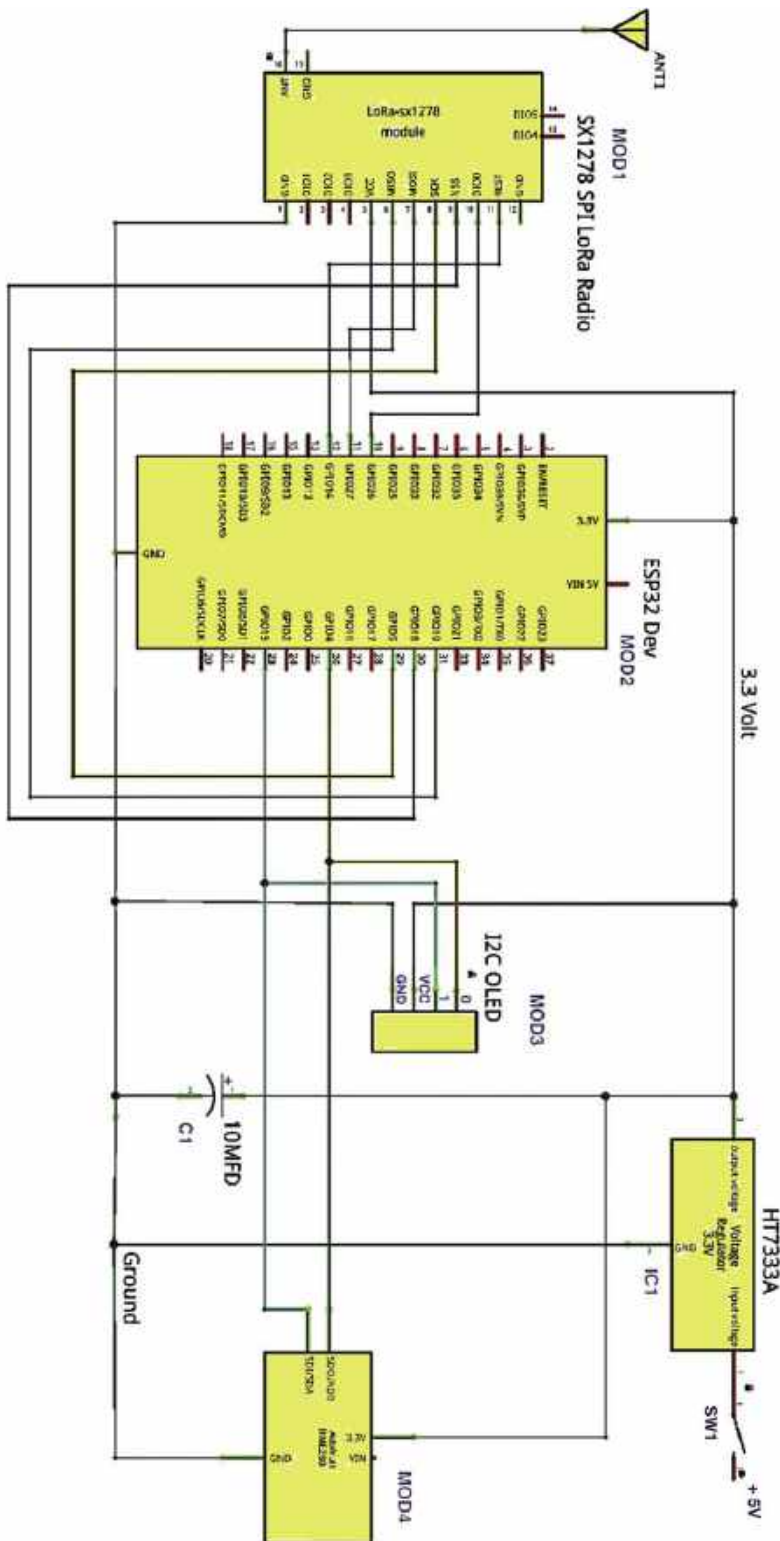
Do skonfigurowania i przetestowania projektu, konieczne jest zainstalowanie aplikacji Meshtastic w telefonie z systemem Android/

iPhone (lub na tablecie). Wersje na Androida i iOS są dostępne w odpowiednich sklepach z oprogramowaniem (Google Play Store/Apple App Store). Do pracy z trzema węzłami, należy wybrać jeden telefon komórkowy (lub tablet) mogący współpracować z każdym węzłem. W ten sposób można uniknąć

wielu kłopotów związanych z konfiguracją. Aplikacja jest darmowa! Na **rysunku 5** przedstawiono ikonę Meshtastic.

Najpierw włączamy zasilanie karty Heltec. Następnie włączamy Bluetooth w telefonie komórkowym i musimy go sparować z kartą Heltec. Na ekranie OLED pojawi się 6-cyfrowy kod parowania. Po pomyślnym sparowaniu otwieramy aplikację Meshtastic, w której zobaczymy, że aplikacja jest połączona (należy szukać znacznika wyboru w prawym górnym rogu) z adresem tablicy Mac. Naciskając trzy kropki (...) w prawym górnym rogu aplikacji, jak pokazano na **rysunku 6** można ustawić nazwę karty w konfiguracji radiowej.

Heltec domowej roboty jest już dostępny. Nazywa się „fabr” (skrót od *fabricated one*). Pozostałe dwa nie działają. Dlatego przy



Rysunek 3. Schemat układu

## //ESHT/ST/C ESP32 web installer



Rysunek 4. Ładowanie oprogramowania firmowego



Rysunek 5. Ikona Meshtastic

### Wykaz elementów:

|                                        |   |
|----------------------------------------|---|
| Heltec ESP32 LoRa/ESP32 MCU (MOD2)     | 3 |
| HT7333A – stabilizator 3,3 V (IC1)     | 1 |
| SX1278 SPI LoRa (MOD1)                 | 3 |
| BME280 (MOD4)                          | 1 |
| MCP9808 – I <sup>2</sup> C OLED (MOD3) | 1 |
| SW1                                    | 1 |
| Czujnik GPS uBLOX z ANT1               | 1 |
| Akumulator/zasilacz 5 V                | 1 |

ich nazwach pojawiają się znaki zapytania. Znacznik wyboru w prawym górnym rogu chmury wskazuje, że ten węzeł jest podłączony.

## Mocowanie czujników

Aktualnie można używać tylko czujników I<sup>2</sup>C (adresy 0x76, 0x77, 0x78, 0x18, 0x40, 0x41, 0x5D, 0x5c, 0x70, 0x44, 0x12, 0x3c są ustalone dla OLED). Więcej czujników zostanie dodanych w najbliższym czasie. Opcja *Radio-configuration* → *Detection Sensor* → *Detection Sensor* powinna być włączona.

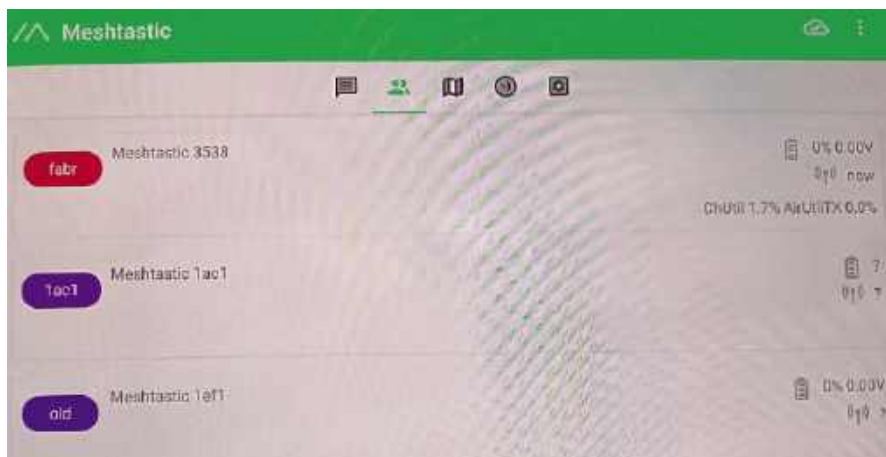
Można podać nazwę czujnika (BME lub Temp\_Ind). W tym samym oknie można również wyznaczyć pin GPIO do monitorowania [wysoki lub niski]. Wykrywanie będzie monitorowane okresowo co 900 sekund (wartość domyślna). Należy pamiętać, że stan tego pinu GPIO jest całkowicie niezależny od wartości czujnika. Na **rysunku 7** przedstawiono interfejs użytkownika aplikacji pokazujący stabilizowany dziennik pomiarów środowiska.

Komendą *Radio Configuration* → *Device* → *Role* (do przesyłania wartości czujnika) nadajemy mu nazwę „SENSOR”. Inne opcje to „CLIENT” lub „ROUTER\_CLIENT”, „REPEATER” itp. Interfejs użytkownika aplikacji pokazujący urządzenia w sieci Mesh pokazano na **rysunku 8**.

