

ELEKTRONIKA

dla wszystkich

2/2022 LUTY • CENA 13,90 zł (w tym 8% VAT)

www.elportal.pl

Podwójnie symetryczny Moduł audio

Pulsująca błyskotka LED

- ▶ Automatyczne nawadnianie także dla Ciebie – Podstawy
- ▶ NanoVNA – Wykonaj precyzyjne pomiary
- ▶ Frezarka CNC
- ▶ Miernik wzmacniaczy operacyjnych
- ▶ Z potrzeby chwili... Zasilacz do laptopa w roli ładowarki akumulatora
- ▶ Elektronika i historia
 - Telewizor bez zasilacza?
- ▶ MPPT – Diody ochronne?
- ▶ MODBUS – Bardzo popularny protokół
- ▶ Klasyczny forward z jednor tranzystorowym kluczem
- ▶ Porady warsztatowe
 - Kilka przydatnych wskazówek
- ▶ Moje pasje
 - Własnej roboty magnetofon szpulowy

Drukarki 3D
filamenty, części zapasowe



sklep.avt.pl

Portale branżowe
AutomatykaB2B.pl
ElektronikaB2B.pl

Miejsca dla
specjalistów



artronic
OPTOELEKTRONIKA
www.artronic.pl

FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

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

www.piekarz.pl

INDEKS 333 62X ISSN 1425-1698



9 771425 116922 1

Przenośna stacja lutownicza KD862 na gorące powietrze



Cyfrowa stacja Hot Air KD 862 umieszczona w kolbie. Kompaktowa forma, łatwa do przenoszenia i transportu - wygodne rozwiązanie dla mobilnych serwisantów. Oprócz typowych zastosowań, nadaje się również do spawania tworzyw sztucznych, obkurczania, usuwania starych powłok z farb, itp.



Sterowanie umieszczone w kolbie: pokrętło do regulacji przepływu powietrza (**max 120l/min**) i przyciski do regulacji temperatury (**od 100°C do 480°C**)

- wyświetlacz LED
- zasilanie 230V
- pobór prądu 650W
- długość całkowita 30.5cm
- system schładzania grzałki
- mocna grzałka wykonana z grubego drutu (większa wytrzymałość i trwałość)
- źródło przepływu powietrza: wentylator z silnikiem bezszczotkowym

173zł

W zestawie:

- kolba
- uchwyt
- instrukcja
- 3 dysze okrągłe
- 1 dysza kwadratowa

kod handlowy: KD862



sklep.avt.pl



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

Nowe kursy on-line w Ulubionym Kiosku!

Chcę zostać programistą PLC



PLCspace

S7-1200

TIA Portal

Chcę zostać programistą PLC
Poziom podstawowy



PLCspace

S7-400

Step7 v5.5

Chcę zostać programistą PLC
Poziom podstawowy



PLCspace

S7-1500

TIA Portal

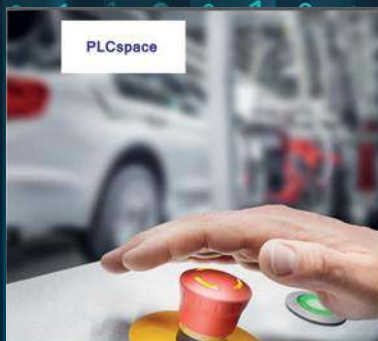
Chcę zostać programistą PLC
Poziom podstawowy



S7-300

TIA Portal

Chcę zostać programistą PLC
Poziom podstawowy



PLCspace

S7-1200F

TIA Portal

Safety
Poziom podstawowy

Z tymi kursami nauczysz się programować tak, jak robią to doświadczeni automatycy!

Zobacz szczegóły i zamów na www.ulubionykiosk.pl/kursy

Firmy
prezentujące
swoje oferty
w niniejszym
wydaniu EdW



ARTRONIC.....1



ELMAX.....29



FERYSTER.....43



PIEKARZ.....1, 21

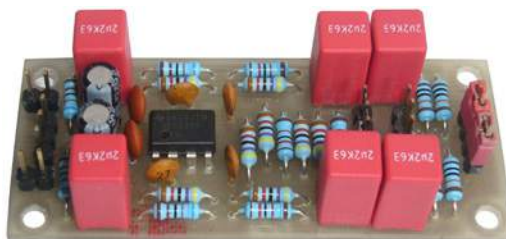


PRODUCENT AUTOMATYKI GRZEWCZEJ

PW KEY.....25



SEMICON.....35



str. 32

Transmisja danych w inteligentnym domu MODBUS – bardzo popularny protokół

MODBUS to wiekowy, ale nadal bardzo popularny protokół, który wykorzystywany jest często w systemach inteligentnego domu. Zapoznaj się ze szczegółami, żeby w pełni wykorzystać czujniki z takim interfejsem.



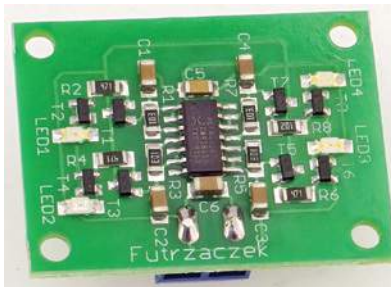
str. 38

NanoVNA

– wykonaj dokładne pomiary

NanoVNA może dokładnie mierzyć także bardzo małe impedancje rzędu miliomów.

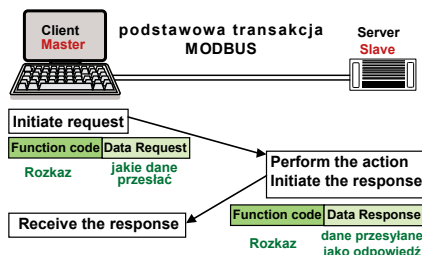
Takie małe oporności trzeba mierzyć z wykorzystaniem obu portów przyrządu, a kluczowe znaczenie ma jak najlepsza kalibracja, którą trzeba wykonać nietypowo.



str. 19

Podwójnie symetryczny moduł audio

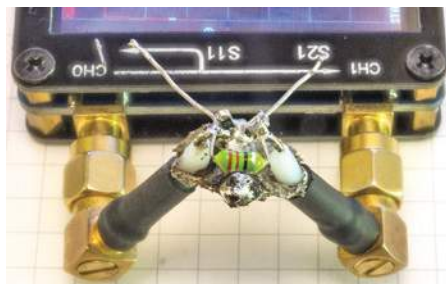
Dlaczego podwójnie symetryczny? Dowiesz się tego podczas lektury artykułu. Powszechnie wiadomo, że połączenia i układy symetryczne są pod wieloma względami lepsze od niesymetrycznych. Miłośnicy audio niewątpliwie zechcą wypróbować proponowany moduł.



str. 34

Automatyczne nawadnianie także dla Ciebie – część 1

Do niedawna systemy automatycznego nawadniania były kosztowną rzadkością. Dziś ceny są przystępne i każdy kto choć trochę ma do czynienia z elektroniką może samodzielnie zrealizować taki system według swoich potrzeb i możliwości.



str. 54

Pulsująca błyskotka LED

Kolorowa, migocząca ozdoba może być urozmaicheniem wielu prezentów od początkującego elektronika. Cztery różnokolorowe diody pulsują płynnie w sobie tylko znanym tempie, co tworzy wrażenie miłego dla oka chaosu.

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”.

Miesięcznik



(12 numerów w roku)
jest wydawany we współpracy
z kilkoma redakcjami
zagranicznymi.

Wydawca:
Wiesław Marciniak

Adres Wydawcy:
AVT-Korporacja sp. z o.o.
ul. Leszczyńska 11
03-197 Warszawa
tel.: (22) 257 84 99
fax: (22) 257 84 00

Redaktor Naczelny:
Piotr Górecki, pg@elportal.pl

Redaktorzy Działów:

Andrzej Janeczek
sp5aht@swiatradio.com.pl

Opracowanie graficzne, skład:

Ewa Górecka-Dudzik

Okladka, zdjęcia, skanowanie:
Piotr Górecki jr

Sekretarz Redakcji

Ewa Górecka-Dudzik
ewa.dudzik@elportal.pl
tel.: (22) 783 00 20
(w godzinach 10:00 – 15:00)

Dział Reklam:

Katarzyna Gugala
katarzyna.gugala@elportal.pl
tel.: (22) 257 84 64

Klasyczne listy i paczki
(projekty i Szkoła Konstruktorów)
prosimy adresować:

AVT – EdW
ul. Leszczyńska 11
03-197 Warszawa
(+dopisek określający zawartość)

Korespondencja elektroniczna:

e-maile do Redakcji EdW:
edw@elportal.pl

e-maile do Szkoły Konstruktorów:
szkola@elportal.pl

rozwiązania konkursów – e-maile:
konkursy@elportal.pl

uwagi do rubryki Errare:
errare@elportal.pl

Prenumerata:

W Wydawnictwie AVT

tel: (22) 257 84 22
(godz. 10:00–14:00)

e-mail: prenumerata@avt.pl

W RUCH S.A.

tel: 801 800 803, (22) 717 59 59
e-mail: prenumerata@ruch.com.pl
www.prenumerata.ruch.com.pl

Stali współpracownicy:

Michał Adamus
Szymon Janek
Krzysztof Kawa
Rafał Orodziński
Michał Pędzimaż
Michał Stach
Szymon Trygar
Adam Sobczyk
Piotr Świerczek
Piotr Wójtowicz

Luty

2 (313)

Projekty

Projekty AVT

Podwójnie symetryczny moduł audio.....	19
Frezarka CNC, część 2.....	24
Miernik wzmacniaczy operacyjnych, część 4.....	27

Elektronika 2000

Pulsująca błyskotka LED.....	54
------------------------------	----

Forum Czytelników

Z potrzeby chwili...	
Zasilacz do laptopa w roli ładowarki akumulatora	56

Szkoła Konstruktorów

Zadanie główne 311

Zaproponuj wykorzystanie elektroniki dla dobra osób z zaburzeniami snu, np. chrapaniem czy bezdechem.....	43
---	----

Rozwiązanie zadania głównego 306

Zaproponuj interesujące,	
najlepiej nietypowe zastosowanie diod LED	44
Druga klasa Szkoły Konstruktorów Co tu nie gra? 311, 306.....	48
Trzecia klasa Szkoły Konstruktorów Policz 311, 306.....	50

Artykuły różne

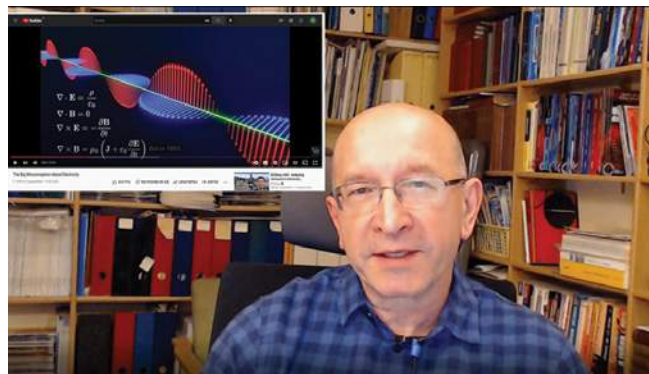
Elektronika i historia	
Telewizor bez zasilacza?, część 2.....	16
MPPT, część 18.....	30
Transmisja danych w inteligentnym domu	
6. MODBUS – bardzo popularny protokół	32
Automatyczne nawadnianie także dla Ciebie – podstawy	34
NanoVNA – wykonaj precyzyjne pomiary, część 2.....	38
Odkrywamy schematy	
Klasyczny forward z jednorozmiarowym kluczem, część 6.....	58
Porady warsztatowe. Kilka przydatnych wskazówek.....	65
Moje pasje – własnej roboty magnetofon szpulowy, część 1.....	66

Rubryki stałe

Nowości, ciekawostki	6
Poczta	8
Prenumerata	13
Skrzynka porad	14
Księgarnia AVT.....	18
Oferta handlowa AVT	70

Konkursy

Co to jest?	62
Krzyżówka	63
Jak to działa?.....	68



Luty

W tym miesiącu na okładce mamy *Podwójnie symetryczny moduł audio*. Miłośnicy audio, żeby nie użyć tu słowa audiofile, poświęcają mnóstwo uwagi kwestiom związanym z symetrią toru i jego poszczególnych bloków. Ponieważ teoretyczne dyskusje z reguły nie prowadzą do konsensusu, naprawdę warto osobiście przekonać się czy, kiedy i na ile pełna symetria ma wpływ na właściwości dźwięku.

W numerze znajdziecie zapowiadaną już wcześniej pierwszą część materiału o samodzielnej budowie systemów automatycznego nawadniania. Oprócz tego oczywiście szereg znanych już cykli i artykułów, w tym jakże ważne wskazówki dotyczące wykorzystania NanoVNA oraz protokołu MODBUS. Do uważnej lektury zapraszam wszystkich Czytelników!

Jednak moim zdaniem w tym numerze najważniejsza jest rubryka Poczta. Zawiera bowiem omówienie i jedno z wyjaśnień zagadki, która w połowie listopada spowodowała burzę i lawinę dyskusji na całym świecie. Okazuje się, że to, czego o prądzie elektrycznym uczy się nas w szkole podstawowej i średniej, nie tylko nie wyjaśnia zagadnienia, ale je zaciemnia, wręcz kieruje na manowce. To jeszcze jeden dowód, jak bardzo potrzebna jest Radiowa Ośła Łączka, czyli przystępne wyjaśnienie podstaw elektroniki i techniki radiowej. Nieco szerzej omawiam tę zagadkę tak jak ostatnio, na YouTube:

www.youtube.com/watch?v=9Vnhet9KbCo

Zachęcam do lektury i do udziału w konkursach.

Jak zawsze serdecznie pozdrawiam

Piotr Górecki



Prenumerata
– naprawdę warto!



ONEPLUS 10 PRO TO MASZYNA DO ZDJĘĆ

W ostatnich kilku dniach przed premierą, o OnePlus 10 Pro dowiedzieliśmy się od samego producenta bardzo wiele, przez co oficjalna premiera straciła na mocy. Jednak odbyła się zgodnie z planem i poznaliśmy nowy sztandarowy model z chińskiej tajni.

10 Pro jest wyposażony w system aparatów Hasselblad drugiej generacji, mocny procesor Qualcomm Snapdragon 8 Gen 1 i przewodowe ładowanie z mocą 80W za pomocą dołączonej do zestawu ładowarki (co nie jest standardem, gdy kolejni producenci, tacy jak Apple i Samsung, z nich rezygnują). Model dostał 6,7-calowy ekran 1440p z adaptacyjną częstotliwością odświeżania od 1 do 120Hz, do 12GB RAM-u LPDDR5 i 25GB pamięci wewnętrznej UFS 3.1. Całość zasilają baterie Li-ion 5000mAh, którą można ładować przewodowo (80W) i bezprzewodowo (50W). Za koordynację wszystkich podzespołów odpowiada system ColorOS 12.1, bazujący na Androidzie 12. Oczywiście poza Chinami będzie to OxygenOS.

Najwięcej uwagi OnePlus poświęcił możliwościom fotograficznym. Zaprezentowano trzy aparaty wchodzące w skład modułu głównego: 48Mpix + 50Mpix + 8Mpix z Dual OIS. Jeden z obiektywów z polem widzenia 150° można wykorzystać do robienia zdjęć w stylu rybiego oka. Kolejny zaś wyposażono w obiektyw 110°. Z pewnością 10-bitowy kolor czy 12-bitowe RAW+ w trybie Hasselblad Pro będą czymś wartym przetestowania. OnePlus 10 Pro jest dostępny w Chinach od 13 stycznia.



RTX 3090 TI WSTRZYMANY?

Na początku października ubiegłego roku pojawiły się pierwsze informacje na temat najbardziej zaawansowanej karty graficznej z generacji Ampere, czyli RTX 3090 Ti. Podczas prezentacji NVIDIA na tegorocznych targach CES 2022, producent zaprezentował tę kartę graficzną, pokazując referencyjny model, jednak obyło się bez prezentacji specyfikacji, a także prezentacji jej możliwości. Marka wcześniej zapowiedziała, że jeszcze w tym miesiącu zadebiutuje GeForce RTX 3090 Ti, ale najnowsze dane sugerują, że premiera może się opóźnić.

To kolejne oznaki kłopotów na rynku kart graficznych, spowodowanych głównie ograniczonymi możliwościami produkcyjnymi i masowym wykorzystaniem kart GPU do kopania kryptowalut.

GeForce RTX 3090 Ti ma być najbardziej zaawansowaną konsumencką kartą graficzną jaka kiedykolwiek powstała. Jest to ulepszona wersja podstawowej 3090, a więc bazuje na tym samym procesorze graficznym Ampere GA102, jednak w pełnej wersji. Układ ma zatem na pokładzie 10752 rdzenie CUDA, 336 Tensor oraz 84 jednostki ray tracing. GPU jest wspierane przez 24GB pamięci wideo na superszybkich kościach GDDR6X o taktowaniu 21Gb/s, która zapewnia przepustowość na poziomie 1018GB/s.

Karta graficzna ma cechować się TDP 450W, a więc w podstawowej referencyjnej wersji NVIDIA będzie wymagała pojedynczej wtyczki 16 pin do zasilania. W przypadku niereferencyjnych modeli GeForce RTX 3090 Ti do poprawnego działania karty mogą być potrzebne nawet 3 wtyczki 8 pin.



TESLA Z PROBLEMAMI

Mimo że rok 2021 był dla Tesli wyjątkowy, a sprzedaż modeli 3 oraz S biła wszelkie możliwe rekordy, to firma boryka się ze znaczącymi kłopotami związanymi z wprowadzeniem kolejnych modeli, których premiera znowu się opóźnia. Zmiana na oficjalnej stronie sugeruje, że nowy pickup producenta pojawi się w... nieokreślonym bliżej terminie. Zapowiedziany w 2019 roku elektryk o kontrowersyjnym wzornictwie miał pojawić się na drogach w ciągu dwóch lat. W sierpniu 2021 roku firma poinformowała, że premiera została przesunięta na 2022 rok. Teraz z oficjalnej strony, gdzie można było składać zamówienia, w ogóle zniknął rok premiery pojazdu.

Zauważono, że wcześniej na stronie widniała następująca informacja: „Konfiguracja będzie możliwa niedługo przed premierą pojazdu w 2022 roku”. Komunikat skrócono teraz jednak do lako-



nicznego tekstu „Konfiguracja będzie możliwa niedługo przed premierą pojazdu”.

Możliwe że Tesla Cybertruck jest kolejną ofiarą obecnie panującej sytuacji, która przekłada się m.in. na problemy dostawą części i podzespołów elektronicznych. Kanciasty wygląd jest mocno nietypowy, ale oprócz wywoływania kontrowersji, wymaga od projektantów nietypowych rozwiązań, które mogą opóźnić ostateczną premierę. Wszyscy pamiętamy konferencję, podczas której „pancerna” szyba niespodziewanie pękła.

Innym powodem opóźnień może być popularność pozostałych samochodów Tesli. W Europie Model 3 był pierwszym autem elektrycznym, które trafiło na szczyt miesięcznej listy najchętniej kupowanych nowych aut. Producent w poprzednim roku sprzedał prawie milion samochodów na całym świecie. Najwyraźniej Cybertruck nie jest priorytetem.

SSD I RAM W JEDNYM

W komputerach i w większości urządzeń przetwarzających informacje wyróżniamy dwa główne rodzaje pamięci – statyczną i dynamiczną, które cechują się różnymi zastosowaniami. Fizycy z brytyjskiego uniwersytetu wrócili do znanego od dawna pomysłu połączenia ich w jeden komponent, mający eliminować wąskie gardła przesyłu danych i zmniejszyć pobór energii. Zarówno SSD, jak i RAM bazują na pamięci flash, czyli błyskawicznego dostępu, lecz – jak widać – w nieco inny sposób. Moduły operacyjne wykorzystywane są do zadań, gdzie potrzeba natychmiastowego dostępu, zaś pamięci półprzewodnikowe w przypadku danych statycznych. Jednak w ostatnich latach wraz z gwałtownym rozwojem półprzewodnikowych dysków SSD, różnice w prędkościach pamięci się zatarły, a interfejs PCI-Express 4.0 pozwala na naprawdę wysokie transfery.

Proponowany przez naukowców UltraRAM to teoretyczny sposób na zunifikowanie pamięci statycznych, które za sprawą odpowiednich magistrali mają być tak szybkie, że będą mogły wykonywać zadania losowego dostępu. Brytyjscy naukowcy z uniwersytetu w Lancaster stworzyli pojęcie UltraRAM, czyli urządzenia, który ma zintegrować RAM i SSD w jeden komponent.

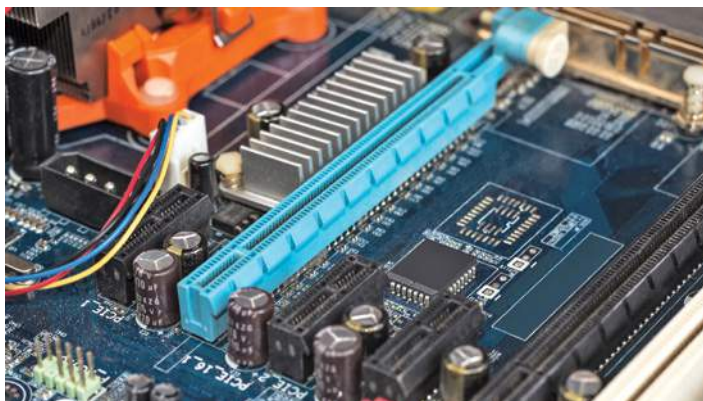
Rozwiązanie to pomogłoby w dalszej miniaturyzacji i poprawianiu wydajności energetycznej – pozwoliłoby zaoszczędzić trochę miejsca w środku urządzenia, a producenci, zamiast tworzyć dwa komponenty, musieliby produkować tylko jeden.



PCI 6.0 JESZCZE SZYBSZY

Nowe komputery PC charakteryzują się coraz nowszymi złączami i interfejsami, które są w stanie obsłużyć jeszcze wydajniejsze komponenty o większych przepustowościach. Obecnie nadal najchętniej wykorzystywana jest magistrala PCI-Express 4.0, na której bazują dyski SSD M.2, a także karty graficzne. Mimo że piąta generacja interfejsu na dobre się nie zadomowiła w naszych pecetach, to wypuszczono już finalną wersję standardu PCIe 6.0, która ma zapewnić jeszcze większe prędkości przesyłu danych.

Standard PCI-Express 4.0 jest coraz powszechniejszy – mimo wysokiego zainteresowania dyskami z trzecią generacją interfejsu. Większość dysków czwartej generacji kosztuje sporo, bowiem



200 MEGAPIKSELI W SMARTFONIE

Pierwsze pytanie – po co? Okazuje się jednak, że w dobie tak zwanej fotografii obliczeniowej i zaawansowanych algorytmów przetwarzaniu obrazu, bardzo wysokie rozdzielności nie są tylko chwytem marketingowym.

Interesującym projektem, zaprezentowanych podczas targów CES 2022, jest sensor Omnivision o rozdzielczości 200Mpix. To, co go wyróżnia, to najmniejszy na świecie rozmiar piksela.

Omnivision OVB0B, bo o nim mowa, z punktu widzenia specyfikacji zapowiada się bardzo interesująco.

Nie jest tajemnicą, że obecnie najwięcej zdjęć robionych jest przy pomocy smartfonów – w końcu te urządzenia mamy zawsze przy sobie. W związku z tym, konsumenci oczekują, że jakość, jaką oferują aparaty w tego typu sprzęcie, będzie coraz lepsza. Naprzeciw tym wymaganiom wychodzi Omnivision – prezentując nowy produkt, który oferuje najwyższą w branży rozdzielczość w niewielkiej obudowie dla smartfonów z najwyższej półki, a także najlepszą w swojej klasie wydajność przy słabym oświetleniu. To ostatnie jest zaskoczeniem, bo małe elementy światłoczułe matrycy, wynikające z jej wysokiej rozdzielczości, nie sprzyjają głębi tonalnej czy dobremu poziomowi szumów przy gorszym oświetleniu.

Sensor ten cechuje się rozdzielczością 200Mpix, a wielkość pojedynczego piksela wynosi 0,61 mikrona. Co ważne, matryca oferuje technologię poczwórnego fazowego autofokusa dla 100% pikseli. Omnivision OVB0B ma zapewnić najwyższą jakość wideo o wielkości 12,5 megapikseli – wszystko za sprawą łączenia aż 16 pikseli w jeden punkt. Jest to, jak twierdzi producent, pierwsza w branży funkcja zlepienia punktów dla wideo w rozdzielczości 4K2K.



topowe półprzewodniki są średnio dwukrotnie droższe od tych przeznaczonych do poprzedniej wersji, ale za wyższą ceną idzie wysoka wydajność. Standard trafia też do nowoczesnych kart graficznych, które cechują się coraz wyższą wydajnością, a także przepustowością transmisji.

Nowoczesne płyty główne (na chipsecie Z690) wspierają magistralę PCIe 5.0, która pozwala na przepustowość do 128GB/s (na 16 liniach; pojedyncza linia oferuje 32GT/s, czyli 4GB/s). Dyski PCI-Express 5.0 x4 teoretycznie mogą rozpędzać się do 16GB/s, natomiast pierwsze testy kontrolera Phison E26 wykazują przepustowość 14GB/s. Ostatnio zaprezentowany interfejs PCIe 6.0 na 16 liniach (x16) będzie gwarantował zarówno odczyt, jak i zapis na poziomie aż 256GB/s, zaś pojedyncza ścieżka 64GT/s – 8GB/s.

W rubryce „Poczta” zamieszczamy fragmenty Waszych listów oraz nasze odpowiedzi i komentarze. Prosimy o listy dotyczące bieżących wydań EdW, a także o listy z Waszymi komentarzami, propozycjami, problemami, pytaniami, oczekiwaniami względem nas,

z propozycjami tematów do opracowania, itp. Autorzy najciekawszych, wartościowych listów otrzymują upominki, najczęściej w postaci drobnych kitów AVT. Piszcie do nas, bardzo cenimy Wasze listy, choć nie wszystkie prośby możemy zrealizować.

UWAGA! UWAGA!

Potwierdzamy otrzymanie każdego e-maila. Zachęcamy do wykorzystywania opcji: *Żądaj potwierdzenia doręczenia*. Jeśli ktoś nie otrzyma potwierdzenia w ciągu tygodnia, proszony jest o wysłanie swojej wiadomości jeszcze raz – do skutku. A gdyby przypuszczalnym powodem skasowania e-maila przez serwery poczty były potencjalnie groźne załączniki (np. typu .exe, bas, itp.), bardzo prosimy wysłać informację o tym bez żadnych załączników.

Do części projektów publikowanych w EdW firma AVT proponuje kompletne zestawy elementów albo tylko płytki drukowane. Na początku i końcu takich artykułów-projektów podana jest informacja o numerze kitu AVT. Jeżeli w artykule numeru kitu nie ma, a Czytelnicy byliby zainteresowani nabyciem zestawów albo samych płytek, jest to możliwe.

AVT uruchomi realizację kitów/płytek, o ile tylko gotowość zakupu wyrazi przynajmniej kilku chętnych. Zgłoszenia i pytania w tej sprawie należy nadsyłać wprost na adres: **kity@avt.pl**

Dzień dobry.

Jestem z wykształcenia elektrykiem, ale dzięki czasopismom takim jak EdW uwielbiam elektronikę. Od wielu lat zaczytuję się EdW, to czasopismo jest biblią elektroniki ze stale dodawanymi księgami.

Cieszę się, że redakcja stara się, aby podążać za nowościami elektroniki oraz próbuje przekazać w prosty sposób „smakoliki” w niej zawarte. Otrzymana wiedza przyjaźnie wdraża się w praktyce.

Oglądałem film na Youtube, podoba mi się taki sposób przemawiania do „wiernych”, ale do rzeczy.

Od chwili wydania pierwszego artykułu do dziś z zapartym tchem śledzę materiały Pana Karola Świerca dotyczące zasilaczy. Powiem szczerze: są wyborne, nigdy nie spotkałem się z tak bardzo rzetelnie i rzeczowo przekazywaną wiedzą. Ogrom zaangażowania i pracy, którą Pan Karol poświęcił zasilaczom, jest godny szacunku i podziwu, każdy skrawek wiedzy poparty jest licznymi badaniami i dogłębną analizą rzeczową. Z mojej strony proszę o dalsze zamieszczanie artykułów o zasilaczach, bo wiadomości te są bezcenne.

Kolejną sprawą, którą chciałbym poruszyć, to artykuły Pana Jerzego Szymańskiego. Jego styl pisanie i sposób przekazywania wiedzy o starszej elektronice jest kojący. Naprawdę miło jest przyjmować wiedzę od tego Wielkiego Człowieka.

Choć nietrudno zauważyć, że w tej dziedzinie sprzęt i wiedza są już niestety bardzo przestarzałe. Jeżeli to możliwe,

proszę, aby Pan Jerzy podzielił się z nami odrobiną wiedzy o samych lampach, na co zwrócić szczególną uwagę przy ich badaniu, jak dobrać lampę do obwodu itp. Przyjazne będzie zamieszczenie przykładowego schematu, zasilania, wzmacniacza, mieszacza itp. wraz z analizą jego działania.

Pozdrawiam.

Tomasz

*Dotyczy: Zapowiedź numeru 1/2022 – superkondensatory
Ja używam: <https://youtu.be/8zMcWfSiHTo>*

Porównanie UPS na „elektrolicie” i SuperCAP:

<https://www.youtube.com/playlist?list=PLdtkbWTUVMmYDYjiCwiDvK3tY4qXv5FJ>

Można kupić gotowy moduł superCap:

https://www.modelmania.com.pl/product_info.php?cPath=2131_212_2105_2888&products_id=124849

w skrócie: <https://bit.ly/3pGUUfK>

Pierwsze co poraża, to cena. Same kondensatory, tak wielkie gabarytowo, to koszt ok. 15zł. Są też dużo mniejsze za ok. 27zł.

Kolejny problem, brak balancera, ze względu na wymiary, w postaci rezystorów.

Największy problem. Norma DCC określa zasilanie w granicach 12–22V. $7 \cdot 2,5 = 17,5V$. Uwzględniając spadek napięcia na mostu prostowniczym jakiś 1V (99% na diodach Schottky’ego z wielu powodów), max. zasilanie to 18,5V. Do 22V daleka droga.

Co przykre, ten drogi chłam produkuje jedna z najlepszych firm od dekoderów do lokomotyw.

Ja stosuję 8 szt. na 2,7V o małych gabarytach i prosty balancer.

Pozdrawiam

SS

Witam!

Zaobserwowałem przypadkiem pewne ciekawe zjawisko / lifehack, którym chciałbym się podzielić w rubryce „Poczta”.

Zrobiłem kiedyś przedłużacz z wyłącznikiem. Prosta rzecz, ale wyłącznik jest na osobnym, metrowym kablu i rozłącza tylko jedną żyłę. Przedłużacz ten zasilał łańcuch „bombek” zrobionych z szeregowo połączonych diod LED (bez żadnego sterownika, choć być może z jakimś rezystorem). Stety lub niestety, przy wyłączonym wyłączniku diody nadal delikatnie świeciły. Nie była to kwestia obecności fazy na diodach i ich pojemności do ściany, na której były zawieszone, bo odwrócenie wtyczki przedłużacza nie rozwiązywało tego problemu. Diody najprawdopodobniej świeciły ze względu na prąd, który płynął przez pojemność kabla prowadzącego do wyłącznika. Zmierzyłem prąd, jaki można pobrać z tego przedłużacza przy wyłączonym wyłączniku i wyszło mi 13

mikroamperów.

Z powodu świąt użyłem tego przedłużacza do zasilania lampki choinkowych z kontrolerem. To typowa konstrukcja, w której kontroler zasilany jest, poprzez mostek prostowniczy i rezystor, bezpośrednio z sieci, i łączyła łańcuchy diod LED poprzez tyrystory. Kontroler ten ma kilka trybów, które zmieniane są przyciskiem. Problemem takich kontrolerów jest to, że po odłączeniu od sieci zapominają ostatnio ustawiony tryb.

Okazało się, że zasilanie kontrolera przez ten przedłużacz rozwiązuje ten problem. 13 mikroamperów, płynące poprzez pojemność kabla, wystarczy, żeby kontroler nie zresetował się i pamiętał ostatnio ustawiony tryb. Pewnie kontroler cały czas pracuje, próbując sterować tyrystorami, ale nie starcza mu już na to prądu.

Circuit Chaos

Czytając październikowy numer EdW, zwróciłem uwagę na felieton pana Adama Sosnowskiego. Bardzo się ucieszyłem, że autor jest sympatykiem elektroniki analogowej. Miło spotkać każdą pokrewną duszę, która kocha dawną technikę, mimo że ze współczesną też jest za pan brat.

Drugi powód do sympatii to ten, że Autor jest również miłośnikiem poezji i to także dawniejszej. Być elektronikiem i kochać poezję to podwójna radość. Cieszy mnie również, że pan Adam z rezerwą odnosi się do cyfryzacji wszystkiego, co możliwe. Powtórzę swoimi słowami za Adamem Asnykiem i panem Sosnowskim: nie poprawiajcie na siłę tego, co jest dobre i co zdało egzamin. Poprawiać też trzeba umieć.

Jerzy Szymański

Jeden ze współpracowników EdW nadesłał krótki e-mail, zawierający tylko link do filmu z kanału RS Elektronika:

<https://www.youtube.com/watch?v=Z99fTUJye-Y>

Jest tam przedstawiona zagadka logiczna dla elektroników (rysunek A): widzimy tu akumulator, żarówkę, wyłącznik oraz dwa odcinki kabla mające razem 300 tysięcy kilometrów długości. Pytanie brzmi: po jakim czasie od zamknię-

cia wyłącznika zaświeci żarówka (zakładamy, że żarówka i wyłącznik nie wykazują żadnych opóźnień, pomijamy nieliniowość włókna żarówki).

W drugim e-mailu ten współpracownik, będący doświadczonym konstruktorem, głównie układów cyfrowych, napisał: Dla mnie RF to MAGIA!

Zapewne podobnie jest w przypadku większości Czytelników EdW. Nieco bliższe zbadanie zagadnienia wykazuje, że jest to wierzchołek góry lodowej. Szerszy kontekst tej zagadki tak naświetlił jeden z Czytelników EdW:

Wszystko zaczęło się od filmiku na kanale Veritasium (rzetelny kanał popularnonaukowy), zatytułowanego „Nie zrozumienie, jak płynie prąd”

<https://www.youtube.com/watch?v=bHlhgxav9LY>

Autor zaproponował eksperyment myślowy: mamy baterię, wyłącznik i żarówkę, połączone ze sobą, lecz w niestandardowy sposób, gdzie przewody od obu terminali żarówki zataczają długą pętlę prawie do Księżyca i z powrotem, cały czas utrzymując odległość jednego metra od siebie. Żarówka od wyłącznika [i akumulatora] również znajduje się w odległości 1 metra [rysunek B].

Pytanie brzmiało, jak długo będziemy czekać od momentu zamknięcia wyłącznika do momentu, aż zaświeci żarówka. Czy będzie to jedna sekunda, dwie sekundy, $1m/c_0 = 3,3ns$, czy może bardzo, bardzo długo?

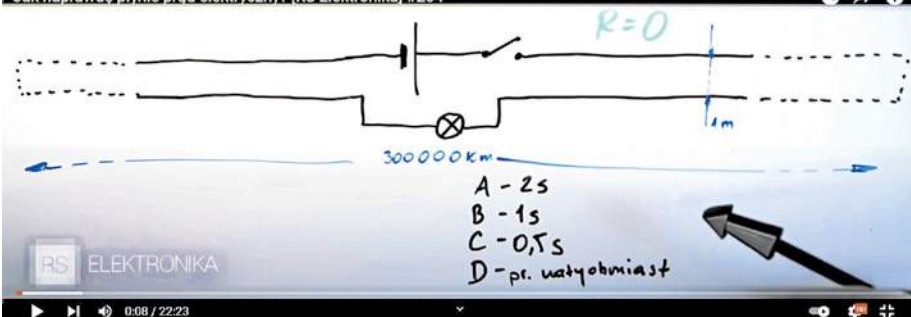
Następnie zaczyna wyjaśniać, iż na bazie twierdzenia Poyntinga można wykazać, że przepływ energii nie wymaga zamkniętego obwodu elektrycznego, i już samo zamknięcie wyłącznika wywoła impuls napięciowy, który przebędzie odległość 1m i wyindukuje prąd w żarówce i ta się zaświeci. Czyli [właściwa] odpowiedź to 3,3ns.

Ze swojej strony od razu powiedziałbym, że żarówka podpięta jest do linii transmisyjnej (lub jak to się na polskich uniwersytetach mówi: linii długiej) i w związku z tym popłynie od razu prąd wyznaczony przez napięcie baterii i impedancję tej linii (jest to linia dwuprzewodowa, równoległa w powietrzu, przy odległości 1m i powiedzmy drucie 12AWG, szacuję, że impedancja to $\sim 750\Omega$, więc popłynie niewielki prąd, ale popłynie). Czyli (...) sygnał biegnie jako

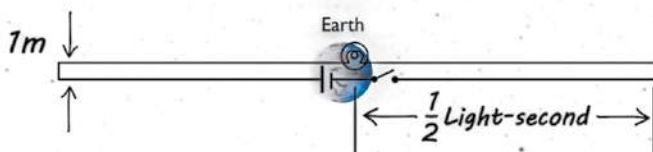
pole EM między dwoma przewodnikami. I nie ma tu znaczenia, czy mówimy o układzie DC, AC, czy RF. W przypadku DC będziemy mieć zamknięcie włącznika, który wywoła zbocze narastające rzędu nanosekund (no dobra, nie wyłącznik, ale klucz tranzystorowy). A skoro tak, to pasmo częstotliwości takiego sygnału to $BW = 0,35/10E-9 \sim 35MHz$, czyli mamy 35MHz biegnących w układzie DC, po prawie nieskończenie długiej linii. W przypadku AC 50Hz, w tej skali odległości będziemy mieć 50 długości fali, więc mamy typowy układ w.cz.

Wracając do przypadku DC, prąd popłynie natychmiast, lecz po osiągnięciu końca linii sygnał zorientuje się, iż jest tam zwarcie na końcu, odbije się i zacznie wracać w kierunku żarówki-baterii. Odbije się po raz kolejny, i kolejny, za każdym razem podnosząc VSWR, aż do momentu, aż ustali się nieskończenie wielki VSWR i po prostu zapanuje napięcie DC baterii. Paradoks jest taki, że gdy

Jak naprawdę płynie prąd elektryczny? [RS Elektronika] #204



- A) 0.5 s
- B) 1 s
- C) 2 s
- D) $1/c$ s
- E) None of the above



DC się ustabilizuje, to prąd elektryczny ustanie prawie do zera, ponieważ taka długość kabla miedzianego da sumaryczną rezystancję ponad $3\text{M}\Omega$. Tyle analizy tego problemu. Tymczasem, jak można się było domyślać, nastąpiła lawina komentarzy pod filmem. Komentarzy ludzi, którzy nie dowierzali lub nie przyjmowali takiej rzeczywistości do wiadomości. Do sprawy zaczęli się odnosić popularni youtuberzy z dziedziny inżynierii i elektroniki, niektórzy bardzo szanowani, co i tak nie uchroniło ich od blamażu. Oraz zaczęła się ujawniać gama ludzi typu „nie znam się, to się wypowiem”. Ci ostatni zaczęli dywagować, że taki mechanizm przeczy szczególnej teorii względności. Pojawiły się analizy, iż nie wytworzy się żadne pole magnetyczne. Dowodem na niesłuszność tego eksperymentu miało być to, iż włókno żarówki ma bezwładność i nie włączy się po 3 nanosekundach. A w ogóle to nie ma takiej żarówki, która działałaby przy dowolnie niskim prądzie, więc to wszystko nieprawda. Ręce opadają! Ostatecznie pojawił się domorosły eksperymentator, który kupił kilometr emaliowanego drucika i po prostu zbudował zestaw pomiarowy, rozciągając równoległe przewody o długości 500m na swojej farmie. I mierzył oscyloskopem spadek napięcia na rezystorze odgrywającym rolę żarówki. Wyszło mu, że natychmiast pojawia się prąd, ale niewielki (mniej więcej zgodny z moim przewidywaniem $\sim 700\Omega$ impedancji), a po czasie, w którym impuls osiąga koniec pętli, napięcie zaczyna rosnąć, i szybko osiąga wartość DC. Konkluzja jest taka, że eksperyment mu nie potwierdził, iż teza z początkowego filmiku jest prawdziwa. Naprawdę ręce opadają!

Honor youtuberów uratował niezastąpiony Dave Jones ze swojego kanału EEVBlog, który na bieżąco oglądał i komentował ten kwestionowany filmik. Jego generalna reakcja była taka: „Żarówka zapali się natychmiast, bo taka jest fizyka, taka jest teoria linii transmisyjnej. Tego się uczy na zaawansowanych kursach elektroniki na uniwersytetach. A ci, którzy to kwestionują, to albo nie uważali na tym wykładzie, albo nie nadają się na bycie inżynierem-elektronikiem”.