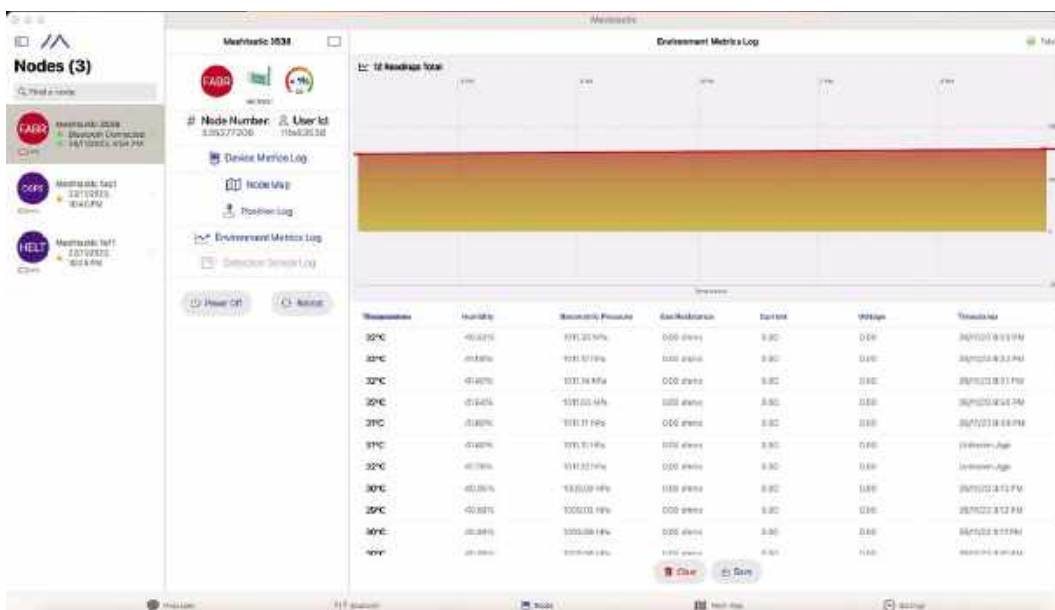
Aby wysyłać sygnały do zdalnego sprzętu, należy najpierw zdefiniować pin GPIO. Jednak w tym przypadku włączony jest tylko cyfrowy odczyt i cyfrowy zapis.

## Powiadomienia zewnętrzne

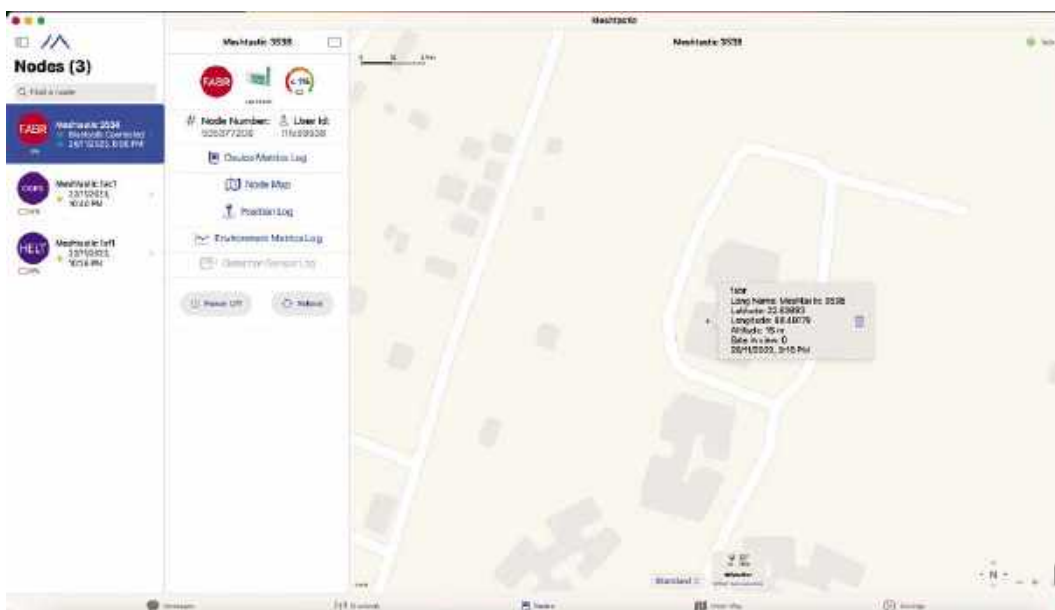
W celu uzyskania powiadomień zewnętrznych, przechodzimy do *Radio Configuration* → *External Notification* → *Enable*. Należy również ustawić wyjściowe piny GPIO, które będą odbierać powiadomienia. Można ustawić diodę LED/brzęczyk alarmowy/wibrator alarmowy



Rysunek 6. Aplikacja Meshtastic



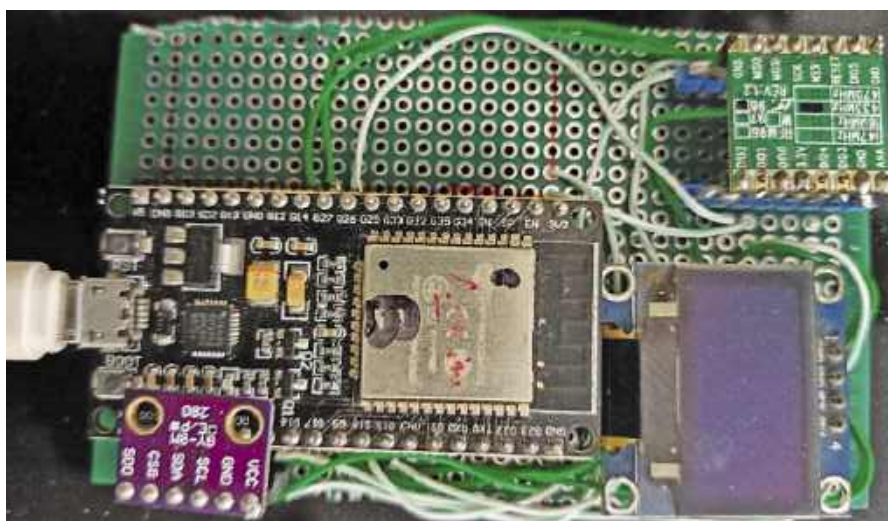
Rysunek 7. Interfejs użytkownika aplikacji pokazujący stabilizowany dziennik pomiarów środowiska



Rysunek 8. Interfejs użytkownika aplikacji pokazujący urządzenia w sieci Mesh



Fotografia 9. Włączony alarm LED



Rysunek 10. Węzeł sieci z czujnikiem

na trzech różnych pinach GPIO wraz z czasem trwania alarmu.

W demonstracyjnym urządzeniu „fab” alarm na GPIO-16 jest ustawiony tak, aby zapalał diodę LED na dwie sekundy. Wartości czujnika pojawią się na ekranie Android/iOS i na ekranie OLED. Dodanie nowego węzła sieci jest bardzo łatwe, wystarczy go skonfigurować.

## Lokalizacja GPS

Jeśli płytkę ma lokalizator GPS, lokalizacja zostanie dostarczona automatycznie. W przeciwnym razie, jeśli urządzenie z systemem Android/iOS ma włączone informacje o lokalizacji, zostanie dostarczona lokalizacja GPS, która będzie widoczna dla tego węzła w oknie aplikacji. Aplikacja ma możliwość dołączenia oddzielnego czujnika GPS, dla którego należy ręcznie ustawić piny GPIO dla Tx i Rx: *Radio configuration* → *Position* (patrz **rysunek 8**, gdzie lokalizacja pochodzi z usługi lokalizacji telefonu). Po ustawieniu czujnika GPS dane lokalizacji (szerokość i długość geograficzna) będą jednak pochodzić z czujnika GPS.

Zgodnie z informacjami na stronie, w tej chwili obsługiwane są tylko czujniki GPS uB-LOX, GLONASS. Stary czujnik GPS autora

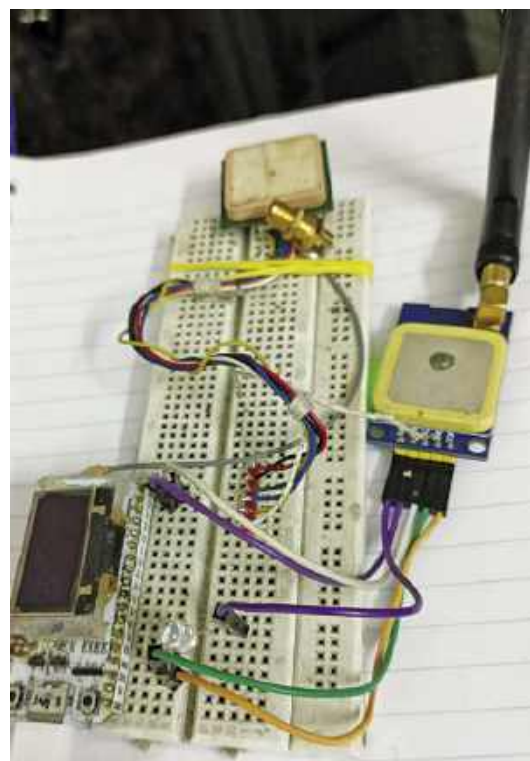
(czujnik GPS V.KEL TTL) po podłączeniu również działał.

## Sieć

Sieć można włączyć klikając: *Radio Configuration* → *Network* → *Enable Wi-Fi* → enter SSID and PSK. Jednak wykonanie tej czynności spowoduje, że węzeł będzie niedostępny przez Bluetooth. Następnie należy uzyskać do niego dostęp za pomocą kabla USB lub protokołu sieciowego za pomocą serwera WWW, co nie jest zbyt wygodne. Odwrócenie tej operacji jest bardzo frustrujące. Konieczne jest ponownie wgranie firmware'u!

## Typowe zastosowania

Sieć LoRa zainstalowana na obszarach leśnych może być używana do monitorowania pożarów lasów. W takich przypadkach każdy węzeł będzie monitorował temperaturę, prędkość i wilgotność powietrza w swojej strefie. W obszarze przemysłowym taka siatka może być bardzo skutecznie wykorzystywana do monitorowania zagrożeń pożarowych. W mieście, zmieniając pojazdy w węzły LoRa, można monitorować prędkość ich przemieszczania się oraz położenie.



Rysunek 11. Finalny prototyp na płytce prototypowej

W miasteczku/kompleksie mieszkaniowym, przy użyciu czujników laserowych V53LOX, można monitorować wszystkie poziomy zbiorników wody. Możliwe jest czatowanie z węzłami, a przy użyciu MQTT możliwe jest nawet za pośrednictwem sieci Meshtastic LoRa dotarcie do innych miast.

*Komenda Radio configuration* → *Module configuration* → *Store & Forward* działa tylko dla ESP32. Tutaj można zapisać kilka rekordów na wypadek rozłączenia z węzłem, jednak Meshtastic nie zaleca jego używania.

## Efekt końcowy

Oprogramowanie Meshtastic pokazało, w jaki sposób za pomocą sieci LoRa można wdrożyć moduły radiowe Lora w celu uzyskania najlepszych cech sieci. Są to: niski pobór mocy, duży zasięg i wysoka odporność na zakłócenia. Używanie sieci to świetna zabawa. Po skonfigurowaniu węzłów można odłączyć urządzenia z systemem Android/iOS, a węzły będą nadal działać automatycznie.

Na **fotografii 9** przedstawiono sytuację, w której został uruchomiony alarm LED, natomiast na **fotografii 10** widzimy węzeł mesh wraz z czujnikiem. Ostateczny prototyp zawierający ESP32, LoRa, ekran OLED i czujnik BME 280 pokazano na **fotografii 11**. ■

**Somnath Bera**

Ten artykuł był wcześniej opublikowany na łamach „EFY”, marzec 2024 (efymag.com)



Lukasz z Wrocławia

**Jak to dobrze, że niektóre przyjemne sprawy mogą trwać dłużej niż wakacje! Na przykład nasze cykliczne „juniorskie” spotkania. Choć od lat prawie dwudziestu, zamiast wakacji, przysługują mi co najwyżej urlopy, a moje ciało młodości ducha nie zdradza, w środku wciąż jestem takim „Juniorem” (jednym z Was) i przyznam się Wam, że strasznie lubię te nasze „juniorskie” spotkania.**

A co u Ciebie? Idę o zakład, że wizyta w sklepie z elektroniką (lub też zakupy w sklepie internetowym) na dobre zapadnie Ci w pamięć. Za młodu uwielbiałem odwiedzać giełdy i stacjonarne sklepy z elektroniką, a każda wizyta, zawsze z przydługą listą części do zakupienia, okazywała się wielką przygodą. Jeśli masz taką możliwość, podziel się z nami wrażeniami (wyślij list do Redakcji: redakcja@elportal.pl) po ostatnim spotkaniu. Napisz, jakiego kompana (tato, ciocia, dziadek, a może starszy brat?) i w jaki sposób udało Ci się namówić do wspólnych zabaw elektroniką. Poszło gładko? Czy też musiałeś się nieco natrudzić? A może nawet odwołać do niewinnego podstępów z fotografiami w tle, który starałem Ci się w ubiegłym miesiącu ukradkiem podsunąć? Jeśli udało Ci się zorganizować sprzęt i „zmontować ekipę” (czytaj: znaleźć opiekuna na czas zabaw z elektroniką) to świetnie, ponieważ dzisiaj złożyście wspólnymi siłami kolejną zabawkę!

Nie wątpię, że lampka, podczas budowy której dużo się w zeszłym miesiącu nauczyłeś, wciąż pięknie rozświetla Twój pokój każdego wieczora, niemniej najwyższa pora na poskładanie kolejnego układu, tym razem intrygującej zabawki! Będzie to „Konsola Audiochaos” (AVTEDU624), która pozwoli Ci poczuć się niczym didżej w królestwie elektroniki i za pomocą konsoli trzech potencjometrów tworzyć zaskakujące efekty dźwiękowe.

## Zasada działania

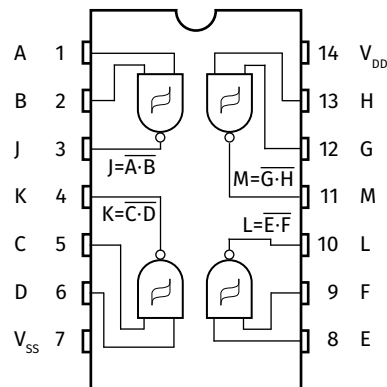
Zerknijmy na schemat widoczny na **rysunku 1**, aby przed przystąpieniem do montowania czegokolwiek, mieć podstawowe zrozumienie zasady działania naszego nowego (po)tworka.

Oprócz komponentów, które być może pamiętasz z poprzedniego odcinka, a więc rezystorów (stawiających opór przepływającemu przez nie prądowi), kondensatorów (gromadzących ładunek w polu elektrycznym) a nawet



Konsola  
Audiochaos  
(AVTEDU624)

potencjometru (rezystora o regulowanej wartości), na schemacie Konsoli Audiochaos znajdziesz dwa nowe elementy. Pierwszym z nich jest głośniczek piezo (buzzer pasywny) oznaczony na schemacie za pomocą desygnotora SP1. Jest to przetwornik sygnałów elektrycznych na dźwięk, a więc pełni rolę analogiczną do klasycznego głośnika. Przetworniki piezoelektryczne oraz głośniki różnią się między sobą budową. Głośniki działają na zasadzie elektromagnetycznej. Składają się z membrany, cewki i magnesu. Przepływ prądu zmiennego przez cewkę powoduje ruch membrany, co generuje

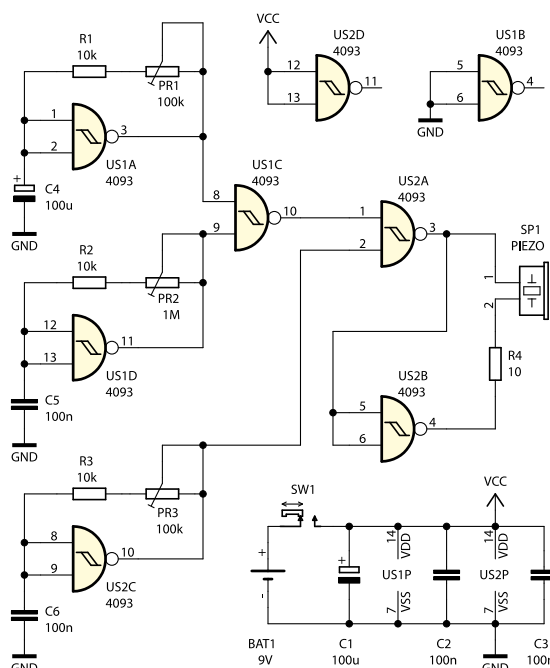


Rysunek 2. Cztery bramki NAND zawarte wewnątrz pojedynczego układu scalonego serii 4093

fale dźwiękowe. Tymczasem buzzer pasywny piezo wykorzystuje zjawisko piezoelektryczne. Jest zbudowany w oparciu o kryształ piezoelektryczny, który odkształca się pod wpływem cyklicznie przykadanego napięcia elektrycznego, generując vibracje a przez to dźwięk. Buzzery mają gorsze tzw. pasmo przenoszenia, co w dużym uproszczeniu oznacza, że nie sprawdzą się one np. w odbiorniku radiowym. Natomiast świetnie zachowują się w odtwarzaczach prostych melodyjek i dźwięków alarmowych, jak również w zabawkach takich jak ta, którą za chwilę zbudujesz.

Drugim nowym symbolem są dla Ciebie z pewnością półkola, z trzema wyprowadzeniami każde, których znajdziesz na schemacie aż osiem. Są to tak zwane bramki logiczne. Nie są one jednak samodzielnymi elementami, które mógłbyś zakupić w sklepie. W sklepie zakupisz natomiast układy scalone, które będą zawierały w swojej obudowie kilka lub kilkanaście takich bramek. W projekcie Konsoli Audiochaos użyto dwa identyczne układy scalone CD4093BE, i każdy z nich zawiera po 4 bramki logiczne, co możesz dostrzec, spoglądając na **rysunek 2**.

W przypadku układu serii 4093 są to bramki typu NAND



Rysunek 1. Schemat ideowy Konsoli Audiochaos AVTEDU624

z przerzutnikiem Schmitta. O funkcji przerzutnika Schmitta i histerezy być może opowie Ci innym razem, przy jakimś bardziej intuicyjnym do tego układzie, teraz natomiast zapraszam Cię do świata matematycznej logiki, by opisać Ci najkrócej jak tylko potrafię, czym w ogóle jest bramka logiczna.

## Z logiką za pan brat

Nie wnikając w detale, które wykraczają poza ramy naszych spotkań, wyobraź sobie, że bramka logiczna to taka czarna skrzynka, z której, po uprzednim dostarczeniu jednej lub kilku danych wejściowych w postaci zer lub jedynek (logicznych prawd lub fałszów) wyciąga się pojedynczą odpowiedź, również w postaci, albo prawdy (logicznej „1”), albo fałszu (logicznego „0”).

Jednym z możliwych rodzajów bramek logicznych jest bramka typu AND (jak wiadomo „and” z angielskiego tłumaczy się na polskie „i”), która posiada dwa wejścia. Żeby na wyjście bramki zwrócona została „prawda” (logiczna „1”), „prawda” (logiczna „1”) musi zostać uprzednio dostarczona na oba wejścia bramki, zarówno na wejście pierwsze, jak „i” na wejście drugie. Jeśli na którymkolwiek z wejść (albo na obu) pojawi się fałsz (czyli logiczne „0”), fałsz pojawi się również na wyjściu. Dlatego też mówi się, że bramka logiczna typu AND zwraca na wyjściu wynik iloczynu logicznego swoich wejść. Iloczyn to mnożenie, i zgodnie z jego prawidłami, gdy dowolny czynnik mnożenia stanie się zerem, wynikiem też będzie zero. Symbol oraz zasadę działania bramki AND pokazano na **rysunku 3**.

Innym (z wielu możliwych) rodzajem bramki logicznej jest bramka typu NOT. Jak wiadomo „NOT” z języka angielskiego tłumaczy się na polskie „nie”. Łatwo się zatem domyślić, że bramka wykonuje operację negacji, czyli zaprzeczenia. Bramka ta posiada tylko jedno wejście. Jeśli na to wejście

dostarczymy „prawdę” (logiczną „1”), wówczas na wyjściu otrzymamy „fałsz” (logiczne 0). Analogicznie, jeśli na wejście dostarczymy „fałsz” (logiczne „0”), wówczas na wyjściu otrzymamy „prawdę” (logiczne 1). Symbol oraz zasadę działania bramki NOT pokazano na **rysunku 4**.