Z mojego wyczucia sytuacji wynika, że tego nie rozumie nikt poza garstką elektroników robiących RF i projektujących PCB do szybkich układów cyfrowych. Nad elektrotechnikami, którzy myślą, że elektrony z elektrowni biegną z prędkością światła poprzez urządzenia domowe, a potem wracają do elektrowni, to już spuszcza zasłonę milczenia. Na szczęście Pan Krzysztof z RS Elektronika też nie dał plamy. Ten przykład pokazuje, jak trudno jest nauczyć RF-ów, i to nie dlatego, że jest to bardzo skomplikowany technicznie i matematycznie materiał (choć jak wszystko, w szczególności jest bardzo skomplikowany), lecz dlatego, iż należy się mierzyć z już ugruntowanymi, acz niepoprawnymi wyobrażeniami, zakorzenionymi w myśleniu o prądzie stałym. Jeśli dla takiej rzeszy praktyków koncepcja linii transmisyjnej jest nieakceptowalna, bo chcą myśleć, że energia elektryczna to elektrony lecące rurką niczym piłeczki pingpongowe, to jak takie osoby nauczyć, że w układzie RF-owym nie da się zrobić zwarcia? Lub że zwyczajna sonda oscyloskopowa $10\times$ jest anteną zbierającą wszystko wokół, tylko akurat nie sygnał z ścieżki, do której ją przykładamy.

Sam ten eksperyment ma też ogromny walor edukacyjny. Po pierwsze uczy, co to linia transmisyjna i charakterystyka

przejściowa (transient) obwodu. Uczy, że sygnał to pola, etc. Po drugie, możemy prowadzić dyskusję, co się dzieje po dojściu impulsu do końca pętli, jak tworzy się sygnał odbity i jak będzie się budował VSWR lub co się stanie, jeśli na końcu będzie rezystor o różnych wartościach.

A po trzecie, co będzie, jeśli koniec tej pętli rozewrzemy tak, by utworzyć dipol i promieniować energię na zewnątrz?

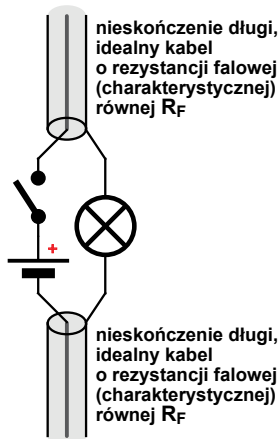
Krzysiek z Krakowa

Pełne i wyczerpujące rozwiązanie zagadki nie może być przedstawione mniej zorientowanym w kilku słowach czy zdaniach. Wymaga bowiem zrozumienia szeregu zagadnień dotyczących techniki w.cz., w tym różnych aspektów zachowania linii długiej. A jak widać, problem z tym mają też osoby, które na pewno zaliczały te tematy na studiach. Jednak do uzyskania odpowiedzi na postawione pytanie o opóźnienie wcale nie jest potrzebna gruntowna wiedza z zakresu techniki w.cz. W EdW nie było wprawdzie systematycznego wyjaśnienia tej tematyki, ale poruszaliśmy kiedyś dziwny problem: jak zmierzyć rezystancję falową kabla za pomocą najzwyklejszego omomierza?

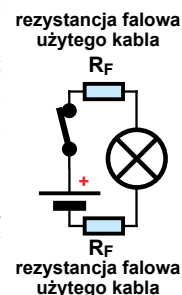
Łatwo rozwiąże zagadkę każdy, kto wie, jak można byłoby zmierzyć rezystancję charakterystyczną kabla za pomocą zwyczajnego omomierza mierzącego rezystancję dla prądu stałego (w ramach eksperymentu myślowego z użyciem idealnego, nieskończonego długiego kabla). Nie trzeba przy tym wcale rozumieć zjawisk falowych, kwestii odbić i dopasowania – wystarczy jedynie rozumieć, co dzieje się po dołączeniu długiego kabla (linii długiej) do źródła napięcia stałego. Otóż **po dołączeniu do akumulatora kabla – linii długiej, natychmiast zaczyna tam płynąć prąd stały, o niezmienniej wartości, wyznaczonej przez napięcie akumulatora i rezystancję falową tego kabla (tej linii długiej)**. Nie wszyscy o tym wiedzą. Bardziej szczegółowo i przystępnie Naczelny wyjaśnia to na stronie:

<https://www.youtube.com/watch?v=9Vnhet9KbCo>.

Dla ułatwienia analizy na początek warto pominąć informację o metrowym przesunięciu między żyłami kabla, a schemat z rysunku A przekreślić o 90 stopni w lewo, do postaci jak na **rysunku C** i na początek założyć, że dwa kable, na przykład koncentryczne, są idealne i nieskończone długie oraz że mają jakąś rezystancję charakterystyczną (falową) R_F .



Wtedy po zwarceniu wyłącznika schemat zastępczy będzie wyglądał jak na **rysunku D**. Przy nieskończonej długości obu kabli taka sytuacja utrzyma się na stałe. Prąd będzie wyznaczony przez sumę rezystancji: żarówki oraz dwóch rezystancji falowych R_F . Jeżeli byłyby to idealne kable współosiowe, to ich rezystancja falowa R_F byłaby rzędu kilkudziesię-



ciu omów. W tej wersji, gdy akumulator, wyłącznik i żarówka będą bardzo blisko siebie, opóźnienie zadziałania żarówki będzie znikome, poniżej 1 nanosekundy, więc w przybliżeniu można uznać, że opóźnienia nie będzie.

I już taka uproszczona analiza odpowiada na pytanie o opóźnienie w sytuacji z rysunku A. O rysunku B za chwilę.

W następnym kroku analizy można założyć, że kabel jest idealny, ale dwa użyte odcinki mają określoną długość, właśnie po 150 tysięcy kilometrów, czyli fala EM przebiegnie je tam i z powrotem w ciągu 1 sekundy (pomijamy fakt, że w kablu z dielektrykiem innym niż próżnia prędkość fali jest mniejsza od prędkości światła nawet o 1/3). W każdym razie warto wtedy rozważyć dwa przypadki. Jeden, gdy na dalekim końcu kabla jest *zwarcie* i drugi, gdy jest tam *rozwarcie*. W obu przypadkach na końcu kabla nastąpi całkowite odbicie fali. Należy dodać, że dokonujemy tu pewnego uproszczenia: rozważamy tylko pierwsze odbicie od końca kabla, a pomijamy dalsze odbicia. Nie ma to jednak żadnego znaczenia w rozważaniach dotyczących opóźnienia włączenia żarówki. Słusznie przyjmujemy, że po jakimś czasie zjawiska falowe zanikną, co jest jak najbardziej prawdziwe dla realnych przypadków. Ogólnie biorąc, zjawiska falowe występują tylko wtedy, gdy mamy do czynienia ze zmieniającym się napięciem (tu z falą napięcia wędrującą w kablu). Gdy w kablu ustabilizuje się już napięcie stałe, zjawiska falowe zanikną i pozostanie zwyczajny obwód prądu stałego. A w omawianym przypadku mamy do czynienia z jednorazowym zwarciem wyłącznika, więc najprościej biorąc, zjawiska falowe nie będą występować ciągle.

Na pewno w obu wspomnianych przypadkach, przez sekundę po zwarceniu wyłącznika sytuacja będzie taka jak na rysunku D. Natomiast później w obu przypadkach będzie inaczej. Dla pierwszego przypadku zwarcia na końcu kabla, podtrzymując wcześniejsze założenie, że kabel jest idealny, czyli że ma zerową rezystancję, uznamy, iż po jakimś dłuższym czasie prąd będzie wyznaczony tylko przez rezystancję żarówki według **rysunku E**. Interesujący jest też przypadek drugi, gdy dwa „(pół)sekundowe” odcinki kabla będą *rozwarne* na dalekich końcach. Wtedy też przez pierwszą sekundę po zwarceniu wyłącznika sytuacja będzie jak na rysunku D, natomiast później ustali się stan stabilny, gdy prąd (stały) w ogóle przestanie płynąć przez żarówkę właśnie z uwagi na rozwarcie na końcu odcinków kabla.

Tylko odrobinę trudniejsza jest analiza przy założeniu, że kabel nie jest idealny. Także wtedy przez pierwszą sekundę po zwarceniu wyłącznika sytuacja będzie jak na rysunku D. W przypadku zwarcia na końcach, po dłuższym czasie nie uzyskamy sytuacji jak na rysunku E, tylko sytuację pokazaną na **rysunku F**, gdzie mamy rezystancję R_{Cu} (zapewne miedzianego drutu kabla, stąd takie oznaczenie).

Troszkę trudniejszy jest przypadek z rysunku B, gdzie dodatkowo akumulator i wyłącznik są odsunięte od żarówki o 1 m i przewody linii długiej też są od siebie oddalone o 1 metr. Nadal w pierwszej chwili po zwarceniu wyłącznika sytuacja będzie jak na rysunku D, tylko dodatkowo wchodzi

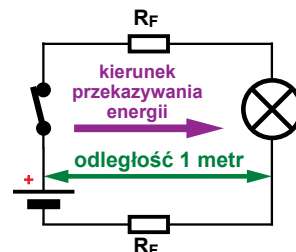
w grę odległość 1 metra według **rysunku G**. Warto zwrócić uwagę na dwa aspekty. Pierwszy, najważniejszy jest taki: najprościej biorąc, według aktualnej wiedzy maksymalna prędkość przesyłania energii to prędkość światła, więc oddalenie żarówki o metr spowoduje opóźnienie o trochę więcej niż 3 nanosekundy. Drugi, mniej istotny aspekt to rezystancja falowa (charakterystyczna) linii długiej R_F o tak dużym rozstawieniu przewodów: będzie ona dużo większa niż linii współosiowych, z którymi mamy do czynienia w praktyce. Najczęściej wykorzystujemy współosiowe kable 50- i 75-omowe. Starsi Czytelnicy pamiętają stary płaski przewód antenowy 300-omowy, a linia z przewodami oddalonymi o metr będzie mieć rezystancję falową rzędu kilkuset omów (750Ω według szacunków *Krzyśka z Krakowa*).

Oto podsumowanie przypadku z rysunku B: Opóźnienie zaświecenia żarówki wyniesie ponad 3ns, potem przez sekundę po zwarceniu wyłącznika sytuacja będzie taka jak na rysunku D, tylko rezystancja falowa R_F będzie większa niż w przypadku kabla współosiowego. Jeszcze później sytuacja ustali się według rysunku F i o prądzie zadecyduje rezystancja między przewodów zastosowanej linii R_{Cu} , która przy długości linii rzędu setek tysięcy kilometrów zapewne byłaby wielokrotnie większa od R_F , czyli prąd w warunkach ustalonych byłby znikomo mały. Natomiast przy założeniu, że idealna linia ma zerową rezystancję przewodów, wrócilibyśmy do rysunku E.

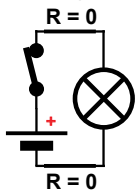
W tym przypadku, do rozważań o samym opóźnieniu zaświecenia żarówki nie jest niezbędna gruntowna wiedza z zakresu techniki radiowej. Jednak taka zagadka znakomicie pokazuje potrzebę przedstawienia przygotowywanej przez Naczelnego Radiowej Oślej Łączki (ROŁ), o której piszemy od miesięcy, a nad którą prace idą powoli, głównie z uwagi na nawał zajęć nad bieżącymi numerami EdW. Ta i pokrewne zagadki będą przekonująco wyjaśnione w ramach ROŁ, tylko nie jest to proste z kilku względów. Jak słusznie wspomniał *Krzysiek z Krakowa*, ogromnym problemem przeszkadzającym w zrozumieniu techniki w.c.z. (RF) są zbyt uproszczone szkolne wyobrażenia na temat prądu elektrycznego, jakie wbito nam do głowy w szkole podstawowej (gimnazjum) i w szkole średniej. Właśnie dlatego kurs ROŁ ma składać się z czterech głównych części, z których pierwsza poświęcona ma być głównie „odkreśleniu” błędnych wyobrażeń. Dopiero w najważniejszej części drugiej mają być wyłumaczone kluczowe kwestie: zjawiska falowe, w szczególności odbicia oraz równie ważne różne aspekty dopasowania. Kolejne części mają dotyczyć zagadnień łatwiejszych do zrozumienia: różnych rodzajów modulacji oraz kwestii przemiany częstotliwości. Wstępnie planowane tytuły tych czterech części to:

1. Radiowa Ośła Łączka,
2. Radiowa Ośła Łączka,
3. Radiowa Ośła Łączka
4. Radiowa Ośła Łączka,

czyli to naprawdę nie jest Twoja wina
czyli dziwny jest ten świat...
Modulacje, czyli majstrowanie przy fali EM
ta *panta rhei* czyli zmiany, zmiany, zmiany.

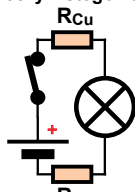


zerowa rezystancja
żył idealnego kabla



zerowa rezystancja
żył idealnego kabla

rezystancja żył
rzeczywistego kabla



rezystancja żył
rzeczywistego kabla

Do redakcyjnej skrzynki pod koniec grudnia wpłynął następujący e-mail skierowany do Naczelnego:

Szanowny Panie Piotrze,
 Tu weteran Szkoły Konstruktorów :)
 Może zaciekać Pana wyniki moich niedzielnych zabaw:
<https://www.eevblog.com/forum/projects/led-blinker-that-operates-below-1mv-hold-my-beer/>
 w skrócie: <https://bit.ly/3exTQUF>
 Z najlepszymi życzeniami z okazji Świąt Bożego Narodzenia i Nowego Roku!

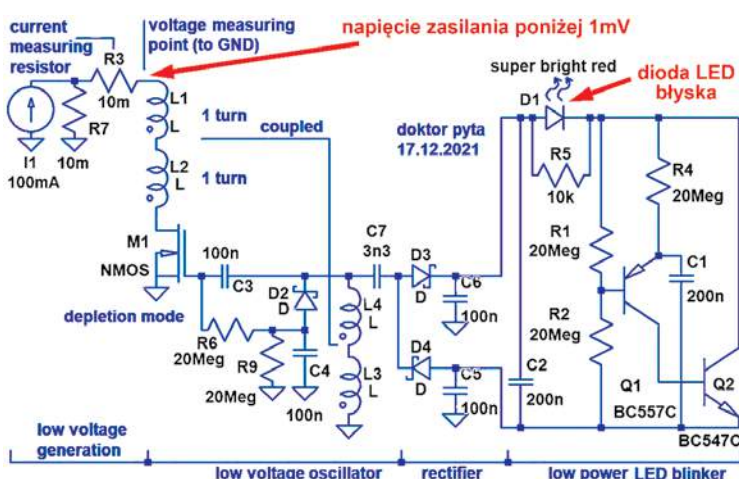
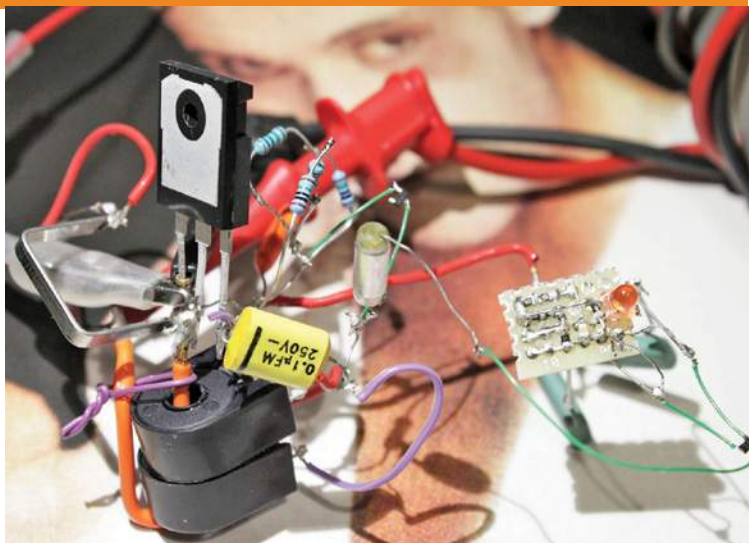
Pozdrawiam
Bartłomiej Radzik

W EdW już kilkakrotnie pisaliśmy o mało znanych zubożonych tranzystorach MOSFET (depletion mode MOSFETs). Były to jednak tylko informacje teoretyczne. Bartek udowodnił, że od teorii nietrudno przejść do praktyki. Zrealizował jak najbardziej działający układ, który zasilany jest napięciem poniżej 1 miliwolta (tak!). W roli transformatora podwyższającego pracują... dwa przekładniki. **Fotografia obok** pokazuje model zrealizowany według nieskomplikowanego schematu.

Do redakcyjnej skrzynki wpłynął ostatnio e-mail:

[Tu] Henryk Cisek z Jeleniej Góry. Szanowny Panie, być może kojarzy Pan mnie, gdyż kilkakrotnie pisałem do Pana na ten adres. Jeżeli nie, to nic złego – zaraz wyluszczę swój temat. Otóż prawie dwa lata temu nabyłem dwie części **Wypraw w świat elektroniki**, a potem **Praktyczny Kurs Elektroniki**. Zabawę z elektroniką rozpocząłem od sumiennego przestudiowania obu wypraw. Nie miałem z tym żadnego kłopotu, gdyż jak Panu pisałem i może pan pamiętać, jestem z wykształcenia inż. mechanikiem, a w latach nastoletnich interesowałem się teoretycznie elektroniką. Życie spowodowało, że zainteresowania z bólem zarzuciłem, ale jak mówią, stara miłość nie rdzewieje. Wróciłem, będąc emerytem, do tematu elektroniki i to w wydaniu praktycznym – zdecydowanie łatwiejszym niż kiedyś (tranzystory i diody w wydaniu miniaturowym zamiast lamp; płytki stykowe; miniaturowe rezystory i kondensatory itd. – dobrze Pan to zna), posiłkując się książkami Pańskiego autorstwa. Praktyka, z uwagi na specyfikę montażu układów na płytce stykowej, szła mi na początku trochę opornie, kilka razy dostałem po nosie, ale ten etap mam już za sobą. Teraz montowane układy częściej od pierwszego razu pracują poprawnie niż było to kiedyś – jakieś 80%/20%. Jestem obecnie w 1/3 wyższego stopnia wtajemniczenia i idzie mi bardzo dobrze. Jak Pan może się domyślać, jestem z Pańskich książek bardzo zadowolony, do tego stopnia zadowolony, że postanowiłem kupić sobie na mikołaja drugi świeżutki zestaw obu części **Wypraw w świat elektroniki**. Pierwsze kupione egzemplarze są egzemplarzami „roboczymi”, w których zapisuję sobie dodatkowe informacje, których Pan nie zamieścił, a są dla mnie ważne i dotarłem do nich w innych źródłach – książki te są nadal w doskonałym stanie. Nowe natomiast będą leżały na półce trochę na „wszelki wypadek”, trochę dla walorów estetycznych. Muszę stwierdzić, że doskonale rozumiem dlaczego **Wyprawy w świat elektroniki** za dwa lata będą obchodziły dwadzieścia lat od pierwszego wydania w 2004 roku – serdeczne gratulacje dla Pana za napisanie tak ciągle „młodej” – choć już dojrzałej książki.

Serdecznie pozdrawiam
Henryk Cisek



A potem jeszcze jeden:

Panie Piotrze, to jeszcze raz ja – Henryk Cisek. Z każdym dniem, z każdą przeczytaną stroną Pańskich Wypraw..., Z każdym z sukcesem wykonanym układem w niej zamieszczonym, Pańskie **Wyprawy w świat elektroniki** podobają mi się coraz bardziej, coraz bardziej je cenię, cieszę się ze słusznego wyboru i ich zakupu – drugi ich zestaw już do mnie dotarł i leży oprawiony w okładki na honorowym miejscu w mojej bibliotece. Absolutnie nie powinien się Pan dziwić, że wyprawy są już dorosłe – mają osiemnaście, a wkrótce będą miały dwadzieścia lat. Widać z tego, że wiele osób docenia je tak samo jak ja – i wieści się rozchodzą – i bardzo dobrze!!! Przyznam się, że na początku kontaktu z nimi miałem pewne problemy wynikające z moich błędów związanych ze specyfiką składania układów elektronicznych na płytkach stykowych. Teraz, gdy poznałem reguły wykorzystywania płytek stykowych, poznałem i przyswoilem sobie podstawowe reguły elektroniki analogowej (np. $U_{BE} = 0,6...0,7V$), układy składam bezbłędnie za pierwszym razem i za Pańskimi radami modyfikuję je, aby poznać ich działanie w nowej konfiguracji. W pierwszym „podejściu” przestudiowałem obie części, a teraz, ponownie czytając, z przyjemnością majsterkuję!! Jeszcze raz wielkie gratulacje i do siego roku :) :) :) :) :) :) :) :))

Henryk Cisek

Upominki za listy do Poczty otrzymują: **Tomasz, Circuit Chaos, Bartłomiej Radzik i Krzysiek z Krakowa.**

**Zaprenumeruj
Elektronikę dla Wszystkich,
a zawsze dostaniesz
najnowszy numer wprost
do Twojej skrzynki!**

**na start
do 6* wydań gratis**

**po 5 latach
nieprzerwanej
prenumeraty
do 12* wydań gratis**



* Cena prenumeraty rocznej **na start** wynosi 152,90 zł. Przy zamówieniu prenumeraty dwuletniej za 250,20 zł oszczędność wynosi równowartość sześciu wydań Elektroniki dla Wszystkich.

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

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

prenumerata@avt.pl
AVT-Korporacja sp. z o.o., ul. Leszczyńska 11, 03-197 Warszawa,
konto 18 1050 1012 1000 0024 3173 1013

Skrzynka Porad

W rubryce przedstawiane są odpowiedzi na pytania nadesłane do Redakcji. Są to sprawy, które, naszym zdaniem, zainteresują szersze grono Czytelników.

Jednocześnie informujemy, że Redakcja nie jest w stanie odpowiedzieć na wszystkie nadesłane pytania, dotyczące różnych drobnych szczegółów.



Niejestem inżynierem, tylko amatorem (...) czy w EdW mógłby się ukazać artykuł, czym tak naprawdę różnią się kable komputerowe klas od 1 do 7, bo chyba wyższych nie ma (...) nie chce mi się wierzyć, że są tak ogromne różnice szybkości danych (...) przecież materiały miedź i izolacja są te same albo podobne (...) ale tak na chłopski rozum (...) [w sposób] zrozumiały także dla mnie (...)

W EdW 4/2021 było rozwiązanie zadania *NieGra296* omawiającego „pseudoskrętkę” zakończoną wtyczką „ethernetową”. Autor powyższego pytania nawiązał do zamieszczonych tam informacji: Dawniej w modemach telekomunikacyjnych szybkości transferu danych była rzędu kilkunastu, potem kilkudziesięciu tysięcy bitów na sekundę (bps). Z upływem lat wprowadzano coraz szybsze sposoby przesyłania informacji cyfrowych. Warto przypomnieć, że w pierwszych lokalnych sieciach internetowych (LAN) prędkość transmisji danych była rzędu 10 milionów bitów na sekundę, czyli 10Mbps, (10Mb/s). Dziś w sieciach LAN wykorzystujemy urządzenia, z których wszystkie mogą przekazywać dane z szybkością 100Mb/s, a wiele z nich z prędkością 1000Mb/s czyli 1Gb/s. I realizują to za pomocą „kablów internetowych”, najczęściej wykorzystując tylko dwie spośród czterech par (skrętek) dostępnych w kablu UTP i podobnych.

Podczas transmisji danych stosuje się ogólnosiłkowe standardy. I rzeczywiście, przez jedne kable można transmitować dane z prędkością wielokrotnie większą niż przez inne, podobne z wyglądu. Uzasadnione jest pytanie, czym w istocie różnią się kable o tak różnych możliwościach?

Próbę wyjaśnienia „na chłopski rozum” należy zacząć od wyliczenia trzech głównych problemów: tłumienia, odbić i zakłóceń.

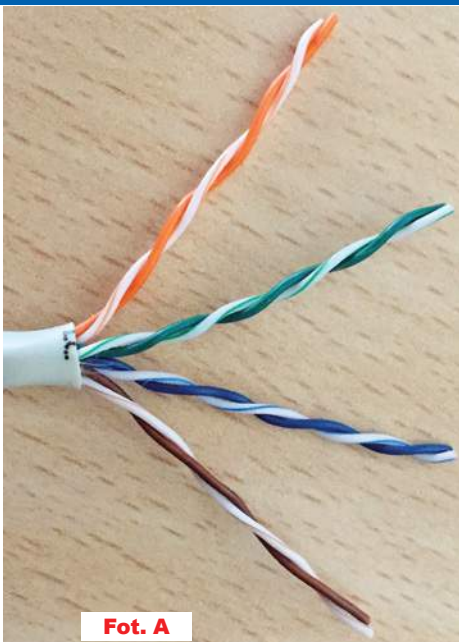
Tłumienie. Jeżeli na jednym końcu kabla dołączymy zasilacz 5-woltowy, a na drugim woltomierz, to woltomierz pokaże 5V, niezależnie od długości kabla. Podobnie jeżeli na jednym końcu kabla dołączymy generator impulsów prostokątnych o amplitudzie 5V i jakiejś małej częstotliwości, np. 100Hz, to oscyloskop dołączony na drugim końcu kabla pokaże na ekranie impulsy prostokątne o amplitudzie 5V, mające ładne, ostre zbocza.

Jeżeli jednak w generatorze będziemy zwiększać częstotliwość tych 5-woltowych impulsów prostokątnych, to zauważymy, że po pierwsze impulsy na ekranie oscyloskopu zaczynają się zniekształcać, a po drugie, stają się coraz mniejsze. Zniekształcanie impulsów ma różne przyczyny, co częściowo omówimy za chwilę. Teraz interesuje nas fakt, że czym wyższa częstotliwość i czym dłuższy kabel, tym silniej tłumione są impulsy. Po przejściu przez kabel stają się one mniejsze. W praktyce nie trochę mniejsze, tylko wielokrotnie mniejsze.

Tu trzeba od razu (bez szerszego uzasadnienia) poinformować mniej zaawansowanych, że przyczyną tłumienia nie jest ani nieunikniona pojemność kabla, ani tym bardziej jego indukcyjność. Indukcyjność i pojemność decydują o tak zwanej rezystancji falowej kabla i mają wpływ na opóźnienie sygnału, a nie na jego tłumienie. Tłumienie kabla zależy od jakości, od dobroci jego pojemności i indukcyjności. Ta dobroć (Q) wiąże się ze stratami cieplnymi, co jest reprezentowane przez rezystancję strat kabla. Ale nie tylko przez rezystancję miedzianych żył, zmierzoną omomierzem, nie przez rezystancję falową, tylko przez zastępczą rezystancję strat. Są to głównie straty w pojemności – w dielektryku (izolatorze) umieszczonym pomiędzy żyłami kabla, a mniej straty w indukcyjności – w miedzianych przewodach. Najprościej biorąc: tanie są materiały izolacyjne o znacznych stratach, a realizacja małostratnego kabla jest wprawdzie możliwa, ale jest kłopotliwa i kosztowna. Dlatego pomijając inne szczegóły, możemy powiedzieć: jeżeli dane kable transmisyjne mają być szeroko wykorzystywane, to muszą być tanie, a więc muszą być budowane z tańszych, gorszych materiałów o większym tłumieniu.

Następna kwestia to **odbicia**. Najprościej biorąc, dowolne sygnały o wysokich częstotliwościach, w tym impulsy wysokiej częstotliwości, ulegają odbiciom na wszelkich niejednorodnościach, jakie napotykają na swej drodze. Główna kwestia to tzw. *dopasowanie falowe* – każdy kabel transmisyjny ma jakąś indukcyjność oraz pojemność, i to one określają jego rezystancję falową. Kable radiowe najczęściej mają rezystancję falową 50Ω albo 75Ω, natomiast popularne kable komputerowe (tzw. skrętka internetowa) mają rezystancję falową około 100...110 omów. Aby uniknąć odbić, konieczne jest dopasowanie rezystancji (impedancji) z obu stron kabla. Prawidłowe dopasowanie falowe na obu końcach kabla to tylko część problemu odbić. Niewielkie odbicia mają źródło we wszelkich nieciągłościach i nieregularnościach kabla. Najprościej biorąc: jeżeli kabel wykonany jest byle jak, w szczególności jeżeli występują zaburzenia jego (symetrycznej) budowy, to występują różne małe odbicia. I suma tych odbić dodaje się do transmitowanego sygnału i w różny sposób go deformuje. Najprościej biorąc – rozmywa. Do problemu tłumienia dochodzi problem odbić, powodujący deformację i rozmycie sygnału.

I jeszcze trzeci problem **zakłóceń**. W naszym otoczeniu występują najróżniejsze zakłócenia impulsowe, pochodzące z wielu źródeł naturalnych (np. wyładowania atmosferyczne) i sztucznych (np. zasilacze impulsowe czy pracujące silniki komutatorowe). Niewątpliwie takie zewnętrzne zakłócenia trzeba tłumić. Są dwie drogi: jedna to stosowanie symetrycznej skrętki, druga to ekranowanie. W przypadku ekranowania jego skuteczność zależy od kilku czynników, między innymi



Fot. A

od grubości i rodzaju materiału ekranującego. Aby skuteczność tłumienia zakłóceń była jak najlepsza, skrętka powinna być idealnie symetryczna – każde odstępstwo (niesymetria, niejednorodność) zwiększa wrażliwość na zewnętrzne zakłócenia.

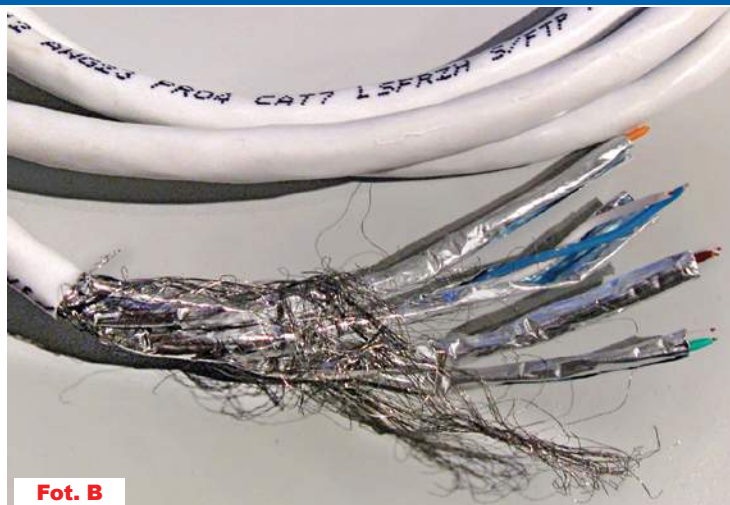
A oto inny aspekt problemu zakłóceń. O ile ze sporadycznie występującymi zakłóceniami zewnętrznymi można skutecznie walczyć, stosując korekcję błędów, ewentualnie też przez powtarzanie

zakłóconej transmisji, o tyle w praktyce największym problemem są zakłócenia powiedzmy „wewnętrzne” – przesłuchy sygnału między sąsiednimi skrętkami (parami). W typowym dziś kablu komputerowym mamy cztery niezależne tory – cztery skrętki. Właśnie dla poprawy separacji między tymi czterema torami każda para ma nieco inny skok skrętki, co widać wyraźnie na **fotografii A**. (z Wikipedii: Johan Braeken CC BY 3.0). Lepsze kable komputerowe, pozwalające na szybszą transmisję, są dodatkowo ekranowane. Przykład na **fotografii B**.

Zależnie od trzech omówionych właśnie głównych czynników, wynikających ze sposobu produkcji, jakości użytych materiałów oraz staranności i precyzji wykonania, można uzyskać kable o lepszych lub gorszych właściwościach. W praktyce wszystko to jest naprawdę bardzo starannie przemyślane, żeby zrównoważyć cenę oraz szybkość i zasięg transmisji.

Autor pytania wspomina o klasach. Istotnie, ogólnosiwiatowa norma ISO/IEC 11801 wymienia szereg klas kabli do przesyłania informacji, które to klasy należy powiązać z szerokością pasma przesyłanych sygnałów. W praktyce, w przypadku kabli, prawie nigdy nie mówimy o *klasach*, tylko o *kategoriach* (w skrócie cat.), które jednak mają ścisły związek z klasami. Poszczególne kable mają swoje kategorie, od 1 do 8. Czym wyższa kategoria, tym szybsza może być prawidłowa transmisja danych.

Kategoria	Max. prędkość Data Rate	Pasma Bandwidth	Max. długość
Kategoria 1	1Mbps	0.4MHz	
Kategoria 2	4Mbps	4MHz	
Kategoria 3	10Mbps	16 MHz	100m
Kategoria 4	16Mbps	20MHz	100m
Kategoria 5	100Mbps	100MHz	100m
Kategoria 5e	1Gbps	100MHz	100m
Kategoria 6	1Gbps	250MHz	100m przy 37m
Kategoria 6a	10Gbps	500MHz	100m
Kategoria 7	10Gbps	600MHz	100m
Kategoria 7a	10Gbps	1000MHz	100m przy 50m
Kategoria 8	25Gbps (Cat8.1) 40Gbps (Cat8.2)	2000MHz	30m



Fot. B

Oto poszczególne klasy:

Klasa A: do 100kHz, wykorzystuje kable kategorii 1.
Klasa B: do 1MHz, wykorzystuje kable kategorii 2.
Klasa C: do 16MHz, wykorzystuje kable kategorii 3.
Klasa D: do 100MHz, wykorzystuje kable kategorii 5e.
Klasa E: do 250MHz, wykorzystuje kable kategorii 6.
Klasa EA: do 500MHz, wykorzystuje kable kategorii 6A.
Klasa F: do 600 MHz, wykorzystuje kable kategorii 7.
Klasa FA: do 1000MHz, wykorzystuje kable kategorii 7A.
Klasa BCT-B: do 1000MHz, wykorzystuje kable koaksjalne.
Klasa I: do 2000 MHz wykorzystuje kable kategorii 8.1.
Klasa II: do 2000 MHz wykorzystuje kable kategorii 8.2.

Mamy więc do dyspozycji kable o kategoriach 1...8, w tym najpopularniejsze dziś kategorii 5e, a także „szybsze” 6A, 7A oraz 8.1 i 8.2. W danym zastosowaniu ze względów ekonomicznych wykorzystuje się kable możliwie najniższej kategorii, byle tylko spełniły swoje zadanie.

Warto dodać, że w praktyce nie spotykamy kabli kategorii 1 i 2. Najstabsza, telefoniczna skrętka ma kategorię 3 i nie nadaje się na współczesne kable ethernetowe, gdzie wymagana jest kategoria co najmniej 5. Są jednak wyjątki, o czym za chwilę

Autor pytania chce wiedzieć, skąd biorą się aż tak duże różnice. Nie jest to łatwa kwestia. Otóż różnice szerokości gwarantowanego pasma przenoszenia w poszczególnych klasach są duże. Dla klasy najniższej to tylko 100kHz, dla najwyższej 2000000kHz, czyli 20 tysięcy razy więcej.

Ale jeszcze większe są różnice prędkości transferu, co pokazuje **tabela obok**. Co ważne, nie ma tu prostej zależności między szerokością pasma i maksymalną prędkością transmisji.

Na pewno omówione trzy problemy: tłumienie, odbicia oraz zakłócenia, mają bezpośredni związek z szerokością pasma przenoszenia, wyrażoną w megahercach. Ale w grę wchodzi też maksymalna długość kabla, dla której gwarantowane są dane parametry. W przypadku Ethernetu w najpopularniejszej wersji 100BASE maksymalna długość kabla to 100m. W innych klasach ta maksymalna przewidziana w danym standardzie (klasie, kategorii) długość może być inna. Na przykład dla najszybszej kategorii 8 maksymalna długość kabla – łączy to tylko 30m.

Ciąg dalszy na stronie 61

Telewizor bez zasilacza?

Poprzedni odcinek zakończyliśmy rysunkiem 1 i pytaniem:

Czy w ogóle jest tu zasilacz?

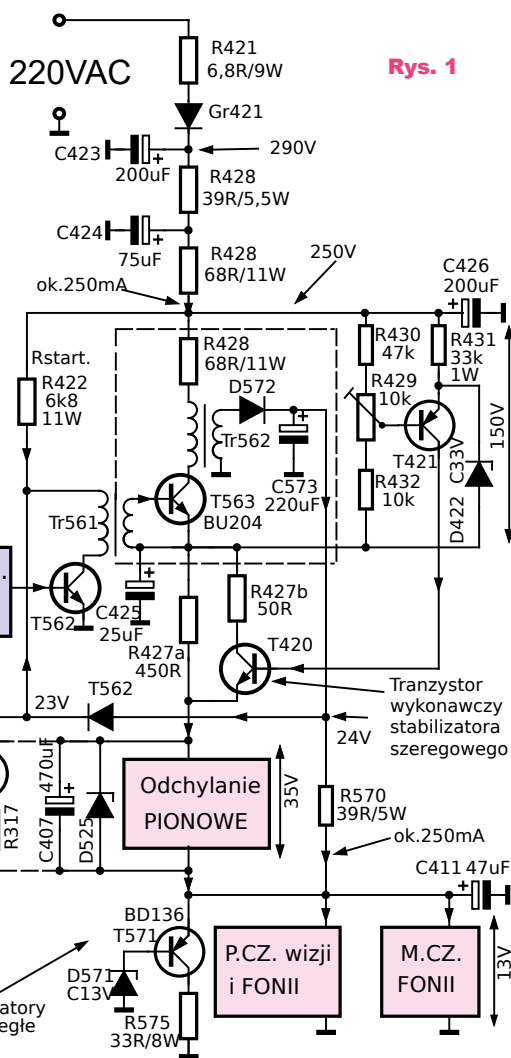
Otóż jako szcztąkowy zasilacz należy traktować obecne tu stabilizatory równoległe plus „pływający” szeregowy.

Rezystory R421, R420 i R428 służyły wstępnemu zbiciu napięcia z ok. 310V do ok. 250V i wraz z istniejącymi tam kondensatorami, wstępnej filtracji. To opory o mocach kolejno 9W, 5,5W i 11W. I faktyczna moc wydzielana na nich była bliska nominalnej. 5,5-watowy R420 miał na sobie bezpiecznik termiczny, który miał przerwać obwód, gdy się przegrzeje. Faktycznie, często „rozpinał”, ale z reguły przed lawinowym uszkodzeniem uchronić nie zdążył. Podobnie jak istniejący tam (w gałęzi prądu stałego), niepokazany na rysunku 1, bezpiecznik topikowy 400mA. W warunkach sprawnego odbiornika prąd gałęzi szeregowego zasilania nie powinien przekraczać 300mA. Zwróćmy uwagę na rezystory przed i za diodą GR421. R421 o mniejszej wartości jest bardziej skuteczny w „zbiciu” nadwyżki napięcia. I pytanie pierwsze: Czy jest tu jakiś zysk pod względem mocy? To znaczy, czy przerywając rezystor przed pierwszym kondensator filtra (C423), uda się zbić napięcie o tę samą wartość z mniejszą stratą mocy?

Analizujemy dalej uproszczony schemat z rysunku 1. Obwód objęty linią przerywaną to stopień końcowy odchyłania poziomego. Tu też mamy rezystor-grzejnik, który można by wyrzucić poza ten obręb, ale z przyczyn, które staną się zaraz jasne, należy wliczyć w stopień końcowy odchylenia linii. Obwód z tranzystorami T420 i T421 utrzymuje stałe napięcie między „górną stroną” R426 i emiterem tranzystora kluczującego linii T563 – BU204. To właśnie szeregowy stabilizator pływający. Kontroluje napięcie, jak powiedziano wcześniej, a tranzystor wykonawczy jest w gałęzi całego zasilania szeregowego. To T420 wyposażony w spory radiator. Podwójny rezystor 25-watowy R427 (lub dwa osobne opory dużej mocy w Neptunie 625) służyły przejściu części

mocy z tranzystora T420 na siebie. Równocześnie ich wartość była tak dobrana, aby zapewnić wystarczająco szeroki zakres regulacji. Szczególnie w Neptunach dwudziestowatowy R971 często się uszkadzał. A odbiornik pracował dalej. Tak długo, aż uszkodził się T420. Ale nawet wtedy pracował dalej, choć już niedługo! Jakie elementy były narażone (na praktycznie nieuniknione) uszkodzenie w tym stanie? To pytanie drugie. Jak poprawić konstrukcję, aby zmniejszyć awaryjność uszkadzania R971, ale równocześnie nie przerywać mocy na tranzystor T420? To pytanie trzecie.