Bramka NAND (takich właśnie użyto do zbudowania Konsoli Audiochaos) jest elementem, który łączy dwie funkcje logiczne: AND i NOT. Nazwa „NAND” pochodzi właśnie od tych dwóch słów (not and). Innymi słowy jest to bramka AND z zanieganowaną odpowiedzią (**rysunek 5**). Zwróć uwagę na „kółeczko” na wyjściu bramki NOT oraz na wyjściu bramki NAND. To kółeczko oznacza właśnie negację (zaprzeczenie).

Symbol oraz zasadę działania bramki NAND pokazano na **rysunku 6**.

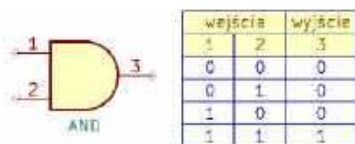
## Powrót do świata elektroniki

„Prawda”, reprezentowana liczbowo przez „1” oraz „fałsz”, reprezentowany liczbowo przez „0”, to tworzy typowo logiczne, z kolei logika jest działem matematyki. W elektronice logiczną jedynkę reprezentuje zazwyczaj napięcie bliskie napięciu zasilania. Ponieważ Konsola Audiochaos zasilana jest z baterii 9 V, logiczna jedynka będzie rozpoznawana jako wystąpienie na wejściu lub wyjściu bramki napięcia bliskiego tej właśnie wartości. „Fałsz”, czyli logiczne „0”, będzie natomiast równoznaczne z pojawieniem się na wejściu lub wyjściu bramki napięcia bliskiego potencjałowi masy (0 V).

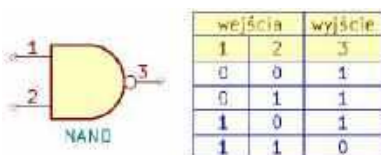
## Powrót do analizy schematu

Nieco wyżej wspomniałem o tym, że bramek logicznych nie da się kupić w postaci pojedynczych komponentów. W projekcie naszej konsoli, potrzebne było sześć bramek NAND, tym samym konieczne było użycie dwóch układów scalonych. Dwóch nadmiarowych bramek, autor projektu nie wykorzystał. Są to bramki US1B oraz US2D. Ich wyjścia nie zostały do niczego podłączone, natomiast wejścia podłączono, w przypadku US1B do napięcia 0 V a w przypadku US2D do napięcia

9 V. Tak naprawdę, nie ma większego znaczenia, do jakich napięć wejścia nieużywanych bramek zostały podłączone. Ważne, że zostały podłączone „gdziekolwiek” przez co napięcia na ich wejściach zostały jednoznacznie ustalone. Gdyby było inaczej, z dużym prawdopodobieństwem wejścia te działałyby jak anteny, zbierając z otoczenia zakłócenia radiowe i elektromagnetyczne, powodując tym samym losowo zmienny przebieg napięć na ich wyjściach i generując potencjalne i nikomu niepotrzebne wtórne zakłócenia na liniach zasilania konsoli. Ponadto pozostawianie niepodłączonych wejść zwiększałoby ryzyko uszkodzenia pojedynczych bramek a nawet całych układów scalonych na skutek wyładowań elektrostatycznych (ESD), o których kiedyś postaram się opowiedzieć nieco więcej. Teraz dla nakreślenia zjawiska ESD wspomnę tylko, że na pewno nie raz dotknąłeś klamki samochodowej, albo nawet ręki drugiej osoby (natychmiast cofając dłoń) poczułeś charakterystyczne i trwające ułamek sekundy „kopnięcie prądem”. Każda taka sytuacja to przepływ wysokiego napięcia na przykład pomiędzy dwiema osobami, albo czyjąś ręką i samochodową klamką na skutek tzw. wyładowania elektrostatycznego (ESD, Electrostatic Discharge). To co dla Ciebie będzie zaledwie chwilowym niemiłym doświadczeniem, dla elektroniki (na przykład dla układu scalonego) potrafi być wręcz „zabójcze”. Dlatego lepiej by tego typu napięcie trafiło do szyn zasilania i zostało przefiltrowane na kondensatorach, niż „przemieściło” układ scalony od środka, trafiając na jedno z wrażliwych a niepodłączonych do niczego wejść. Z tego samego powodu, przed przystąpieniem do prac z elektroniką warto zawsze usunąć zgromadzone w ciele ładunki elektryczne, dotykając dłonią, na przykład, niemalowanej części metalowego kaloryfera, przy założeniu, że metalowa (miedziana bądź stalowa) instalacja centralnego ogrzewania, będzie wystarczając dobrym sposobem na odprowadzenie zgromadzonych w ciele ładunków. Dodatkowo na schemacie znajdziesz kondensatory C2 i C3, które, zgodnie z zasadami projektowania płytek PCB powinny zostać wpięte możliwie najbliżej wejść zasilania przy każdym z układów scalonych. Pamiętaj jak działają kondensatory? Magazynują one ładunek elektryczny w polu elektrycznym. Te miniaturowe „magazyny energii” pozwalają natychmiast uzupełnić wszystkie „niedobory napięcia” na liniach zasilania. Tym samym przeciwstawiają się niewielkim skokom i wahaniom napięcia i przez to skutecznie eliminują wszystkie pojawiające się tam zakłócenia i szumy. W naszym



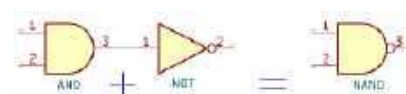
Rysunek 3. Symbol oraz tabela prawdy bramki logicznej typu AND



Rysunek 6. Symbol oraz tabela prawdy bramki logicznej typu NAND



Rysunek 4. Symbol oraz tabela prawdy bramki logicznej typu NOT



Rysunek 5. Bramka NAND jako połączenie bramek AND oraz NOT

układzie gwarantują układom scalonym możliwość najlepszej, stabilnej pracy. O ile niskie pojemności, rzędu 100 nF (stu nanofaradów) pozwalają skutecznie walczyć z zakłóceniami o wysokich częstotliwościach, o tyle większe pojemności pozwalają skutecznie przeciwdziałać tętnieniom napięcia o niższych częstotliwościach. Taki właśnie kondensator o nieco większej pojemności 100 µF (stu mikrofaradów) i oznaczeniu C1 został umieszczony na płytce PCB w pobliżu wejścia zasilania. Niemniej jednak, przy zasilaniu układu z „czystego” napięcia wprost z baterii, jego rola będzie tu mocno ograniczona. Łuźno „wiszące” na schemacie pary wyprowadzeń, opisane jako US1P i US2P i podłączone do VCC i GND to właśnie wejścia zasilania układów scalonych, odpowiednio US1 i US2 (napięcia te rozprowadzane są wewnątrz układów scalonych do wszystkich pojedynczych bramek). W ten sposób oznacza się właśnie na schematach miejsce doprowadzenia zasilania do układów scalonych, bez konieczności domalowywania wejść zasilania przy każdej z ośmiu widocznych na naszym schemacie bramek logicznych. Całkiem wygodne, nieprawdaż?

Nie wiem, czy zwróciłeś na to uwagę, ale na schemacie konsoli Audiochaos większość bramek ma trwale połączone ze sobą oba swoje wejścia. Niepołączone razem wejścia pozostawiono jedynie dla bramek US1C i US2A. W praktyce oznacza to tylko tyle,

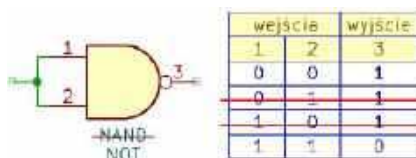
**Fotografia 2. Obcinanie nadmiaru wyprowadzeń po przyłutowaniu rezystorów do PCB. Na zdjęciu Tymek (Wrocław, koło zainteresowań Młodych Entuzjastów Elektroniki)**



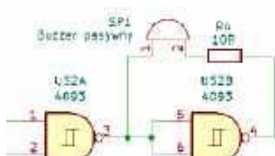
że wszystkie pozostałe bramki mają już tylko pojedyncze wejście (logiczna „1” albo „logiczne „0” będą pojawiały się na obu wejściach jednocześnie). W ten sposób bramki NAND funkcjonalnie zredukowane zostały do prostych bramek NOT, co pokazano na **rysunku 7**.

Teraz, jeśli dobrze się przyjrzyj schematowi z rysunku 1, zauważysz, że bramka logiczna US2B to bramka typu NOT (bramka NAND zredukowana do pełnienia funkcji NOT). Specjalnie dla Ciebie obwód sterowania pasywnym buzzerem narysowałem trochę inaczej (**rysunek 8**).

Jeśli na wyjściu bramki US2A pojawi się logiczne „1”, czyli napięcie bliskie dodatniemu potencjałowi zasilania (w przypadku zasilania układu z baterii 9 V będzie to właśnie 9 V), to na wyjściu bramki US2B pojawi się potencjał bliski GND (czyli 0 V). I odwrotnie: jeśli na wyjściu bramki US2A pojawi się logiczne „0” (czyli napięcie bliskie 0 V), to na wyjściu bramki US2B pojawi się napięcie bliskie 9 V. Innymi słowy polaryzacja napięcia płynącego przez buzzer pasywny piezo i rezystor R4 będzie się cyklicznie zmieniała, co z kolei będzie cyklicznie odkształcało element piezoelektryczny, raz to w jedną, raz w drugą stronę, generując drgania (docelowo dźwięk) w zakresie częstotliwości ustawionej za pomocą generatorów, o których za moment Ci opowiem. To co dobrze jest zauważyć, to to, że buzzer pasywny sterowany jest tutaj nie dostarczaniem i przerwą w dostawie prądu, ale cyklicznym przepływem prądu, raz w jednym kierunku, raz w drugim. Dzięki temu kryształ piezoelektryczny, nie jest poddawany wyłącznie naprzemiennemu odgięciu (w jednym kierunku) i spoczynkowi, ale jest odkształcany cyklicznie w obydwu kierunkach, raz w jedną stronę, a raz w drugą. Dzięki powyższemu odkształceniu są intensywniejsze a dźwięk głośniejszy.