Kolejnymi blokami zasilanymi kaskadowo (jeden nad drugim) był układ odchyłania pionowego i obwody małosygnałowe pośredniej częstotliwości wizji i fonii, łącznie ze wzmacniaczem m.c. fonii. Ten ostatni charakteryzował zmienny pobór prądu, był wykonany na układzie scalonym pracującym w klasie B. Stabilizację napięcia ok. 13V zapewniał tu stabilizator równoległy z tranzystorem T571, BD136 wyposażony w radiator. Na rysunku 1 nie uwzględniono jedynie zasilania głowicy w.c. i wzmacniacza wizji. Przypomnijmy tylko, że w starszych odbiornikach tej rodziny głowica wymagała ujemnego napięcia -12V, w nowszych +12V. Mały pobór mocy usprawiedliwia nieujęcie tych bloków w strukturze logiki zasilania odbiornika. Osobnym wątkiem, wartym ewentualnego omówienia, jest zasilanie głowicy tzw. napięciem warikapowym. To zawsze było realizowane z obwodu wysokiego napięcia przez duży rezystor, co oznaczało, z najprościej wykonanego źródła prądowego, z bardzo niską sprawnością. Mimo to, w całym bilansie energetycznym odbiornika, było to praktycznie bez znaczenia. Pierwsze głowice pracowały z tzw. syntezą napięciową, czyli – bez syntezy częstotliwości. Na



Rys. 1

dodatek nie miały nawet obwodu ARCz (automatycznej regulacji częstotliwości), obecnej nawet w prostych odbiornikach radiowych. Wymagały w związku z tym bardzo stabilnego napięcia w celu przestrajania pojemności obwodu heterodyny i obwodów w.c. Obwód zasilania zawierał układ scalony będący skompensowaną termicznie diodą Zenera. Mimo to aż dziwne, że nie było większych kłopotów (bo mniejsze były) z odstawianiem się odbiornika od zaprogramowanej stacji. Wracamy do głównego tematu.

Wyższa kaskada, którym był tu układ odchyłania pionowego, także miała swój stabilizator równoległy.

W odbiornikach, gdzie „ramka” pracowała z układem scalonym, był to stabilizator równoległy tej samej budowy, co kaskada najniższa. W odbiornikach z odchyleniem pionowym, wykonanym na dyskretnych tranzystorach, była tu jedynie dioda Zenera. O napięciu 47V, praktycznie w charakterze zabezpieczenia. W jednych i drugich była interakcja między regulacją szerokości i wysokości rastra obrazu. W pierwszych, powiększając szerokość, obserwowano obniżanie wysokości, w drugich obserwowano się także wzrost wymiarów pionowych. I tu pytanie czwarte: dlaczego?

Wypada kończyć to opowiadanie, ale koniecznie trzeba jeszcze dopowiedzieć parę uwag do schematu na rysunku 1. Główną trudnością tego projektu jest takie zaprojektowanie wszystkich bloków funkcjonalnych telewizora, aby zasilanie szeregowe, kaskadowe miało sens. Trzeba różnicować napięcie, ale zbilansować – wyrównać pobór prądu. Aby stabilizatory nie musiały dużo „wyrównywać”. Ale pobór prądu przez najniższą kaskadę jest większy i to zapewne za sprawą wzmacniacza m.cz. fonii. W przypadku stabilizacji równoległej trzeba przewidzieć maksymalny pobór mocy (maksymalna siła głosu, jak wtedy to nazywano). Rezystor R570 (39Ω, pięciowatowy) doprowadza dodatkowy prąd do najniższej kaskady, pobierając energię z uzwojenia dodatkowego transformatora wysokiego napięcia, co się odbija też w prądzie gałęzi głównej (jednak tu energia przetwarzana jest z dużą sprawnością). Korekta nie jest marginalna. Nietrudno policzyć, znając napięcia i wartość rezystora, że to ok. 250mA. Czyli drugie tyle, ile do tej (najniższej) kaskady wpływa z szeregowego zasilania tego odbiornika. Około połowa mocy pobranej z uzwojenia 6-5 transformatora i tak jest wydzielana w rezystorze R570. A przy niskim poziomie głosu cała nadwyżka grzeje tranzystor T571 i opór w jego kolektorze, który musi być tak dobrany, aby w żadnych warunkach nie nastąpiło nasycenie T571. Dobór tej rezystancji wymaga kilku kompromisów i trudno w nich uwzględnić także sytuacje awaryjne. Napięcie pozyskane z dodatkowego uzwojenia 6-5 transformatora służy do jeszcze jednego celu. To napięcie ok. 24V zasila również stopień sterujący odchyleniem poziomego i po zbiegu do 9V także jego generator.

Ten skróty opis wyczerpuje najważniejsze aspekty szeregowego zasilania

zastosowanego w tym odbiorniku. Poza jednym faktem: tak odbiornik działał nie będzie! Chciałoby się zostawić to jako kolejne pytanie dla Czytelników, chcących dogłębnie zrozumieć zastosowaną tu ideę i logikę. Ale trzeba to wyjaśnić: otóż układ wymaga startu! Rezystorem startowym jest tu R422. Rezystor 11-watowy i nietrudno policzyć, że wydzielą się na nim w nominalnych warunkach ok. 8W mocy. Można by go po starcie jakimś tranzystorem odłączyć, ale kto by się tam wtedy przejmował kilkoma watami mocy, podobnie jak „niezdrowym” dla sieci obciążeniem spowodowanym jednopółprzewodnikowym prostowaniem.

Widzimy, że mimo sporych starań i zabiegów, marnotrawstwo energii jest duże. Zastosowanie zasilacza liniowego z dużym i ciężkim transformatorem sieciowym niewiele sytuację tę poprawi. To przykład ze starych schematów, bo jakże może być inaczej? W nowych takiego nie znajdziemy, pomijając drobny fakt, że teraz nie ma schematów! Przetwornice tak zdominowały całą branżę obwodów zasilania, że ich stosowanie wydaje się nam oczywiste i nikt w żadnym projekcie nie będzie się tak gimnastykował, aby pogodzić sprawność, prostotę i zrealizować zasilanie z pominięciem ciężkiego i drogiego transformatora sieciowego. Konstrukcje przetwornic też są coraz doskonalsze i jak lubią podawać materiały reklamowe, z coraz wyższym stosunkiem mocy na sześcienny cal.

Temat bieżącego artykułu jest dobrym przykładem, jaka była presja, aby opracować sprawne zasilacze impulsowe. W tym miejscu aż się prosi o przykład z branży RTV z wykorzystaniem odchylenia tyrystorowego. Tu brak zasilacza jest dodatkowym bonusem. Zdaniem autora jest on wart co najmniej jednego odcinka tego cyklu. Ale dobór tematów dostosujemy do opinii Czytelników.

Ostatnim OTV czarno-białym na naszym rynku był Hermes T400/T600. Równocześnie był to pierwszy polski odbiornik z przetwornicą napięcia. Wyprzedził go radziecki Foton 234D, także wyposażony w porządną zasilacz impulsowy. Odbiornik też o interesującej konstrukcji, ale to temat na inne opowiadanie.

Na koniec tego opracowania chciałbym zadać jeszcze jedno pytanie. Odbiornik, rodzina odbiorników OTV z zasilaniem szeregowym, była zapro-

jektowana na napięcie sieci 220VAC. Widać wyraźnie, że cały problem i kłopot miał swoje źródło w tym, że to wartość niewygodnie wysoka. A jakie są konsekwencje pracy tego telewizora na napięciu 230VAC? Jak się zmieni pobierana moc i na których elementach wydzieli się jej nadwyżka? Czy jest tu jakiś prosty sposób zaradczy, gdyby te odbiorniki były nadal w użyciu?

I na koniec jeszcze kilka pytań dla szczególnie dociekliwych Czytelników.

Zalóżmy, że napięcie sieci jest stałe, zadane. Powiedzmy nominalne 220VAC. A dociążamy jedno z uzwojeń wtórnych transformatora wysokiego napięcia. Na przykład uzwojenie WN poprzez zmianę prądu anodowego kineskopu (jasność ekranu). Ta dodatkowa energia przetransformowana jest z uzwojenia pierwotnego, głównego tego transformatora. Ale napięcie zasilania tego stopnia jest stabilizowane. Na których elementach wydzieli się nadwyżka mocy? I jaka orientacyjnie będzie ta nadwyżka w stosunku do mocy użytecznej pobranej z uzwojenia wtórnego (dodatkowego) transformatora, zakładając, że sprawność w transformatorze jest stuprocentowa? Pytanie kolejne, trudne: jaka jest funkcja mocy wydzielanej w tranzystorze szeregowego regulatora pływającego (T420) od powyższych zależności? I kolejne pytanie, z częściową odpowiedzią. Stabilizator ten musi nadążać za tętnieniami o częstotliwości sieci 50Hz. A którą zamyka się składowa prądów o częstotliwości linii 15625Hz? Czy stabilizator ten ją „widzi” i próbuje za nią nadążać?

I jeszcze ostatnie, też trudne: jakie znaczenie ma fakt, że obwodem sprzężenia zwrotnego – stabilizacji napięcia na stopniu końcowym odchylenia linii, objęto też rezystor R426? To opór dużej wartości 68Ω/11W. Co by się zmieniło, gdyby opór ten wyrzucić poza pętlę sprzężenia zwrotnego?

Powyższy artykuł można potraktować jedynie jako ciekawostkę techniczną. Można też potraktować jako wstęp do dobrodziejstw, jakie przyniosły przetwornice. Czekamy na Wasze e-maile!



Karol Świerc
rtv@silnet.pl

Wybrane książki dla Czytelników „Elektroniki dla Wszystkich”

Encyklopedia elementów elektronicznych. Tom 1. Rezystory, kondensatory, cewki indukcyjne, przełączniki, enkodery, przekaźniki i tranzystory

Autor: Charles Platt; Stron: 296; Oprawa: miękka; Kod: KS-210200

To książka przeznaczona dla początkujących i zaawansowanych elektroników, zarówno inżynierów, jak i hobbystów. Zawiera starannie zebrane, skompletowane, uporządkowane, a co najważniejsze, sprawdzone i potwierdzone informacje o elementach elektronicznych. Pierwszy z trzech tomów obejmuje informacje o podstawowych elementach, wykorzystywanych chyba we wszystkich projektach.

Rezystory, kondensatory, cewki indukcyjne, przełączniki, enkodery, przekaźniki i tranzystory. Dokładne informacje o każdym komponencie: funkcja, działanie, rodzaje, wartości, stosowanie, możliwe błędy.

Absolutny niezbędnik każdego elektronika: wiarygodny, kompletny, wyczerpujący!



Encyklopedia elementów elektronicznych. Tom 2. Tyrystory, układy scalone, układy logiczne, wyświetlacze, LED-y i przetworniki akustyczne

Autor: Charles Platt i Fredrik Jansson; Stron: 304; Oprawa miękka; Kod: KS-210202

Drugi tom niezwyklej encyklopedii przeznaczonej dla praktyków elektroniki. Podobnie jak w pierwszym, tak i tutaj znalazły się skompletowane, uporządkowane, a co najważniejsze - sprawdzone i potwierdzone informacje o elementach elektronicznych. Drugi z trzech tomów jest poświęcony układom scalonym, tyrystorom, źródłom światła i dźwięku, wskaźnikom oraz wyświetlaczom - ich opisy zostały uzupełnione licznymi fotografiami, schematami i wykresami. Dowiesz się, do czego służy każdy z prezentowanych podzespołów, jak działa, kiedy jest najbardziej przydatny i w jakich odmianach występuje. Oto prawdziwa pomoc dla praktyków, którzy chcą szybko uzyskać wskazówki potrzebne do pracy!

Absolutny niezbędnik każdego elektronika: wiarygodny, kompletny, wyczerpujący!



Lutowanie od podstaw

Autor: Witold Wrotek; Stron: 160; Oprawa miękka; Kod: KS-201000

Jeśli chcesz poznać technikę lutowania i nauczyć się prawidłowo stosować ją w praktyce, sięgnij po odpowiednie źródło wiedzy! Książka Lutowanie od podstaw krok po kroku wprowadzi Cię w tajniki sztuki łączenia elementów, przedstawi niezbędne narzędzia i dobre praktyki, nauczy unikać typowych błędów popełnianych przez początkujących oraz pokaże najlepsze sposoby lutowania różnych elementów elektrycznych i elektronicznych.

Nauczysz się też dzięki niej, jak wykonać proste prace elektryczne w swoim domu, a nawet jak naprawić typowe usterki występujące w urządzeniach AGD.

Zostań prawdziwym mistrzem lutownicy!

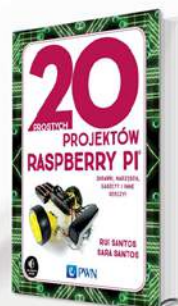


20 prostych projektów Raspberry Pi

Autorzy: Rui Santos, Sara Santos; Stron: 276; Oprawa miękka; Kod: KS-210401

Książka krok po kroku uczy, jak realizować interaktywne projekty z wykorzystaniem Raspberry Pi – małego i niedrogiego komputera – takie jak np. cyfrowy zestaw perkusyjny, robot kontrolowany przez WiFi, gra Pong, alarm antywłamaniowy wysyłający powiadomienia e-mail, domowa kamera do monitoringu, detektor wycieku gazu, stacja pogodowa czy gadżety Internetu Rzeczy (IoT) sterujące elektroniką w całym domu. W trakcie lektury czytelnik pracuje z podstawowymi komponentami, takimi jak diody LED, ekrany LCD, kamery i czujniki oraz gry i zabawki. Uczy się, jak skonfigurować własny serwer WWW, stworzyć pierwszą stronę internetową czy napisać prostą grę komputerową.

Każdy projekt zawiera instrukcję krok po kroku, kolorowe zdjęcia i diagramy, a także kompletny kod, dzięki któremu czytelnik ożywi swoje projekty.

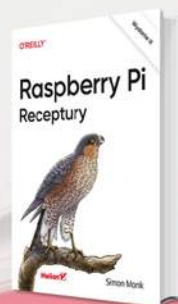


Raspberry Pi. Receptury. Wydanie III

Autor: Simon Monk; Stron: 528; Oprawa miękka; Kod: KS-200901

Zaktualizowane wydanie znakomitego zbioru receptur ułatwiających wykorzystanie potencjału Raspberry Pi. Uwzględniono tu nowe modele tego komputera, a także zmiany i ulepszenia systemu operacyjnego Raspbian. Dodano rozdziały traktujące o dźwięku i automatyce domowej. Te receptury bez trudu wykorzystasz dla zwiększenia wygody we własnym domu. Dzięki lekturze poznasz podstawowe reguły tej technologii, aby łatwiej zrozumieć zagadnienia dotyczące konkretnej płytki czy kodu. Z tej pozycji możesz korzystać podobnie jak z książki kucharskiej: przeczytać od deski do deski albo skupić się na rozwiązaniu jednego, konkretnego problemu. Być może docenisz, że w recepturach dotyczących sprzętu uwzględniono przede wszystkim rozwiązania niewymagające lutowania obwodów.

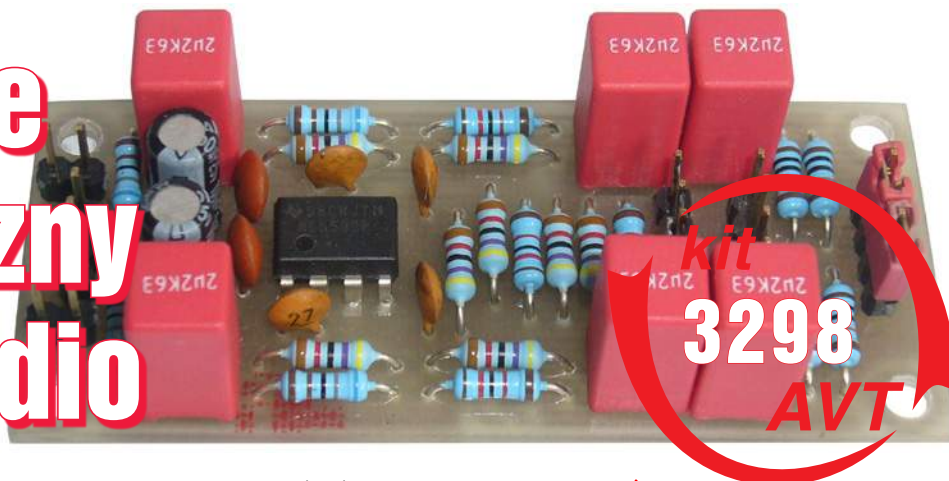
Raspberry Pi: morze możliwości dla inżyniera z pasją!



sklep.avt.pl

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

Podwójnie symetryczny moduł audio



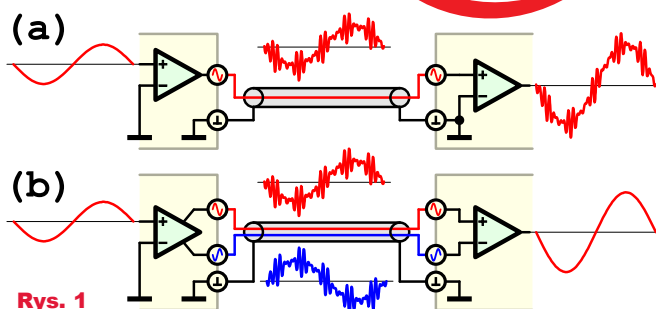
Symetryczne wejścia i wyjścia sygnału liniowego są cechą charakterystyczną profesjonalnych urządzeń audio. Urządzenie audio, doposażone w opisywany układ (moduł), zyskuje symetryczne wejście lub wyjście liniowe.

Połączenia kablowe do przesyłania analogowych sygnałów audio między urządzeniami są szczególnie narażone na przenikanie zakłóceń, co w uproszczeniu przedstawione jest na **rysunku 1**. Urządzenia audio powszechnego użytku (tzw. konsumenckie) wyposażone są z reguły w wyjścia i wejścia niesymetryczne (ang. unbalanced) (**1a**) i są łączone ze sobą za pomocą kabla z jedną żyłą w ekranie. Choć ekranowanie przewodu sygnałowego zmniejsza podatność na zakłócenia wnikające przez pojemności (zmienną pole elektryczne), to nie jest skuteczną ochroną przed zakłóceniami, których źródłem jest zmienną pole magnetyczne. Wnikające przez szkodliwą indukcyjną niesymetrycznego połączenia kablowego zakłócenia (np. przydźwięk, brum sieciowy) są w jakimś stopniu obecne w sygnale wyjściowym. Połączenie takie nie sprawdza się przy większych odległościach, dlatego w zastosowaniach profesjonalnych preferowany jest standard symetrycznego (ang. balanced) (**1b**) przesyłania sygnału. Sygnał symetryczny przesyłany jest dwoma identycznymi żyłami (zwykle w postaci skrętki w ekranie) w przeciwfazie, tj. wartości napięć zmiennych w obu żyłach różnią się znakiem. Zakłócające pola o zmiennym natężeniu (elektryczne, magnetyczne) powodują indukowanie napięć (i pra-

dów) o zbliżonych wartościach (i takich samych znakach) w obu żyłach kabla. Gdy oba wejścia odbiornika sygnału charakteryzują się identycznymi właściwościami (impedancja, pojemność wejściowa), indukowane zakłócenia są niemal jednakowe dla każdej z żył.

Odbiornik to nic innego jak wzmacniacz różnicowy, który tłumi sygnały wspólne (tj. zakłócenia) i na jego wyjściu pojawia się tylko sygnał będący wzmocnioną różnicą potencjałów na obu jego wejściach (sygnał użytkowy).

Opisywany układ jest nietypowym, dwuwyjściowym, wzmacniaczem różnicowym. Uniwersalność nieczęsto stosowanej topologii sprawia, że układ może przetwarzać sygnał symetryczny na niesymetryczny i niesymetryczny na symetryczny, tzn. może równie dobrze pracować jako odbiornik bądź nadajnik sygnału symetrycznego. Moduł przeznaczony jest do montażu w posiadanym lub budowanym urządzeniu audio, dzięki czemu zyska ono wejście (lub i wyjście) liniowe dla sygnałów audio w standardzie symetrycznym. Układ charakteryzuje się stosunkowo niską impedancją wejściową, co w przypadku pracy w roli wejścia liniowego audio może być zaletą ze względu na mniejszą podatność na przenikanie zakłóceń. Niewątpliwie najlepsze właściwości

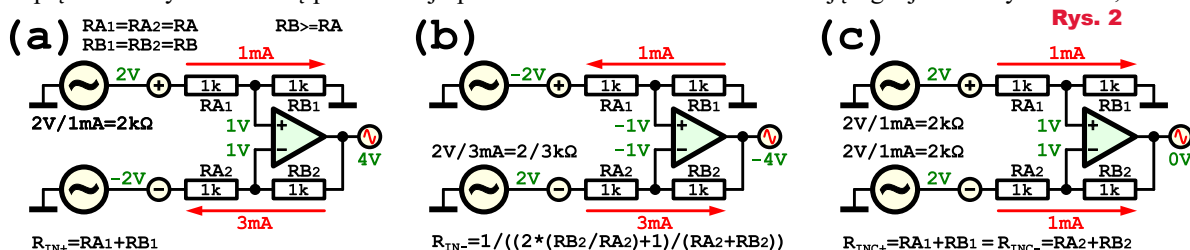


Rys. 1

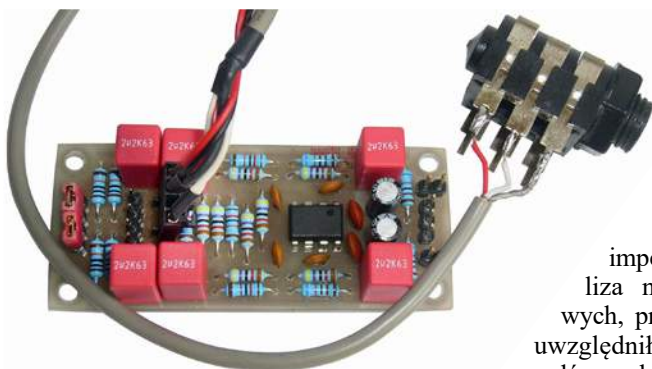
wejścia różnicowego można uzyskać, stosując specjalizowany do tego celu układ scalony lub wzmacniacz nazywany pomiarowym, najczęściej realizowany z użyciem trzech wzmacniaczy operacyjnych. Niemniej przy pracy z sygnałami audio o poziomach liniowych, gdzie mamy stosunkowo duży odstęp sygnału od zakłóceń, prezentowany układ podwójnego sprzężonego wzmacniacza różnicowego jest dobrą, mniej kosztowną alternatywą.

Opis układu

Podstawowa aplikacja wzmacniacza operacyjnego OA (ang. Operational Amplifier) pracującego jako wzmacniacz różnicowy widoczna jest na **rysunku 2**. Na pierwszy rzut oka taki układ jest połączeniem wzmacniacza odwracającego z nieodwracającym. We wzmacniaczu odwracającym napięcie na wejściu nieodwracającym OA jest równe zeru (tzw. wirtualna masa). Gdy potencjał na wejściu wzmacniacza odwracającego jest różny od zera,



Rys. 2



z/do wyjścia OA płynie prąd wywołujący spadek napięcia na rezystorze RB_2 w pętli ujemnego sprzężenia zwrotnego tak, by napięcie na wejściu nieodwracającym OA było równe zero (wirtualna masa). Innymi słowy, napięcie na wyjściu wzmacniacza jest równe potencjałowi wejścia pomnożonego przez współczynnik wzmocnienia (RB_2/RA_2), lecz o przeciwnym znaku. Impedancja wejściowa wzmacniacza odwracającego jest równa RA_2 . We wzmacniaczu różnicowym jest inaczej, napięcie na wejściu nieodwracającym OA nie jest stałe i wyznacza je dzielnik RA_1 i RB_1 , zależnie od potencjału na wejściu nieodwracającym wzmacniacza różnicowego. Na wejściach wzmacniacza różnicowego, przy dołączonym sygnale różnicowym (symetrycznym względem masy) panują napięcia równe co do wartości, ale różniące się znakiem. Zależnie od polaryzacji wejść, z/do wyjścia OA płynie prąd wywołujący spadek napięcia na rezystorze RB_2 równający potencjał wejścia odwracającego OA, z potencjałem na wejściu nieodwracającym OA. Ponieważ napięcie na wejściu nieodwracającym nie jest równe zero, prąd wejścia odwracającego wzmacniacza różnicowego (płynący z/do wejścia OA przez RB_2) jest większy, niż wynikałoby to z ilorazu wartości napięcia wejściowego do rezystancji RA_2 . W takich warunkach pracy impedancja wejścia odwracającego wzmacniacza różnicowego na pewno nie jest równa rezystancji rezystora wejściowego RA_2 , jak ma to miejsce we wzmacniaczu odwracającym. Gdy potencjały obu wejść wzmacniacza różnicowego są identyczne (sygnał wspólny), jego wyjście ma potencjał równy zero. Podsumowując, na wyjściu takiego wzmacniacza pojawia się tylko napięcie będące wzmocnioną róż-

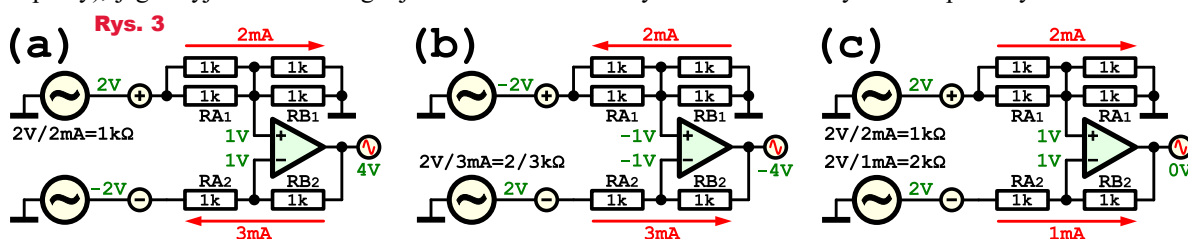
nicą napięć wejściowych, a sygnały jednakowe co do wartości i znaku, tj. wspólne (np. zakłócenia, składowa stała itp.), są tłumione.

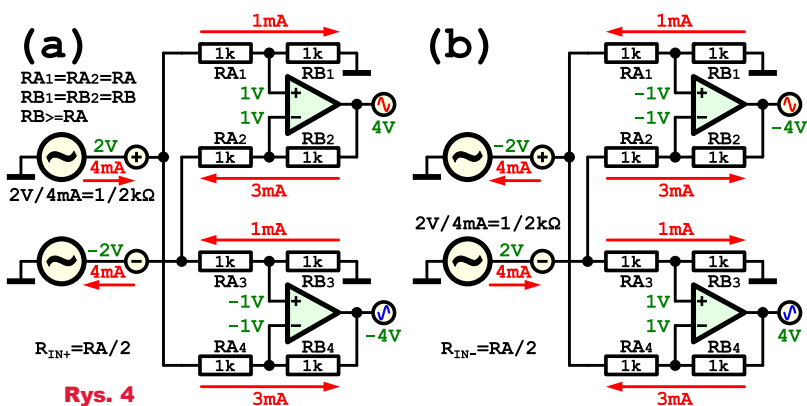
Wyznaczenie wartości impedancji wejść ułatwia analiza napięć i prądów wejściowych, przy czym na rysunkach nie uwzględniłem pomijalnych wartości prądów polaryzacji OA. Przy zastosowaniu jednakowej wartości wszystkich rezystorów ($R=RA_X=RB_X$) impedancja wejścia odwracającego dla sygnału różnicowego (**2a**) wynosi: $R_{IN-}=2/3 \cdot R$, natomiast impedancja wejścia nieodwracającego (**2b**) wynosi: $R_{IN+}=2 \cdot R$. Jak widać, wadą tego układu jest niejednakowość dla obu wejść impedancja wejściowa dla sygnału różnicowego ($R_{IN+} \neq R_{IN-}$). Nie jest to istotne, gdy wejścia układu są sterowane ze źródeł o niskiej impedancji wyjściowej, np. jak ma to miejsce we wzmacniaczu różnicowym, nazywanym pomiarowym (gdzie oba wejścia sterowane są z wejściowych OA). Nierówność impedancji wejściowych nabiera negatywnego znaczenia przy sterowaniu ze źródeł o znaczącej rezystancji poprzez kable połączeniowe. Co bardziej istotne (korzystne), impedancje dla tłumionego sygnału wspólnego (**2c**) są jednakowe dla obu wejść, i wyraża je formuła: $R_{INC+}=RA_1+RB_1=R_{INC-}=RA_2+RB_2$. Przy jednakowej wartości wszystkich rezystorów wzmocnienie wzmacniacza wynosi: $K=2$ (6dB). Wzmocnienie układu ogólnie określa wzór: $K=2 \cdot (RB/RA)$, i można je zwiększyć kosztem pogorszenia współczynnika tłumienia sygnału wspólnego CMRR (ang. Common Mode Rejection Ratio), zmieniając proporcję rezystorów wchodzących w skład dwóch grup tj. $RA=RA_1=RA_2 < RB=RB_1=RB_2$. Wtedy to dla sygnału różnicowego impedancję wejścia odwracającego określa wzór: $R_{IN-}=1/((2 \cdot (RB_2/RA_2)+1)/(RA_2+RB_2))$, natomiast impedancję wejścia nieodwracającego wyznacza wyrażenie: $R_{IN+}=RA_1+RB_1$.

Analizowanie działania wzmacniacza różnicowego jako dwóch niezależnych

wzmacniaczy (odwracającego i nieodwracającego) może prowadzić do błędnych wniosków. **Rysunek 3** ilustruje przypadek próby poprawy parametrów wzmacniacza z jednakowymi wartościami rezystorów ($R=RA_X=RB_X$) przez dodanie rezystorów (równoległe do RA_1 , RB_1) mających na celu „zrównanie” impedancji wejścia nieodwracającego z impedancją wejścia odwracającego, przez co impedancja wejścia nieodwracającego (**3a**) zmniejsza się z $R_{IN+}=2 \cdot R$ do $R_{IN+}=R$. Jednak przynosi to więcej szkody niż pożytku. Co prawda taki zabieg zmniejsza impedancję wejścia nieodwracającego dla sygnału różnicowego, ale i tak nie jest ona równa impedancji wejścia odwracającego (**3b**) ($R \neq 2/3 \cdot R$). Co gorsza, ze względu na powstałą różnicę ($1R \neq 2R$) impedancji wejść dla sygnału wspólnego (**3c**), popsuty jest współczynnik tłumienia sygnału wspólnego CMRR, który jest gorszy o $\approx 3\text{dB}$ niż w podstawowym układzie bez dodatkowych rezystorów.

Opisywany w artykule układ mający dwa wyjścia tj. w fazie (0°) i w przeciwfazie (180°), widoczny na **rysunku 4**, stanowi połączenie dwóch podstawowych wzmacniaczy różnicowych z wejściami połączonymi „krzyżowo”. Według niektórych źródeł taka topologia wzmacniacza została wymyślona i rozpropagowana przez Teda Fletchera. Każde z wejść pierwszego wzmacniacza jest połączone z „przeciwnym” wejściem drugiego wzmacniacza, tj. wejście odwracające do nieodwracającego i na odwrót. Dzięki takiemu połączeniu układ charakteryzuje się jednakową wartością impedancji dla obu wejść, a to dlatego, że impedancja pojedynczego wejścia jest równa równoległemu połączeniu impedancji wejścia odwracającego pierwszego OA z impedancją nieodwracającego wejścia drugiego OA. Innymi słowy, oba wejścia wzmacniacza różnicowego uzyskanego z pary tak sprzężonych wzmacniaczy różnicowych charakteryzują się jednakowymi właściwościami. Do prawidłowego działania wymagane jest, by wartości rezystorów spełniały warunek:



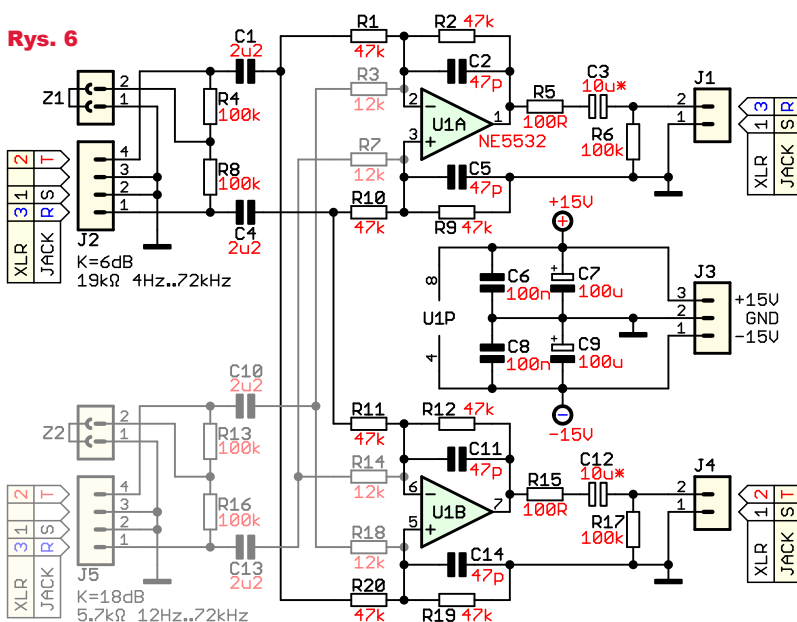
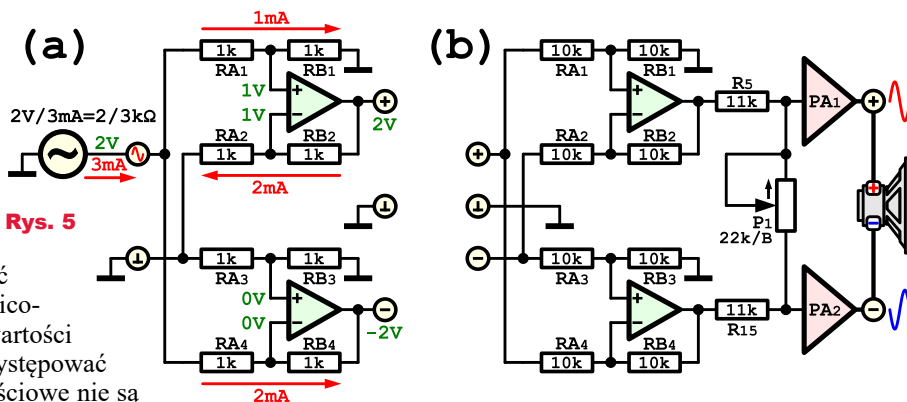


$RA=RA_1=RA_2=RA_3=RA_4 \leq RB=RB_1=RB_2=RB_3=RB_4$. Równoległe połączenie wejść OA jest przyczyną stosunkowo niskiej wartości wypadkowej impedancji wejściowej. Wypadkowa impedancja wejściowa pojedynczego wejścia dla sygnału różnicowego (**4a**, **4b**) równa jest $R_{IN} = RA/2$, a dla sygnału wspólnego (**4c**) wynosi $R_{INC} = RA$. W praktyce różnica między impedancją dla sygnału różnicowego i impedancją dla sygnału wspólnego ($1/2RA \neq RA$) nie jest tak istotna, jak nierówność impedancji w układzie z rysunku 2. Ze względu na dołączone do masy rezystory RB_1 i RB_3 (polaryzujące wejścia OA), symetryczny sygnał dołączony do wejść może być tzw. pływającym (połączenie dwuprzewodowe bez połączenia masy).

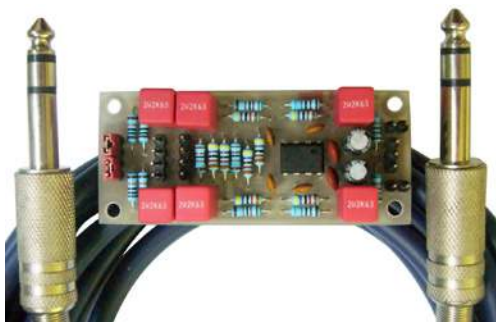
Obecność dwóch wyjść tj. w fazie (0°) i w przeciwfazie (180°) sprawia, że układ nadaje się do konwersji sygnału niesymetrycznego na symetryczny, co obrazuje **rysunek 5**. Układ może pełnić funkcję nadajnika sygnału symetrycznego (**5a**), przy czym rezystor RB_3 może być niemontowany, a rezystor RA_3 może być zastąpiony zworą. Zaletą takiego układu jest jednakowy czas „wytwarzania” obu sygnałów wyjściowych 0° i 180° (co oznacza relatywnie mały błąd fazy), ale i tak jakość tak uzyskanego wyjściowego sygnału różnicowego w głównej mierze zależy od rozrzutu wartości rezystorów RA , RB , przez co zawsze będą występować małe błędy amplitudy tak, że oba sygnały wyjściowe nie są dokładnymi „odbiciami” względem siebie (nie sumują się idealnie do zera). Innym zastosowaniem takiego układu jest praca w roli stopnia wejściowego (z wejściem symetrycznym lub/i niesymetrycznym) sterującego dwoma identycznymi, niemożliwymi końcówkami mocy PA (**5b**), umożliwiając ich pracę w konfiguracji przeciwsoonej tzw. mostkowej BTL (ang. Bridge Tied Load), gdzie regulacja głośności może być dokonywana jednym potencjometrem o charakterystyce wykładniczej audio. Teoretycznie, aby zapewnić maksymalne tłumienie sygnałów, gdy P_1 jest ustawiony na minimum, rzeczywiste wartości R_5 i R_{15} powinny być równe. W praktyce lepiej jest zastosować R_5 o nieco mniejszej wartości od R_{15} , z szeregowo włączonym potencjometrem montażowym umożliwiającym kalibrację maksymalnego tłumienia dla ustawionej, minimalnej rezystancji P_1 .

Schemat ideowy układu, którego zasada działania została właśnie opisana, widoczny jest na **rysunku 6**. Układ wymaga zasilania napięciem symetrycznym $\pm 15V$ doprowadzanym do złącza J_3 , które jest odprze-

powstawanie przesunięcia stałonapięciowego w układzie, w przypadku różnicy potencjału mas, w zakresie wytrzymałości napięciowej elementów wejściowych RC, źródła sygnału różnicowego i układu. Gdy wejście jest niepodłączone, rezystory R_4 , R_8 polaryzują C_{IN} potencjałem masy (założone są zworniki $Z1$, $Z2$). Ponieważ R_4 , R_8 są dołączone równoległe do impedancji wejść, ich obecność wpływa na wypadkową wartość impedancji wejściowych wzmacniacza. Nie mają one jednak wpływu na dolną częstotliwość graniczną wejścia, dlatego wzór na jej obliczenie wygląda następująco: $f_{GP} = 1/(2\pi \cdot RA/2 \cdot C_{IN})$. Kondensatory C_2 , C_5 , (C_{DP}) zapewniają ograniczenie pasma przenoszenia od góry, przy czym rozrzut wartości (C_{DP}) ma nie mniej istotne znaczenie,



gane kondensatorami $C_6...$ C_9 . Elementem aktywnym jest podwójny wzmacniacz operacyjny $U1$. Odcinające składową stałą kondensatory wejściowe C_1 , C_4 , (C_{IN}) uniemożliwiają



jak rozrzut wartości RA i RB na charakterystykę tłumienia CMRR, tym razem w funkcji częstotliwości (zafalowania charakterystyki). Górną częstotliwość graniczną wzmacniacza obliczamy ze wzoru: $f_{DP} = 1/(2\pi \cdot RB \cdot C_{DP})$. Rezystory R5, R15 stanowią zabezpieczenie wyjść U1 przed zwarciami. Rolę elementów: R6, R17, C3, C12, opisałam w następnym śródtytuł.

Przy wartościach elementów: $RA=RB=47k\Omega$, $C_{IN}=2,2\mu F$, $C_{DP}=47pF$, czyli jak na schemacie, wzmacniacz charakteryzuje się wzmocnieniem: $K=2V/V$ (+6dB), impedancją wejściową dla sygnału różnicowego odniesionego do masy: $R_{IN}=19k\Omega$ (tj. $38k\Omega$ między wejściami, dla wejścia pływającego), impedancją pojedynczego wejścia dla sygnału wspólnego odniesionego do masy: $R_{IN}=32k\Omega$ (tj. $64k\Omega$ między wejściami, dla wejścia pływającego), dolną częstotliwością graniczną: $f_{GP}=1/(2\pi \cdot 23,5k\Omega \cdot 2,2\mu F)=3Hz$, i górną równą: $f_{DP}=72kHz$.

Opcjonalne elementy (rysowane szarym kolorem) umożliwiają uzyskanie dodatkowego wejścia różnicowego na złączu J5, wtedy to układ pełni funkcję dwu-wejściowego różnicowego sumatora, miksera. Sygnał na wyjściach układu jest sumą wzmocnionych sygnałów różnicowych z obu wejść J2 i J5. Podane na schemacie wartości ($12k\Omega$) rezystorów RA wyznaczają parametry dodatkowego wejścia, które charakteryzuje się właściwościami: $K=7,8V/V$ (+18dB), $R_{IN}=5,7k\Omega$ ($11,4k\Omega$), $R_{INC}=10,7k\Omega$ ($21,4k\Omega$), $f_{GP}=12Hz$, $f_{DP}=72kHz$. Większa czułość toru J5 sprawia, że współczynnik tłumienia sygnału wspólnego CMRR dla tego wejścia jest gorszy niż dla toru J2. Oczywiście zastosowanie identycznych wartości elementów w obwodach obu torów (J2, J5) skutkuje uzyskaniem dwóch wejść różnicowych o zbliżonych parametrach.