**Rysunek 7. Bramka NAND funkcjonalnie zredukowana do roli bramki NOT przez trwałe elektryczne połączenie ze sobą obu wejść**



**Rysunek 8. Układ zmieniający polaryzację buzzera na podstawie stanu logicznego panującego na wyjściu bramki US2A**



**Fotografia 1. Lutowanie rezystorów do płytki PCB. Na zdjęciu Tymek (Wrocław, koło zainteresowań Młodych Entuzjastów Elektroniki)**

Kilka zdań wcześniej wspomniałem o generatorach częstotliwości. Trzy niezależne, indywidualnie regulowane generatory stanowią „elektroniczne serce” naszej wspianej konsoli.

Pierwszy generator częstotliwości zbudowany jest z użyciem bramki NAND US1A, kondensatora C4 oraz rezystancji złożonej z szeregowo połączonych rezystora R1 oraz potencjometru PR1. Drugi generator częstotliwości zbudowany jest w sposób analogiczny, z użyciem bramki NAND US1D, kondensatora C5 oraz rezystancji złożonej z szeregowo połączonych rezystora R2 oraz potencjometru PR2. Oba te generatory w sposób bardzo losowy i chaotyczny, w oparciu o losowe nastawy potencjometrów PR1 i PR2, sterują wejściami bramki NAND US1C, na której dokonywana jest dodatkowo operacja NOT AND. Przebieg z wyjścia bramki US1C trafia na... jeszcze jedną bramkę NAND, którą można by rzec, miksuje tak uzyskany sygnał z przebiegiem uzyskiwanym z trzeciego już generatora, zbudowanego znów w sposób analogiczny tym razem z użyciem bramki NAND US2C, kondensatora C6 oraz rezystancji złożonej z szeregowo połączonych rezystora R3 oraz potencjometru PR3. Elementy stanowiące pojemności i rezystancje dla każdego z generatorów zostały dobrane w taki sposób, by zapewnić nieco odmienne pasmo częstotliwości, dostrajane oczywiście w dosyć szerokim zakresie za pomocą potencjometrów. Na koniec sygnał w specyficzny sposób „zmiksowany” za pomocą bramki NAND US2A trafia na pierwsze wyprowadzenie pasywnego buzzera piezo, a do jego drugiego wyprowadzenia trafia napięcie logicznie odwrócone przez bramkę NOT (bramkę NAND zredukowaną do funkcji logicznej NOT). Tym sposobem wymuszany zostaje cykliczny przepływ prądu przez buzzer pasywny raz w jednym kierunku, raz w drugim. Zmiany napięć na wyprowadzeniach buzzera SP1 prowadzą do powstania drgań elementu piezoelektrycznego i w konsekwencji do powstania dźwięku, o czym opowiedziałem na samym początku. Tyle „czarnej magii” na dzisiaj, a skoro wiesz już, jak urządzenie



**Fotografia 3.** Po prawej stronie zdjęcia widoczne jest pótokrągłe zagłębienie w materiale układu scalonego, podstawki oraz nadruk na warstwie opisowej płytki PCB pokazujący właściwy kierunek montażu



**Fotografia 4.** Niektóre z pinów układu scalonego nie przeszły na drugą stronę płytki PCB. Uczestnik zajęć tego nie zauważył i przylutował podstawkę, tworząc trudną do wykrycia usterkę



**Fotografia 5.** Podstawka po wylutowaniu z płytki. Dwa wyprowadzenia były zagięte i schowane pomiędzy podstawką a płytką drukowaną. Połączenie elektryczne nie mogło zaistnieć

z tego odcinka jest skonstruowane możecie wraz z Twoim kompanem lub kompanką z czystym sumieniem przystąpić do budowania układu.

## Montaż urządzenia

Pora na część najprzyjemniejszą, czyli montaż. Lutować nauczyłeś się w ubiegłym miesiącu. Teraz pozostaje Ci z tych umiejętności skorzystać. Jednocześnie to znowu ten moment, w którym powinieneś poprosić o uczestnictwo towarzysza Twych zabaw, a przynajmniej zadbać o to, by był w pobliżu. Dla porządku przypomnę tylko, że będziesz trzymał się dwóch zasad: „od najmniejszych do największych” i „nie wszystko na raz”. Innymi słowy proponuję Ci włożenie do płytki a następnie przylutowanie do niej kolejnych grup komponentów. W tym zestawie najniższe są rezystory zatem włóż je do płytki na odpowiednich pozycjach, zgodnie z listą komponentów dostępną w instrukcji dołączonej do zestawu. Rezystory w zestawie posiadają dwie różne wartości. R1, R2 i R3 to oporniki o rezystancji 10 kΩ (czytaj: „kilo-omów”, a więc są to rezystory o wartości dziesięciu tysięcy omów każdy, bo kilo znaczy tysiąc). Rezystor R4 ma wartość 10 Ω (czytaj: „omów”). Wartości rezystancji opisywane są na obudowach rezystorów za pomocą kodu złożonego z kolorowych pasków, a lista elementów podaje układ kolorów dla danej wartości i jednoznacznie wskazuje, która wartość na jakich desygntorach powinna być zamontowana. Po włożeniu rezystorów do płytki rozegnij ich wyprowadzenia pod lekkim kątem tak, żeby nie wypadły gdy odwrócisz płytkę do góry stroną lutowania. Następnie przylutuj wszystkie rezystory do płytki.

Oczywiście rezystory są komponentami symetrycznymi, czyli takimi, które posiadają identyczne właściwości elektryczne w obu kierunkach przepływu prądu.

Oznacza to, że ich działanie nie zmienia się w zależności od kierunku, w którym prąd przez nie przepływa, a to z kolei oznacza, że można je montować w dowolnym kierunku. Po ich przylutowaniu pozbądź się za pomocą obcinaczek nadmiaru wyprowadzeń (przypomnij sobie co mówiłem na naszym poprzednim spotkaniu odnośnie posługiwania się obcinaczkami, ponieważ nieodpowiednio je używając, pady lutownicze łatwo oderwać od płytki i ścieżek, a tego zdecydowanie nie chcemy).

Następnie zamontuj dwie podstawki pod układy scalone na pozycjach oznaczonych desygntorami US1 oraz US2. Jest szansa, że będą one niższe niż kondensatory, które trzeba będzie zamontować w ich pobliżu, zatem wygodniej najpierw zamocować podstawki. W przypadku podstawek, kierunek nie jest obojętny. Nie dlatego, że elektrycznie cokolwiek on zmieni, ale dlatego, że podstawki mają na swojej obudowie wskazówkę: znacznik kierunku dla osoby, która później będzie montowała w niej układ scalony. Układy scalone to zdecydowanie komponenty niesymetryczne, wymagające właściwego kierunku montażu. Zarówno podstawka, jak i później osadzany w niej układ scalony muszą zostać zamontowane zgodnie z kierunkiem zaznaczonym na warstwie opisowej płytki drukowanej, jak pokazano na **fotografii 3**.

Zamontowanie układu w odwrotnym kierunku i podłączenie zasilania w większości przypadków spowodują natychmiastowe uszkodzenie układu scalonego oraz bardzo prawdopodobną jego eksplozję i wystrzał fragmentów obudowy we wszystkich kierunkach (**uwaga na oczy!**). Dlatego zawsze proszę chłopaków, by podczas podłączania do zmontowanego układu zasilania mieli założone okulary ochronne.

Podstawki posiadają zazwyczaj metalowe wyprowadzenia zrobione z bardzo cienkiej blachy. Łatwo je zgiąć lub wykrzywić. Przed

próbą włożenia ich do płytki, proszę upewnij się, że wszystkie wyprowadzenia są proste i prostopadłe do podstawy komponentu. Jeśli znajdzie taka potrzeba, powykrzywiane piny ostrożnie i z wyczuciem naprostuj i wyrównaj, na przykład przy użyciu śrubokręta. Podczas wkładania podstawki do płytki nigdy się nie spiesz. Zanim zdecydujesz się włożyć podstawkę do płytki, upewnij się, że wszystkie otwory mają prześwit na drugą stronę, i czy żaden jakimkolwiek sposobem nie został zalutowany (zatkany cyną). Zanim dociśniesz ciało komponentu do płytki PCB, upewnij się, czy wszystkie wyprowadzenia trafiły w otwory, albowiem lubią one zahaczać o ścianki otworów i zamiast przechodzić przez otwór, potrafią się zagiąć i schować pomiędzy płytką i komponentem lub też, pod wpływem działającej na nie siły, wysunąć się z plastikowej ramy podstawki. Gdybyś przylutował podstawkę nie zauważając zagiętego pinu ciężko byłoby taką usterkę wykryć i jeszcze trudniej naprawić. Mówię Ci o tym, ponieważ moim chłopakom zdarza się to co jakiś czas i naprawa jest zawsze żmudna i czasochłonna.

**Fotografia 4 i fotografia 5** przedstawiają taki „przykład z życia”. Po zmontowaniu jednego z zestawów układ nie działał. Po odesaniu cyny z nóżek układu scalonego okazało się, że dwa piny podstawki w ogóle nie przeszły na drugą stronę płytki. Po odlutowaniu i wyciągnięciu podstawki z płytki „zagubione” nóżki „odnalazły” się. Były schowane między

**Fotografia 6.** Lutowanie podstawek pod układy scalone do płytki. Na zdjęciu Łukasz (Wrocław, koło zainteresowań Młodych Entuzjastów Elektroniki)





**Fotografia 7. Obszar poprawnie przylutowanej podstawki układu scalonego do płytki PCB (przy założeniu, że podstawka została włożona do płytki zgodnie z oznaczeniem kierunku na warstwie opisowej)**

komponentem i płytką drukowaną. To nie miało prawa działać.

Po włożeniu podstawek do płytki, upewnieniu się, czy każdy z pinów w całości przeszedł przez własny otwór, oraz zagięciu skrajnych ich wyprowadzeń można bezpiecznie odwrócić płytkę stroną lutowania do góry i przylutować oba układy scalone (**fotografia 6**).