Wartości elementów na schemacie nie są przypadkowe, wynikają z ogólnie przyjętych poziomów sygnałów liniowych audio, profesjonalnego +4dBu ($1,228V$, $1,736V_{pp}$) i konsumenckiego

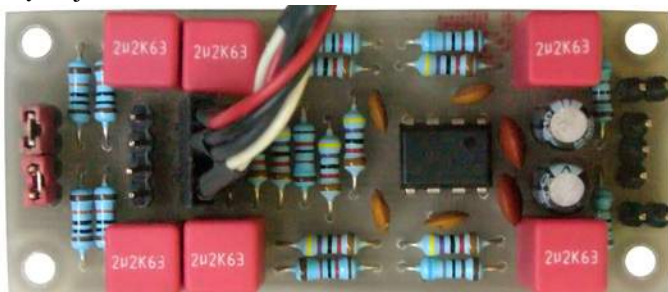
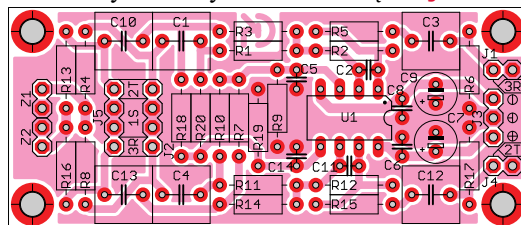
–10dBV ($0,316V$, $0,447V_{pp}$). Różnica czułości torów wejściowych wynika z potrzeby uzyskania zbliżonego poziomu sygnału na wyjściu układu, dla sygnału o poziomie +4dBu dołączonego do J2, i dla sygnału o poziomie –10dBV dołączonego do złącza J5. Montując jedynie R1, R2, R14, R18 bez (R3, R7, R11, R20) możliwe jest uzyskanie modułu (stereo) z dwoma niezależnymi wzmacniaczami różnicowymi, według topologii z rysunku 2. Wtedy to, by zachować jednakowe fazy sygnałów wyjściowych obu wzmacniaczy, należy zamienić między sobą sygnały wejściowe w jednym ze złączy wejściowych (J2 lub J5).

Montaż i uruchomienie

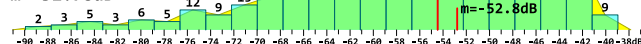
Lutowanie elementów w, uprzednio sprawdzoną pod względem zwarc i braku ciągłości ścieżek, płytkę drukowaną, widoczną na **rysunku 7**, należy przeprowadzić, stosując kryterium gabarytowe, tj. montując je w kolejności od najmniejszych do największych. W zasadzie jako U1 można zastosować dowolny podwójny OA przeznaczony do pracy w układach audio np. NE5532, NE5532A, TL072 albo lepszy. W zależności od potrzeb, wartości elementów można zmieniać w szerokich granicach. Na przykład, przy użyciu rezystorów RA, RB, R4, R8 o wartości równej $10k\Omega$, kondensatorach: $C_{IN}=2,2\mu F$, $C_{DP}=220pF$, uzyskujemy wzmacniacz różnicowy o wzmocnieniu $K=2V/V$ (6dB), częstotliwościach granicznych odpowiednio dolnej i górnej: $f_{GP}=14Hz$, $f_{DP}=72kHz$, impedancji wejść dla sygnału różnicowego odniesionego do masy: $R_{IN}=3,3k\Omega$ tj. $6,6k\Omega$ między wejściami, dla wejścia pływającego oraz impedancji wejść dla sygnału wspólnego między wejściami: $R_{INC}=10k\Omega$ tj. $5k\Omega$ dla pojedynczego wejścia odniesionego do masy. W większości przypadków zwory Z1, Z2 mogą być wlutowane na stałe. Z uwagi na symetryczne zasilanie układu kondensatory C3, C12 najczęściej nie

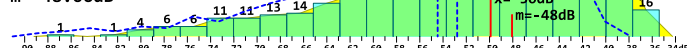
będą potrzebne i mogą być zastąpione zworami. Gdy ich obecność jest konieczna, np. z powodu występowania składowej stałej, to w przypadku montażu kondensatorów elektrolitycznych należy zwrócić uwagę na ich

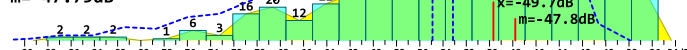
biegunowość, która powinna być zgodna z polaryzacją stałoprądową współpracujących obwodów. Oczywiście biegunowość kondensatorów nie ma znaczenia w przypadku typów foliowych oraz elektrolitycznych bipolarnych. Pojemność ($C3=C12=C_X$) będzie od impedancji (R) wejściowej obwodu, do którego kierowany jest sygnał wyjściowy układu, i złożonej dolnej częstotliwości granicznej (f_{DP}), zgodnie ze wzorem: $C_X=1/(2\pi \cdot R \cdot f_{DP})$. Rezystory R6, R17 zapewniają potencjał masy na wyjściach, gdy C3, C12 są wlutowane, a wyjścia nigdzie niepodłączone. Podsumowując, montaż i wartości elementów C3, C12, R6, R17 zależne są od docelowej roli modułu, tj. nadajnik lub odbiornik sygnału różnicowego, sterownik końcówek mocy BTL. Chociaż CMRR współczesnych OA zwykle jest lepszy od –70dB (NE5532 $\approx$ –100dB, TL072 $\approx$ –76dB) i słabnie wraz ze wzrostem częstotliwości od kilkudziesięciu kHz, to w rzeczywistym układzie jest on z reguły słabszy. W praktyce CMRR jest w dużej mierze zależny od rozrzutu wartości zastosowanych rezystorów w ramach grup RA, RB. Im mniej różnią się od siebie rzeczywiste wartości rezystorów RA oraz im mniejsze różnice rzeczywistych wartości między RB, tym lepszy CMRR (większe tłumienie). Wskazują na to symulowane w LTspice XVII (dla NE5532) analizy n-krotnego ($n=1000$) pseudolosowego doboru rezystorów (o tolerancji 1%) metodami Monte-Carlo i Gaussa, których wynikowe $k = \sqrt{n/2} \approx 23$ klasowe histogramy liczebności widoczne są na **rysunku 8**. Na wykresach średnią oznaczono x, natomiast medianę oznaczono m. Wyniki uzyskane metodą **Rys. 7**



(a) NE5532 (Dotyczy Rys.2)
 RA=RB=47kΩ±1% K=6dB

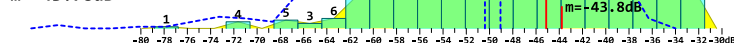
 CMRR @50Hz Gauss: n=1000 k=23
 CMRR_{MAX}=-89.93dB
 CMRR_{MIN}=-38.95dB
 R=50.98dB
 R/k=2.22dB
 x=-54.46dB
 m=-52.78dB

(b) NE5532 (Dotyczy Rys.2)
 RA=RB=47kΩ±1% K=6dB

 CMRR @50Hz Monte-Carlo: n=1000 k=23
 CMRR_{MAX}=-88.05dB
 CMRR_{MIN}=-35.40dB
 R=52.65dB
 R/k=2.29dB
 x=-49.91dB
 m=-48.06dB

(c) 2xNE5532 (Dotyczy Rys.4)
 RA=RB=47kΩ±1% K=6dB

 CMRR @50Hz Monte-Carlo: n=1000 k=23
 CMRR_{MAX}=-88.10dB
 CMRR_{MIN}=-35.50dB
 R=52.60dB
 R/k=2.29dB
 x=-49.67dB
 m=-47.79dB

(d) NE5532 (Dotyczy Rys.2)
 RA=12kΩ±1% RB=47kΩ±1% K=18dB

 CMRR @50Hz Monte-Carlo: n=1000 k=23
 CMRR_{MAX}=-77.47dB
 CMRR_{MIN}=-31.45dB
 R=46.02dB
 R/k=2.00dB
 x=-45.56dB
 m=-44.02dB

(e) 2xNE5532 (Dotyczy Rys.4)
 RA=12kΩ±1% RB=47kΩ±1% K=18dB

 CMRR @50Hz Monte-Carlo: n=1000 k=23
 CMRR_{MAX}=-78.84dB
 CMRR_{MIN}=-31.49dB
 R=47.35dB
 R/k=2.06dB
 x=-45.13dB
 m=-43.78dB


Rys. 8

CMRR w stosunku do wzmacniacza z jednym OA. Jak widać na (8d) i (8e), zwiększenie wzmocnienia w obu typach wzmacniacza do $K \approx 18\text{dB}$, przez zmianę wartości $RA=12\text{k}\Omega$, powoduje degradację CMRR o około 4dB ($1,6\times$) w stosunku do poprzednich symulowanych wariantów o wzmocnieniu $K = 6\text{dB}$. Prawdopodobieństwo wystąpienia najbardziej niekorzystnego wariantu rozrzutu wartości rezystorów 1%, przy którym to CMRR wynosi $\approx -34\text{dB}$ dla $K=6\text{dB}$ i $\approx -29\text{dB}$ dla $K=18\text{dB}$, jest znikome, choć nie niemożliwe. Przy zastosowaniu dobrej jakości 1% rezystorów metalizowanych, od renomowanego producenta z jednej serii produkcyjnej, osiągnięta wartość CMRR najczęściej będzie się plasować w zakresie najbardziej licznych klas histogramu. Dlatego też wydaje się, że w znacznej większości przypadków dla opisywanych wyżej układów, w zakresie wzmocnień $K=6\dots 18\text{dB}$ możliwe jest uzyskanie $\text{CMRR} \leq -40\text{dB}$ bez specjalnego dobiegania rezystorów 1%. Dla układu pracującego z sygnałem różnicowym o poziomie liniowym, 40dB tłumienie sygnału wspólnego jest zadowalającym rezultatem. Występujące na tle 1V sygnału różnicowego, 100mV sygnały wspólne (zakłócenia), na wyjściu wzmacniacza są tłumione do 1mV i stanowią znikomą wartość w stosunku do wzmocnionego do 2V, przy $K=6\text{dB}$ sygnału różnicowego. Jest to redukcja z 10% na wejściu do 0,05% na wyjściu.

Zwykle układy audio powinny być zasilane napięciami symetrycznymi, jednak w przypadku projektowania identycznego funkcjonalnie układu, który miałby być zasilany napięciem niesymetrycznym, obwód wytwarzający sztuczną masę powinien charakteryzować się niską impedancją (np. wtórnik zrealizowany na OA). Często stosowany do tego celu rezystorowy dzielnik napięcia z kondensatorem odsprężającym nie spełni dostatecznie zadania, ponieważ powoduje znaczną degradację tłumienia sygnałów wspólnych, tj. pogarsza CMRR.

Zależnie od zastosowania układu, zwykle będzie on wbudowany w urządzenie audio, używane wejścia/wyjścia układu należy połączyć przewodami sygnałowymi do współpracujących obwodów urządzenia oraz gniazd wejściowych/wyjściowych XLR, jack TRS, według oznaczeń na schemacie ideowym. Przy czym ekran przewodów powinien być tak połączony, by nie powstawały szkodliwe pętle masy. Oczywiście układ powinien mieć zapewnione symetryczne napięcia zasilania $\pm 15\text{V}$.

Zmontowany ze sprawnych elementów moduł powinien działać od pierwszego włączenia napięć zasilających. Układ nie wymaga czynności regulacyjnych, a jego sprawdzenie polega na kontroli sygnału na wyjściach (oscylloskop, końcówka mocy itp.) przy dołączonym do wejścia, symetrycznym lub niesymetrycznym sygnale audio.

Cyprian Kamil Kowalski
 c4v2@o2.pl

Literatura:

Balanced I/O <https://sound-au.com/articles/balanced-io.htm>
 Balanced Interfaces <https://sound-au.com/articles/balanced-2.htm>

Gaussa dla pojedynczego OA z rysunku 2, z rezystorami $RA=RB$ ($K=6\text{dB}$), widoczne są na histogramie (8a), gdzie w najgorszym przypadku minimalny $\text{CMRR} \approx -39\text{dB}$. Dla tego samego wzmacniacza wyniki uzyskane metodą Monte-Carlo są mniej optymistyczne (8b), gdzie minimalny $\text{CMRR} \approx -35\text{dB}$.

W celu porównania wyników w zależności od metody obliczeń, na kolejne wykresy uzyskane metodą Monte-Carlo naniosłem wielobok liczebności uzyskany metodą Gaussa. Histogram (8c) wykonany dla opisywanego podwójnego wzmacniacza różnicowego z rysunku 4 (dla jego pojedynczego wyjścia), przy identycznych jak poprzednio rezystorach (wzmocnieniu $K=6\text{dB}$), uwiadamia znikome, nieistotne pogorszenie

Wykaz elementów

R5,R15		100kΩ
R3,R7,R14,R18		12kΩ
R1,R2,R9,R10,R11,R12,R19,R20		47kΩ
R4,R6,R8,R13,R16,R17		100kΩ
C2,C5,C11,C14		47pF ceramiczny
C6,C8		100nF ceramiczny
C1,C4,C10,C13		2,2μF foliowy
C3*,C12*		10μF/25V bipolarny
C7,C9		100μF/25V
U1		NE5532, NE5532A, TL072
J1,J4		Złącze grzebieniowe 2,54mm 1×2 pin
J2,J5		Złącze grzebieniowe 2,54mm 1×4 pin
J3		Złącze grzebieniowe 2,54mm 1×3 pin
Z1,Z2		Złącze grzebieniowe 2,54mm 1×2 pin + Żwornik

Komplet podzespołów z płytą jest dostępny
 w Sklepie AVT jako zestaw AVT3298

Frezarka CNC

część 2

W artykule opisana jest prosta obrabiarka sterowana numerycznie, której wykonanie nie przekracza możliwości większości elektroników – hobbystów.

W pierwszej części artykułu przedstawione były ogólne informacje na temat frezarki CNC własnej konstrukcji. Oto dalszy opis.

Wrzeciono

Zadaniem wrzeciona CNC (*ang. CNC spindle*) jest usuwanie nadmiaru materiału za pomocą obracającego się frezu. Teoretycznie mogłoby się wydawać, że funkcję wrzeciona może pełnić np. zwykła wiertarka albo uniwersalne narzędzie popularnie zwane dremelkiem od nazwy jednego z producentów tego typu urządzeń. Wbrew pozorom, nie jest to dobre rozwiązanie. Moc narzędzia uniwersalnego nie przekracza zwykle 130W, co często jest wartością zbyt małą dla naszych potrzeb. Dodatkową wadą, zarówno narzędzia uniwersalnego, jak i wiertarki, jest fakt, że ich konstrukcja nie jest przeznaczona do przenoszenia obciążeń bocznych, które występują podczas frezowania. Nie są one również przeznaczone do długotrwałej pracy. Użycie do frezowania wiertarki czy narzędzia uniwersalnego powoduje ich szybkie niszczenie. Wrzeciono ma inny system łożyskowania niż wiertarka. Większa moc wrzeciona umożliwia osiągnięcie większych szybkości frezowania oraz frezowanie bardziej „opornych” materiałów. Zastosowane przez autora wrzeciono zasilane jest napięciem stałym 48V i ma moc 400W. Na górze wrzeciona umieszczony jest mały wentylator wymuszający przepływ powietrza wkoło wrzeciona, a tym samym poprawiający jego chłodzenie. Do niektórych wrzecion dostępny jest uźebrowanie poprawiające chłodzenie, pełniące funkcję radiatora. Prędkość wrzeciona zasilanego napięciem stałym można regulować za pomocą regulacji współczynnika wypełnienia (PWM) napięcia zasilającego wrzeciono. Szybkość obrotu wrzeciona musi być dostosowana do obrabianego materiału. Zbyt szybko obracające się wrzeciono pod-

czas obrabiania tworzyw sztucznych będzie powodować nadtapianie materiału, natomiast zbyt wolno obracające się wrzeciono wydłuży czas frezowania. W przypadku urządzeń profesjonalnych stosuje się wrzeciona zasilane napięciem przemiennym jednofazowym, a w przypadku wrzecion dużej mocy – trójfazowym. Ich prędkość reguluje się za pomocą falownika. Do redukcji prędkości obrotowej wrzeciona zasilanego prądem zmiennym nie stosuje się sterowania fazowego ze względu na fakt, że sterowane fazowo wrzeciono ze spadkiem prędkości coraz łatwiej jest zatrzymać – traci ono moment obrotowy. Zastosowanie falownika pozwala na zredukowanie prędkości pracy wrzeciona z zachowaniem jego momentu obrotowego. W przypadku dużej mocy wrzecion pracujących długotrwale stosuje się chłodzenie wodne. Kupując wrzeciono, najlepiej kierować się opiniami o nim na stronach poświęconych CNC. Kupując wyrób nieznanego producenta, można kupić albo wrzeciono bardzo słabe, albo całkiem przyzwoite, za rozsądną cenę – nie ma reguły. Do niektórych wrzecion dostępny jest odpowiedni uchwyt mocujący. Jeśli taka opcja jest przewidziana, to należy go zamówić w komplecie, gdyż bardzo ułatwia mocowanie wrzeciona. Na końcu wrzeciona znajduje się uchwyt na frez. Uchwyty na frez mają różne oznaczenia w zależności od swojej średnicy.

Oś Z

Zespół sterowania w osi Z jest najtrudniejszym do wykonania elementem opisanej frezarki. Gotową mechanikę osi Z można kupić na zagranicznych portalach aukcyjnych, w wyszukiwarkach należy wpisać frazę „Z axis”. Koszt kupna mechaniki osi Z będzie jednak wyraźnie większy niż wykonanie jej samodzielnie z zakupionych podzespołów. Zadaniem osi Z jest podnoszenie wrzeciona w górę i opuszczanie w dół. Pokazana na **fotografiach 1, 2** oś Z mogłaby być znacznie krótsza. Jest niepotrzebnie długa – przez co ma tendencję do uginania się, szczególnie przy większych prędkościach frezowania. Docelowo zostanie ona znacznie

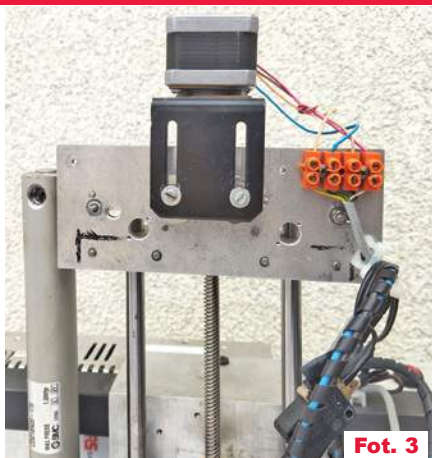


Fot. 1

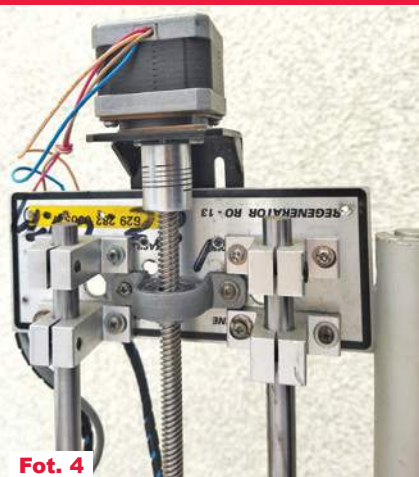


Fot. 2

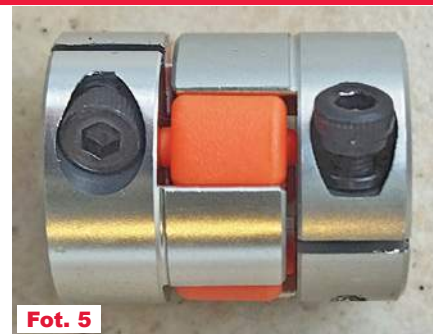
skrócona, co wymagać będzie również obniżenia bramy. Wykonując oś Z, wszystkie otwory mocujące w materiale należy wykonywać bardzo starannie, gdyż nawet niewielkie przesunięcia otworów powodują problemy w dokładnym spasowaniu elementów osi Z. Zastosowanie blach o kształcie pro-



Fot. 3



Fot. 4



Fot. 5

stokąta do poszczególnych bloków osi Z bardzo ułatwia nanoszenie pozycji otworów. Wymiary boków materiału należy kontrolować za pomocą suwmiarki i porównać je ze sobą. Linie wierceń na materiale najprościej zaznaczyć za pomocą suwmiarki z metalowymi szczękami. Jedną szczękę przykładamy do krawędzi materiału, drugą do płaszczyzny materiału i przesuwając ją względem krawędzi, zarysowujemy lekko materiał. W ten sposób uzyskamy linię równoległą do płaszczyzny materiału.

Autor przed wierceniem otworów montażowych nabijał je punktiakiem, następnie rozwiercał wiertłem o średnicy 2mm, a dopiero potem wiercił do właściwej średnicy za pomocą odpowiedniego wiertła. Wszystkie wiercenia wykonywano za pomocą wiertarki kolumnowej. W przypadku wykonywania otworów wstępnych bardzo ważne jest dobre oświetlenie stanowiska pracy. Przy odpowiednio dużym wgłębieniu w materiale wykonanym przez punktak i przy niskich obrotach wiertarki, przesuwając powoli materiał pod wiertłem, wiertło „samo” się naprowadzi na wgłębienie.

Oś Z zbudowana jest z trzech bloków przy wykorzystaniu prowadnic liniowych (linear shaft rod) w postaci dwóch niepodpartych wałków stalowych o średnicy 12mm. Wszystkie bloki osadzone są na profilach z blachy aluminiowej o grubości 3mm.

W pierwszym bloku – **fotografie 3 i 4**, na kawałku aluminium o kształcie prostokąta zamocowany jest silnik krokowy za pomocą uchwytu silnika. Mocowanie w oryginale wykonane jest za pomocą dwóch śrub i czeka na przeróbkę, polegającą na zastosowaniu dwóch dodatkowych śrub. Silnik zawiera amortyzator drgań. Do aluminium przymocowana jest kostka elektryczna, do której doprowadzone są sygnały sterujące silnikiem krokowym i włącznik krańcowy osi Z. Z lewej i prawej strony krótszego boku umieszczone są w jednej linii dwie podpory wałka, w odległości od siebie kilku cm. Rozwiązanie takie poprawia sztywność mechaniczną konstrukcji.

Napęd z silnika przenosi śruba trapezowa o średnicy 8mm. Śruba trapezowa sprzęgnięta jest z silnikiem za pomocą sprzęgła elastycznego, identycznie jak w przypadku osi X i Y. Lepszym, lecz nieco droższym rozwiązaniem byłoby zastosowanie w tym miejscu sprzęgła kłowego – **fotografia 5**. Śruba trapezowa podparta jest na początku i na końcu za pomocą łożyska samonastawnego w aluminiowej obudowie – KP08 8mm.

Drugi z bloków – **fotografie 6, 7**, zamocowany jest na profilu o kształcie litery L. Krótszy z boków profilu

bearing), przez które przechodzą prowadnice. Łożyska liniowe mają kształt walca i umieszczone są w obudowie zamkniętej, ułatwiającej ich mocowanie. Łożyska te zapobiegają możliwości wykonania ruchu obrotowego przez wrzeciono, dzięki nim ruch obrotowy silnika zamieniany jest w ruch liniowy.

Średnica łożysk równa jest średnicy wałków liniowych. Również w wypadku łożysk ustawiono dwa łożyska liniowe w jednej linii, co poprawiło sztywność całej konstrukcji. W przypadku stosowania wałków podpartych należy zastosować łożyska liniowe otwarte, a ich obudowy powinny być na tyle mocne, aby zapobiegać rozginaniu się. Do tego samego profilu przymocowana jest podpora końcowa wałka o średnicy 12mm, w której zamocowana

R E K L A M A

KEY

PRODUCENT AUTOMATYKI GRZEWCZEJ

11-200 Bartoszyce ul. Bohaterów Warszawy 67
tel. (89)7635050 fax (89)7635051

pwkey@onet.pl

TANIE REGULATORY

DO KOTŁÓW WĘGLOWYCH I NA DREWNO

z wbudowanym termostatem pokojowym
zapewniającym komfort i oszczędność



REGULATORY DO KOTŁÓW Z PODAJNIKIEM

REGULATORY POGODOWE

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

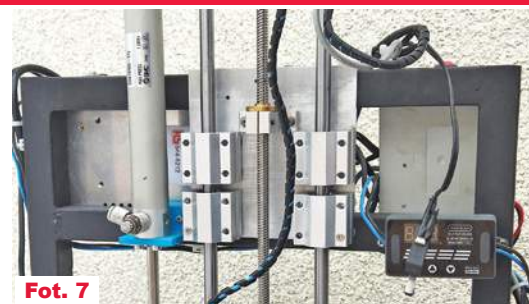
www.pwkey.pl

jest nakrętka trapezowa, przez którą przechodzi śruba trapezowa przenosząca napęd na oś Z. Znacznie lepszym rozwiązaniem, niż śruba trapezowa, jest zastosowanie śruby kulowej z nakrętką. Jest to jednak rozwiązanie stosowane w urządzeniach dużo wyższej klasy.

Trzeci blok – **fotografie 8, 9**, wykonany jest na kawałku aluminium o kształcie prostokąta. Umieszczone są na nim: dwie podpory wałka końcowego, szeregowo w niewielkiej odległości od siebie, analogicznie jak w przypadku pierwszego bloku, łożysko poparte podpierające śrubę trapezową, mocowanie wrzeciona i kostka elektryczna, do której doprowadzane jest napięcie zasilające wrzeciono. Wrzeciono stanowi najcięższy element osi Z i w przypadku braku wspomaganie pneumatycznego, po zabranii sygnałów sterujących pracą silnika, oś Z ma tendencję do opadania pod ciężarem wrzeciona. Aby zapobiec opadaniu wrzeciona pod własnym ciężarem, do ramy i bloku trzeciego przymocowany jest siłownik pneumatyczny, działający jako wspomaganie pneumatyczne. Siłownik pneumatyczny powoduje, że wrzeciono nie opada pod wpływem swego ciężaru w dół, a jednocześnie nie utrudnia w sposób istotny jego podnoszenia. Siłowniki pneumatyczne stosuje się np. w meblach do zapobiegania gwałtownemu opadaniu ciężkich elementów np. drzwiczek, leży łóżek czy w samochodach do zapobiegania gwałtownemu opadaniu klapy. Autor dobrał siłownik eksperymentalnie z kilku posiadanych egzemplarzy. Mocowanie siłownika wykonano na drukarce 3D, jednak można je również wykonać z blachy. W ostatecznej wersji wrzeciono zostanie opuszczone maksymalnie do dołu, blok 1 zostanie przymocowany do górnej poprzeczki ramy, blok 3 do dolnej poprzeczki ramy. Do bloku drugiego zostanie przymocowane wrzeciono, przy czym łożyska będą zamocowanie z odwrotnej strony bloku, a samo wrzeciono zostanie maksymalnie opuszczone. Rozwiązanie takie zredukuje znacznie możliwość podnoszenia osi Z, ale uczyni ją bardzo sztywną.



Fot. 6



Fot. 7

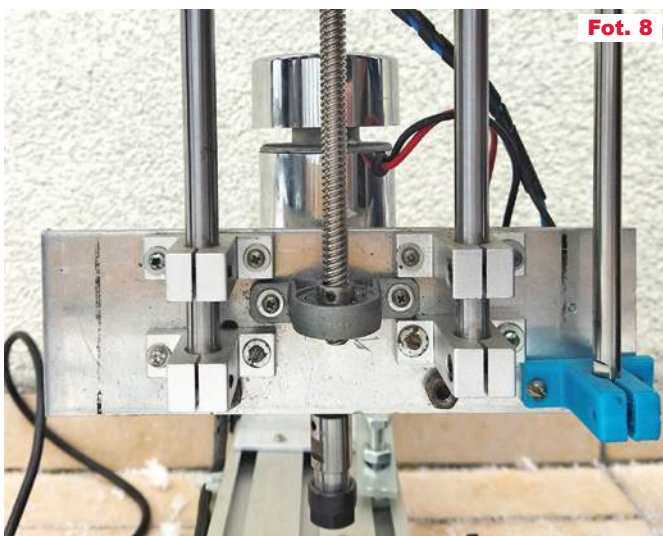
Aby usztywnić konstrukcję, można zastosować również większą liczbę przewodów liniowych.

Odsysacz

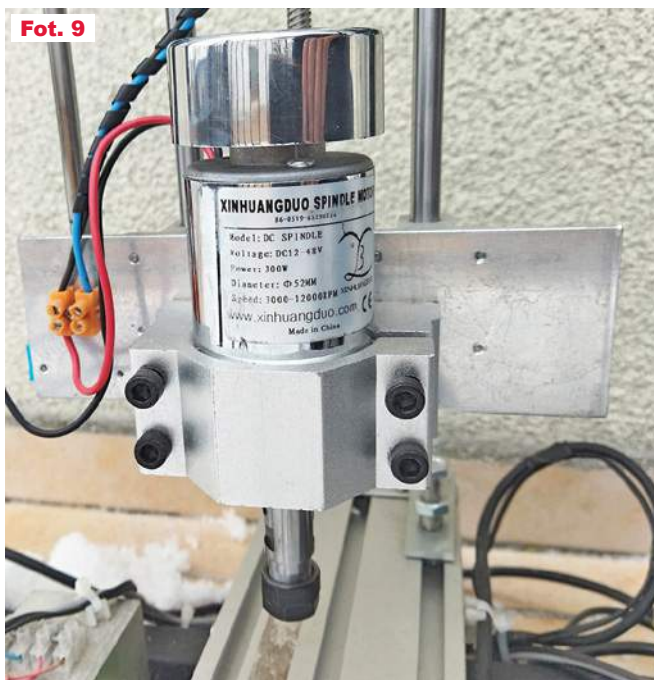
Wszelkie pyły, wiórki i skrawki obrabianego materiału powodują silne zanieczyszczenie miejsca pracy, a w przypadku dostania się na prowadnice czy do łożysk powodują zwiększenie oporów ruchu. Budując frezarkę, warto więc wykonać również odsysacz, który byłby podłączany do odkurzacza i zbierał zanieczyszczenia. W opisanym układzie funkcję odsysacza może pełnić zwykła ssawka odkurzacza przymocowana na stałe do ramy. W przypadku ruchomej osi Y odsysacz najłatwiej wykonać na drukarce 3D i przymocować do wrzeciona. Podczas pracy frezarki najlepiej umieścić w zamkniętym pomieszczeniu, gdzie nie będzie przeszkadzać powstający pył lub wykonać minibudkę, w której znajdzie się frezarka. Autor jako obudowę frezarki zastosował miniszklarnię kupioną na jednym z portali aukcyjnych, co było zdecydowanie najtańszym rozwiązaniem.

W następnym odcinku cyklu zostanie opisane podłączenie modułów elektronicznych. Na zakończenie artykułu autor chce podziękować Waldkowi 3Z6AEF za uwagi do tego tekstu. W materiałach dodatkowych umieszczonych na Elportalu jest szereg zdjęć w wysokiej rozdzielczości, pokazujących szczegóły wykonania układu.

Jerzy Wilczewski
Rafał Orodziński
sq4avs@gmail.com



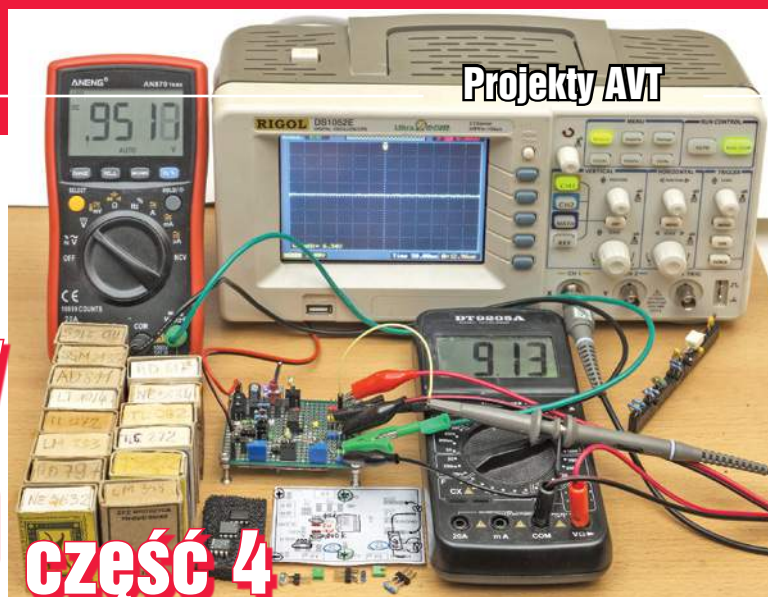
Fot. 8



Fot. 9

Miernik wzmacniaczy operacyjnych

część 4



Zgodnie z zapowiedzią, w czwartej części artykułu przedstawione są informacje, jak powstawały – schemat oraz mój przyrząd.

Rozważania projektowe

Pomysł, a właściwie potrzeba realizacji miernika parametrów wzmacniaczy operacyjnych wiąże się ściśle z cyklem *Kuchnia Konstruktora* – *Taki zwyczajny zasilacz*, a także z innymi układami planowanymi do publikacji w EdW. Otóż bardzo pożądane, a wręcz niezbędne okazało się sprawdzenie faktycznych parametrów niektórych posiadanych wzmacniaczy operacyjnych. Z uwagi na niską cenę zakupu na chińskim portalu zachodziła bowiem obawa, że ich rzeczywiste właściwości nie odpowiadają parametrom wynikającym z oznaczeń na obudowie.

W przypadku wzmacniaczy operacyjnych, przeznaczonych do wykorzystania w konstruowanym zasilaczu warsztatowym, podstawowym wymaganiem jest prawidłowa praca przy zasilaniu pojedynczym napięciem 5...5,5V i możliwość pracy wejść na poziomie ujemnego napięcia zasilania – masy. Aby umożliwić wykonanie jak najbardziej precyzyjnej wersji zasilacza, potrzebowałem wzmacniaczy operacyjnych o jak najmniejszym napięciu niezerównoważenia oraz jak najmniejszym dryfcie termicznym. Trzeba było to jakoś sprawdzić, czy posiadane wzmacniacze „trzymają parametry”.

Przeszukałem dostępne w Internecie materiały dotyczące pomiaru wzmacniaczy operacyjnych, w szczególności pomiaru napięcia niezerównoważenia i parametrów pokrewnych. Różnych takich materiałów jest sporo, jednak większość informacji powtarza się. Poza tym najczęściej są to rozważania

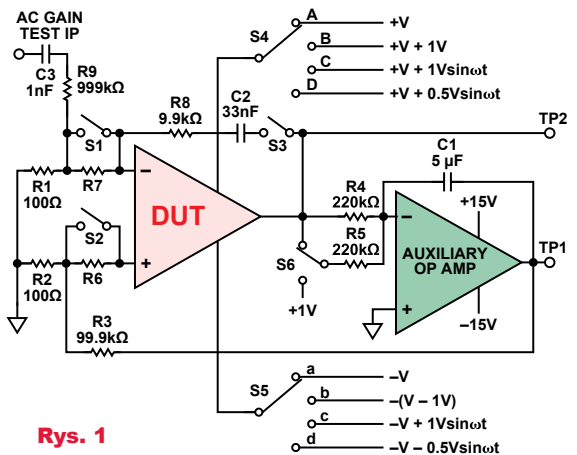
teoretyczne, dotyczące kilku standardowych metod i układów pomiarowych. Bardzo mało jest informacji o praktycznych rozwiązaniach oraz o różnych problemach, które trzeba rozwiązać przy realizacji tego rodzaju mierników.

Moją uwagę przykuł przede wszystkim artykuł *Simple Op Amp Measurements*, którego autorem jest James Bryant z Analog Devices. Artykuł jest dostępny na stronie: www.analog.com/en/analog-dialogue/articles/simple-op-amp-measurements.html

a wersja PDF pod adresem:

www.analog.com/media/en/analog-dialogue/volume-45/number-2/articles/simple-op-amp-measurements.pdf
w skrócie: <https://bit.ly/39YhJ3w>.

Bardzo obiecująco wygląda zamieszczony tam schemat, który w częściowo spolszczonej wersji pokazany jest na **rysunku 1**. Pozwala on zmierzyć szereg parametrów stałoprądowych i zmiennoprądowych wzmacniacza oznaczonego DUT (Device Under Test). W literaturze można też znaleźć schematy bez dodatkowego wzmacniacza pomocniczego oraz z dwoma wzmacniaczami pomocniczymi, na przykład według **rysunku 2**.



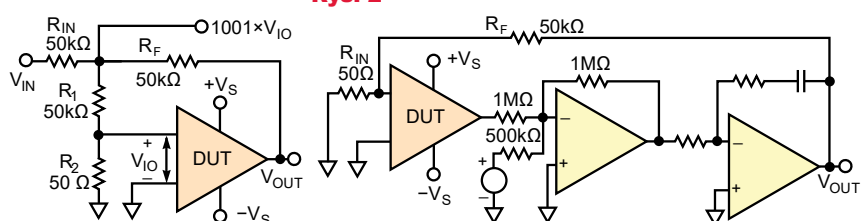
Rys. 1

	S1	S2	S3	S4	S5	S6
napięcie niezerównoważenia U_{OS}	1	1	0	A	a	0
prądy wejściowe I_B	0/1	0/1	0	A	a	0
wzmocnienie napięciowe DC	1	1	0	A	a	0/1
wzmocnienie napięciowe AC	1	1	0	A	a	0
tłumienie sygnału wspólnego DC CMRR	1	1	0	A/B	a/b	0
tłumienie tętnień zasilania DC PSRR	1	1	0	A/B	a/b	0
tłumienie sygnału wspólnego AC CMRR	1	1	1	C	c	0
tłumienie tętnień zasilania AC PSRR	1	1	1	D	d	0

Ja zdecydowałem się na koncepcję z rysunku 1. Oferuje ona też możliwość pomiaru prądów polaryzacji wejść przez rozwarcie styków S1, S2.

Gdy są one zwarte, mierzone jest napięcie niezerównoważenia wzmacniacza DUT. Rozwarcie jednego z tych styków spowoduje, że prąd wejściowy wywoła spadek napięcia na rezystancji R7 lub R6. Spadek ten doda się do zmierzonego wcześniej napięcia niezerównoważenia. Można w ten sposób zmierzyć prądy polaryzacji obu wejść (I_{B+} , I_{B-}), a także ich

Rys. 2



różnicę, czyli wejściowy prąd nierównoważenia (I_{OS}).

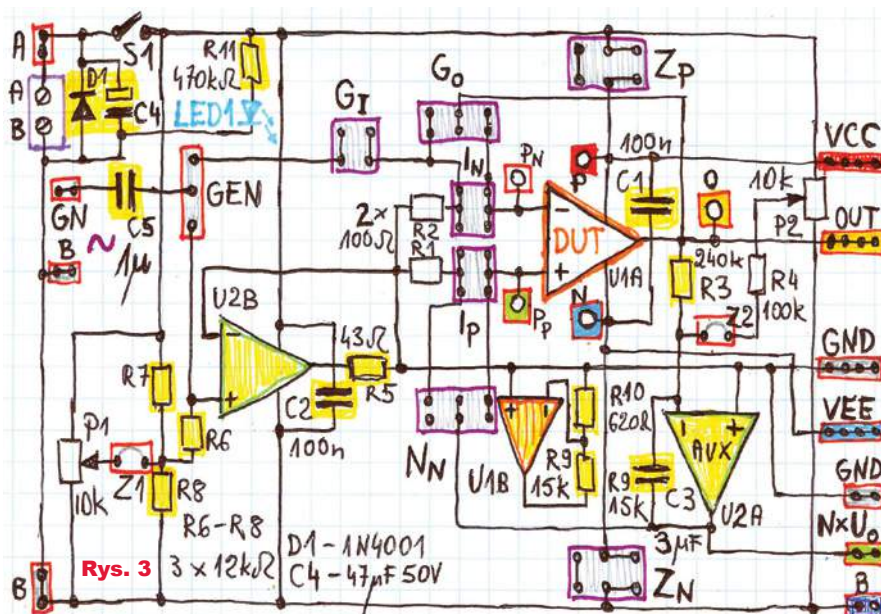
Dalsza analiza wymaga porównania rysunku 1 z moim schematem, przypominanym na **rysunku 3**.

W układzie z rysunku 1 napięcie wyjściowe można zmienić za pomocą S6 tylko o 1 volt. Ja zamiast S6 dałem zworkę (Z2) i chciałem uzyskać szerszy zakres zmian. Potrzebowałem też możliwości sprawdzenia zakresu użytecznych napięć **wyjściowych** przy różnym obciążeniu wyjścia. Dlatego dodałem w tym obwodzie potencjometr (P2), który pozwala dowolnie zmieniać napięcie na wyjściu wzmacniacza DUT. W układzie z rysunku 1 nie ma możliwości sprawdzania zakresu napięć **wyjściowych**, jakie może wytworzyć badany wzmacniacz. U mnie umożliwia to właśnie potencjometr P2.

Liczne współczesne wzmacniacze operacyjne mogą pracować przy napięciach **wejściowych** wykraczających poza napięcia zasilania i dobry układ testujący powinien mieć możliwość sprawdzenia także takiej nietypowej pracy.

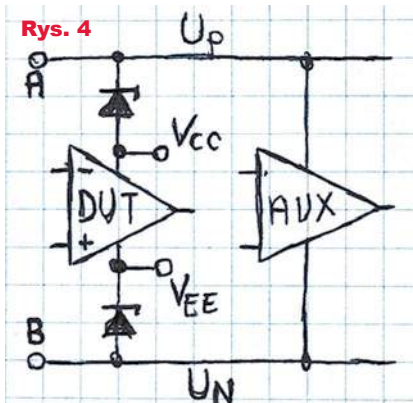
Na rysunku 1 nie są pokazane szczegółowo obwody zasilania badanego wzmacniacza DUT. Widać tylko, że DUT jest zasilany innym napięciem niż pomocniczy wzmacniacz AUX. I tak musi być. Dawniej wzmacniacze operacyjne zasilane były napięciem symetrycznym, najczęściej $\pm 15V$. Później napięciami niższymi, także niesymetrycznymi. Dziś wzmacniacze operacyjne zwykle są zasilane napięciem pojedynczym, a wiele ma maksymalne zalecane napięcie zasilania 5,5V i maksymalne dopuszczalne około 7V. Aby przyrząd pomiarowy był uniwersalny, trzeba przewidzieć różne możliwości zasilania, ale także konfiguracji masy

Miernik laboratoryjny mógłby być zasilany z kilku oddzielnych zasilaczy. Dla hobbysty lepszym, prostszym rozwiązaniem wydaje się wykorzystanie pojedynczego napięcia zasilania wzmacniacza DUT i płynnego regulowania potencjału (sztucznej) masy za pomocą potencjometru P1 względem zasilania. Aby to jednak działało prawidłowo, napięcie zasilania potencjometrów P1, P2 oraz wzmacniacza pomocniczego AUX powinno być wyższe niż napięcie zasilania badanego wzmacniacza DUT (między punktami VCC, VEE). Można



to zrobić na różne sposoby, na przykład zastosować oddzielne źródła zasilania, jednak najprostsze okazuje się zasilanie całości z jednego źródła. Można to łatwo zrealizować przez zastosowanie dwóch diod Zenera według idei z **rysunku 4**. Przy takim rozwiązaniu można też sprawdzić, czy wejścia wzmacniacza DUT mogą pracować w całym zakresie jego napięć zasilania ($V_{CC} - V_{EE}$), a nawet nieco poza nim, co w wielu wzmacniaczach jest możliwe. W moim rozwiązaniu funkcję diod Zenera pełnią wkładane w gniazda Z_P, Z_N diody LED, które jednocześnie są też dobrymi, wizualnymi wskaźnikami poboru prądu.

Dla rozszerzenia możliwości pomiarowych warto też przewidzieć możliwość pracy DUT w roli wzmacniacza sygnałów zmiennych. Potrzebne są do tego elementy sprzężenia zwrotnego w obwodzie wejścia odwracającego. W wersji z rysunku 1 przewidziane są do tego rezystory R8, R9 włączone w szereg z kondensatorami i wyłącznikiem S3. Ja zamiast tego przewidziałem gniazda G₁, G₀, w które można wetknąć potrzebne elementy sprzężenia zwrotnego, a dodatkowo wprowadziłem możliwość podania sygnału zmiennego na obwód sztucznej masy za pomocą przełącznika oznaczonego GEN. Na wejście oznaczone u mnie GN można podać przebieg zmienny z zewnętrznego generatora, co pozwoli zmierzyć kolejne pożyteczne parametry. Skonfigurowanie DUT do roli wzmacniacza sygnałów zmiennych pozwoli zmierzyć nie tylko zmiennoprądowe wzmocnienie



nie z otwartą pętlą, PSRR i CMRR, ale też umożliwia pomiar na wyjściu takich sygnałów zmiennych jak szumy. A to otwiera też drogę do pomiaru gęstości szumów napięciowych i prądowych. Pozwoli także zmierzyć szybkość zmian napięcia wyjściowego (SR) mniej szybkich wzmacniaczy.

Co ciekawe, w takim przyrządzie można zastosować rozwiązanie niedopuszczalne w typowych aplikacjach – mianowicie w pętli ujemnego sprzężenia zwrotnego można włączyć szeregowy kondensator, który na rysunku 1 oznaczony jest C2. Wyjście wzmacniacza DUT nie nasyci się, ponieważ nie dopuści do tego wzmacniacz AUX, który nietypowo realizuje stałoprądową pętlę ujemnego sprzężenia zwrotnego.

Możliwości zmian

Podane właśnie informacje pokazują w dużym skrócie tylko kluczowe etapy projektowania i realizacji mier-



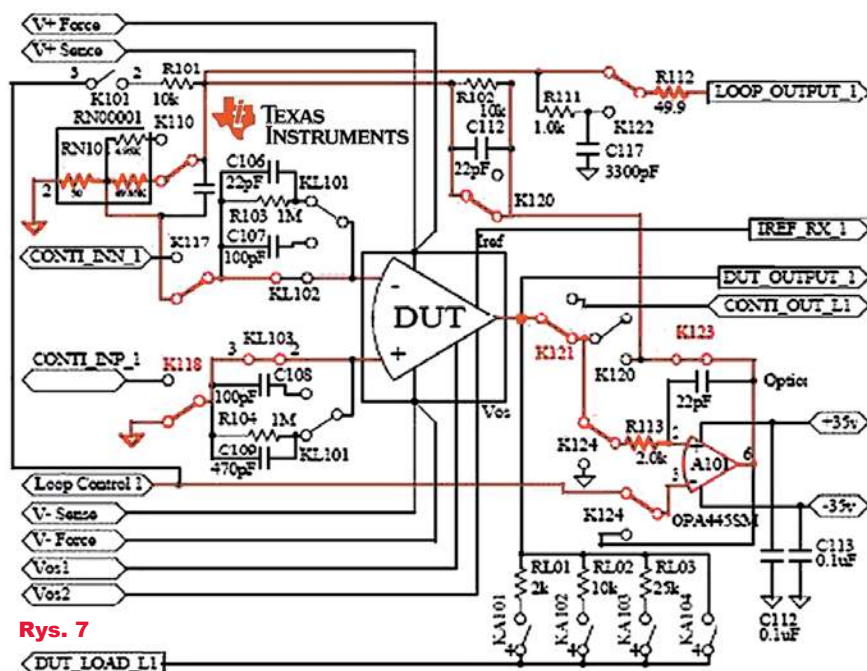
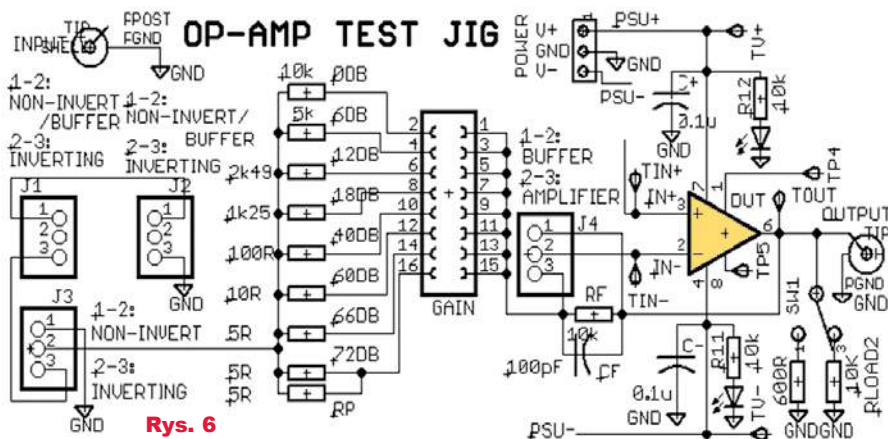
nika oraz przeprowadzania pomiarów. Omawiamy te szczegóły nie tylko dlatego, żeby zaspokoić ciekawość, ale także dlatego, że niektórzy Czytelnicy zapewne nie potrzebują mierzyć wszystkich wymienionych parametrów i zechcą zrealizować wersję uproszczoną przyrządu, mierząc jedynie niektóre parametry, na przykład tylko napięcie niezrównoważenia. Układ do pomiaru wzmacniaczy operacyjnych można bowiem zrealizować na wiele sposobów.

Na stronie <http://fivefishaudio.com/diy/opampstestjig/> za 17 dolarów oferowany jest kit niepomiarowo prostszego układu do pomiaru tylko nielicznych parametrów. Na fotografii 5 pokazana jest płytka pomiarowa z wetkniętym w podstawkę pomiarową wzmacniaczem operacyjnym z elementami dyskretnymi, wskazanym czerwoną strzałką (niektórzy uważają, że takie zamienniki scalonych wzmacniaczy operacyjnych dają w układach audio lepszy dźwięk). Schemat układu pomiarowego pokazany jest na rysunku 6.

Oczywiście opisywany w artykule miernik wzmacniaczy można rozbudować i zautomatyzować, stosując mikroprocesor z przetwornikami ADC i DAC, wyświetlacz oraz szereg przekaźników. Wtedy naciśnięcie przycisku lub dotykowego ekranu skonfiguruje system do konkretnego pomiaru, a na wyświetlaczu od razu będzie można odczytać wynik pomiaru wybranego parametru.

W materiałach TI można znaleźć dość obszerny opis takiego miernika, którego kluczowa część schematu pokazana jest na rysunku 7.

Ja wybrałem koncepcję niejako przeciwną: prosty schemat z rysunku 1 stał się bazą i punktem wyjścia. Dodałem obwody rozszerzające funkcjonalność, jednocześnie maksymalnie upraszczając układ. Nie była to



szybka i jednorazowa decyzja. Najpierw na papierze powstało kilkanaście różnych wersji schematu, a potem po wybraniu prostej wersji, podczas budowy i testów modelu konieczne lub pożądane okazały się dalsze modyfikacje, które doprowadziły do modelu według rysunku 3.

Rozumiejąc działanie całości, każdy Czytelnik może system uproszczyć albo rozbudować i zrealizować tego rodzaju miernik param-

trów wzmacniaczy operacyjnych inaczej, zgodnie ze swoimi upodobaniami oraz potrzebami.

Piotr Górecki

Certyfikat Underwriters Laboratories

94V-0 E480148 TYPE 1

Zakład produkcyjny:

05-260 Marki ul. Duża 1
tel. 22 781 63 95
22 761 95 80
fax. 22 781 63 95 w 23
www.elmax.waw.pl
elmax@elmax.waw.pl

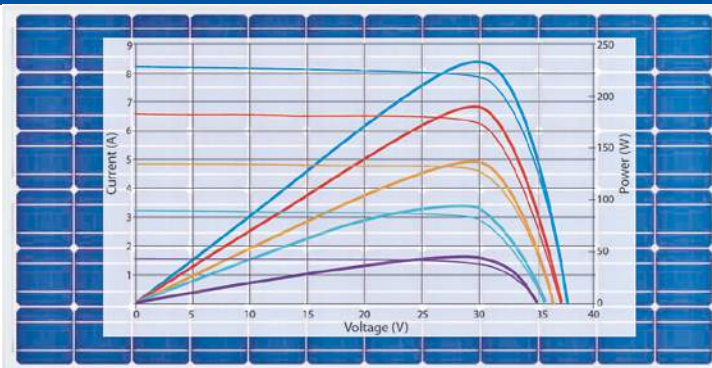
OBWODY DRUKOWANE

Produkcja, Projektowanie, Montaż

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

MPPT

część 18



Zgodnie z zapowiedzią, wracamy do kwestii sprawności energetycznej.

Straty w przewodach

Jeżeli zależy nam na wykorzystaniu całej energii z panelu FV, musimy też minimalizować straty w przewodach. Hobbysta, nie mając doświadczenia i przyjmując błędną koncepcję, niepotrzebnie zwiększy straty. Na przykład popularne są akumulatory 12-woltowe i popularne są panele FV tzw. 12-woltowe. Hobbyście 12-woltowy system FV może się więc wydawać czymś wręcz oczywistym, optymalnym.

Zacznijmy od przypomnienia: bardzo rzadko spotykane są instalacje z panelami o napięciu niższym od napięcia akumulatora i nie jest to przypadek. Zdecydowanie korzystniejsze jest wykorzystanie paneli o napięciu wyższym niż ma 12-woltowy akumulator.

Prąd z umieszczonych na zewnątrz paneli trzeba przesłać przewodami do kontrolera i akumulatorów (albo do inwertera), które zapewne umieszczone są w pomieszczeniu. A czym większy prąd, tym większe są straty w przewodach ($P_{str} = I^2 R$). Przy danej mocy paneli, czym niższe napięcie, tym większy prąd, bo przecież $P = U \cdot I$.

W instalacjach 12-woltowych o mocach rzędu kilkuset watów i większych prądy byłyby ogromne. Przykładowo instalacja z panelami „12-woltowymi” zawierająca kilka metrów kwadratowych paneli o mocy elektrycznej 1000 watów i o napięciu roboczym 19 woltów musiałaby dostarczać do (potężnego zestawu) akumulatorów prąd ponad 52A!

Podana moc nominalna (szczytowa) paneli 1000W nie powinna szokować, ponieważ dziś powszechnie mamy do czynienia z domowymi instalacjami FV o mocach kilku kilowatów. Tylko nie są to instalacje wyspowe *off-grid* z akumulatorami, tylko instalacje *on-grid*, wykorzystujące inwerter i sieć 230V w roli akumulatora, co nie zmienia podstawowych zasad.

Wracamy do strat w przewodach: popularne przewody miedziane do domowych instalacji elektrycznych o przekroju 1,5mm² i 2,5mm² w ogóle nie nadają się do pracy przy prądzie rzędu 50A, bo stałyby się bardzo gorące i izolacja z tworzywa po prostu by się stopiła. Stopienie izolacji nie powinno nastąpić na wolnym powietrzu w pojedynczych przewodach o przekroju 4mm² i grubszych.

Akumulatory zwykle są umieszczone wewnątrz budynku i są oddalone od paneli o kilka, a nawet o kilkadziesiąt metrów. Załóżmy, że odległość wynosi 20 metrów, co daje sumaryczną długość przewodów 40m. Przewód o przekroju 4mm² ma rezystancję około 5 omów na kilometr, więc 40 metrów będzie mieć rezystancję 0,2 oma. Niewiele, jak dla elektronika, ale prąd 52A płynąc przez nią, wywoła spadek napięcia...

tak: 10,4 wolta, a napięcie z paneli to 19V. Ponad połowę mocy stracilibyśmy w przewodach!

Do obliczeń wzięliśmy gruby jak dla elektronika przewód o przekroju 4mm². Jednak według tabel „elektrycznych” dla prądu 52A potrzebny jest miedziany przewód o przekroju co najmniej 10mm². Ma on rezystancję około 2Ω/km, co dla długości 40m metrów daje 0,08 oma. Przy prądzie 52A spadek napięcia na nim wyniesie ponad 4V. W energetycznej instalacji 230V zmniejszenie napięcia o 4V to tylko redukcja o 1,7%. Ale przy napięciu z panelu 19V, strata 4 woltów to redukcja mocy o 21%, czyli nadal straty będą duże. W takiej instalacji FV trzeba byłoby zastosować kosztowny miedziany przewód o niewyobrażalnie dużym dla elektronika przekroju 25mm² albo 35mm².

Problem się zmniejsza, jeśli panele mają wyższe napięcie. Oczywiście w systemie mogą pracować panele o niskim napięciu, ale połączone w szereg, by napięcie przesyłane do kontrolera, zawierającego przetwornicę indukcyjną, było możliwie duże. Tak,

ale ze względów bezpieczeństwa nie może być jednak zbyt duże i w praktyce wynosi kilkadziesiąt woltów.

Nawet przy napięciu paneli rzędu 70 woltów trzeba brać pod uwagę spadki napięć i straty mocy w rezystancji przewodów.

I tu wracamy do przetwornic indukcyjnych i MPPT. Wcześniej analizowaliśmy sprawność systemu FV z 12-woltowym akumulatorem i „12-woltowymi” panelami. Zastosowanie paneli o znacznie wyższym napięciu, przetwornicy indukcyjnej i zestawu akumulatorów 12- albo 24-woltowych daje szansę na zmniejszenie strat w przewodach, a więc zwiększenie całkowitej sprawności systemu FV. W każdym razie w instalacjach FV trzeba stosować przewody grubsze, niż wymagają tego reguły „energetyki 230V”. Widzimy też potrzebę, a wręcz konieczność wykorzystania przetwornic indukcyjnych i to o jak najwyższej sprawności.

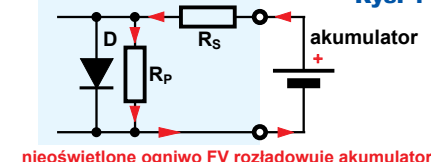
W szczególności nie będziemy się zagłębiać – podane informacje mają tylko uczulić każdego, kto chce zająć się tematyką MPPT, by starannie przeanalizował wszystkie wchodzące w grę czynniki.

A teraz kolejny zapowiadany temat.

Diody w systemie FV

W dotychczasowych rozważaniach wielokrotnie podkreślaliśmy konieczność zapobiegania „nocnemu cofaniu prądu”. Przypomnijmy, że fotodiody w panelach FV mają dużą powierzchnię, a to oznacza znaczne prądy upływu przez zastępczą rezystancję R_p pokazaną na **rysunku 1**. W sumie chodzi o to, żeby akumulator w systemie FV nie był rozładowywany nocą prądem płynącym przez rezystancję upływu fotodiod. Może temu zapobiec dioda lub tranzystor MOSFET.

Rys. 1



Jednak w panelach FV stosowane są też inne diody, zwane „diodami bajpasowymi”, obejściowymi (bypass diodes). W trzecim odcinku niniejszego cyklu wspomnieliśmy o diodach montowanych w panelach. Nie są to diody blokujące, tylko właśnie „bajpasowe”.

W wielu źródłach informacje uzasadniające potrzebę, a wręcz konieczność ich stosowania nie są jasne. W uzasadnieniach mówi się o zapobieganiu powstawania „hotspotów”, czyli gorących punktów i wiąże się to z problemem cieni lub zanieczyszczeń poszczególnych ogniw panelu, co rzeczywiście prowadzi do zmniejszania możliwości wytwarzania energii przez poszczególne ogniwa i cały panel. A to osobom mniej zorientowanym nasuwa myśl, że „diodы bajpasowe” przy niesprawności niektórych ogniw umożliwiają prawidłową pracę pozostałych.

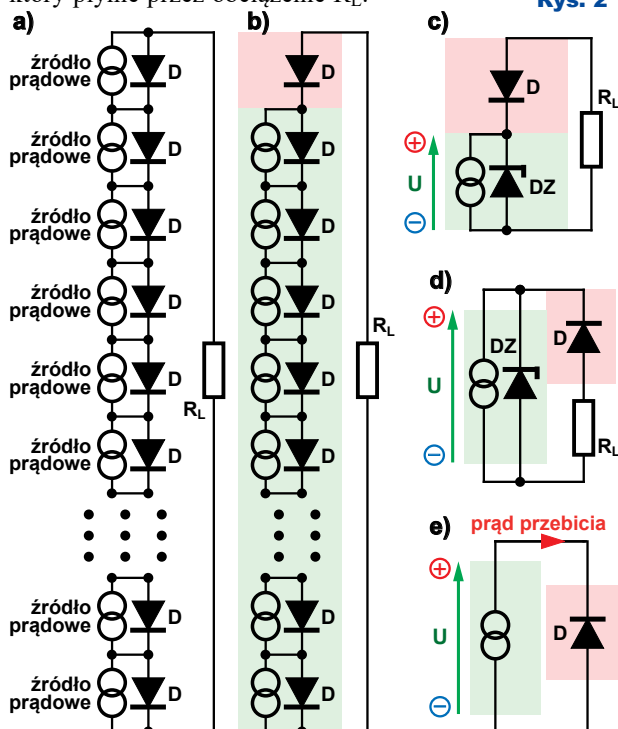
Po części słusznie, ale podstawowa prawda jest następująca:

brak diod „bajpasowych” może doprowadzić do pożaru!

Jest to zilustrowane między innymi na stronie: www.hioki.com/en/products/detail/?product_key=6415

w skrócie: <https://bit.ly/2NhjVgD>

Aby zrozumieć istotę problemu, przypomnijmy, że w panelu mamy wiele fotodiod połączonych szeregowo, a ich działanie w pewnym uproszczeniu pokazuje **rysunek 2a**. W idealnym przypadku wszystkie te fotodiody składowe zgodnie wytwarzają prąd elektryczny, który płynie przez obciążenie R_L .



Rys. 2

Problem powstaje, gdy fotodiody działają w niejednakowym stopniu. Dla uproszczenia założmy, że tylko jedno ogniwo, jedna fotodioda została całkowicie zasłonięta. Nie wytwarza więc prądu, a w układzie jest reprezentowana jako pojedyncza zwyczajna dioda. Na **rysunku 2b** została wyróżniona różową podkładką. Ponieważ pozostałe fotodiody próbują nadal zgodnie pracować, wytwarzają znaczne napięcie U , co można zilustrować jak na **rysunku 2c**. Można też układ przerysować jak na **rysunku 2d**.

Sprawne fotoogniwa wytwarzają prąd, ale ten prąd nie może płynąć przez obciążenie R_L , bo przeszkadza temu zasłonięta fotodioda, która jak widać, jest włączona w kierunku zaporowym. Prąd wytwarzany przez sprawne fotodiody marnuje się wewnątrz panelu, płynąc przez struktury diodowe poszczególnych ogniw, co reprezentowane jest przez diodę Zenera.

Jak już wiemy, takie „marnowanie prądu” nie grozi niczym złym. Tak zachowuje się w pełni sprawne ogniwo bez obciążenia. Problem dotyczy natomiast niesprawnego ogniwa, fotodiody, która nie wytwarza prądu. Dioda ta jest włączona zaporowo i powinna blokować przepływ prądu przez rezystancję obciążenia. Owszem, tylko jest pewien problem. Otóż okazuje się, że fotodiody z ogniw FV z kilku względów mają dość niskie napięcie przebicia. Jeżeli więc panel zawiera wiele fotodiod i wytwarza wysokie napięcie, to nastąpi przebicie nieoświeconej fotodiody, co w uproszczeniu pokazuje **rysunek 2e**.

Okazuje się, że napięcie przebicia fotodiod w ogniwach słonecznych jest rzędu kilkunastu woltów, często około 13V. A panele są różne, niektóre wytwarzają napięcia rzędu 40V i więcej.

Gdy nastąpi przebicie takiej nieczynnej fotodiody według **rysunku 2e**, popłynie przez nią duży prąd, wystąpi na niej znaczne napięcie, właśnie rzędu kilkunastu woltów, a to spowoduje wydzielanie się w tej fotodiodzie dużej ilości ciepła – powstanie „gorący punkt” – *hot spot*.

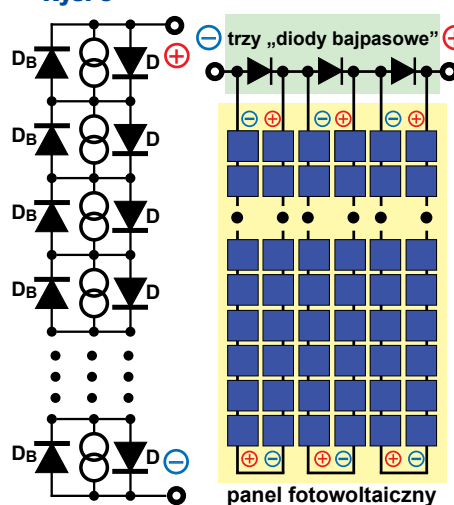
Duże panele FV wytwarzają prądy rzędu kilkunastu amperów lub więcej, co pomnożone przez kilkanaście woltów daje moc strat rzędu 200 watów lub więcej. Taka ilość ciepła silnie rozgrzeje nieczynne ogniwo i nie tylko może je uszkodzić: może spowodować pożar!

Można temu zapobiec w prosty sposób. Także z innych, nieomówionych względów najlepszym sposobem byłoby włączenie równoległe do każdej z fotodiod „diodы obejściowej”, jak widać z lewej strony **rysunku 3**. Najlepiej gdyby to były diody o jak najmniejszym napięciu przewodzenia, w praktyce diody Schottky'ego.

Jednak nie robi się tego, bo byłoby to zbyt kosztowne: każda taka dioda musiałaby mieć prąd pracy równy maksymalnemu prądowi wytwarzanemu przez panel. W praktyce stosuje się rozwiązanie uproszczone. Otóż napięcie przebicia elementarnej fotodiody to kilkanaście woltów. Jedna fotodioda może wytworzyć napięcie 0,5...0,6V. Dwadzieścia szeregowo połączonych fotodiod da napięcie co najwyżej 12V, bliskie napięciu przebicia jednej takiej fotodiody. Wystarczy więc „długi” panel FV podzielić na sekcje po kilkanaście ogniw i między te sekcje włączyć diody obejściowe, jak pokazuje prawa część **rysunku 3**. Muszą to być diody o odpowiednim prądzie i mocy strat albo specjalizowane diodowe układy scalone (np. Smart Bypass Diode SM74611).

Kończymy cykl dotyczący technicznych aspektów wykorzystania paneli fotowoltaicznych. Zupełnie inną kwestią są aspekty ekonomiczne oraz związane z tym nadużycia oraz wprowadzanie klientów w błąd. To jednak odrębny temat, który nie mieści się w ramach tego cyklu.

Rys. 3



Piotr Górecki

panel fotowoltaiczny

Transmisja danych w inteligentnym domu

6. MODBUS – bardzo popularny protokół

Wielu elektroników chciałoby samodzielnie zbudować elementy lub cały system inteligentnego domu. Jednym z kluczowych problemów okazuje się dobranie odpowiednich sposobów transmisji danych. Wybory w tym zakresie decydują o finalnym sukcesie lub porażce. W cyklu artykułów omawiamy rozmaite aspekty tego bardzo ważnego zagadnienia.

W poprzednim odcinku dość obszernie omawialiśmy łącze RS-485. Trzeba uściślić, że RS-485, czyli Recommended Standard 485, określa warstwę fizyczną, a więc „kwestie drucikowe”. W tym odcinku chcemy omówić MODBUS, czyli protokół, a więc zbiór reguł transmisji, który nieprzypadkowo wielu osobom nierozłącznie kojarzy się z RS-485. Protokół ten powstał pod koniec lat 70. w firmie Modicon do obsługi urządzeń przemysłowych, a przez obecnego właściciela Schneider Electric jest bezpłatnie udostępniony dla wszystkich chętnych.

Konieczne trzeba podkreślić, że MODBUS, do tej pory słusznie zresztą kojarzony z RS-485, może także wykorzystywać łącze RS-232 oraz internetowe łącza TCP/IP. Dlatego omawianie zaczniemy od rysunku 1, gdzie mamy pełniący funkcję zarządcy (*master*) komputer z łączem RS-232 i dołączoną „czarną skrzynką” oznaczoną *slave*.



Rys. 1

Na razie nie jest istotne, jakie funkcje pełni ta czarna skrzynka slave. Ważne jest tylko, jak kontaktujemy się z nią za pomocą protokołu MODBUS w trybie master-slave. A to oznacza, że wymianę danych zawsze inicjuje master. Urządzenie slave nie ma takiej możliwości – ono zawsze wykonuje jedynie rozkazy mastera. Ma to zawsze postać: zapytanie – odpowiedź, a więc według koncepcji klient – serwer, master jest klientem, a slave jest serwerem świadczącym usługi.

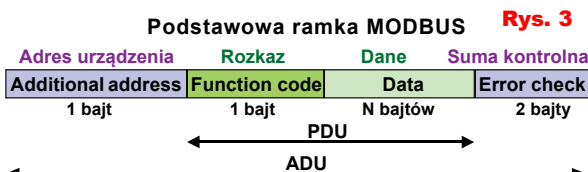
W przypadku z rysunku 1 nie byłby potrzebny żaden adres urządzenia slave, bo łącze RS-232 może obsłużyć tylko jedno takie urządzenie. Teoretycznie w takim przypadku master mógłby wysłać tylko rozkaz i ewentualnie

jakiś „informacje dodatkowe”. Rozkaz może być „jednokierunkowy” w tym sensie, że master nakazuje urządzeniu slave wykonanie jakiegoś prostego zadania. Ale rozkaz może być „dwukierunkowy”, gdy nakazuje urządzeniu slave odesłanie jakichś danych, do mastera.

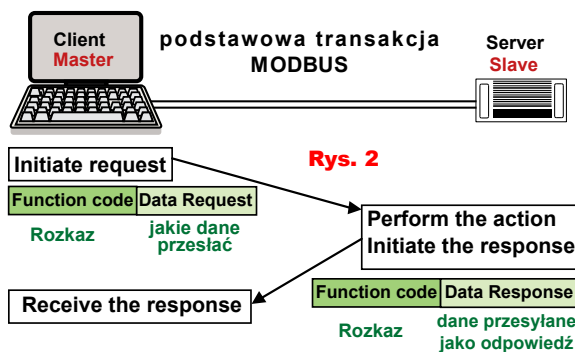
W systemie MODBUS nawet przy rozkazach „jednokierunkowych”, komunikacja zawsze polega na tym, że master wysyła rozkaz, a slave potwierdza wykonanie rozkazu i ewentualnie dostarcza informacje, których zażądał master. Rysunek 2 pokazuje w uproszczeniu podstawową, typową transakcję wymiany „podstawowych danych”, które nazywane są PDU (Protocol Data Unit). PDU zawsze zawiera rozkaz (Function code) i jakieś dodatkowe informacje. W przypadku łącza RS-232 to wystarczy. Ale protokół MODBUS może wykorzystywać łącze RS-485 lub TCP/IP, a w systemie może być wiele różnych urządzeń slave (od jednego do 247 urządzeń Slave o adresach 1...247).

Dlatego w systemach z łączem RS-485 oprócz „podstawowych danych”, czyli PDU, na początku wysyłany jest zawsze także adres urządzenia slave (jeden bajt) oraz pozwalająca wykryć błędy transmisji suma kontrolna (dwa bajty).

Rysunek 3 pokazuje „podstawową paczkę danych”, zwaną ADU, przy wykorzystaniu RS-485. Do realizacji urządzenia slave MODBUS z RS-485 wystarczy nawet mały mikrokontroler z kilkoma kilobajtami pamięci programu.

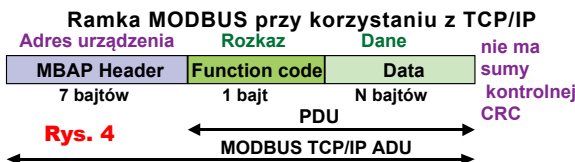


Rys. 3



Rys. 2

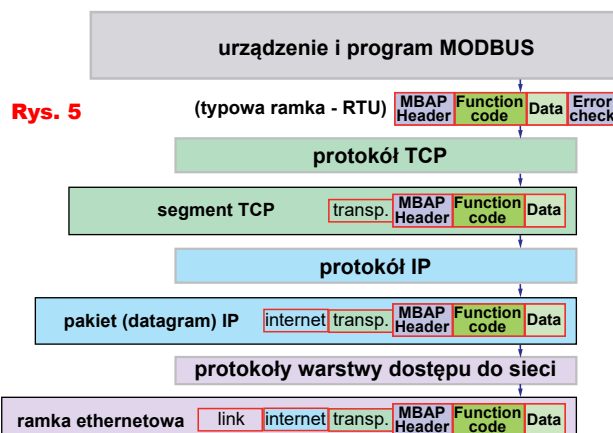
Przy wykorzystaniu TCP/IP paczka ADU ma nieco inną strukturę, pokazaną na rysunku 4.



Rys. 4

Nagłówek MBAP zawiera więcej informacji – 7 bajtów. Nie ma sumy kontrolnej, ponieważ kontrolę błędów zrealizuje, najprościej biorąc, protokół TCP/IP.

I taka właśnie zawartość jest enkapsulowana w ramki TCP i w kolejne warstwy stosu według rysunku 5. Wersja MODBUS over TCP/IP (z zarejestrowanym stałym numerem portu 502) jest coraz częściej wykorzystywana. Jeśli chcesz, poszukaj dodatkowych informacji samodzielnie, jednak zwykle nie jest to potrzebne, bo sieć internetową wykorzystuje się w typowy sposób. Znacznie ważniejsze jest co innego.



Rys. 5

Wymiana informacji

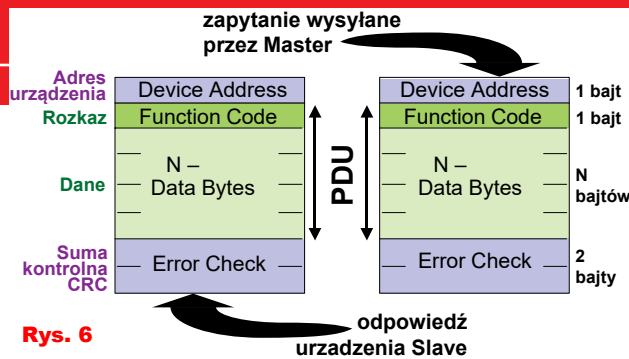
Przeanalizujmy podstawy, opierając się na rysunkach 2, 3 oraz **rysunku 6**, który pokazuje, że zapytanie i odpowiedź mają taką samą strukturę.

Mamy cztery podstawowe „składniki”: adres urządzenia slave, rozkaz, jakieś dane oraz sumę kontrolną. Adres zawsze zajmuje jeden bajt, rozkaz też jeden bajt, dane – różnie, nawet do 252 bajtów, a suma kontrolna – dwa bajty. Osobom samodzielnie realizującym i wykorzystującym elementy inteligentnego domu opierające się na protokole MODBUS, nie są potrzebne wszystkie szczegóły. Zgodnie z rysunkiem 2 i 6, master musi wysłać polecenie, zwane czasem *telegramem* albo nazywane *query*, czyli raczej zapytaniem. Na początku oczywiście ma być adres urządzenia slave, potem kod rozkazu, następnie związane z tym rozkazem dane. W praktyce nie trzeba martwić się o sumę kontrolną, bo oblicza ją program bez udziału użytkownika.

W ten sposób zapytane urządzenie slave odsyła swoją odpowiedź (Response – potwierdzenie) w bardzo podobnej postaci. Mianowicie w pierwszym bajcie wysyła swój adres, numer wykonanego rozkazu i jakieś „informacje dodatkowe”. Nawet gdy rozkaz nie dotyczy przesłania danych ze slave, master otrzymawszy takie „echo” wysłanego rozkazu, wie, że polecenie zostało zrealizowane. Często jednak slave przesyła do mastera jakieś dane, o czym za chwilę. W ramach zapoznawania się jedynie z podstawami protokołu MODBUS nie rozpatrujemy sytuacji, gdy nastąpi jakiś błąd i nie wgłębiamy się w kwestie sumy kontrolnej. Hobbysta interesujący się inteligentnym domem powinien tylko wiedzieć, jaka jest struktura komunikatów – telegramów MODBUS. Może przyjąć, że oprogramowanie w modułach master i slave potrafi poradzić sobie z błędami.

Adresy urządzeń slave

Adres jest 1-bajtowy, więc do dyspozycji jest 256 adresów, ale adres 0 nie jest przydzielany, tylko oznacza, iż następujący dalej rozkaz jest skierowany do wszystkich urządzeń slave w systemie. Zasadniczo adresy 248...255 są zarezerwowane, więc urządzenia slave w danym systemie mogą mieć adresy 1...247. Urządzenie master nie ma adresu.



Rys. 6

Powinniśmy teraz omówić rozkazy, ale lepiej zacznijmy od określenia, czym jest urządzenie slave.

Slave i rodzaje rejestrów

Ogólnie biorąc, slave może być niemal dowolnym, bardziej czy mniej skomplikowanym urządzeniem wykonawczym lub czujnikiem.

W latach 70. w urządzeniach automatyki przemysłowej powszechnie stosowano przyciski, przełączniki oraz przekładniki. Rozkaz mastera wysłany do slave mógł polegać na włączeniu przekładnika (o określonym numerze). Ponieważ rozkaz włączenia jest w pewnym sensie „chwilowy”, więc w urządzeniu slave dla każdego przekładnika musiał istnieć jakiś (jednobitowy) element pamiętający – komórka pamięci, inaczej *jednobitowy rejestr*. Urządzenia slave zawierały więc szereg jednobitowych rejestrów, sterujących cewkami przekładników (a cewka po angielsku to *Coil*). Do takiej komórki – rejestru master może wpisać jednobitową informację (stan przekładnika: 1 – włączony, 0 – wyłączony), ale master może też wysłać rozkaz odczytu takiego rejestru, by poznać aktualny stan danego przekładnika.

Master mógł też zająć, żeby urządzenie slave przysłało mu informacje o stanie, ale nie przekładnika, tylko jakiegoś przełącznika czy przycisku (o określonym numerze). Zasadniczo tu niepotrzebne są rejestry, czyli elementy pamiętające, bo stan przycisku można odczytywać „na żywo”. Jednak możemy i powinniśmy sobie wyobrazić, że w urządzeniu slave mamy też szereg jednobitowych rejestrów „pojedynczych wejść” (*Discrete Inputs*), ale w przeciwieństwie do przekładników i ich cewek (*Coils*), stan tych rejestrów można tylko odczytywać.

Dziś nie jest ważne, czy wszystko to ma jeszcze jakkolwiek związek z przekładnikami i przyciskami. Ważne jest, że za pomocą protokołu MOD-

BUS, master może obsługiwać cztery rodzaje rejestrów zawartych w urządzeniach slave:

A – jednobitowe rejestry „przekładnikowe”, które można zapisywać i odczytywać, nazywane *Coils*, co dosłownie oznacza cewki (przekładniki).

B – jednobitowe rejestry „przycisków”, nazywane *Discrete Inputs*, które można tylko odczytywać.

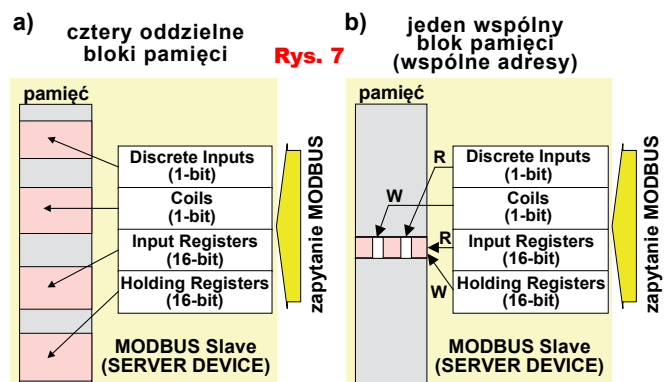
C – 16-bitowe rejestry nazywane *Input registers*, które można tylko odczytywać. Możemy sobie wyobrazić, że na przykład w rejestrach tych otrzymujemy wyniki pomiarów z przetworników ADC.

D – 16-bitowe rejestry, które można i zapisywać, i odczytywać, nazywane *Holding Registers*. Mogą to być rejestry czymś sterujące, na przykład prędkością dołączonego silnika. Ale mogą to być rejestry „czysto informatyczne”, przechowujące jakieś dane cyfrowe, na przykład adres tego urządzenia slave.

W danym urządzeniu slave mogą być zrealizowane wszystkie cztery rodzaje rejestrów albo tylko niektóre. W skrajnym przypadku może to być tylko jeden rejestr jednego rodzaju. Natomiast maksymalna liczba rejestrów każdego z czterech rodzajów zasadniczo wynosi 9999, co w sumie daje prawie 40 tysięcy (w niektórych urządzeniach jeszcze więcej). Poszczególne urządzenia slave mogą pełnić różne funkcje i mieć najrozmaitszą budowę. Nie warto wgłębiać się w szczegóły. Warto tylko zaznaczyć, że te logiczne rejestry mogą być w realnym urządzeniu slave oddzielnymi obszarami pamięci, co ilustruje **rysunek 7a**. Ale oddzielne logiczne rejestry MODBUS mogą być w fizycznym urządzeniu wspólne (**rysunek 7b**), czyli do tych samych komórek i danych można mieć dostęp przez różne rejestry MODBUS.

W następnym odcinku omówimy rozkazy dostępne w standardzie – protokole MODBUS.