Luty powinny być wolne od zwarcie pomiędzy sąsiednimi padami, cyna powinna obficie rozplwać się wokół pól lutowniczych i wyprowadzeń, a same luty powinny być błyszczące. Obszar każdego z układów scalonych po zakończeniu lutowania powinien wyglądać mniej więcej tak, jak pokazano na **fotografii 7**.

W przypadku podstawek, po ich przylutowaniu nie obcinaj nadmiaru wyprowadzeń. Wyprowadzenia podstawek są tak krótkie, że po pierwsze, nie ma czego obcinać, po drugie, wzrosłoby ryzyko uszkodzenia lutów i zerwania padów, a tego przecież nie chcemy.

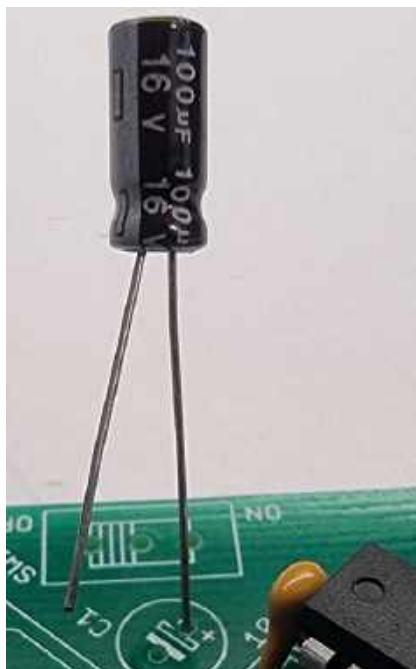
Pora na zamontowanie kondensatorów stałych o pojemności 100 nF (czytaj sto „nanofaradów”). Rozpoznaś je po oznakowaniu na obudowie w postaci napisu 100 nF, albo 0,1  $\mu$ F (100 nanofaradów to 0,1 mikrofarada) albo też nadruku 104 (10 i cztery zera w pikofaradach, czyli 100000 pF (100000 pikofaradów to 100 nanofaradów). Jak wspomniałem, są to kondensatory stałe (niespolaryzowane, symetryczne) więc kierunek ich montażu jest dowolny. Zamontuj je na pozycjach C2, C3, C5, C6, następnie rozegnij lekko ich nóżki, tak by nie wypadły po odwróceniu płytki a następnie przylutuj.



**Fotografia 8. Montaż pasywnego buzzera piezo**

Kierując się polityką zwiększania wysokości montowanych elementów przylutuj teraz buzzer pasywny piezo. Powinien on zostać zamontowany na pozycji oznaczonej na warstwie opisowej PCB desygnatorem SP1 (**fotografia 8**).

W układzie wykorzystano buzzer pasywny, czyli taki bez zaimplementowanej elektroniki, która po podaniu napięcia generowałaby jakiś sygnał dźwiękowy. Generowaniem przebiegów audio zajmie się zbudowana przez Ciebie Konsola Audiochaos. Buzzer pasywny (w przeciwieństwie do buzzerów aktywnych z wbudowaną elektroniką) jest elementem niespolaryzowanym i może być zamontowany w dowolnym kierunku. Wysokość pasywnego buzzera dołączonego do zestawu była podobna do wysokości podstawek pod układy scalone, więc jeśli zachowałeś sugerowaną kolejność montażu elementów, po włożeniu buzzera do płytki możesz przytrzymać go palcem tak, by nie wypadł (jego twarde wyprowadzenia mogą być trudniejsze do zagięcia) i położyć płytkę na stole. Odwrócona do góry stroną lutowania płytka powinna teraz leżeć na podstawkach i buzzerze, wystarczająco dobrze go unieruchamiając. Przylutuj wyprowadzenia buzzera do pól lutowniczych na płytce. Czasami buzzery piezo zabezpieczone są na czas montażu i czyszczenia elektroniki w taki sposób, że specjalna naklejka zakrywa otwór przez który emitowany jest dźwięk. Jeśli widzisz



**Fotografia 9. Montaż kondensatora elektrolitycznego, zgodny z oznaczeniami na warstwie opisowej płytki PCB. Dłuższe wyprowadzenie trafia do otworu ze znakiem „+”. Krótsze wyprowadzenie to „-” co zaznaczone jest również podłużnym białym nadrukiem ze znakiem „-” na korpusie kondensatora**

u siebie taką naklejkę, nie zapomnij jej odkleić, w przeciwnym razie dźwięk będzie bardzo skutecznie tłumiony (o ile w ogóle słyszalny).

Następnym elementem do zamontowania jest włącznik zasilania i należy go umieścić w pozycji opisanej jako SW1. Przełącznik jest komponentem symetrycznym. Niezależnie od kierunku zamontowania, będzie działał dokładnie tak samo. Po włożeniu przełącznika do płytki odwróć ją, przytrzymując włącznik palcem w taki sposób, aby nie wypadł, i albo unieruchom przycisk pod ciężarem płytki pomiędzy stołem a płytką, lub też przytrzymaj przełącznik palcem na czas lutowania. W obu przypadkach przylutuj najpierw pin środkowy, a następnie odwróć płytkę, i sprawdź, czy włącznik przylutował się równolegle do powierzchni płytki. Jeśli wszystko jest równe, przylutuj do płytki dwa pozostałe jego wyprowadzenia.

Jako następne przylutuj dwa kondensatory elektrolityczne na pozycjach C1 i C4. Oba mają w tym konkretnym zestawie tę samą wartość 100  $\mu$ F (czytaj: sto „mikrofaradów”). Należy jedynie pamiętać o tym, że kondensatory elektrolityczne montujemy zgodnie z polaryzacją zdefiniowaną na obudowie kondensatora. Białe paski wskazujące ujemne wyprowadzenie kondensatora zazwyczaj jest bardzo wyraźny. Naprawdę trudno byłoby takiego oznaczenia nie zauważyć. Ponadto nowe (nie przecięte jeszcze) kondensatory elektrolityczne mają jedno wyprowadzenie dłuższe a drugie krótsze. Dłuższe to „plus” a krótsze to „minus”. To już druga cecha po której można rozpoznać polaryzację kondensatora. Pozostaje więc poprawnie (dłuższym wyprowadzeniem do namalowanego na warstwie opisowej znaczka „+”) włożyć je do płytki (**fotografia 9**), odgiąć (jak zwykle) wyprowadzenia a następnie przylutować.

Po przylutowaniu obu kondensatorów elektrolitycznych nadeszła pora na zamocowanie potencjometrów. Dwa z tych potencjometrów mają wartość 100 k $\Omega$  (czytaj: sto „kilo-omów”) i te muszą zostać zamontowane na pozycjach PR1 i PR3. Trzeci z potencjometrów ma wartość 1 M $\Omega$  (czytaj: jeden „megaom”, mega znaczy milion, innymi słowy



**Fotografia 10. Potencjometry z wyprowadzeniami wygiętymi w stronę tulei regulacyjnych**

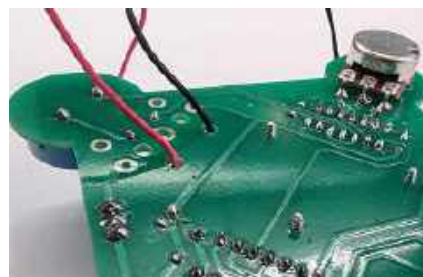


**Fotografia 11. Przykręcenie potencjometru do płytki PCB (a), wycentrowanie potencjometru względem pól lutowniczych (b) oraz lutowanie wyprowadzeń uprzednio przykręconego do płytki i wycentrowanego względem padów potencjometru (c)**

jeden megaom to tyle co milion omów) i ten powinien zostać zamontowany na pozycji PR2. W przypadku potencjometrów (inna nazwa to rezystory nastawne) chodzi oczywiście o maksymalne możliwe do nastawienia opory, a ich możliwa do uzyskania wartość mieści się zawsze w przedziale od zera omów do wspomnianej przed chwilą wartości maksymalnej.

Najpierw wyginamy piny tych potencjometrów pod kątem prostym w stronę ich tulei regulacyjnych. Najłatwiej uczynić to w sposób odpowiedni i z wycuciem dociskając odpowiednio wyprowadzenia potencjometrów do blatu stołu roboczego (odpowiednio zabezpieczając go jakąś deseczką, czy twardą tekturą, żeby sobie nie porysować biurka). Tak przygotowane potencjometry (**fotografia 10**) montujemy jeden po drugim na odpowiednich pozycjach na PCB.

W tym celu każdy z potencjometrów przewlekamy tuleją regulacyjną (ale jeszcze bez założonej, plastikowej gałki) przez duży otwór na PCB, centrujemy wyprowadzenia potencjometru względem pól lutowniczych (**fotografia 11b**), od strony górnej nakładamy podkładkę i przykręcamy nakrętkę (**fotografia 11a**). Czynności te powtarzamy dla każdego z potencjometrów. Po solidnym przykręceniu potencjometrów do płytki i upewnieniu się, że wszystkie wyprowadzenia są dobrze wycentrowane względem pól lutowniczych, lutujemy wszystkie wyprowadzenia do płytki, obficie zalewając je cyną (**fotografia 11c**).



**Fotografia 12. Na zdjęciu widać gwinty śrub wkręconych w płytkę. Kabelki koszyeczka baterii zostały przewleczone przez płytkę przed jego zamontowaniem**

Na końcu zamontuj koszyczek na baterię 9 V. Najpierw za pomocą małego, precyzyjnego śrubokręta krzyżakowego przykręć koszyczek do płytki PCB. Koszyczek należy przymocować od strony tulei regulacyjnych potencjometrów po uprzednim przewleczeniu kabelków koszyeczka baterii przez otwory (**fotografia 12**) znajdujące się pod koszyczkiem (po przykręceniu koszyeczka nie będzie do nich dostępu, zatem w pierwszej kolejności należy przewlec przez otwory płytki wspomniane kable). Pola lutownicze służące do przylutowania kabelków baterii nie są w żaden sposób na płytce PCB oznakowane, dlatego też pozostaje Ci kierować się fotografiami zamieszczonymi poniżej.