Piotr Górecki



Rys. 7

Automatyczne nawadnianie także dla Ciebie – podstawy

część 1

W niniejszym artykule zajmiemy się *systemami automatycznego nawadniania ze szczególnym uwzględnieniem strony elektronicznej przedsięwzięcia.*

EdW niewątpliwie jest czasopiśmie dla elektroników, ale współczesnym elektronik powinien mieć szersze horyzonty, żeby mieć przynajmniej ogólne pojęcie o różnych dziedzinach życia, w których wykorzystywana jest elektronika. Między innymi o systemach automatycznego nawadniania (podlewania) ogrodów i ogródków, które już przestały być luksusem dostępnym tylko dla najbogatszych. Ceny sprzętu nie są wygórowane, a **największą barierą okazuje się... brak informacji.** Niektórzy nie interesują się tematem, ponieważ nie mają wiedzy o sytuacji rynkowej i możliwościach.

Dobra wiadomość jest taka, że elementy systemu nawadniającego nie są drogie, a praktycznie wszystkie prace można wykonać samodzielnie. Współczesny elektronik może jeszcze bardziej obniżyć koszty, samodzielnie wykonując albo cały system sterowania, albo przynajmniej jego elementy. System nawadniania można realizować stopniowo, co rozkłada koszty na dłuższy czas. Najłatwiej jest, gdy instalacja zawiera tylko linie kroplujące („przeciekające węże”).

Połączenia rur i kształtek instalacji wodnej wykonuje się szybko i łatwo. Nie trzeba do tego doświadczenia, wystarczy krótka instrukcja sprzedawcy i odrobina testów praktycznych. Nie są potrzebne specjalizowane narzędzia – większość złączek dociska (skręca) się ręcznie. Co istotne, w przypadku sieci rur zakopanych w ziemi (na głębokości około 30cm) ewentualne drobne przecieki nie są problemem. Bez względu na szczelność muszą być tylko fragmenty instalacji umieszczone przed zaworami, gdzie stale występuje ciśnienie wody, w szczególności te umieszczone w piwnicy (można też wykonać instalację „całkowicie zewnętrzną”).

W przypadku elektronika jest jeszcze jedna bardzo dobra wiadomość: **można najpierw zrealizować tańszą wersję podstawową**, zawierającą elektrozawory, sterowane w najprostszy sposób, nawet ręcznie za pomocą przełączników lub jakichkolwiek timerów. W kolejnych etapach można dodać automatyczny sterownik, a **potem ulepszać system oraz eksperymentować z różnymi sposobami sterowania i z różnymi czujnikami** (deszczu, wilgotności gleby, korzystaniem z SMS-ów i Internetu, itp.).

Podstawy

W telegraficznym skrócie: do realizacji automatycznego systemu nawadniania potrzebne są trzy główne „składniki”: (1) źródło wody o znacznej wydajności, (2) system plastikowych rur doprowadzających wodę do rur kroplujących i zraszaczy, a do tego (3) układ sterujący z elektrozaworami i mniej lub bardziej rozbudowanym sterownikiem (zazwyczaj mikroprocesorowym).

Najłatwiejsze do wykorzystania są tzw. linie kroplujące („dziurawe węże”), mające maleńkie otworki (ściślej wtopione w wąż emiter), przez które pomału przesącza się woda. Przykładowo typowa **linia kroplująca** o średnicy 16mm (**fotografia 1**) ma kroplowniki – emiter w odległości co 33cm. W pierwszym przybliżeniu można przyjąć **nominalny** wydatek wody każdego takiego kroplownika około 2 litry na godzinę, czyli 6 litrów na metr linii na godzinę. Linia kroplująca o długości 20 metrów ma wydatek wody 120 litrów na godzinę, co daje 2 litry na minutę, czyli drobnutkie 33 mililitry na sekundę. Oczywiście zmiana ciśnienia, długości linii oraz nierówności terenu powodują zmianę średniej wydajności.

Z liniami kroplującymi o długościach do 50 metrów w płaskich terenach nie ma problemu – nawadnianie jest jednolite na całej długości

linii. Aby dostosować się do ciśnienia w instalacji, można łatwo zmierzyć wydatek wody w określonym czasie i określić potrzebny czas nawadniania stosownie do potrzeb roślin lub przewidywanych kosztów.

W niewielkich instalacjach stosuje się też tzw. **mikrozraszacze** oraz **indywidualne kroplowniki**, też łatwe do wykorzystania. Znacznie trudniej jest z podlewaniem trawnika. Na niezbyt dużych domowych trawnikach z reguły stosuje się **zraszacze wynurzalne** (**fotografia 2**), dołączane do grubszych rur zasilających (25mm) zazwyczaj za pomocą łatwych do zamontowania obejm. Pojawienie się ciśnienia wody powoduje wypchnięcie ze zraszacza ruchomego tłoka w górę, a woda tryska z umieszczonej na wierzchu głowicy z odpowiednio dobraną dyszą. Zraszacze wynurzalne są atrakcyjne, ponieważ w spoczynku ich głowice są na poziomie gruntu i nie przeszkadzają w koszeniu trawnika.

Oprócz tego do nawadniania trawników stosowane są też **zraszacze obrotowe** (rotacyjne: młoteczkowe i turbino-

Fot. 2



Fot. 1



we) oraz ogólnie biorąc „stałe” – statyczne. Poszczególne zraszacze i głowice mają różne właściwości i parametry, choćby tylko zasięg, kąt i wydatek wody, zależne też od ciśnienia, które z kolei zależy od kilku czynników.

Oczywiście wszystko podlega ścisłym prawom fizyki. Ogólnie sytuacja jest taka, jak w obwodzie elektrycznym: mamy źródło wody o jakimś ciśnieniu (napięciu) i jakiejś wydajności (oporze wewnętrznym). Punktem wyjścia zawsze muszą być parametry źródła wody. Jeden podstawowy parametr to ciśnienie „spoczynkowe”, przy zamkniętym kranie (mierzone manometrem, wyrażone w atmosferach lub barach, odpowiednik napięcia SEM baterii), drugi to wydajność przy pełnym otwarciu zaworu (wyrażona w litrach na sekundę albo w metrach sześciennych na godzinę np. przez pomiar czasu napełniania wiadra o znanej pojemności, odpowiednik prądu zwarciaowego baterii). Wydajność realnego źródła wody – sieci wodociągowej albo lokalnej pompy – jest ograniczona. Dlatego większy system nawodnienia konieczne trzeba podzielić na sekcje. Czym mniejsza wydajność źródła wody (i niższe ciśnienie), tym więcej musi być mniejszych sekcji.

W systemie nawadniającym występują różne opory przepływu, więc przy przepływie wody następują spadki ciśnienia we wszystkich elementach systemu: w rurach, złączkach, filtrze, zaworach. Często na wejściu systemu nawadniania umieszczany jest filtr, który m.in. zmniejsza ryzyko błędnego działania elektrozaworów membranowych. W niektórych systemach zasilanych z sieci wodociągowej o wysokim ciśnieniu sensowne może się okazać zastosowanie reduktora ciśnienia. Przy ciśnieniu niezbyt dużym kluczowe może się okazać dobranie elementów, nie tylko rur, ale też zaworów, kształtek i licznika o odpowiednio dużym przekroju. Liczba zraszaczy, ich rozmieszczenie oraz dobór dysz mają zapewnić równomierne nawodnienie trawnika. Finalnie na głowicy-dyszy zraszaczej powinno wystąpić takie ciśnienie, żeby uzyskać potrzebny zasięg zraszania.

Zasady są takie same, jak w obwodach elektrycznych, ale problemem jest konieczność uwzględnienia dużej liczby czynników. Największego doświadczenia wymaga zaprojektowanie systemu nawadniania trawników, które z reguły

mają nieregularne kształty. Wprawdzie są głowice (dysze) na różne zakresy ciśnienia i mają możliwość regulacji kąta zraszania i zasięgu, niemniej **dla osób nieorientowanych zaprojektowanie nawadniania trawnika jest zdecydowanie zadaniem zbyt trudnym. Dlatego projekt instalacji wodnej oraz dobór zraszaczy wynurzalnych i dysz należy zostawić profesjonalistom.**

Do projektu potrzebny jest plan ogrodu. Dobra wiadomość jest taka, że większość firm/sklepów handlujących sprzętem nawadniającym robi taki projekt za 100...300zł, a potem... tę sumę odejmie od ceny zakupu sprzętu. Właśnie dlatego optymalne są zakupy nie przez Internet, tylko w pobliskiej firmie, która gotowa byłaby pomóc nie tylko przy projekcie, ale ewentualnie też przy uruchomieniu i finalnej regulacji dysz w spryskiwaczach.

Koszty

Koszty materiałów nie są duże. Rury z polietylenu kosztują, zależnie od średnicy, 1...3zł/metr. Popularna linia kroplująca (16mm) kosztuje około 1zł/metr. Złączki z tworzywa kupuje się po kilka złotych. Popularne elektrozawory 1-calowe to koszt 60...100 złotych za sztukę. Zraszacze, dysze i potrzebne złączki to sumaryczny koszt mniej więcej 100zł na zraszacz.

Sterowniki mają bardzo różne ceny, zależnie od możliwości. Elektroniki może różne sterowniki wykonać samodzielnie, na początek w prościutkiej wersji bez zaawansowanych funkcji.

W sumie ceny materiałów są przystępne. Natomiast dla wielu mniej zamożnych hobbystów **przykrym zaskoczeniem mogą się okazać koszty użytkowania**, o ile źródłem wody ma być publiczna sieć wodociągowa. Nawet przy założeniu podlicznika (by nie płacić za odprowadzenie ścieków, a tylko za samą poborą z sieci wodę) rachunki mogą znacząco wzrosnąć przy nawadnianiu według ogólnie przyjętych zaleceń. Według takich zaleceń codziennie na 1 metr kwadratowy trawnika należałoby wylać 6 litrów wody. Jeśli przykładowo ktoś ma trawnik o powierzchni 300m², to dziennie powinien zużyć około 1800 litrów wody, czyli prawie 2 metry sześciennie. Jeżeli 1m³ wody kosztuje 5zł (bez ścieków), dziennie da to koszt 10 zł, miesięcznie 300zł, a przez cały sezon zapewne znacznie ponad 1000

złotych. Do tego dojdzie zużycie wody przez linie kroplujące.

Jeśli koszty uzyskania wody są decydujące, atrakcyjna może okazać się eksploatacja własnego ujęcia – istniejącej, a już niewykorzystywanej studni/pompy. Natomiast wykonanie nowej studni, zakup pompy i zbiornika buforowego (hydroforu) oraz filtra to wydatek co najmniej kilku tysięcy złotych i trzeba się dobrze zastanowić, kiedy się taka inwestycja zamortyzuje.

Nie ma co liczyć na zbieranie dużych ilości deszczówki – jej zasoby są zwykle nieduże, uzysk nieregularny, a koszty pozyskania – wbrew pozorom duże. Średnio można przyjąć roczną ilość opadów w Polsce 500mm (0,5m³/m²), z czego tylko część uda się wykorzystać w sezonie. Przyjmując połowę tej wartości, z dachu o powierzchni 100m² (10×10m) rocznie można byłoby uzyskać 25m³ cennej miękkiej deszczówki. W zasadzie dużo, ale koszty porządnego zbiornika na deszczówkę są rzędu 1000 złotych na każdy metr sześcienny (1000 litrów). Tymczasem 1m³ wody z wodociągu kosztuje około 5zł (bez kanalizacji), więc inwestycja w porządną system zbierania deszczówki okazuje się zaskakująco kosztowna.

Oczywiście należy też podkreślić, że wspomniane „ogólnie przyjęte zalecenia” dotyczą obfitego podlewania, żeby rośliny miały optymalne warunki wegetacji. Z uwagi na koszty, wielu użytkowników zmniejsza dawki wody.

Ponadto rzeczwiście zapotrzebowanie na wodę nie jest stałe w ciągu sezonu. Użytkownik może znacząco ograniczyć zużycie wody, w skrajnym przypadku na przykład żeby jedynie zapobiec całkowitemu wyschnięciu, powodującemu zniszczenie trawnika.

W każdym razie naprawdę warto zainteresować się możliwościami optymalizacji nawadniania. Ścisłej biorąc, miarą jest tu parametr znany jako **ewapotranspiracja** (patrz np. www.nawadnianie.inhort.pl/slownik), obejmujący zarówno parowanie powierzchni gleby, jak też „wyciąganie” wody przez rośliny, co zależy od temperatury i innych czynników. W grę wchodzi opady deszczu, nie tylko z ostatnich dni, ale też przewidywane, żeby nie podlewać tuż przed nadejściem frontu z obfitymi deszczami.

Dostępne na rynku sterowniki mają różne możliwości. Do niektórych można dołączyć czujniki wilgotności

gleby lub elementy stacji meteo, w tym deszczomierz. Inne mogą pobierać dane meteorologiczno-nawodnieniowe z Internetu. Oczywiście zaawansowane sterowniki fabryczne są odpowiednio drogie. W przypadku elektronika optymalnym rozwiązaniem może okazać się samodzielne realizowanie i stopniowe ulepszanie sterowników nawadniania, począwszy od najprostszych. Tym bardziej że można wtedy wykorzystać specyficzne i nietypowe rozwiązania, z zastosowaniem interesujących czujników, o których jeszcze wspomnimy. A teraz aspekty elektroniczne systemu nawadniania.

Elektrozawory

W nawet najmniejszym systemie automatycznego nawadniania potrzebne są elektrozawory. Są to stosunkowo duże (zwykle 1-calowe) elektrozawory membranowe przeznaczone do nawadniania o niecodziennej zasadzie pracy. **Fotografia 3** pokazuje popularny 1-calowy elektrozawór Hunter PGV. Ponieważ na forach można znaleźć wiele błędnych lub nieścisłych informacji, trzeba wyraźnie podkreślić, że do takich elektrozaworów stosowanych w kraju można wkręcić albo jeden, albo drugi rodzaj cewki-elektromagnesu: najczęściej na prąd zmienny (24VAC) jak na fotografii 3, rzadziej na prąd stały (9VDC).

Fot. 3



Te dwa rodzaje wymagają zupełnie innego sposobu sterowania. Zawór jest otwarty, gdy na cewkę 24VAC podane jest napięcie. Inaczej działa cewka-elektromagnes 9VDC: cewka taka musi być sterowana krótkimi impulsami, ponieważ **działa jak przekaźnik bistabilny jednocewkowy**. Podanie impulsu o jednej biegunowości powoduje przyciągnięcie tłoka i przytrzymanie go w tej pozycji. Podanie impulsu o przeciwnej biegunowości cofa tłok do położenia spoczynkowego. Elektrozawór z cewką DC jest więc bardzo energooszczędny: wymaga tylko króciutkich impulsów o czasie rzędu kilkadziesiąt milisekund. To wyjaśnia tajemnicę, dlaczego niektóre sterowniki nawadniania są zasilane albo małą baterią 9V (błoczek), która wystarcza na cały sezon pracy, albo 2...3 ogniwami AA. Takie energooszczędne, bistabilne cewki-elektromagnesy DC są mało popularne w amatorskich instalacjach. Jednym z praktycznych problemów jest to, że standardowy elektrozawór zawiera cewkę 24VAC i taki komplet można kupić już za około 70zł. A cewkę 9VDC trzeba zakupić oddzielnie i zwykle kosztuje więcej niż cały elektrozawór z cewką 24VAC. Pewne zamieszanie dotyczy możliwości pracy cewek 24VAC przy napięciu stałym. Otóż zasadniczo **są to elektromagnesy prądu zmiennego, ale pod pewnymi warunkami mogą być sterowane prądem stałym!**

To jest bardzo ważna wiadomość dla wielu elektroników, którzy realizują własne sterowniki nawadniania i chcieliby zasiląć je napięciem stałym. Zagadnienie to zostanie szerzej opisane w oddzielnym artykule.

Elektrozawory kulowe

Do nawadniania wykorzystuje się standardowo omówione właśnie duże elektrozawory membranowe. Ich awarie zdarzają się rzadko, ale niestety nie można ich wykluczyć – skutkiem może być np. zalanie piwnicy lub nieplanowane zwiększenie rachunku za wodę o kilka tysięcy złotych. Gdy system nawadniania jest pod stałym nadzorem domowników, ryzyko jest niewielkie. Ale liczne instalacje są obsługiwane zdalnie: w domkach letniskowych albo w domach mieszkalnych podczas urlopu.

Prawdopodobieństwo awarii i niekontrolowanego wypływu wody jest wprawdzie nieduże, ale pełna eliminacja ryzyka i spokojny sen są warte poniesienia dodatkowych kosztów. Dlatego **konieczny jest dodatkowy elektrozawór główny, umieszczony na wejściu do systemu nawadniania**. Ale nie powinien to być kolejny zawór „nawodnieniowy”, tylko sterowany silnikiem porządnym, niezawodnym zawór kulowy.

Fotografia 4 pokazuje tego rodzaju elektrozawór w wersji z pokrętką, pozwalającym też na sterowanie ręczne.

Zawiera on mały silniczek prądu stałego na 12V z przekładnią oraz dwa mikrowyłączniki, sygnalizujące osiągnięcie krańcowych pozycji zaworu otwarte/zamknięte. Zasada pracy pod-

Fot. 4



stawowych wersji jest prosta: podanie napięcia na silnik powoduje obrót kuli z otworem. Po kilku sekundach, gdy nastąpi obrót o 90 stopni, czyli zawór zostanie w pełni otwarty, zadziała jeden z mikrowyłączników i silnik ma zostać wyłączony. Taki elektrozawór pobiera energię tylko w czasie zmiany stanu, gdy pracuje silniczek

Fot. 7



Po dobrej analizie oferty rynkowej takie zawory można kupić już za kilkadziesiąt złotych. **Uwaga! Na rynku dostępnych jest szereg wersji tego rodzaju elektrozaworów, różniących się sposobami sterowania** (niektóre zawierają dodatkowe elementy). Przy zakupie trzeba na to zwrócić uwagę i kupić wersję dostosowaną do własnych potrzeb. To też będzie szerzej omówione w oddzielnym artykule.

Sterowniki nawadniania
Oferta sterowników nawadniania jest bardzo szeroka. Dostępne są też proste sterowniki programowane mechanicznie (fotografia 5). Sterowniki elektroniczne można podzielić na grupy. Najprostsze są jednoobwodowe – przykład wyrobu znanego polskiego producenta na fotografii 6.

Sterowniki nawadniania

Popularne są niedroge sterowniki kilkukanałowe, przeznaczone do współpracy z elektrozaworami zawierającymi cewki 24VAC. Droższe wersje mają

możliwość zdalnego sterowania, albo przez telefon, albo przez Internet. Przykład na fotografii 7 to sterownik znanej firmy Rain Bird.

Bardziej zaawansowane i droższe sterowniki mogą współpracować z czujnikami deszczu, więc uwzględniają aktualną sytuację. Tylko najdroższe wersje mają czujniki wilgotności gleby lub mogą z nimi współpracować. Niektóre przekazują i uzyskują dodatkowe informacje za pomocą Internetu.

Ceny sterowników w stosunku do możliwości są stosunkowo niskie, jednak dla elektronika hobbysty jeszcze bardziej interesujące są rozwiązania DIY (zrób to sam). Można zacząć od części hydraulicznej z elektrozaworami i prostym

sterownikiem ręcznym, a potem stopniowo ulepszać sterowanie. Zaletą jest możliwość nieograniczonej rozbudowy, eksperymentów i testowania różnych opcji. Ogromnie cenna jest satysfakcja z wyników samodzielnej pracy, ale rozbudowane konstrukcje finalnie nie będą wcale tańsze od fabrycznych.

Sterowniki nawadniania można zrealizować w najróżniejszy sposób, choćby przy użyciu Arduino, RaspberryPi, ESP32. W Internecie można znaleźć szereg przykładów realizacji, często z dokumentacją. Warto wiedzieć, że w Internecie od dawna dostępna jest pełna dokumentacja projektu OpenSprinkler (<https://opensprinkler.com/>) i gotowe urządzenia (rysunek 8). W ramach niniejszego artykułu nie będziemy zagłębiać się w szczegóły

A w następnym odcinku omówimy inne elementy o charakterze elektronicznym, które mogą się znaleźć w systemie nawadniania.

Piotr Górecki



Fot. 5



Rys. 8

Fot. 6



NanoVNA – wykonaj precyzyjne pomiar



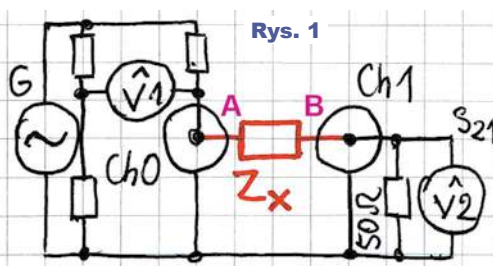
W poprzednim odcinku omówiliśmy sposoby pomiaru z pomocą VNA dużych i małych impedancji. Pomiar dużych impedancji, powyżej 50Ω , należy przeprowadzać według **rysunku 1**. Do pomiaru małych impedancji, poniżej 50Ω , ma służyć konfiguracja *shunt-through* według **rysunku 2**.

Wejście Ch1 VNA mierzy napięcie na mierzonej impedancji Z_x i teoretycznie dolny zakres mierzonych wartości tej impedancji jest wyznaczony przez poziom szumów własnych użytego analizatora VNA. W porównaniu z kosztownymi, profesjonalnymi VNA, w NanoVNA ten poziom szumów jest stosunkowo wysoki i dynamika jest niezbyt duża. Ogranicza to od dołu zakres mierzonych impedancji. Ale to nie wszystko: także w przypadku taniutkiego NanoVNA trzeba omówić pewien istotny problem, dotyczący wszystkich VNA. Jest to bardzo ważne w praktyce, ponieważ często chcemy w szerokim zakresie częstotliwości mierzyć małą impedancję obwodów zasilania (PDN – Power Distribution Networks) i ich elementów składowych, w szczególności stosowanych tam kondensatorów odsprężających.

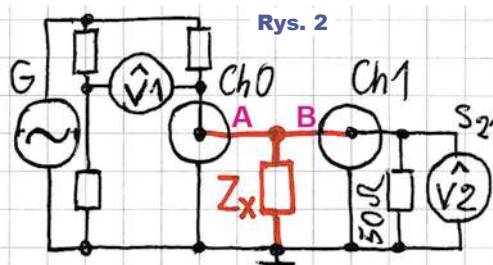
Pełna kalibracja

Gdy wykorzystywaliśmy tylko gniazdo Ch0 i pomiar S_{11} , wystarczyła trzystopniowa kalibracja **SOL** (*Short, Open, Load*).

Jeżeli jednak chcemy wykorzystywać konfiguracje pomiarowe według rysunków 1, 2, czyli mierzyć także parametr S_{21} , powinniśmy, a wręcz musimy przeprowadzić pełną kalibrację pięciostopniową **SOLIT**.



Rys. 1



Rys. 2

W programie *NanoVNA-Saver*, wykorzystując *Calibration assistant*, po trzech krokach nie klikniemy *Apply*, tylko *Yes*, a wtedy wyświetlą się kolejno dwa komunikaty pokazane na **rysunku 3**.

Pierwszy dotyczy kalibracji **I** – *Isolation*, na czas której na gniazdo wyjściowe Ch0 należy nakręcić nakrętkę *Load* (50 omów), a jeżeli mamy drugą 50-omową, to także na wejście Ch1. Wtedy wejście Ch1 jest

izolowane od wyjścia Ch0 i kalibrujemy jakość tej izolacji.

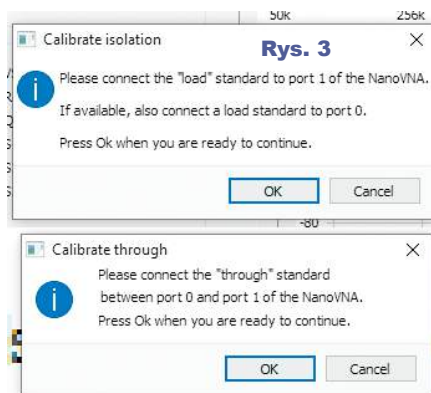
W następnym kroku **T** mamy kalibrować *przejście* – *Through* i na ten czas Ch0 należy bezpośrednio połączyć z wejściem Ch1.

Na koniec taką pełną kalibrację trzeba zastosować – wprowadzić do programu (*Apply*), ewentualnie także zapisać na dysku w postaci pliku **.cal** po kliknięciu *Save calibration*.

W zasadzie procedura pełnej, pięciostopniowej kalibracji jest dziecinnie prosta, jednak jak mówi przysłowie, diabeł tkwi w szczegółach. Przekonamy się o tym niebawem.

W każdym razie dopiero pełna kalibracja **SOLIT** pozwala wykorzystać konfiguracje pomiarowe z rysunków 1 i 2, które genialnie rozszerzają zakres pomiarów impedancji w porównaniu z wykorzystaniem jedynie gniazda Ch0.

Ja do dokładniejszych pomiarów dwuportowych wykorzystałem pokazany na **fotografii 4** 10-centymetrowy (4 cale) półsztywny kabel RG402 z dwoma męskimi kątowymi wtykami SMA, zakupiony kiedyś pod adresem: <https://bit.ly/3F7SAnj>.



Rys. 3

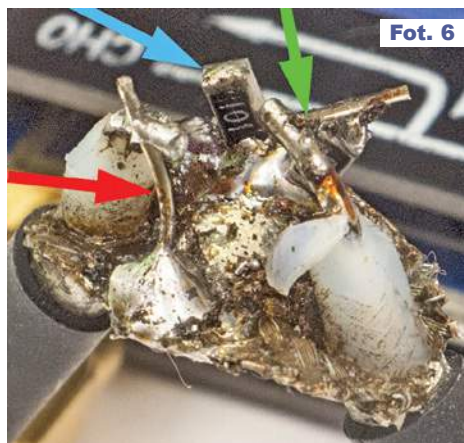


Fot. 4

Rozciąłem ekran kabla, przeciąłem wewnętrzną żyłę, a do ekranu (masy) przylutowałem na stałe dwa rezystory 50Ω (równolegle po dwa 100Ω SMD), które są potrzebne do kalibracji. W ten sposób do pomiarów według rysunków 1, 2 wykonałem adapter pokazany na **fotografii 5**.

Do kalibracji **S** (short – zwarcie) zwieram końcówkę **A** połączoną z gniazdem Ch0 za pomocą kawałeczka drutu, co na **fotografii 6** pokazuje czerwona strzałka. Jednocześnie drugą końcówkę pomiarową **B**, połączoną z Ch1, dołączam do rezystancji 50 omów, jak pokazuje zielona strzałka (co ma niewielkie znaczenie, ale jest zalecane). Jeden z rezystorów 50Ω pozostaje niepodłączony (niebieska strzałka).

W następnym kroku kalibracji **O** (open – rozwarcie) końcówka **A** jest odłączona, czyli „wisi w powietrzu”, a końcówka **B** nadal jest dołączona do rezystora 50Ω.

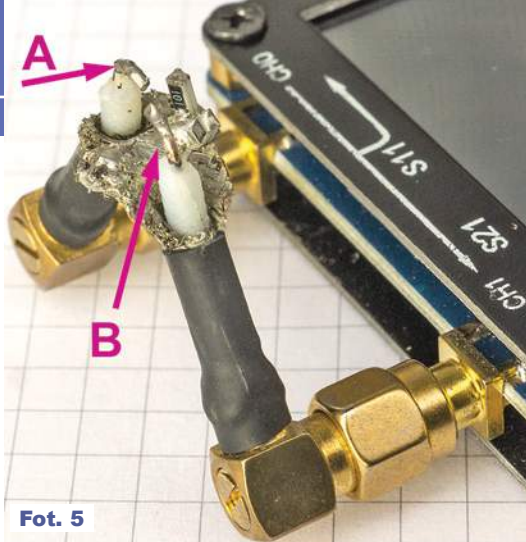


Fot. 6

W następnym kroku kalibracji **L** (load – obciążenie) dołączam końcówkę pomiarową **A** do niepodłączonego wcześniej rezystora 50Ω według **fotografii 7**. Kolejny krok kalibracji to **I** (isolation – izolacja między portami Ch0, Ch1). Układ połączeń jest taki



Fot. 7



Fot. 5



Fot. 8

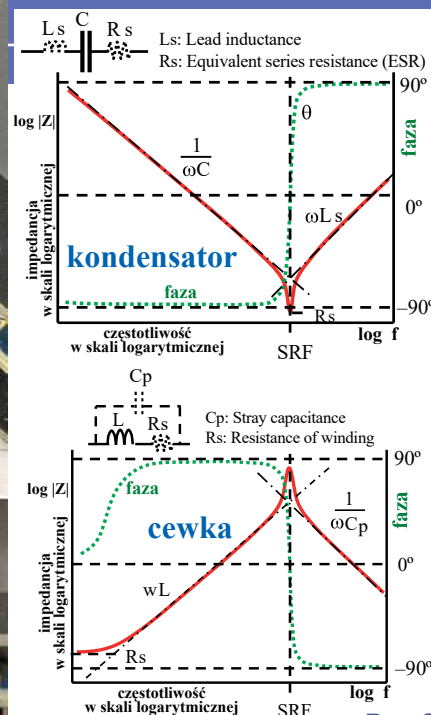
sam jak podczas poprzedniego kroku **L** – do obu gniazd dołączone są rezystancje 50Ω.

Ostatni krok pełnej kalibracji to **T** (through – przejście z Ch0 do Ch1). Dwa rezystory pomocnicze zostają odłączone, a punkty **A**, **B** – zwarte, jak pokazuje **fotografia 8**.



W procesie takiej pełnej kalibracji celowo nie wykorzystuję nakrętek kalibracyjnych z zestawu NanoVNA, tylko zwyczajne rezystory SMD wielkości 0805, jednak w zakresie częstotliwości do kilkuset megaherców taki sposób z rezystorami SMD jest wystarczający do naszych celów edukacyjnych.

Efekty pomiarów z użyciem obu portów po takiej właśnie kalibracji już widziałeś m.in. w poprzednim odcinku na rysunkach 1, 2, które pokazują wyniki pomiarów rezystorów (22kΩ i 470kΩ) w konfiguracji z zamiesz-

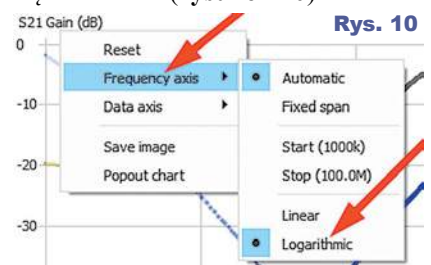


Rys. 9

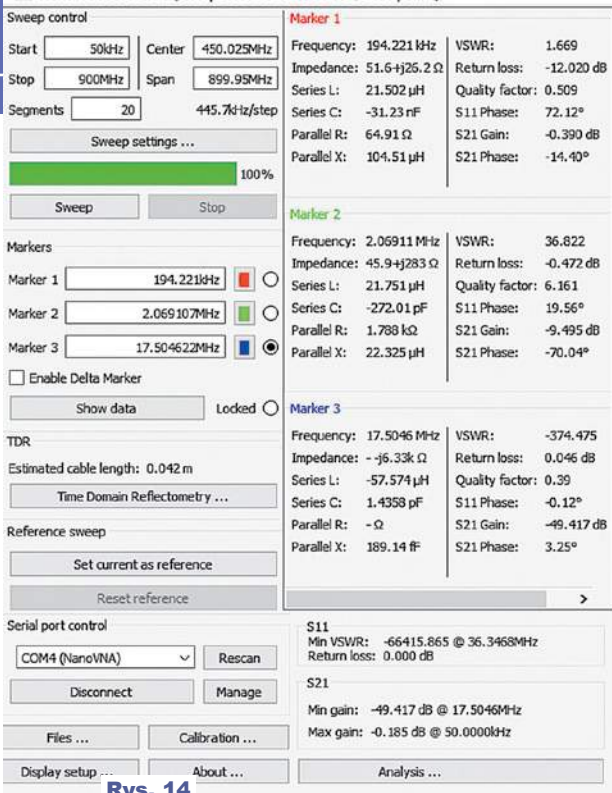
zonego tam rysunku 6. Pomiar z wykorzystaniem tylko portu Ch0 i parametry S11 mają istotne wady, nie tylko wąski zakres pomiarowy, w pobliżu 50 omów. Jest też kłopot ze zobrazowaniem wyników, a mianowicie niemożliwość przedstawienia modułu impedancji $|Z|$ obliczonej na podstawie S11 w skali logarytmicznej (a przynajmniej ja nie potrafię tego zrobić). Przy pomiarach z użyciem dwóch portów mamy dwie konfiguracje dla dużych i małych impedancji (rysunki 1, 2), a wartość zmierzonej impedancji można przedstawić w skali podwójnie logarytmicznej.

Właśnie w skali podwójnie logarytmicznej przedstawiane są charakterystyki częstotliwościowe impedancji elementów, jak pokazuje przykład z **rysunku 9** (z materiałów Keysight). Wtedy bardzo dobrze widać, że reakcja pojemnościowa maleje ze wzrostem częstotliwości, a indukcyjna – rośnie. Sensownie przedstawiony jest też wtedy rezonans własny (SRF) wynikający z niedoskonałości kondensatorów i cewek elementów.

W programie NanoVNA-Saver klikając „prawą myszką” na wykresach można ustawić logarytmiczną skalę częstotliwości (**rysunek 10**).



Rys. 10

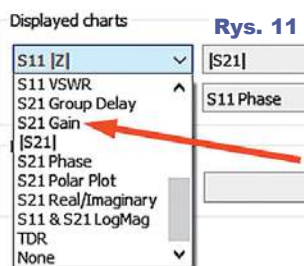


Rys. 14

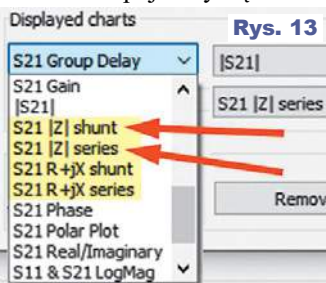
Wartość modułu impedancji też powinna być przedstawiona w skali logarytmicznej – wtedy uzyskamy piękne, podręcznikowe charakterystyki.

Jak pokazuje **rysunek 11**, w wersji 0.3.8 NanoVNA-Saver po kliknięciu **Display setup** była możliwość wyboru logarytmicznego zobrazowania tylko wartości **S21 Gain**, która rzeczywiście niesie informacje o badanej impedancji, ale nie pokazuje jej w omach. Na pionowej osi na wykresie mamy decybele, ale nie są to „decybeloomy” i najprawdopodobniej należałoby uwzględnić rezystancję charakterystyczną 50Ω.

Przy pomiarach dużych rezystancji według **rysunku 1**, na wykresie otrzymamy charakterystykę odwróconą, czego przykład mamy na **rysunku 12**, będącego rysunkiem 9c z EdW 8/2021 str. 19. To był dość istotny kłopot, ale na szczęście w wersji 0.3.9 NanoVNA-Saver pojawiły się nowe możliwości zobrazowania wyników. Autorzy programu dodali cztery nowe sposoby zobrazowania potrzebne właśnie do pomiarów impedancji według **rysunków 1 (series)**, **2 (shunt)**. Czerwone strzałki na **rysunku 13** pokazują wybór modułu impedancji |Z|.

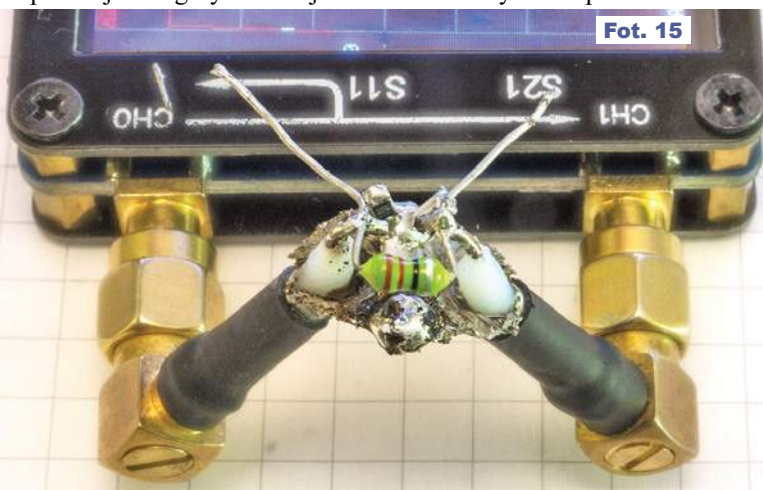


Rys. 11

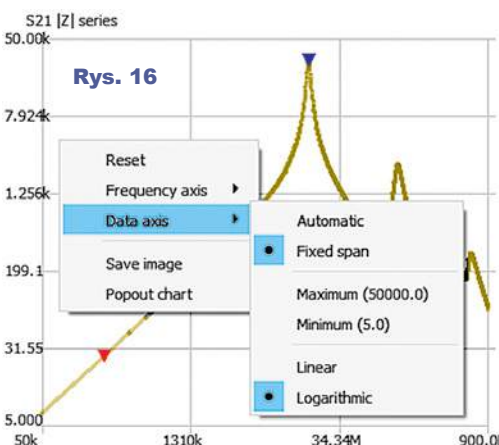


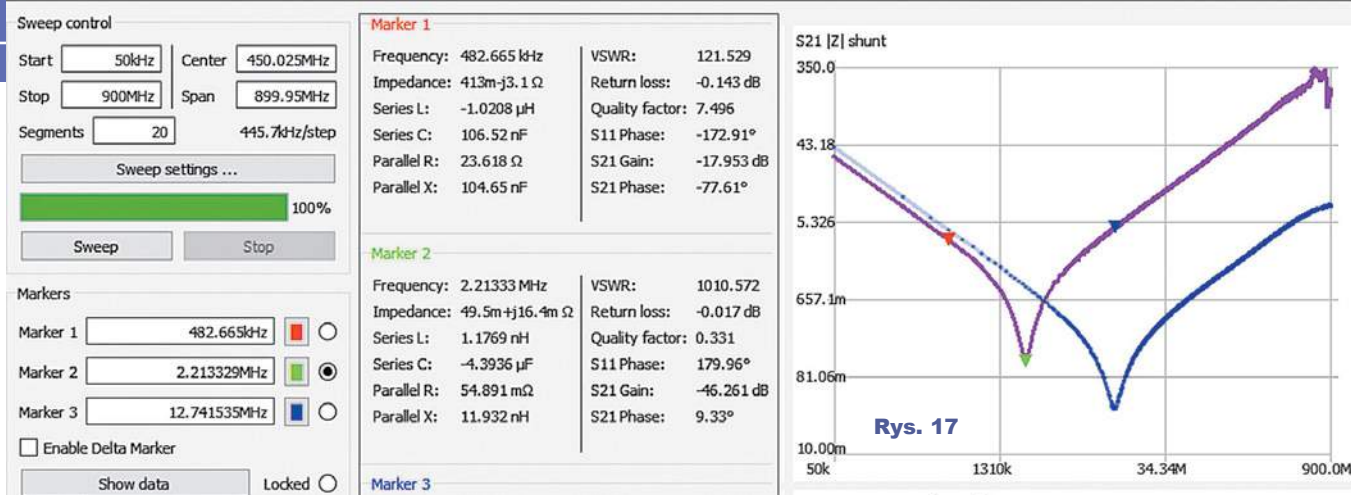
Rys. 13

I ten moduł impedancji można zobrazować w skali liniowej albo logarytmicznej. **Rysunek 14** to wyniki pomiaru tego samego dławika 22uH (**fotografia 15**) w wersji programu 0.3.9. Wartość od razu podana jest w omach. W dwóch górnych oknach moduł impedancji (**S21 |Z| series**) przedstawiony w skali podwójnie liniowej i ten sam wynik w skali podwójnie logarytmicznej. Komentarza chyba nie potrzeba.



Skalę zobrazowania wartości modułu impedancji można zmienić, klikając „prawą myszką” na wykresie według **rysunku 16**. Ja na potrzeby artykułu wybrałem zakres **Fixed span** i dobrałem wartości skrajne. Jednak najczęściej korzystamy z opcji **Automatic**.





Rys. 17

Na rysunku 14 w dolnym lewym oknie zobrazowane są oddzielnie: składowa rzeczywista (rezystancyjna – R; skala z lewej strony) oraz składowa urojona (reaktancyjna – X; skala z prawej strony). Przy takim zobrazowaniu nie można ustawić logarytmicznej skali oporności, a tylko logarytmiczne zobrazowanie częstotliwości. Ma to uzasadnienie, bo skala logarytmiczna nie obejmuje zera ani wartości ujemnych, a wartość reaktancji może być dodatnia lub ujemna. Nie jest to więc niedoróbka. Tym bardziej że mało kto czuje intuicyjnie taki sposób zobrazowania impedancji.

Pomimo ulepszeń wersji 0.3.9, nadal jest kłopot z odwzorowaniem kąta przesunięcia (fazy) między prądem i napięciem w mierzonej impedancji. Nawet w nowszej wersji programu nie przewidziano odpowiedniego wykresu. Na rysunku 14 w dolnym prawym oknie przedstawiona jest faza parametru S21. Wykres taki okazuje się jak najbardziej użyteczny, ale nie jest to dokładne, a jedynie przybliżone odwzorowanie przesunięcia między prądem i napięciem w badanej impedancji.

Z lewej strony rysunku 14 warto też zwrócić uwagę na informacje o trzech kolorowych markerach, które można dowolnie ustawić na skali częstotliwości. Można tam odczytać, co algorytmy zastosowane w programie NanoVNA-Saver wyliczyły ze zmierzonych wartości S11, S21. Wyliczają między innymi indukcyjność i podają wynik bliski 22uH. Niestety, w wersji 0,3,9 nie wyliczają wartości modułu impedancji, co byłoby bardzo pożyteczne przy takich pomiarach. Wartości modułu impedancji trzeba odczytać z rysunku. W razie potrzeby można to zrobić dokładniej, korzystając z opcji „Fixed span” i ustawiając według potrzeb wartości pokazywanej oporności: minimalną i maksymalną.

Rysunek 17 analogicznie pokazuje wyniki pomiaru dwóch kondensatorów 100nF w konfiguracji *shunt* według rysunku 2. Krzywa fioletowa to przebieg modułu impedancji starego, dużego kondensatora MKSE 100nF 630V. Krzywa niebieska to zapamiętana (za pomocą polecenia *Set current as reference*) charakterystyka małego kondensatora SMD 0805 o nominalnie 100nF, który jest wskazany czerwoną strzałką na **fotografii 18**. Fotografia ta pokazuje pomiar *shunt* dużego kondensatora foliowego, a małego kondensatora SMD jest niepodłączony, tylko jedną końcówką przylutowany do ekranu kabla.

Widzimy wyraźnie, że najpierw ze wzrostem częstotliwości impedancja (reaktancja pojemnościowa) maleje, ale tylko do częstotliwości rezonansu własnego. Powyżej dominuje reaktancja indukcyjna szkodliwej indukcyjności. Co najważniejsze, na wykresach 14 i 17 pionowa skala podaje wartość modułu impedancji wprost w omach. Jest to szczególnie ważne w przypadku badania kondensatorów, bowiem po niezbędnej prawidłowej kalibracji uzyskujemy informację o wartości ESR, która generalnie powinna być jak najmniejsza. Wartość ESR określa najniższy punkt wykresu na rysunku 17. A ogólnie biorąc, wysokość nachylonej charakterystyki pokazuje, jaka jest reaktancja (pomijając niewielkie szkodliwe

rezystancje). Czym wyżej przebiega nachylna charakterystyka, tym mniejsza pojemność albo większa indukcyjność, które zresztą można obliczyć dla dowolnego punktu wykresu.

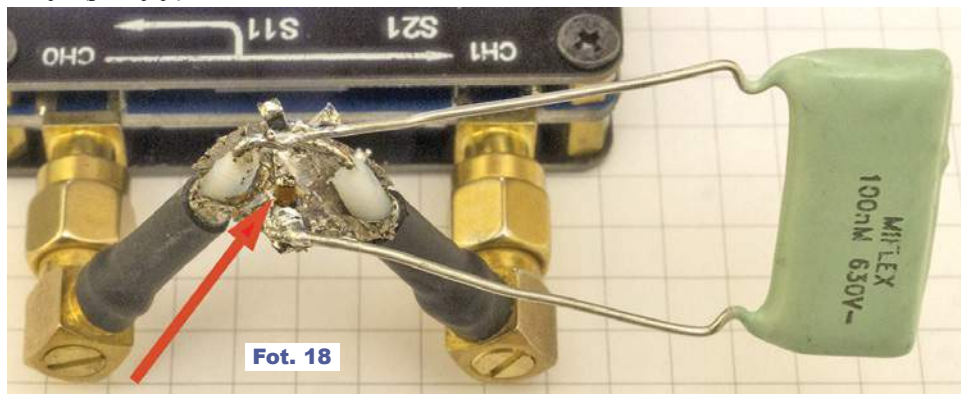
Na rysunku 17 też warto zwrócić uwagę na markery. Jak pokazuje marker czerwony, z wyników pomiaru przy częstotliwości około 500kHz program prawidłowo wyliczył wartość szeregowej pojemności bliską 100nF.

Przypominam, że zasadniczo NanoVNA mierzy parametry tylko przy 101 częstotliwościach rozmieszczonych liniowo. Dla uzyskania większej dokładności i możliwości sensownego zobrazowania w logarytmicznej skali częstotliwości trzeba przeprowadzić wiele pomiarów – ja i przy pomiarach, i przy kalibracji ustawiłem 20 segmentów, więc miernik dokonuje ponad 2000 pomiarów, co trwa dość długo. Pojedynczy pomiar trwa prawie 40 sekund, ale za to wykres jest dokładniejszy.

W tym odcinku nauczyliśmy się mierzyć impedancję w szerokim zakresie wartości. **Zachęcam do samodzielnego pomiaru impedancji różnych elementów i obwodów!**

Podczas takich pomiarów zapewne pojawią się dziwne wyniki i różne zagadkowe zjawiska. Dlatego w następnych odcinkach w temat będziemy wniknąć głębiej.

Piotr Górecki



Fot. 18

Szkoła Konstruktorów



W Szkole Konstruktorów może wziąć udział każdy Czytelnik EdW, także i Ty!

Możesz zostać stałym uczestnikiem Szkoły, ale możesz tylko jednorazowo nadesłać pojedyncze rozwiązanie jednego zadania, które Cię najbardziej zainteresowało. Nie trzeba się zapisywać, nie ma żadnych zobowiązań – można tylko zyskać. Co miesiąc przydzielane są punkty, upominki, nagrody i kupony do Sklepu AVT, a raz na rok najaktywniejsi uczestnicy Szkoły Konstruktorów są nagradzani dodatkowo. W każdym numerze zamieszczane są zadania trzech klas (*Zadanie główne*, *Co tu nie gra?* oraz *Policz*).

W terminie dwóch miesięcy możesz więc nadesłać e-mailem na adres: szkola@elportal.pl (szkola, a nie szkoła), rozwiązanie jednego, dwóch albo wszystkich trzech zadań Szkoły z danego numeru.

Potwierdzam otrzymanie rozwiązań, nadsyłanych e-mailem. Jeśli w terminie dwóch tygodni nie otrzymasz mojego potwierdzenia, prześlij rozwiązanie jeszcze raz (o przyczynach ewentualnych kłopotów przeczytasz na początku rubryki *Pocztą* na stronie 10).

Bardzo proszę: dla ułatwienia segregacji niech tytuł Twojego e-maila (i nazwa każdego ewentualnego załącznika), oprócz **nazwy konkursu** oraz **numeru zadania**, zawiera też **Twoje nazwisko** (najlepiej bez typowo polskich liter), na przykład: *Szko311Kowalski*, *Policz311Zielinski*, *NieGra311Malinowski*, *Jak02Krzyzanowski*. Chodzi o to, żeby w tytule e-maila i w nazwach wszystkich załączników była zarówno informacja o zadaniu, jak i o Autorze. Bardzo też proszę, żeby jeden Twój e-mail zawierał rozwiązanie tylko jednego konkursu, a nie kilku, co znacznie mi ułatwi segregowanie poczty.

Do wysyłki nagród i upominków potrzebny jest Twój adres pocztowy. Oszczędzisz mi sporo niepotrzebnej pracy, jeśli podasz go w jednej linii: **Imię Nazwisko ulica nr domu kod pocztowy Miejscowość e-mail**

Jeśli na łamach czasopisma nie chcesz ujawniać imienia i nazwiska – napisz, a zachowam dyskrecję, podając albo pseudonim, albo imię i pierwszą literę nazwiska, ewentualnie miejscowość zamieszkania. Jeśli nadeślesz rozwiązanie zadania głównego, możesz dołączyć swoją fotografię (portret), która będzie zamieszczona przy rozwiązaniu zadania. Zachęcam też do podawania **roku urodzenia**, a w przypadku uczniów i studentów także informacji o szkole/klasie lub uczelni. Jest to pomocne przy opracowywaniu i ocenie rozwiązań (Twoje dane nie są nigdzie przekazywane, tylko wykorzystywane w redakcji EdW wyłącznie w związku z oceną prac i przydzielanymi nagrodami).

Najbardziej cieszę się z krótkich i zwięzłych rozwiązań, bo to ułatwia ich opracowanie. Ale jeżeli Twoje rozwiązanie będzie obszerniejsze, mam prośbę dotyczącą kwestii technicznych: Nie umieszczaj ilustracji w tekście! Wszystkie ilustracje (fotografie i rysunki) prześlij w e-mailu jako oddzielne pliki – załączniki. Bardzo proszę też o przysyłanie schematów, projektów płytek i wszelkich innych rysunków w popularnych formatach, na przykład PDF, SVG, JPG, GIF czy PNG, i to także wtedy, gdy przysyłasz oryginalny, źródłowy plik z danego programu projektowego (.sch, .pcb, .brd, .ddb, itp.).

Jeżeli w ramach zadania głównego zrealizujesz rozwiązanie praktyczne, czyli zbudujesz konkretny układ-model, mam następujące wskazówki i prośby:

Nie przysyłaj modelu do redakcji! Nie ma też potrzeby nadsyłania papierowych wydruków, płyty CD/DVD, ani modelu – całkowicie wystarczy załączone do e-maila pliki i fotografie zrobione przez Ciebie.

Przygotowując opis **skorzystaj z szablonu** dostępnego pod adresem: www.elportal.pl/szablon.

Więcej wskazówek na temat przygotowania materiałów i prawidłowego fotografowania modeli znajdziesz w Elportalu na stronie: <https://elportal.pl/zostan-wspolautorem-elektroniki-dla-wszystkich/>.

Twoje praktyczne rozwiązanie głównego zadania Szkoły może być później opublikowane jako artykuł w EdW, za który otrzymasz honorarium. Dlatego w treści e-maila umieść wtedy tekst: *Oświadczam, że materiał, który przesyłam w tym e-mailu do redakcji „Elektroniki dla Wszystkich”, jest moim osobistym opracowaniem i nie był wcześniej nigdzie publikowany.*



Zadanie główne 311

Temat zadania głównego 311 ma źródło w rozmowie, jaką na początku grudnia przeprowadziłem z jednym z Autorów piszących do EdW. Otóż pomysłodawcą zadania jest dobrze znany naszym Czytelnikom **Karol Świerc**, specjalista od serwisu i wszelkich zasilaczy impulsowych, który w luźnej rozmowie na różne tematy wspomniiał o problemie wielu osób, jakim są kłopoty ze snem. Jedną z najczęstszych przyczyn jest stres i wiele terapii zaburzeń snu polega na próbie walki ze stresem. Ale nie zawsze tak jest. Zaburzenia snu są dość częste i mają bardzo różne przyczyny. We wspomnianej rozmowie **Karol Świerc** wspomniiał o swoich bardzo dawnych planach budowy monitora aktywności mózgu (wykrywania, zobrazowania, analizy fal EEG). Później napisał też: (...) *Wątpię, żebym coś sensownego wymyślił, skoro nie wymyśliłem przez 45 lat. Są dwa główne problemy (...) do przeskoczenia. Sygnał EEG jest bardzo słaby. W typowym mieszkaniu to sygnał poniżej szumu. (...) ekranowanie łóżka [raczej] jest nierealne. Pomieszczenia, w których wykonuje się rejestracje EEG mają w ścianach siatki [ekranujące]. Nie miałem okazji i możliwości zapoznania się z jakimś schematem (...) rejestratora encefalograficznego. (...) Druga sprawa to „problem mechaniczny”: umocowanie elektrod. Nie*

mogą to być tak płaskie elektrody jak przy EKG. [Wtedy] doszedłem do wniosku, że pomysł nie ma szans na realną realizację (...) I... niestety, poszedłem w farmakologię. Może teraz są znacznie lepsze elementy, niż wtedy (...) Są wzmacniacze operacyjne (...) Gdyby wrócić to tematu, trzeba by się wzorować na jakichś rozwiązaniach aparatury medycznej. Nie ma szans, żebym wymyślił coś lepszego, a przynajmniej niewiele gorszego. (...) Tak w tej chwili o tym myślę.

Po 45 latach mamy do dyspozycji różnicowe wzmacniacze pomiarowe o CMRR ponad 150dB i mnóstwo informacji w Internecie (*EEG schematic*). Zadanie jest więc realne! Warto się tym zainteresować! Oczywiście EEG to tylko jeden z kierunków badań oraz prób diagnozy i leczenia. Problem zaburzeń snu często związany jest z chrapaniem i bezdechem. A tu otwiera się szerokie pole do popisu dla współczesnego elektronika. W każdym razie mamy bardzo praktyczny i ważny dla wielu Czytelników i ich bliskich temat zadania 311:

Zaproponuj wykorzystanie elektroniki dla dobra osób z zaburzeniami snu, np. chrapaniem czy bezdechem.

W grę wchodzi zarówno sposoby diagnozowania i wykrywania problemu, jak też

niosące ulgę. Dość proste jest to w przypadku chrapania, bo można wykorzystać mikrofon (kamerę?) oraz układ lub program, który wykryje dźwięk chrapania i da sygnał do układu wykonawczego.

Uwaga!

Każdy Autor, nadsyłając rozwiązanie zadania głównego, może dołączyć też swoją fotografię (portret). Fotografia zostanie opublikowana w artykule, omawiającym nadesłane rozwiązania.

Ten układ wykonawczy ma za zadanie delikatnie skłonić chrapiącego, żeby zmienił pozycję (przekręcił się na bok). Mógłby to być sygnalizator akustyczny (brzęczyk), wibrator, a może lepiej mocowany na ramieniu generator niezbyt silnych impulsów elektrycznych? Analogicznie jest w przypadku bezdechu, tylko trudniejsze jest wykrycie bezdechu, a raczej poprzedzających go anomalii rytmu oddychania.

Proponuję, żeby osoby, których problem dotyczy, po pierwsze poszukały dodatkowych informacji oraz by przeprowadziły testy. Czekam też na wszelkie rozwiązania teoretyczne. Zachęcam do udziału w tym specyficznym, ale jakże praktycznym zadaniu!

Piotr Górecki

Nadsyłajcie propozycje zadań!

Autorzy propozycji zadań, które zostaną wykorzystane w Szkole, otrzymują jako nagrodę kupon 100zł na zakupy w sklepie AVT:

www.sklep.avt.pl

Koszty przesyłki pokrywa AVT.

Dobra propozycja nie powinna być ani zbyt trudna, ani zbyt ogólna, ani zbyt wąsko ukierunkowana.

Dobre zadanie Szkoły powinno mieć na tyle szeroki zakres, żeby mogli w nim wziąć udział zarówno doświadczeni elektronicy, jak i początkujący, w tym najmłodsi.

Zachęcam do nadsyłania propozycji następnych zadań Szkoły!

UWAGA! UWAGA! UWAGA! UWAGA! UWAGA! UWAGA! UWAGA! UWAGA! UWAGA!

Zachęcamy także Ciebie, drogi Czytelniku, żebyś w ramach działu „Wokół Arduino”

opublikował swoją realizację projektu lub artykułu związanego z platformą Arduino.

Chętnie zaprezentujemy na łamach EdW Twój własny projekt albo Twoją realizację projektu z Internetu, wykorzystującego dowolne moduły lub moduły rozszerzeń Arduino,

a także wartościowe artykuły, pokazujące rozmaite aspekty korzystania z tej interesującej platformy.

Bliższe informacje: www.elportal.pl/arduino, a w razie pytań i wątpliwości śmiało pisz: edw@elportal.pl



Rozwiązanie zadania głównego 306

Temat wrześnieowego zadania 306 brzmiał: **Zaproponuj interesujące, najlepiej nietypowe zastosowanie diod LED.**

Temat diod LED jest zawsze aktualny, ponieważ elementy te na naszych oczach przeżywają okres gwałtownego i spektakularnego rozwoju. Przede wszystkim oświetlają nasze mieszkania. Warto też zastanawiać się nad nietypowymi sposobami ich wykorzystania. Oto przegląd nadesłanych prac.

Jacek Konieczny z Poznania napisał: (...) Jeśli chodzi o nietypowe wykorzystanie diod LED, to miałem takie doświadczenie z przełomu lat 80. i 90. ubiegłego wieku. (...) Pod swoją pieczę miałem między innymi magnetofony, zarówno szpulowe (głównie Grundig z serii ZK...), jak i kasetowe (głównie MK232). W celu szybkiego diagnozowania niektórych uszkodzeń stosowałem diody LED. Między innymi boczniowałem bezpieczniki w tych magnetofonach diodami LED (z odpowiednimi opornikami redukcyjnymi). W razie przepalenia się bezpiecznika (...) dioda LED (...) sygnalizowała stan awaryjny. (...) Oczywiście w przypadku bezpieczników sieciowych (zmiennoprądowych) dioda LED była zbocznikowana diodą prostowniczą o przeciwniej polaryzacji. [Spośród] nowych pomysłów (...) [sygnalowe łącze optyczne] (...) Nie chodzi o kable optyczne czy światłowodowe (TOSLINK), ale o zwyczajne, sygnałowe kable elektryczne (np. koncentryczne), ale zakończone diodą LED. W gniazdku „odbiorczym” umieszczona byłaby fotodioda lub fototranzystor. Ponieważ dioda LED jest urządzeniem odwracalnym, tj. może pracować również jako fotodioda ([foto] zestaw linków na ten temat):

<https://bit.ly/3yQdep5>

<https://bit.ly/3pcUFYV>

<https://bit.ly/3qaIcEy>

<https://bit.ly/3yJe7jg>

(...) zarówno we wtyczce, jak i w gniazdku można byłoby umieszczać diody LED i cały proponowany system mógłby dwukierunkowo przysyłać sygnały. Proponowane rozwiązanie (...) eliminowałoby kłopoty z „konfliktem mas” (...) W przypadku pojawienia się jakiegoś przepięcia w torze nadawczym, zniszczeniu uległaby dioda, a samo przepię-

cie nie przeniosłoby się dalej (tj. poza gniazdko); dioda LED byłaby zatem elementem „poświęcanym” (...)

Proponowałbym również wykorzystanie diod LED do sprawdzania gniazdek typu „duży jack”. Chodziłoby o sprawdzenie, czy jest to gniazdko MONO, czy raczej gniazdko STEREO. Głowica próbna (wsuwana do testowanego gniazdzka) zawierałaby zarówno diodę LED, jak i fotodiodę. Fotodioda byłaby ukierunkowana tak, aby odbierać światło odbite od metalowej tulei gniazdzka. Specjalny układ zliczałby przerwy w odbijanym świetle, tj. zliczałby pierścienie izolujące (zazwyczaj wykonane z ciemniejszego niż metal materiału izolacyjnego) obecne wewnątrz gniazdzka. (...) Odpowiednia duża liczba zliczeń wskazywałaby na gniazdko typu STEREO. (...)

jeszcze jedno (...) również chodzi o bolec i gniazdo, a ściślej – o sposób ich wzajemnego „centrowania”. Skąd się wziął taki pomysł? Otóż niedawno dostałem w prezencie licznik rowerowy. Podczas dłuższej wycieczki rowerowej wokół pewnego jeziora ów czujnik przestał działać, tj. przestał pokazywać aktualną prędkość i przebytą drogę. Prawdopodobnie wskutek wstrząsów obsunął się nieco wzdłuż szprychy magnes generujący impulsy czujnika. W używanym przeze mnie liczniku rowerowym nie istnieje żaden sposób wzajemnego centrowania magnesu i czujnika. Założmy, że w magnecie wydrążony jest niewielki otwór, a raczej – niewielka „wnęka” i w tej wnęce umieszczona jest dioda świecąca (pełniąca funkcję „bolca”). Układ odbiorczy zamontowany jest wraz z czujnikiem impulsów. Proponowana konstrukcja wykorzystuje fakt, że stosuje się „zwykłą” diodę LED, a nie diodę laserową, a zatem emitowana wiązka światła jest dostatecznie rozbieżna. Wiązka ta jest na tyle wąska, że wpada przez metalowy pierścień (pełniący funkcję zwierciadła „pierścieniowego”) do wnętrza układu odbiorczego. Układ odbiorczy stanowią dwa fotorezystory współpracujące z odpowiednio ukształtowanym układem optycznym. Na ten układ optyczny składałyby się: „wewnętrzny” odbiornik światła oraz „zewnątrzny” odbiornik światła. Owym odbiornikiem wewnętrznym byłaby nie-

duża soczewka, umieszczona centralnie wewnątrz otworu (wnęki) czujnika, skupiająca zebrane światło diody LED na fotorezystorze „wewnętrznym”. Układ „zewnątrzny” miałby podobną konstrukcję z tym, że dysponowałby dużo większą soczewką i byłby umieszczony za układem „wewnętrznym”. Oznacza to, że środek układu „zewnątrznego” byłby zasłonięty przez układ „wewnętrzny”, a obrzeże dużej soczewki układu „zewnątrznego” kierowałoby najbardziej odchylone „na zewnątrz” promienie diody LED i skupiałoby je na fotorezystorze „zewnątrznym”. Fotorezystory: wewnętrzny i zewnętrzny połączone byłyby w układ mostka Wheatstone’a. Kiedy dioda LED byłaby umieszczona „centrycznie”, tj. dokładnie na wprost środka układu odbiorczego, to promienie światła tej diody nie odbijałyby się od wewnętrznej powierzchni pierścienia „zwierciadlanego” i wpadałyby wyłącznie do czujnika „wewnętrznego”; mostek oporowy (zawierający oba fotorezystory) byłby wówczas w skrajnej nierównowadze i byłby to stan (jak najbardziej) „prawidłowy”. Kiedy dioda świecąca LED nie byłaby umieszczona „centrycznie” albo gdyby znajdowała się zbyt daleko od układu (pary) czujników fotoelektrycznych, to światło diody odbijałoby się również od wewnętrznej powierzchni „zwierciadła pierścieniowego” i oświetlałoby również „zewnątrzny” czujnik fotoelektryczny. Mostek oporowy, zawierający oba fotorezystory, znajdowałby się wówczas w stanie „zbliżonym do równowagi”. Skrajną nierównowagę owego mostka można byłoby wykrywać, łącząc przekątną owego mostka z przetrzutnikiem (bramką) Schmitta. Zatem proponowany układ wykrywałby dwie sytuacje niekorzystne dla licznika rowerowego, tj. zarówno „niecentryczne” usytuowanie magnesu, jak i zbyt dużą odległość magnesu od głowicy czujnikowej. Pozdrawiam.

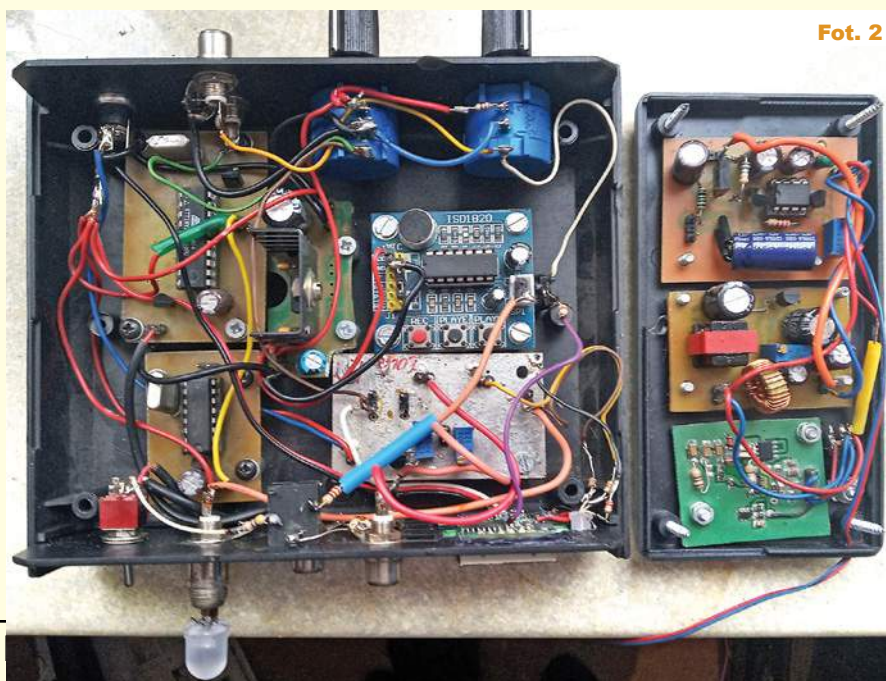
Rafał Orodziński zaproponował: (...) Nietypowe zastosowanie diod LED. (...) **Uprawa roślin.** Na parapecie przygotowane na okres zimowy stanowisko doświetlania sadzonek roślin (...) Stelaż wykonany z prętów zbrojeniowych. Konstrukcja wymaga pomalowania (na niepomalowanej konstrukcji zaczęły poja-



Fot. 1

wiać się ślady rdzy i co znacznie gorzej – także na parapecie) (...) Do doświetlania zastosowano lampy do uprawy roślin emitujące światło czerwone i niebieskie [fotografia 1] (...) Lampy sterowane są za pomocą sterownika czasowego. (...) Lepiej jest stosować układ składający się z większej liczby diod mniejszej mocy niż z kilku diod dużej mocy ze względu na większą równomierność oświetlenia. (...) **Telekomunikacja.** (...) Diody LED mogą pracować jako zwykłe fotodiody oraz jako połączenie elementu detekcyjnego i wzmacniającego, czyli jako fotodioda lawinowa. Fotodetektory ze wzmocnieniem lawinowym (...) są najczulszymi detektorami półprzewodnikowymi. Aby dioda LED wykazywała zjawisko wzmocnienia lawinowego, musi zostać do niej przyłożone napięcie wsteczne od kilkudziesięciu do stu kilkudziesięciu woltów w zależności od diody LED. Przetwornica do zasilania diod LED wykorzystuje bardzo tani i popularny układ scalony MC34063. Na uwagę zasługuje wykorzystanie transformatora podwyższającego. Transformator taki uzyskano z zasilacza telefonu komórkowego, ale działa on w tym układzie „w drugą stronę”. Układ taki ma znacznie większą sprawność niż przetwornice dławikowe. Po przetwornicy następuje stabilizator ciągły w celu redukcji tęt-

nień na wyjściu przetwornicy. W tej samej obudowie zawarta jest przetwornica dławikowa do 50V – prawa strona [fotografia 2]. Korzystną cechą diod LED jako fotodetektorów jest (...) [selektywność], przez co jest ona niewrażliwa na światło o innej barwie – częstotliwości. Taka dioda działa jako filtr pasmowo-przepustowy. W łącznościach na większe odległości stosuje się światło czerwone, gdyż znacznie mniej się rozprasa niż światło o krótszej długości fali, co jest bardzo korzystne w przypadku łączności na duże odległości. Bardzo ważne jest, by dioda LED użyta w nadajniku generowała światło o mniejszej długości fali [o większej energii] niż LED użyta jako fotodioda. Jako diodę odbiorczą wykorzystałem czerwone diody LED o długości fali światła 660nm (diody LED stosowane w uprawie roślin), a jako nadajnik diody LED o długości fali 630nm (typowe czerwone diody LED). Na fotografii 2 pokazany został mikronadajnik z diodą LED do strojenia odborników, umożliwia on modulację diody sygnałami audio oraz video (lewa strona). Obecnie w opracowaniu jest układ modulatora dużej mocy pracującego liniowo (regulowane źródło prądowe) i PWM z wykorzystaniem układu TL494 – fotografia 3. Część nadawcza będzie wykorzystywała dziewięć diod



Fot. 2

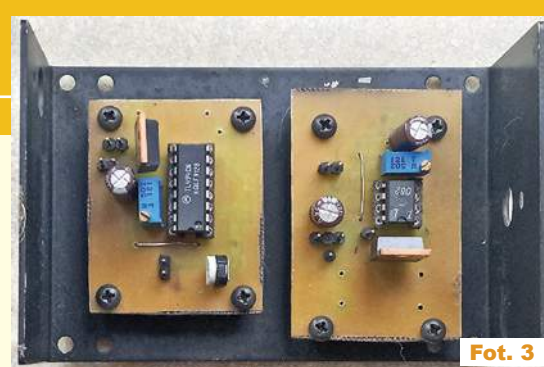


Co z tą energią?

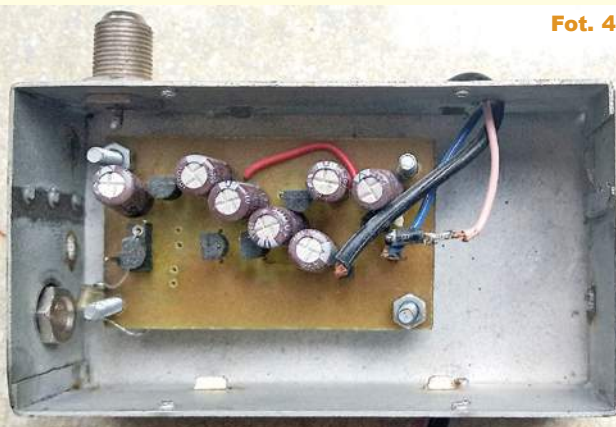
Sen o mocy – czystej i taniej

100 albo nawet 150 bilionów dolarów globalnie kosztować może w ciągu najbliższych trzech dekad przejście z tradycyjnej, opartej na emisjach związków węgla gospodarki do ekonomii klimatycznie neutralnej. Tymczasem w świecie alternatywnych, odnawialnych źródeł energii pełno jest pytań, wyzwań, ślepych zaułków i pułapek. Są też szanse, z których jakoś nie chce się korzystać.

Nowy numer już w sprzedaży
www.ulubionykiosk.pl
 Koniecznie odwiedź serwis:
mlodytechnik.pl



LED 3W wyposażonych w kolimatory służące do skupienia wiązki. Jeden z prototypowych odbiorników pokazany jest na **fotografii 4**. Na świetle wykonywano łączności za pomocą diod LED na ponad 200km z wykorzystaniem odbicia światła od chmur, a także transmitowano sygnał telewizyjny na ponad 50 kilometrów. Mogłbym przedstawić w EdW cały cykl artykułów z opisaniem łączności wykorzystującej fale świetlne (...)



Zygmunt Flisak z Opola zaproponował następujące nietypowe rozwiązanie: *Dzień dobry*. Interującym rozwiązaniem stanowiącym połączenie XX-wiecznej techniki próżniowej i bardziej współczesnych (choć także z ubiegłego wieku) osiągnięć technologii ciała stałego jest wykorzystanie diody LED do polaryzacji siatki lampy elektronowej (triody) w przedwzmacniaczu. Pomysł przedstawiono w książce „Designing High-Fidelity Tube Preamps” autorstwa M. Blencowe opublikowanej w 2016 r. Wydaje mi się, że warto go przytoczyć w kontekście nietypowych zastosowań diod świecących. Zasadniczo trioda powinna pracować przy ujemnym napięciu siatki względem katody. Kiedy przez lampę płynie prąd, to rezystor o odpowiednio dobranej wartości umieszczony w obwodzie katody powoduje podniesienie potencjału tej elektrody względem masy, do wartości równej spadkowi napięcia na rezystorze. Tym samym zbędne staje się dodatkowe źródło ujemnego napięcia polaryzującego siatkę. Takie postępowanie nosi nazwę polaryzacji automatycznej.

Punktacja Szkoły Konstruktorów



Sławomir Węgrzyn Dziekanówce	92	Sebastian Jarmosiewicz Motwica	50	Lukasz Nowak Warszawa	33
Daniel Turbasa Kraków	88	Michał Pędziński Stara Słupia	48	Jacek Konieczny Poznań	26
Michał Stach Kamionka	88	Lukasz Olszok Tarnowskie Góry	45	Piotr Grzegorz Siedlce	25
Lukasz Dachowski Cymbark	72	Krzysztof Kawa Lubcza	44	Andrzej Nowicki Warszawa	24
Artur Beret Barcin	69	Dawid Placha Rdzawa	44	Marian Gabrowski Polkowice	23
Aleksander Bernaczek Magnuszowice	69	Szymon Czepiel Piszczowice	43	Roman Braumberger Bytom	21
Krzysztof Smoliński Poznań	68	Piotr Gajdosz Grybów	41	Jakub Gajda Kraków	20
Szymon Trygar Szczecin	66	Maciej Zieliński Kraków	41	Jacek Ręčka Polonia	20
Radosław Smalec Zabrze	64	Circuit Chaos Warszawa	41	Marian Caruk Luban	17
Paweł Hoffmann Wrocław	62	Rafał Rówiak Słaboszów	40	Bogdan Kosiński Szczecin	16
Robert Szolc Bytom	58	Rafał Orodziński Białystok	36	Lukasz Kojro Gdańsk	15
Andrzej Herbut Siekierzyn	52	Teodor Woźniak Łódź	35	Marcin Malich Wodzisław Śl.	13
Adam Ples Jaworzno	51	Tomasz Zaorski Kalinówka	34	Paweł Sablik Piszczowice	13
Adam Sobczyk Warszawa	50	Jarosław Węgliński Warszawa	34	Piotr Wyderski Wrocław	13

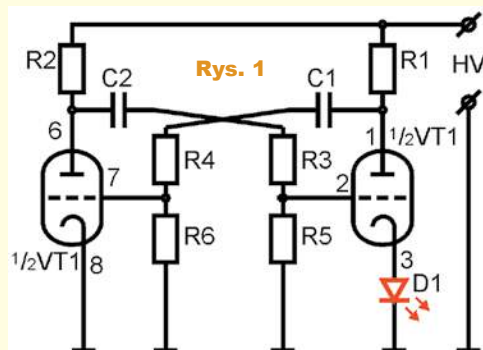
Tymczasem zamiast rezystora katodowego można użyć diody świecącej. W tym przypadku dodatni potencjał katody względem masy będzie równy napięciu przewodzenia elementu półprzewodnikowego. W zależności od długości fali emitowanego promieniowania (lub bardziej precyzyjnie: rodzaju półprzewodnika) mamy do dyspozycji napięcia od 1,0 do 4,0V dla jednej diody LED. Ponieważ wiele lamp odbiorczych pracuje przy prądzie katody rzędu kilkunastu miliamperów, ryzyko przekroczenia dopuszczalnego prądu diody i jej uszkodzenia jest znikome. Dodatkowo atutem tego rozwiązania jest świetlna sygnalizacja przepływu prądu przez triodę i czasem tylko z tego powodu warto je rozważyć.

W tym kontekście z powodzeniem wykorzystalem czerwoną diodę LED do sygnalizacji działania prototypu multiwibratora astabilnego zbudowanego w oparciu o lampę ECC83 (**rysunek 1**). Podobny schemat przedstawiłem kilka lat temu na forum portalu trioda.pl, lecz wtedy dyskusja dotyczyła ograniczenia maksymalnego ujemnego napięcia na siatce. Podczas eksperymentowania należy mieć świadomość, że w układach lampowych z reguły panują napięcia niebezpieczne dla zdrowia i życia. Pozdrawiam.

Andrzej Nowicki z Warszawy napisał: *Dzień dobry*. Dziś rozwiązaniem z serii miniaturowych. (...) Warto zauważyć, że charakterystyka prądowo-napięciowa diody LED w obszarze przewodzenia trochę przypomina charakterystykę diody Zenera (...) Dioda LED może więc służyć jako element obniżający napięcie w obwodzie, w którym jest włączona, o wartość zależną od koloru świecenia – typowo 1,4V do 3V. Prościutkie zastosowanie

– tester baterii 1,5V. Składa się ze wskaźnika poziomu od jakiegoś starego magnetofonu, pojemnika na baterię i czerwonej diody LED włączonej w kierunku przewodzenia, dobranej ze starych zapasów tak, aby prawie wyladowana bateria (1,3V) powodowała minimalne wskazanie miernika, prawie dobra (1,45V) w okolicach zera miernika, a zupełnie nowa powodowała wychylenie wskazówki aż do końca, co widać na **fotografii 5**. Sprawdzenie nawet kilkudziesięciu baterii zajmuje tylko kilka minut. (...) Z najlepszymi pozdrowieniami dla Redakcji.

Michał Stach z Kamionki Małej napisał: (...) *Dzień dobry*. Przesyłam pierwszą część SzK 306 nietypowe wykorzystanie LED. W opisywanym przypadku będzie raczej typowe – [wykorzystanie] do oświetlenia ulicy/ogrodu. (...) lampę charakteryzuje naprawdę dobre wykonanie. Zastosowane materiały: plastik zbliżony do ABS, panel fotowoltaiczny przykryty szkłem i diody LED COB zasilane prądem stałym (bez efektu migotania). Do tego czujnik ruchu i tryby pracy nastawiane pilotem zachęciły do zakupu wielu klientów. Niestety w eksploatacji lampy nie spisuje się tak rewelacyjnie.





Wielu użytkowników skarży się na małą moc świetlną i krótką pracę po zmierzchu. Przyczyna jest dosyć oczywista: zasilanie lampy w całości pochodzi z wewnętrznego akumulatora LiFePO4 ładowanego panelem PV. (...) moc zamontowanego ogniwa PV to maksymalnie 4Wp, a lampa sprzedawana jest jako źródło światła 100W!!!. Do tego energia z panelu marnowana jest po części przez brak obwodu MPPT. Akumulator został podpięty wprost do ogniwa przez diodę zapobiegającą cofaniu się prądu, co sprawia, że zamiast optymalnego napięcia 5,7V, na ogniwie PV mamy ~3,6V, zmniejszając tym samym użyteczną moc o 30%. Zadanie [jest] w trakcie realizacji (...) Pozdrawiam.

Niestety, do chwili oddania materiałów do druku nie udało się zakończyć analizy błędów i przeróbki. Mam nadzieję, że Michał upora się z niełatwym zadaniem i przedstawi efekty w następnym numerze.

A na koniec zaległości. **Paweł Hoffmann** z Wrocławia napisał tak: (...) z opóźnieniem chciałbym przedstawić rozwiązanie zadania dotyczące superkondensatorów. 3 sztuki wykorzystałem do naprawy pokazanego na **fotografii 6** termometru zaokiennego, zasilanego [umieszczonym na tylnej ściance] panelem słonecznym. Ten gadżet reklamowy kupiłem na giełdzie elektronicznej za drobne, zapewne 5 zł. Ponieważ termometr powinien działać nie tylko, gdy panel solarny jest oświetlony promieniami słonecznymi, a właściwie to szczególnie wtedy, gdy nie jest w słońcu, bo wtedy jego wskazania są zawyżone, wymaga akumulatora. W tym egzemplarzu akumulator był uszkodzony, wylany. Było to już wiele lat temu, więc nie pamiętam, jakiego typu był to akumulator, ale najprawdopodobniej typu pastylkowego. Postanowiłem go zastąpić czymś bardziej odpornym, co nie zawiedzie po wielu cyklach ładowania i rozładowania. Czyli właśnie superkondensatorem. Zmierzyłem napięcie wyjściowe panelu fotowoltaicznego w pełnym słońcu. Wynosi ono 6,2V, czyli zastosowanie pojedynczego kondensatora odpada. Napięcie dwóch połączonych szeregowo sztuk też niestety jest za małe. Dopiero 3 sztuki połączone szeregowo dają napięcie 8,1V, czyli odpowiednie, z dużym zapasem. Nie mierzyłem poboru prądu termometru. Pojemność 1F została wybrana „na oko” i żeby nie inwestować za dużo w termometr, który sam był bardzo tani. Oczywiście przy połączeniu szeregowym pojemność wypadkowa wynosi około 0,3F. Jest ona w zupełności wystarczająca. Nawet w krótkie zimowe dni ilość zgromadzonej energii jest wystarczająca, by termometr działał do rana. Napięcie nie spada na tyle, by wyświetlacz LCD tracił

Fot. 6



Publikacja	Nagroda	Talon AVT PLN	Imię	Nazwisko	Miejscowość	Punkty
-	-	100	Karol Świerc za pomysł zadania 311			-
-	U	-	Jacek	Konieczny	Poznań	2
?	-	200	Rafał	Orodziński	Białystok	8
-	U	-	Zygmunt	Flisak	Opole	5
-	U	-	Andrzej	Nowicki	Warszawa	5
?	-	-	Michał	Stach	Kamionka Mała	?
-	-	-	Paweł	Hoffmann	Wrocław	4

kontrast. **Fotografia 7** pokazuje kondensatory zamontowane we wnętrzu termometru. (...) spędził za oknem już wiele lat. Naprawa była skuteczna i działa cały czas niezawodnie. (...) wykorzystałem zaletę będącą przewagą nad akumulatorami – wytrzymałość na bardzo wiele cykli ładowania i rozładowania. Przy codziennym cyklu i wielu latach użytkowania, powtórzeń było już mocno ponad tysiąc albo i dwa. Co dla akumulatorów byłoby już wartością graniczną, a dla kondensatora zupełnie marginalną. Pozdrawiam.