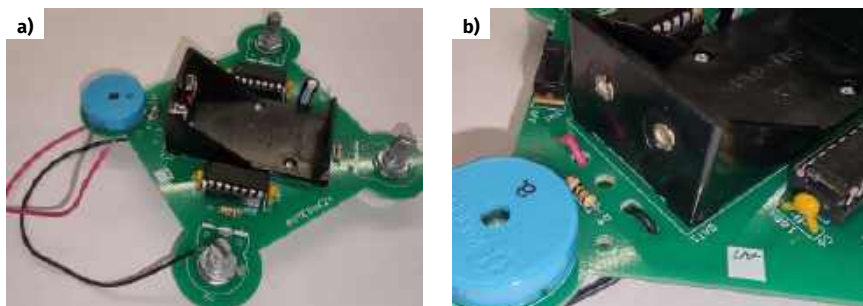
Po przeprowadzeniu kabelków przez pary otworów, zgodnie z **fotografiami 13a i 14a**, można uciąć ich nadmiar (**fotografia 14b**), a następnie usunąć izolację (**fotografia 14c**). Najwygodniej będzie tego dokonać przy użyciu obcinaczek. Po ich przyłożeniu do płaszczyzny płytki, należy zaciśnąć je na szerokości izolacji, ale delikatnie i z wyczuciem, żeby nie uciąć miedzianego rdzenia przewodu. Takimi nie do końca zaciśniętymi obcinaczkami można podeprzeć się i odbić od powierzchni płytki i tym sposobem, jak za pomocą dźwigni, pozbyć się izolacji z końcówek kabelków. Na koniec przylutuj obrane z izolacji przewody do pól lutowniczych na płytce (**fotografia 14d**).

Na koniec zamontuj układy scalone w podstawkach. Postępuj podobnie jak w przypadku

wkładania podstawek do płytki, upewniając się, że wkładasz te układy do podstawek we właściwym kierunku, zgodnie z oznaczeniem namalowanym na płytce drukowanej. Przypilnuj, by wszystkie wyprowadzenia układów scalonych trafiły do swoich gniazd, a następnie dociśnij układy do podstawek.

Na wszystkie trzy potencjometry nałóż plastikowe gałki regulacyjne. Rozkład kółców według uznania (Ty decydujesz).

Po wizualnym sprawdzeniu poprawności i jakości montażu pod kątem obłożenia płytki komponentami, zachowania polaryzacji dla komponentów niesymetrycznych, ewentualnego usunięcia zwarcia czy poprawy lutów wątpliwej jakości przyszedł czas na to, by do koszyeczka włożyć baterię i ustawić włącznik SW1 w pozycji „ON”. Zanim to zrobisz, załóż okulary ochronne (jeśli jakimś sposobem nie masz ich w tej chwili na nosie) na wypadek gdyby któryś układ scalony lub kondensator miał wystrzelić na skutek ewentualnych błędów montażu. Dodatkowo możesz znacznie zminimalizować ryzyko uszkodzenia układów scalonych wyjmując je z podstawek przed podłączeniem zasilania. Wówczas po włączeniu zasilania bez tych układów, będziesz mógł za pomocą multimetru ustawionego w tryb pomiaru napięcia DC w zakresie do 20 V sprawdzić, czy w podstawkach, na pinach docelowo zasilających układy scalone, pojawia się poprawne napięcie. W tym celu, po załączeniu zasilania (bateria włożona, SW1 w pozycji ON) przyłóż sondę czarną multimetru do pinu



**Fotografie 13. Koszyczek na baterię przymocowany do płytki, widok od strony tulei regulacyjnych potencjometrów, a) widok całości, b) sposób prowadzenia kabelków koszyeczka baterii**



**Fotografia 14. Prowadzenie i montaż kabelków koszyeczka na baterię: a) przewleczenie kabelków przez dwie pary otworów, b) docięcie kabelków na odpowiednią długość, c) zdjęcie izolacji, d) przyłutowanie kabelków do płytki PCB**

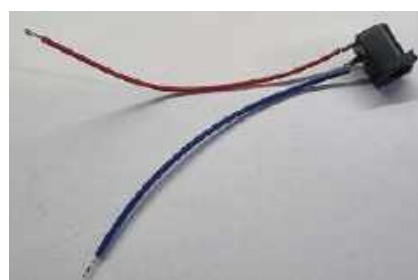
7 a sondę czerwoną do pinu 14 układu scalonego US1. Na multimetrze powinienesz zobaczyć wskazanie około 9 V, koniecznie bez znaku „-” na początku. Jeśli w tej sytuacji widzisz minus na pierwszej pozycji, oznacza to, że masz odwróconą polaryzację (np. kabelki baterii podłączone na odwrot). Jeśli natomiast zmierzylesz napięcie „9 V” lub „+9 V” (bez minusa na początku wskazania) powtórz ten pomiar w analogiczny sposób na podstawie układu US2. Jeśli dla obu układów zmierzylesz napięcie dodatnie 9 V (bez minusa) możesz odłączyć zasilanie, ponownie włożyć w podstawki (zgodnie z polaryzacją!) układy US1 i US2 i mając wciąż na nosie okulary ochronne, spokojnie już ponownie podłączyć zasilanie.

Mam nadzieję, że po włączeniu konsoli usłyszysz dźwięk, który będzie się zmieniał w odpowiedzi na zmianę nastaw każdego z potencjometrów. Układ jest dosyć prosty, więc szansa na to, że zadziała od razu jest całkiem spora. Gdyby jednak urządzenie nie funkcjonowało prawidłowo, na przykład gdyby nie reagowało na regulację któregoś z potencjometrów, wyłącz zasilanie (najlepiej wyciągnij baterię) i wykonaj kilka kroków jak poniżej.

Jeśli nie działa jeden z potencjometrów, podważając uprzednio układy scalone śrubokrętem, możesz je zamienić miejscami (oba są identyczne), wykluczając w ten sposób bądź potwierdzając (jeśli urządzenie

zaczęło działać inaczej) usterkę któregoś z nich. Zwarcia na płytce, gdyby istniały, z pewnością dostrzeżesz gołym okiem, natomiast brak połączeń jest całkiem prawdopodobny, zatem, przy odłączonym zasilaniu konsoli, prześledź miernikiem ustawionym w tryb pomiaru ciągłości, czy jest przejście (czy miernik wydaje sygnał dźwiękowy) gdy jego sondy przykładasz do każdej możliwej pary połączeń lutowanych (pół lutowniczych) pomiędzy którymi widzisz miedziane ścieżki (jaśniejsze zielone linie na ciemnozielonym tle). Płytkę jest jednostronna, zatem w otworach nie ma metalizacji i z tego powodu o zerwanie padu podczas lutowania lub obcinania nadmiaru wyprowadzeń komponentów po ich przyłutowaniu wcale nie trudno.

W przypadku naszych zajęć stacjonarnych, u wszystkich chłopaków układ finalnie zadziałał, choć nie obyło się bez przygód i niespodzianek. Jedną z nich był brak ciągłości pomiędzy pewnymi dwoma punktami lutowniczymi. Brak delikatności przy obcinaniu nogi komponentu po jego przyłutowaniu spowodował oderwanie padu od ścieżki. Usterka niewidoczna gołym okiem ale szybka do wyłapania przy sprawnym posługiwaniu się miernikiem w trybie pomiaru ciągłości. Drugą przygodą na tej samej płytce była noga podstawki, która nie przeszła przez otwór płytki podczas montażu, która to noga zagięła się i utknęła gdzieś pomiędzy płytką a podstawką. W takiej sytuacji najpewniej



**Fotografia 15. Zamiast koszyeczka z baterią, z uwagi na zastosowanie układów scalonych wykonanych w technologii CMOS, do płytki można przyłutować gniazdko i zasilić konsolę z typowego zasilacza napięcia stałego 12 V. Należy tylko pamiętać o zachowaniu prawidłowej polaryzacji zasilania**

byłoby podstawkę odlutować za pomocą odsysacza (opowiem o nim w przyszłości) i lutownicy, a najwygodniej z użyciem stacji rozlutowniczej z kompresorem. Ewentualnie z użyciem plecionki lutowniczej (o niej też kiedyś opowiem). Czasem jednak stosując drobny trik można sobie zaoszczędzić sporo pracy. Może nie jest to optymalna naprawa, ale czasem warto spróbować. Otóż rozgrzaaliśmy pole lutownicze po czym w miejsce brakującego pinu dopchnęliśmy ucięte wyprowadzenie jakiegoś rezystora. Szczęśliwie wyprowadzenie złapało dobry kontakt elektryczny z pinem, który utknął między podstawką i płytką, tym samym zabawkę udało się naprawić bardzo niewielkim nakładem pracy.

Jako ciekawostkę warto dodać, że ponieważ w urządzeniu zastosowano układy scalone zawierające bramki logiczne wykonane w technologii CMOS (seria układów 40xx), urządzenie bez żadnych przeróbek można śmiało zasilić również napięciem stałym 12 V z typowego zasilacza wtyczkowego. W takim przypadku można zrezygnować z montowania koszyeczka i kupowania baterii. Zamiast tego, pamiętając o zachowaniu polaryzacji, można przyłutować kabelki zakończone gniazdem pasującym do posiadanego zasilacza (**fotografia 15**).