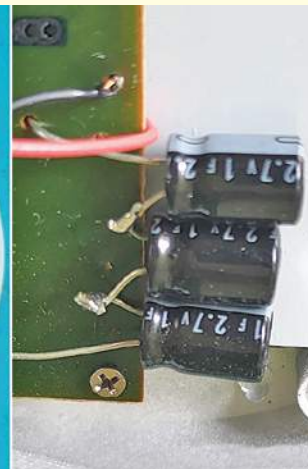
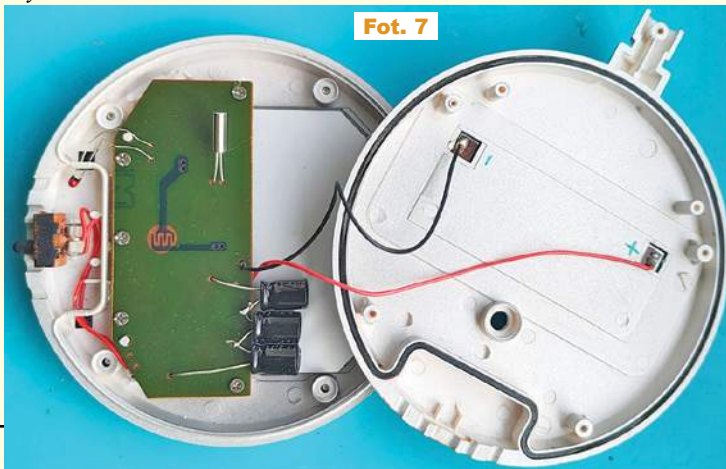
Jak widać, nie ma tu żadnego balansera, a mimo to zestaw trzech kondensatorów pracuje bezawaryjnie kilka lat. Jest to możliwe, jeżeli parametry wszystkich trzech są jednakowe. Przy znaczących różnicach, nie tyle pojemności, co upływności, bez balansera łatwo o niesymetrię i uszkodzenie. Warto co jakiś czas zmierzyć napięcie na kondensatorach i podładować ten o najniższym napięciu.

Aktualne informacje o punktacji oraz rozdziale nagród, upominków i kuponów podane są w tabelkach. Znak zapytania oznacza, że ewentualna publikacja nastąpi dopiero po nadesłaniu ostatecznych materiałów. Osoby nagrodzone kuponami otrzymują z naszej redakcji stosowny e-mail z informacją i wskazówkami, a dopiero potem zamawiają w sklepie AVT (wrzucają do koszyka pod adresem www.sklep.avt.pl) towary za przydzieloną sumę, a w uwagach piszą, że jest to kupon ze Szkoły Konstruktorów. Kupon z zadania z kolejnych miesięcy można sumować, by kupić sprzęt o większej wartości. Istnieje też możliwość dopłaty różnicy cen w przypadku zamówienia na sumę większą niż przydzielony kupon. Ale **uwaga: kupon ważny jest tylko 12 miesięcy – po tym terminie traci ważność i przepada.**

Serdecznie zapraszam do udziału w zadaniu głównym 311, a także w drugiej i trzeciej klasie naszej Szkoły Konstruktorów! Zachęcam uczestników, żeby praktyczne rozwiązania zadań Szkoły przygotowywali według Szablonu ze strony: <http://elportal.pl/zostan-wspolautorem-elektroniki-dla-wszystkich/>

Piotr Górecki

Fot. 7



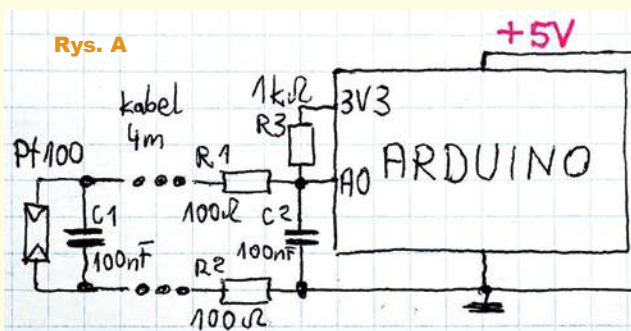


Co tu nie gra? Zadanie 311

Na **rysunku A** przedstawiony jest schemat układu pomiarowego.
Jak zwykle pytanie brzmi:

Co tu nie gra?

Nawet gdy w układzie jest kilka usterek, możesz zgłosić tylko jedną. Bardzo proszę o możliwie krótkie odpowiedzi.



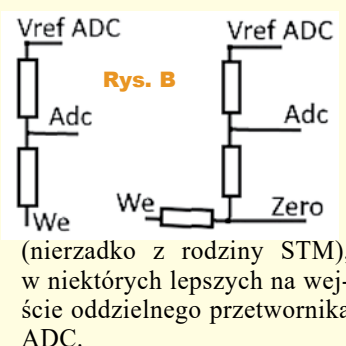
Odpowiedź oznacz **NieGra311** i nadeślij w terminie 60 dni od ukazania się tego numeru EdW. Od razu podaj też swój adres pocztowy, żebym nie musiał pytać, gdy przydzielę upominek. Możesz jeszcze przysłać rozwiązanie zadania *NieGra* z poprzedniego miesiąca. Uczestnicy konkursu otrzymują upominki, a najaktywniejsi uczestnicy są co rok nagradzani bezpłatnymi prenumeratami EdW lub innego wybranego czasopisma AVT.

Co tu nie gra? Rozwiązanie zadania 306

Na **rysunku B** pokazany jest schemat, zamieszczony w EdW 9/2021, który był tam opatrzone komentarzem: *Po artykułach o modułowych miernikach napięcia stałego jeden z Czytelników napisał: (...) Zastanawia mnie, dlaczego nie jest stosowane proste rozwiązanie [według rysunku B]: nie ma problemu z pomiarem napięć bliskich zera, a dodatkowo można mierzyć napięcia ujemne. Można też prosto zrobić autokalibrację zera przez zwarcie wyprowadzenia ZERO do masy. Wystarczy, aby ZERO było podłączone do wyjścia uP typu OpenDrain (...)*

Na standardowe pytanie *Co tu nie gra?* odpowiedziało niezbyt wielu Czytelników. Tym razem nie ja byłem autorem schematu, więc nie wprowadziłem do schematu kilku rozmaitych usterek, jak to bywało we wcześniejszych zadaniach. Dlatego w tym zadaniu należało się skoncentrować na jednym tylko zagadnieniu.

Kluczową sprawą było stwierdzenie, że prezentowane schematy pojawiły się po opublikowaniu w EdW artykułów o modułowych miernikach prądu stałego, które zasilane są pojedynczym napięciem (+3,5...+12V) i z reguły mierzą tylko napięcia dodatnie względem masy. Na wejściach woltomierzy tego typu z reguły umieszczony jest typowy rezystorowy dzielnik napięcia. A osłabiony sygnał z dzielnika podawany jest na wejście przetwornika ADC. W tańszych wersjach jest to przetwornik ADC wbudowany w mikrokontroler

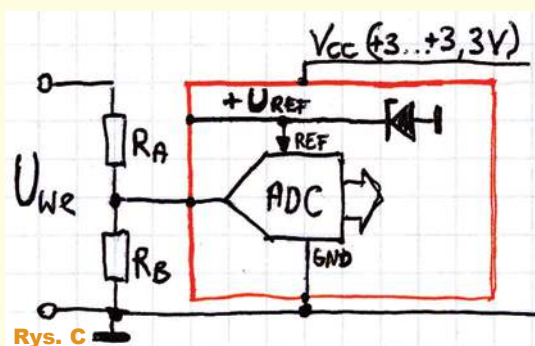


(nierzadko z rodziny STM), w niektórych lepszych na wejście oddzielnego przetwornika ADC.

Ogólnie biorąc, konfiguracja jest taka, jak pokazuje **rysunek C**. Przetworniki ADC zawarte w mikrokontrolerach mierzą tylko napięcia dodatnie, od zera (masa) do napięcia odniesienia U_{REF} . Oznacza to, że tak zbudowany przyrząd może mierzyć tylko napięcia dodatnie o takiej wartości maksymalnej, która na wyjściu dzielnika R_A , R_B da napięcie równe U_{REF} .

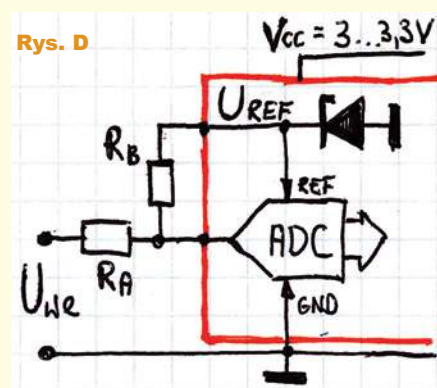
Napięcie U_{REF} w najtańszych miernikach tego rodzaju może być po prostu napięciem zasilania (3V albo 3,3V), ale może też być pobierane z wewnętrznego źródła napięcia odniesienia, o ile mikrokontroler je ma i o ile tak zdecyduje twórca modułu woltomierza.

W każdym razie w klasycznym rozwiązaniu według **rysunku C** nie można mierzyć napięć ujemnych, tylko dodatnie, do wartości maksymalnej, wyznaczonej głównie przez dzielnik rezystorowy. Tą wartością maksymalną może być przykładowo napięcie +100V lub jeszcze wyższe.



Pomysłodawca zadania *NieGra306* pisze: *(...) Zastanawia mnie, dlaczego nie jest stosowane proste rozwiązanie [według rysunku B]: nie ma problemu z pomiarem napięć bliskich zera, a dodatkowo można mierzyć napięcia ujemne (...)*

Z tych słów wynika, że z lewej strony **rysunku B** pokazana jest propozycja lepszego rozwiązania – lepszej koncepcji. Koncepcji według **rysunku D** albo



według **rysunku E**, gdy nie ma wyprowadzonego na zewnątrz napięcia U_{REF} .

Nie innej, odmiennej, tylko koncepcji ulepszonej, która ma mieć dodatkowe zalety względem klasycznej z rysunku C. Dodatkową zaletą ma tu być możliwość pomiaru napięć ujemnych. A możemy się domyślać, że zalety rozwiązania klasycznego według rysunku C zostaną zachowane.

Otóż nie!

W Internecie bez problemu można znaleźć propozycje rozwiązań według rysunku D czy E. Koncepcja ta jest znana i...

bardzo rzadko wykorzystywana.

I tu dochodzimy do odpowiedzi: **co tu nie gra?**: otóż w rozwiązaniu według rysunków B, D, E można wprawdzie mierzyć niemal dowolnie duże napięcia ujemne (zależnie od stopnia podziału dzielnika), ale **nie można mierzyć napięć dodatnich większych niż U_{REF}** .

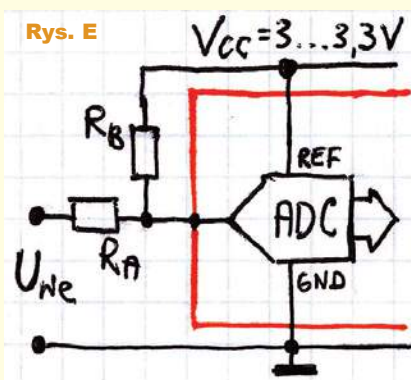
Autor rysunku B zastanawiał się, dlaczego nie jest stosowane zaproponowane przez niego rozwiązanie. Teraz już znamy odpowiedź: Napięcie U_{REF} w niektórych mikrokontrolerach może być rzędu 1V, a na pewno nie będzie wyższe od napięcia zasilania V_{CC} , wynoszącego co najwyżej 3,3V, więc **woltomierz zbudowany według koncepcji z rysunku D na pewno nie może mierzyć napięć dodatnich wyższych niż 3,3V**, a to powoduje, że **jest bezużyteczny w większości zastosowań**.

I to jest podstawowy problem.

Niektórzy Koledzy zwrócili też uwagę na inne kwestie.

Ale tylko jeden z uczestników, który pierwszy raz wziął udział w konkursie, napisał: *Witam (...) szukamy błędów (...) moim zdaniem nie jest to duży błąd, ale w spoczynku na wejściu woltomierza będzie występować napięcie V_{REF} (...) po dołączeniu do układu woltomierz będzie wstrzykiwał prąd do [mierzonego obwodu] (...) Prąd ten zależy od wartości rezystorów dzielnika (...) nie jest duży, aczkolwiek w zwykłych woltomierzach tego nie ma. Wejście woltomierza nie wstrzykuje prądu, tylko jest rezystancją, w cyfrowych mulimetrach dużą 1 albo 10 megomów (...)*

Jeden ze stałych uczestników napisał: *(...) Rozwiązanie z rysunku A nie jest stosowane w prostych cyfrowych woltomierzach DC ze względu na stosowane mikroprocesory, które nie mają wyprowadzonego pinu V_{REF} napięcia odniesienia. W procesorach*



tych napięcie zasilania jest używane jako napięcie odniesienia. Są to niedoskonałe elementy, obwody pracujące niestabilnie (...)

Inny Kolega stwierdził: *(...) Pomiar napięć mniejszych niż 0V jest możliwy. Tego typu rozwiązania wiążą się jednak ze zbyt niską opornością wejściową przyrządu (rzędu kOhm), co wpływa na duże błędy pomiaru. Tego typu rozwiązania od lat wykorzystywane są w woltomierzach opartych na mikrokontrolerach AVR. Wykorzystuje się tu proste zależności wynikające z praw Kirchhoffa, Ohma, oraz czego nie ukrywam, dosyć skomplikowane obliczenia pozwalające dobrać optymalną dla ADC wartość podzespółów w dzielniku (...)*.

Wspominaliście też, bardzo słusznie zresztą, że koncepcja, pokazana z lewej strony rysunku B, jest niepełna, fragmentaryczna i nie obejmuje modułowych woltomierzy, które nie wykorzystują wewnętrznego ADC, wbudowanego w mikrokontroler, tylko zawierają ADC w postaci oddzielnego układu scalonego.

Dokładniejsze przetworniki ADC mają i różnicowe wejście sygnału mierzonego, i różnicowe wejście napięcia odniesienia U_{REF} , co ogromnie rozszerza możliwości zmian i rzeczywiście ulepszeń.

W praktyce w lepszych modułach do pomiaru napięć i prądów stałych są wykorzystywane albo chińskie przetworniki ADC o oznaczeniu BX5815, do których dokumentacji nie sposób dotrzeć, albo 18-bitowe przetworniki MCP3421, które tajemnic nie mają.

MCP3421 mają różnicowe wejście mierzonego napięcia, ale nie mają wyprowadzonego wewnętrznego napięcia odniesienia (2,048V).

Na koniec warto poświęcić parę słów komentarza schematowi z prawej strony rysunku B. Otóż we wszel-

kich miernikach bardzo istotna jest kwestia „stabilności zera”. Z różnych względów nie jest oczywiste, że przy zerowym napięciu wejściowym wynikiem pomiaru będzie zero. To naprawdę szeroki temat, wykraczający daleko poza konkurs *Co tu nie gra?* Konstruktorzy przetworników ADC i kompletnych mierników (woltomierzy) poświęcają dużo uwagi tej kwestii. A ma ona kilka aspektów.

Ogólnie biorąc, standardowe przetworniki (SAR) wbudowane w mikroprocesory mają gorsze właściwości na obu krańcach zakresu pomiarowego. Prosty woltomierz według rysunku C ma więc pogorszone parametry przy napięciach bliskich zero. Pod tym względem lepsza jest wersja z rysunków D, E, ponieważ przy zerowym napięciu U_{we} przetwornik ADC nie pracuje na samym końcu zakresu.

Z jednej strony jest to korzystne, ale w przypadku najtańszych rozwiązań wypadałoby też zwrócić uwagę na stabilność wskazań zera. I właśnie z prawej strony rysunku B mamy koncepcję, pozwalającą procesorowi okresowo wykonać kalibrację zera.

Kwestia „stabilności zera” ma jeszcze inne oblicze w przetwornikach ADC z wejściem różnicowym, bowiem tam napięcie wejściowe równe zero to nie kraniec, tylko połowa zakresu. Ten wątek też nie mieści się w ramach konkursu *NieGra*, ale po części zajmujemy się tym w artykułach opisujących modułowe mierniki zawierające przetwornik MCP3421.

Podsumowując zadanie *NieGra306*, należy stwierdzić, że podstawowym problemem jest to, że w proponowanej na rysunku B koncepcji nie można mierzyć napięć dodatnich wyższych od napięcia odniesienia, a drugą godną odnotowania niedoskonałością jest to, że w spoczynku na wejściu występuje napięcie, a przy małych napięciach mierzonych do badanego obwodu pompowany jest (niewielki) prąd.

Nagrody-upominki za zadanie *NieGra306* otrzymują:

Emil C. – Poznań,

Tadeusz Suszał – Warszawa,

Łukasz Kamiński – Jemiołowo.

Wszystkich uczestników dopisuję do listy kandydatów na bezpłatne prenumeraty.

Piotr Górecki



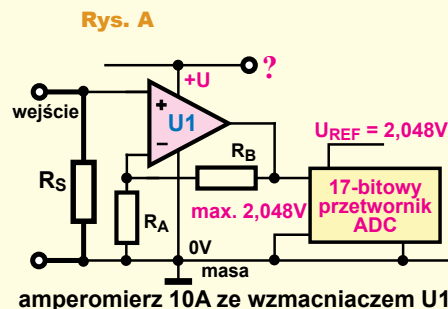
Policz – zadanie 311

Według rysunku A planujemy budowę dokładnego cyfrowego amperomierza z 17-bitowym przetwornikiem ADC o napięciu odniesienia 2,048V. Cyfrowy sygnał z przetwornika podamy na jakiś mikroprocesor, który przeliczy i przedstawi wynik na wyświetlaczu. Zakres pomiarowy projektowanego amperomierza ma wynosić 10A.

W ramach zadania *Policz311* należy:

- zaproponować wartości rezystorów, typ wzmacniacza operacyjnego U1 oraz wartość napięcia zasilania +U.

Zapraszam do udziału zarówno elektroników doświadczonych, jak i początkujących, którzy jeszcze nie potrafią przeanalizować wszystkich subtelności układu. Z uwagi na specyfikę zadania proszę o podawanie swojego wieku oraz miejsca nauki czy pracy.



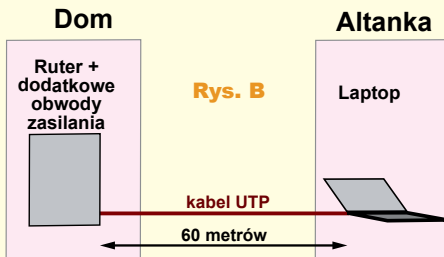
Odpowiedź nadesłaj w terminie 60 dni od ukazania się tego numeru EdW. Tytuł e-maila powinien zawierać nazwę konkursu i numer zadania oraz Twoje nazwisko (**Policz311_Nazwisko**). *Jeżeli chcesz uczestniczyć w podziale upominków, w e-mailu podaj od razu swój adres pocztowy.* Możesz też jeszcze przysłać rozwiązanie zadania *Policz* z poprzedniego miesiąca.

Policz – rozwiązanie zadania 306

W EdW 9/2021 przedstawione było zadanie *Policz306*, które brzmiało tak: *Z uwagi na nieskuteczność Wi-Fi, między budynkiem mieszkalnym a altanką w ogrodzie, gdzie nie ma zasilania 230V, położony został kabel UTP o długości 60m według rysunku B.* W kablu są dwie wolne pary (cztery żyły). W ramach zadania *Policz306* należy: – *wstępnie oszacować możliwość zasilania laptopa przez te wolne żyły kabla UTP.*

Nadesłane przez Was rozwiązania pokazały, że jak najbardziej warto było postawić takie właśnie zadanie. Pytanie, sformułowane w lakoniczny sposób, dotyczyło możliwości zasilania, więc odpowiedź w zasadzie powinna brzmieć: **tak** albo **nie**. Owszem, ale w grę wchodzi dodatkowe czynniki i tak jak to często bywa, odpowiedź brzmi: **to zależy**. Zależy od kilku czynników.

Nieprzypadkowo należało **wstępnie oszacować możliwość**, ponieważ dokładniejszych obliczeń przeprowadzić nie można. Nie znamy bowiem parametrów laptopa, a konkretnie nie znamy poboru mocy, ani średniego, ani maksymalnego. Nie zastanawiamy się też nad realnymi możliwościami obniżenia poboru mocy, choćby nad tym, czy w grę wchodzi zmniejszenie jasności ekranu w jasno oświetlonej altance podczas słonecznego lata. Jest też kilka innych ważnych szczegółów, które należałoby uwzględnić.



Niewątpliwie do wstępnego oszacowania trzeba było przyjąć konkretną, w miarę sensowną wartość mocy, jaką podczas pracy pobiera laptop. To była kluczowa kwestia, od której zależały wyniki dalszych rozważań.

Niektórzy z uczestników przyjęli moc laptopa, a inni – maksymalny pobierany prąd. Niektórych przyjęte założenia oraz rozważania szybko doprowadziły do wniosku, że **nie można**. Oto przykład: (...) Witam. Przewody skrętki komputerowej mają średnicę 0,5mm. Te 4 żyły (po 2 połączone równolegle) dają na długości 60m łącznie 5,256 oma. Gdybyśmy chcieli wykorzystać standardowy zasilacz do laptopa o napięciu wyjściowym 19V i o mocy ok. 60W, to maksymalny prąd 3A dałby spadek na przewodach prawie 16V i na wyjściu zostałoby parę woltów. Oczywiście przy mniejszym płynącym prądzie odpowiedni mniejszy spadek napięcia. Ogólnie, proponowane przewody są za małej średnicy, żeby móc przesyłać z małymi stratami napięciowymi, wymagany prąd zasilający laptop.

Większość uczestników obliczała rezystancję drutu oraz rozważała, jaki spadek napięcia wystąpi na niej pod wpływem płynącego prądu. Kable UTP klasy 5 (5e) mają średnicę około 0,5mm (AWG26... AWG23), co daje przekrój 0,2mm² oraz rezystancję jednej żyły i całego kabla zasilającego nieco ponad 5 omów. Każdy amper płynącego prądu wywoła więc na tej rezystancji spadek napięcia ponad 5V.

To nie jest zły kierunek analizy, ale nie uwzględnia problemu strat ciepłych w kablu i wzrostu temperatury.

Nieliczni Koledzy podeszli do zadania z innej strony. Jeden zastanawiał się, gdzie znaleźć *tabele, pokazujące relację prądu i wzrostu temperatury izolacji kabla*.

Pewien stały uczestnik napisał: (...) wskazówką jest maksymalna moc i prąd PoE. Kiedyś to było 350mA, teraz jest 600mA na parę. Jak na razie moc PoE 602.3bt dochodzi do 100W, ale tylko przy użyciu wszystkich czterech par kabla UTP. Przy dwóch parach jest 50W. (...) nie wiadomo, ile dokładnie pobiera laptop (...) jest na granicy (...) nie można jednoznacznie stwierdzić (...) może tak, a może nie (...)

Inny z uczestników zauważył: (...) Szacowana oporność pojedynczego przewodu skrętki o średnicy 0,5 i długości 60m to (...) około 5,14 oma. Do obliczeń przyjmę 5 omów (...) łączymy żyły po dwie równolegle, czyli wypad-

kowa oporność pętli zasilania wynosi 5 omów. Niby nie za wiele. Zasilacz do laptopa to od 3 do 5A maksymalnego prądu. (...) Zasilanie to standardowo 19V (...), czyli spadek napięcia na przewodach $I \times R$ to 25V. Robi się ciekawie. Na wejście trzeba podać 19+25V, czyli 44V, a widziałem skrętki z izolacją na 60V. Robi się jeszcze ciekawiej. Napięcie na końcu linii będzie się zmieniać, zależnie od prądu pobieranego przez laptopa, czyli praktycznie na 100% w końcu podamy więcej niż nominalne 19V i ugotujemy laptopa. Oczywiście można na końcu linii zastosować przetwornicę, która będzie stabilizowała napięcie dla laptopa. Pozostaje jeszcze kwestia mocy strat w linii zasilającej czyli $I^2 R$, czyli $25 \times 5 = 125W$. Daje to 2W na każdy metr linii. Do tej pory zakładałem, że skrętka ma rzeczywiście 0,5mm i jest z prawdziwej miedzi, a nie z jakiegoś wynalazku o oporności bliżej nieznannej, ale na bank większej niż miedzi. W takim przypadku straty w linii będą jeszcze większe (...) Najlepiej byłoby, kładąc skrętkę, równocześnie położyć przewód [energetyczny] $3 \times 2,5$. Mamy pewne zasilanie, a w razie draki zawsze można podłączyć oświetlenie albo grilla elektrycznego. Pozdrawiam.

Osąd w kwestii można / nie można zależy nie tylko od przyjętej mocy oraz prądu laptopa, ale też od innych założeń. Otóż niektórzy uczestnicy przyjęli, że umieszczonymi w domu dodatkowymi obwodami zasilania będzie zwyczajny zasilacz o jakimś ustalonym napięciu wyjściowym, natomiast napięcie na odległym końcu kabla ma zgodnie z **rysunkiem C** wynosić 19V, bo takie jest typowe napięcie zasilania laptopów. Przyjęcie takiej koncepcji natychmiast daje odpowiedź: **nie można tak zasilac laptopa!**

Już choćby dlatego, że pobór prądu przez laptop zmienia się w szerokim zakresie. Maksymalnie może wynosić kilka amperów i wtedy spadki napięć na przewodach ($U_{P1} + U_{P2}$) mogą nawet przekroczyć 20V, co wymagało-

by napięcia domowego zasilacza U_Z rzędu 40V. A w успіniu i przy naładowanej baterii zmniejszy się niemal do zera i wtedy na gnieździe zasilania wystąpi całe napięcie zasilacza i gdyby to było 40V, najprawdopodobniej uszkodzi wewnętrzne obwody zasilania

laptopa. Koncepcja z rysunku C jest błędna i nieakceptowalna. Można jednak zastosować rozwiązanie według ogólnego **rysunku D** ze stabilizatorem impulsowym umieszczonym przy laptopie. Napięcie na wejściu stabilizatora będzie się zmieniać, zależnie od płynącego prądu, ale na wyjściu utrzyma on potrzebne napięcie 19V. Aby zmniejszyć straty napięcia na rezystancji kabla, należy zmniejszyć prąd, czyli zwiększyć napięcie domowego zasilacza. Na ile zwiększyć? I jakie to ma być napięcie?

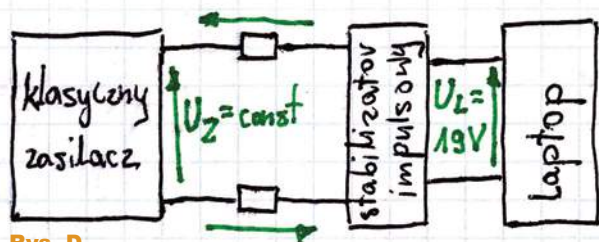
Absolutnie niedopuszczalne byłoby przesyłanie przez wolne pary kabla UTP napięcia sieci 230V! To wcześniej czy później zaprowadziłoby prosto na ławę oskarżonych z zarzutem zabójstwa z zamiarem ewentualnym! Zamiar ewentualny oznacza, najprościej biorąc, że zabójca mógł przewidzieć konsekwencje swojego czynu.

Napięcie w kablu UTP powinno być bezpieczne, a w przypadku napięcia stałego za bezwzględnie bezpieczne uznaje się co najwyżej 60V. Na odległym końcu kabla należy umieścić indukcyjną przetwornicę obniżającą o napięciu wyjściowym 19V i jak najwyższej sprawności energetycznej.

Można byłoby też wykorzystać bezpieczne napięcie zmienne, ale wtedy w grę wchodziłoby 24V, najwyżej 48V.

Koncepcja z rysunku D daje szansę na prawidłową pracę laptopa, ale jej nie gwarantuje. Oddajmy głos uczestnikom zadania, którzy przedstawili różne godne zastanowienia szczegółów.

I tak jeden z Kolegów napisał: (...) W terenie kabel albo wisi, albo jest zakopany. W obu przypadkach jest narażony na wilgoć, więc należałoby użyć kabla żelowanego. Sprawdziłem (...) skrętki (...) wychodzi 84mΩ/m. Znalazłem też skrętkę (...) 95mΩ/m. (...) Wiele laptopów ma też w gniazdku zasilającym trzeci pin oprócz plusa i minusa. Służy on do



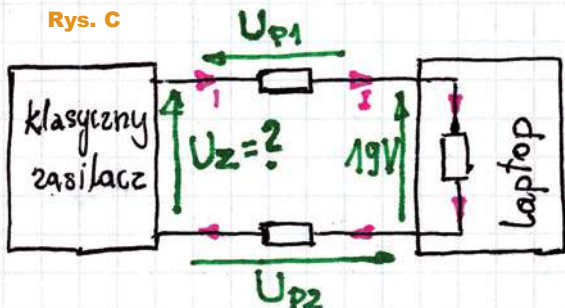
Rys. D

wykrywania, czy zasilacz jest oryginalny. Jeżeli weryfikacja się nie powiedzie, to obniżane jest taktowanie CPU, wyłączane jest ładowanie akumulatora albo laptop w ogóle przestaje być zasilany. Moim zdaniem to nieuczciwe praktyki rynkowe typu vendor lock-in, ale trzeba je brać pod uwagę, budując taki przedłużacz. (...) Ponieważ rezystancja skrętki jest stała, można zbudować kompensator rezystancji przewodów, który był prezentowany w zadaniu Jak7 w EdW. Możemy też puścić skrętką wyższe napięcie (PoE używa do 57V) i zamontować przetwornicę obniżającą przy laptopie. (...) Jeżeli nasz laptop nie działa bez oryginalnego zasilacza, należy przenieść układ producenta z oryginalnego zasilacza do stworzonego przez nas zasilacza znajdującego się w altanie. Pozdrawiam.

Jeden ze stałych uczestników dokładniej zbadał sprawę kabli: (...) dla kategorii 5e średnica żyły waha się od 0,48mm (AWG26) do 0,50mm (AWG24), odpowiednio rezystancja 188mΩ/m do 260mΩ/m (...) Ale jest jeszcze kategoria 7 (...) żyły jednodrutowe miedziane o średnicy Ø 0,56mm (23 wg AWG) o rezystancji 94mΩ/m. Nawiasem mówiąc, te dane katalogowe budzą wątpliwości. Przewód ma nazwę „CCA”, czyli chińska miedź, w tłumaczeniu na polski: aluminium pokryte miedzią, a w opisie „żyły jednodrutowe miedziane”. Natomiast producent drutów DNE podaje rezystancje: $\phi = 0,50mm - 86,6m\Omega/m$, $\phi = 0,56mm - 69,0m\Omega/m$.

Przy okazji, jak trzeba uważać przy przeglądaniu katalogów – impedancja kabla cat. 5e podana jest dla całej linii (1 para przewodów) „tam i z powrotem”, impedancja kabla cat.7 tylko „tam”. I od razu widać, że kabel jest z chińskiej miedzi, potwierdza to symbol na otoczce kabla, mimo zapewnień (a raczej kłamstw) w opisie. (...) I jeszcze kilka uwag praktycznych: jedna para skrętek to zasilanie +, druga para skrętek to zasilanie -. Pary [cat. 7]

Rys. C





są indywidualnie ekranowane, cały kabel ma dodatkowy ekran. Te ekrany należy połączyć ze sobą i połączyć z masą zasilacza/rutera, drugie końce ekranów (w altance) MUSZĄ zostać wolne, z niczym niepołączone!

Gorąco przestrzegam przed „wynalazkami” typu: wykorzystanie ekranu jako jednego z przewodów zasilania lub łączeniem ich z masą zasilacza/rutera na jednym końcu oraz z masą laptopa w altance.

Podsumowując, uważam, że problem da się rozwiązać.

Inny stały uczestnik naszych konkursów napisał: (...) Najczęściej użytkownicy wybierają komputery przenośne, dla których 65W to najczęściej spotykana moc. Laptopy wykorzystywane do pracy biurowej i podstawowych czynności, takich jak na przykład odbieranie i wysyłanie poczty elektronicznej, mają zasilacz o mocy ok. 65W, a wykorzystuje ok. 40W. Mocny laptop do gier wyposażony jest zazwyczaj w zasilacz o mocy ok. 150W. Pod obciążeniem taki sprzęt może pobierać średnio 110W. (...) Pobór prądu jest zmienny, zależy od wykorzystywanych aplikacji, programów (...) Dla 24AWG mamy przekrój 0,205mm², obciążalność prądowa 0,58A. Dla dwóch par przewodów obciążalność wyniesie 1,16A. Dla 23AWG mamy przekrój 0,259mm², obciążalność prądowa 0,73A. Dla dwóch par przewodów obciążalność wyniesie 1,46A. (...) Oprócz skrętki musimy uwzględnić również obciążalność prądową wtyku RJ45. Dane katalogowe podają tę wartość jako 1A przy 50°C dla napięcia 60V. (...) należy wspomnieć, że istnieje technologia zasilania przez Ethernet, powszechnie znana jako PoE (Power over Ethernet). Standard ten ewoluował już do PoE++ (IEEE 802.3bt). W tym jak również w jego wcześniejszych wersjach określono zakres napięć na urządzeniu zasilającym i na zasilanym oraz prąd maksymalny, moc dostępną na urządzeniu zasilanym oraz moc dostarczaną przez urządzenie zasilające. Najnowszy standard PoE++ przewiduje co prawda moc dostępną na urządzeniu zasilanym równą 71W, ale w trybie 4 pary (960mA na parę). Dla trybu 2 par jest to 51W (600mA na parę). Daje to możliwość zasilania urządzeń sieciowych (np. kamery, telefony, punkty dostępowe), ale raczej wyklucza laptopy, których pobór mocy przewyż-

sza dopuszczalną moc wg powyższych standardów. Należy również zwrócić uwagę na fakt, że urządzenia zasilające i zasilane powinny być w standardzie PoE. Pozdrawiam.

W jeszcze innym rozwiązaniu przeczytamy: (...) Najłabsze zasilacze jakie spotkałem, miały zaledwie 45W, przeciętnie jest to 60–90W, a najmocniejszy notebook, jaki miałem, wymagał zasilacza 150W. (...) oczywiście wydaje się wykorzystanie idei zasilania, jaka jest zastosowana w standardzie PoE. Zasilanie napięciem sieciowym 230V/50Hz nie wchodzi w rachubę ze względów bezpieczeństwa, gdyż kabel UTP nie jest przystosowany do takiego napięcia. Najprostsze byłoby zastosowanie dwóch transformatorów na obu końcach kabla, jeden obniżający napięcie na kablu do bezpiecznej wartości i drugi podwyższający do 230V na wyjściu. Jednakże takie rozwiązanie mogłoby doprowadzić do zakłóceń uniemożliwiających prawidłowe przesyłanie danych w parach sygnałowych. Pozostaje zasilanie prądem stałym. Założyłem, że na wejściu zastosujemy inżektor PoE własnej konstrukcji dostarczający prąd o napięciu wg przyjętych założeń. Na końcu kabla poprzez adapter do dwóch par będzie podłączona przetwornica napięcia DC/AC dostarczająca na wyjściu napięcie 230V/50Hz lub prostsza przetwornica DC/DC, jeśli zasilacz może być zasilany napięciem stałym (co jest bardzo prawdopodobne), sprawność również na poziomie 80%. (...) trzy opcje:
1 – ściśle trzymanie się standardu POE wg typu energetycznego 4 „4PPoE” jako zapewniającego największą moc (do 70W na urządzeniu zasilanym),
2 – wykorzystanie maksymalnego potencjału przeciętnego kabla UTP,
3 – wykorzystanie „najmocniejszego” kabla wysokiej jakości, jaki uda się znaleźć, wytypowałem kabel U/UTP firmy BiTLAN.

(...) Napięcie zasilające na wejściu: standard 4PPoE przewiduje maksymalne napięcie zasilające 57V, przeciętny kabel UTP ma dopuszczalne napięcie pracy 70–72V, najlepszy kabel, jaki znalazłem, to kabel firmy BiTLAN wytrzymujący 150V. (...) Maksymalny prąd jednej pary – producenci nie podają maksymalnej obciążalności prądowej w swoich materiałach, więc przyjąłem 960mA na jedną parę wg standardu 4PPoE, ograniczenie to

wynika z dopuszczalnej temperatury przewodów. (...) uwzględniłem spadek napięcia na przewodach (i związane z tym straty mocy) oraz sprawność przetwornicy wyjściowej (80%). (...) Z obliczeń wynika, że trzymając się wytycznych standardu 4PPoE, będziemy mogli uzyskać moc na wyjściu przetwornicy nieco poniżej 73W (czyli w przybliżeniu tyle, co przewiduje standard), co powinno wystarczyć do zasilenia notebooka o najłabszych parametrach.

Wykorzystując do maksimum potencjał typowego kabla UTP, uzyskamy na wyjściu moc ok. 95W, co pozwoli już zasilić większość współczesnych notebooków. A stosując „najmocniejszy” kabel [BiTLAN U/UTP], możemy zasilac chyba każdy dostępny notebook, gdyż moc, jaką będziemy dysponować, to ponad 200W. Dodatkowo przy tym kablu, jeśli zasilacz notebooka jest tzw. „światowym”, czyli może być zasilany napięciami z zakresu 95–250V/50–60Hz i da się go zasilac prądem stałym, to na wyjściu naszej instalacji nie potrzebujemy przetwornicy napięcia DC/AC, gdyż napięcie stałe, jakie uzyskamy na końcu kabla UTP z uwzględnieniem spadków napięcia, to około 140V, co jest w zupełności wystarczające do poprawnej pracy takiego zasilacza. Jednocześnie pozbywając się drugiej przetwornicy, zwiększamy sprawność układu i uzyskujemy dodatkową moc na wyjściu. Można też wykonać przetwornicę dostarczającą dokładnie takiego napięcia, jakiego wymaga notebook i zrezygnować z jego zasilacza, co również podniesie sprawność instalacji (...)

Szczegółowe obliczenia dostępne są w umieszczonym w Elportalu pliku PDF.

Pomimo, że przedstawione finalne opinie (można / nie można) były różne, wszystkie nadesłane rozwiązania śmiało mogę uznać za prawidłowe, ponieważ wszystko zależało od przyjętych założeń.

Nagrody-upominki za zadanie **Policz306** otrzymują:

Teodor Woźniak – Łódź,
Tomasz Sukiennik – Kraków,
Mariusz Hejto – Łowczówek.

Wszystkich uczestników dopisuję do listy kandydatów na bezpłatne prenumeraty.

Piotr Górecki

Profesjonalny telefon monterski MT-8001

Przeznaczony do wszechstronnej kontroli stanu analogowych linii telefonicznych oraz tymczasowej komunikacji.



Wywołaniu numeru z pamięci



Zapis do pamięci

Wskaźnik polaryzacji

Duży, wielofunkcyjny panel

Przycisk pauzy

Przełącznik wybierania: tonowo lub impulsowo

Powtórzenie ostatnio wybranego numeru

Mocny, wytrzymały uchwyt



Krokodylki z kołcami



Kształt dopasowany do ramienia



Pro'sKit®



Telefon monterski MT-8006B

Przeznaczony dla techników i instalatorów do wszechstronnej kontroli stanu analogowych linii telefonicznych oraz tymczasowej komunikacji.



Wskaźnik polaryzacji

Podświetlany wyświetlacz LCD

Pamięć 12 numerów telefonów



Port RJ11



Gniazdo słuchawkowe

sklep.avt.pl



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

Pro'sKit®





Pulsująca błyskotka LED

Kolorowa, migocząca ozdoba może być urozmaicheniem wielu prezentów od początku elektroniki. Cztery różnokolorowe diody pulsują płynnie w sobie tylko znanym tempie, co tworzy wrażenie miłego dla oka chaosu. Jasność każdej z nich nigdy nie spada do zera, więc gadżet świeci przez cały czas!

Do czego to służy?

Zadaniem tego układu jest płynne rozjaśnianie i ściemnianie czterech diod LED w różnych kolorach: czerwonym, zielonym, żółtym i niebieskim. Sterowanie ich jasnością odbywa się poprzez cykliczną zmianę natężenia ich prądu, do czego służy obwód z tzw. lustrem prądowym. Prosty, efektowny, różnokolorowy gadżet, który można wykorzystać jako część składową zabawki, pozytywki albo ozdoby choinkowej.

Diody i inne podzespoły znajdują się na niewielkiej płytce drukowanej, którą może polutować nawet początkujący użytkownik lutownicy. Wprawdzie wszystkie elementy montuje się powierzchniowo (SMD), ale są one na tyle duże, że nie powinno to sprawić większych problemów. Można ten układ potraktować jako ćwiczenie tej techniki montażu.

Niniejsze urządzenie jest zasilane napięciem stałym, niekoniecznie stabilizowanym, o wartości rzędu 5...6V. Pobiera w trakcie pracy około 15mA.

Jak to działa?

Schemat układu można zobaczyć na **rysunku 1**. Składa się z czterech (prawie) identycznych bloków, więc poprzestanę ma omówieniu tylko jednego z nich – na przykład tego, który zawiera diodę LED1 (czerwoną). Zasilanie do nich doprowadza złącze J1, z niego też jest zasilany układ