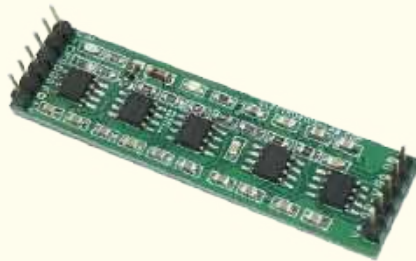
Chociaż dzisiejsze spotkanie dobiega końca, jak zawsze gorąco zachęcam Cię do podzielenia się z pozostałymi młodymi Entuzjastami Elektroniki swoimi wrażeniami i przygodami podczas montażu dzisiaj omówionego zestawu. Poczta elektroniczna i tradycyjna pozostają do Twojej dyspozycji. A już za miesiąc zaproszę Cię do zbudowania kolejnej zabawki z serii AVTEDU, która z pewnością zapewni kolejną porcję niecodziennych, może wręcz kosmicznych wrażeń. Tymczasem przybijam Ci piątkę i, już tęskniąc nieco, żegnam się z Tobą aż do przyszłego miesiąca. Do zobaczenia! ■

**Mariusz Ciszewski**

Przedstawiamy początkowe fragmenty dwóch projektów ze zbioru kilkudziesięciu projektów dostępnych wyłącznie dla prenumeratorów EdW w rubryce **DIY PLUS** na [www.elportal.pl](http://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](http://www.elportal.pl) z niezwykle inspirującymi zasobami rubryki **DIY PLUS**.

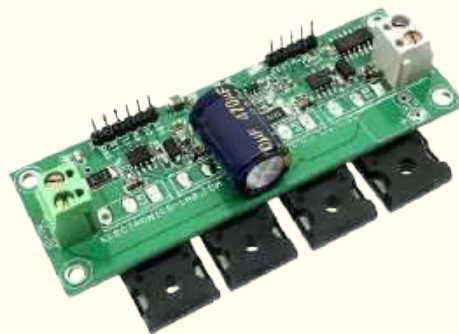
## Pojemnościowy czujnik wilgotności do konwertera wyjścia analogowego

Przedstawiony tutaj projekt to płytki konwertera analogowego pojemnościowego czujnika wilgotności odpowiednia do różnych zadań związanych z pomiarem wilgotności. Obwód pochodzi z noty aplikacyjnej Texas Instruments (patrz poniżej). Większość konwencjonalnych pojemnościowych czujników zbliżeniowych wytwarza sygnał cyfrowy, ale ten obwód wytwarza napięcie wyjściowe DC, które jest funkcją wilgotności względnej otoczenia. Wejście oscylatora dla tego obwodu jest źródłem sygnału sinusoidalnego o częstotliwości 1 kHz (1 Vpp) podawanego na wejście FQ. Niska częstotliwość i przebieg sinusoidalny ograniczają problemy RFI do minimum. Ten konwerter wyjściowy czujnika może okazać się trudny do skonstruowania ze względu na małe pojemności, pasożytnicze pojemności w obwodzie czujnika i wychwytywanie szumów. Obwód wymaga zewnętrznej fali sinusoidalnej 1 kHz z sygnałem 1 Vpp i działa z podwójnym zasilaniem 15 V DC.



## Mostek H dla wysokiej mocy szczotkowego silnika prądu stałego z czujnikiem prądu

Ta płytki z mostkiem H to łatwy w użyciu sterownik szczotkowego silnika prądu stałego, który może obsługiwać duże silniki, takie jak silniki wózków inwalidzkich i wciągarek. Mostek H to obwód elektroniczny, który przełącza polaryzację napięcia przyłożonego do obciążenia. Obwód został zbudowany przy użyciu 4 tranzystorów MOSFET typu N, sterownika bramki IR2104 i obwodu logicznego. Kierunek obrotów zależy od polaryzacji przyłożonego napięcia. Odwrócenie napięcia powoduje zmianę kierunku obrotów.



Niektóre projekty aktualnie dostępne tylko dla prenumeratorów EdW w rubryce **DIY PLUS** na [www.elportal.pl](http://www.elportal.pl):

1. Przetwornica DC-DC buck 12...75 V na 10 V na wyjściu
2. Czujnik prądu low-side 10 µA...10 mA
3. Kontroler ramienia robota z bezprzewodowym pilotem PS3
4. Termiczny czujnik masowego przepływu powietrza – anemometr stałotemperaturowy
5. Precyzyjny wzmacniacz transimpedancjny z przetwarzaniem integratorem
6. Kontroler pełnego mostka z przesunięciem fazowym i prostowaniem synchronicznym wykorzystujący UCC28950
7. Wysokowydajny monofoniczny wzmacniacz audio klasy D o mocy 20 W
8. Monitorowanie poziomu cieczy za pomocą czujnika ciśnienia – wyświetlacz słupkowy
9. Sterowanie silnikiem DC za pomocą joysticka
10. 16-kanalowy sterownik serwo mechanizmów RC z interfejsem I<sup>2</sup>C
11. Programowalny kondycjoner sygnału z czujnika rezystancyjnego mostkowego
12. Choinka z Arduino i pikselowymi diodami
13. 20-segmentowy wyświetlacz słupkowy w rozmiarze jumbo
14. Stacja pogodowa i lilo tygo t5-4.7 z wyświetlaczem typu e-papier
15. Półprzewodnikowy przełącznik mocy DC z prądowym sprzężeniem zwrotnym
16. Wyłącznik nadprądowy – przełącznik wyłączający nadprądowy
17. Uniwersalny konwerter napięcia AC – wyjście 18 V DC z wejścia 85...265 V AC
18. Moduł procesora echa głosu – urządzenie opóźniające do efektów dźwiękowych, echo, reverb
19. Najlepszy sposób na próbkowanie dźwięku za pomocą ESP32
20. Sterownik silnika krokowego z joystickiem
21. RPI – stacja pogodowa IoT
22. Niskobudżetowy monitor jakości powietrza IoT oparty o RaspberryPi 4
23. Automatyczny system ogrodniczy z NodeMCU i Blynk, ArduFarmBot 2
24. TinyML – Rozpoznawanie ruchu przy pomocy RPI Pico
25. Wzmacniacz piezoelektryczny do gitary i skrzypiec
26. Wysokowydajny i niezawodny sterownik bipolarnego silnika krokowego
27. Sonarowy theremin MIDI
28. Sterownik silnika prądu stałego z wykorzystaniem przełącznika i mosfetu – interfejs Arduino
29. Przedwzmacniacz do mikrofonu MEMS
30. Super prosty czuły wykrywacz metali
31. Stymulator czaszkowy Arduino (Bio-BrainTuner)
32. Generator sygnałów AD9833
33. Obserwacja charakterystyk tranzystora
34. Wyświetlacz EKG z użyciem Arduino
35. Łatwy do zbudowania robot kroczący
36. Zamek elektroniczny na kod
37. Prosty tester tranzystorów
38. Zegar binarny z użyciem Microbit
39. Izolowany obwód wykrywania napięcia 250 V AC z pojedynczym wyjściem (wejście 250 V prądu przemiennego, wyjście 5 V)

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](mailto:avt@avt.pl)

**Wydawca:**  
Wiesław Marciniak

**Redaktor naczelny:**  
Mariusz Ciszewski  
[mariusz.ciszewski@elportal.pl](mailto:mariusz.ciszewski@elportal.pl)

**Adres redakcji:**  
03-197 Warszawa, ul. Leszczyńska 11  
e-mail: [edw@elportal.pl](mailto:edw@elportal.pl), [www.elportal.pl](http://www.elportal.pl)

**Redaktor merytoryczny:**  
[Paweł Sujko](mailto:Pawel.Sujko@elportal.pl)

**Dział reklamy:**  
Katarzyna Gugała  
[katarzyna.gugala@elportal.pl](mailto:katarzyna.gugala@elportal.pl), tel. 22 257 84 64

**Szef Pracowni Konstrukcyjnej:**  
Jakub Sobański  
[jakub.sobanski@elportal.pl](mailto:jakub.sobanski@elportal.pl)

**Sekretarz redakcji:**  
Dariusz Welik  
[dariusz.welik@elportal.pl](mailto:dariusz.welik@elportal.pl)

Copyright AVT-Korporacja Sp. z o.o., Warszawa, ul. Leszczyńska 11. Projekty publikowane w „Elektronice dla Wszystkich” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki dla Wszystkich”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice dla Wszystkich” jest dozwolone wyłącznie po uzyskaniu pisemnej zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice dla Wszystkich”.

**DTP, okładka,**  
**Redakcja strony internetowej [www.elportal.pl](http://www.elportal.pl):**  
MAD Sp. z o.o.

**Prenumerata:**  
W Wydawnictwie AVT, e-mail: [prenumerata@avt.pl](mailto:prenumerata@avt.pl)  
tel. 22 257 84 22, (godz. 10:00–14:00)  
[www.ulubionykiosk.pl](http://www.ulubionykiosk.pl)

W RUCH S.A., e-mail: [prenumerata@ruch.com.pl](mailto:prenumerata@ruch.com.pl)  
tel. 801 800 803, 22 717 59 59, [www.prenumerata.ruch.com.pl](http://www.prenumerata.ruch.com.pl)

# Arrow MultiSolution Day Poland

Warsaw,  
September 10th, 2024

**ARROW** | arrow.com



| Time        | Track 1                            | Track 2    |
|-------------|------------------------------------|------------|
| 08.30-09.00 | Registration                       |            |
| 09.15-10.00 | Opening – Realising AI at the Edge |            |
| 10.00-10.45 | Infineon                           | Zephyr OS. |
| 10.45-11.00 | Coffee Break                       |            |
| 11.00-11.45 | Silabs.                            | FPGA       |
| 11.45-12.00 | Coffee Break                       |            |
| 12.00-12.45 | TE                                 | NXP MCX    |
| 12.45-13.45 | Lunch Break                        |            |
| 13.45-14.30 | Wolfspeed                          | Nvidia AI  |
| 14.30-14.45 | Coffee Break                       |            |
| 14.45-15.30 | Onsemi                             | BMS        |
| 15.30       | Grill                              |            |

## Alm the Power

Już 10 września od 9.00 zapraszamy  
na tegoroczne wydarzenie AMS,  
podczas którego:

- 2× ścieżka szkoleniowa
- showcase 40 producentów
- grill dla gości i producentów



Info z pełną agendą  
wydarzenia + rejestracja:

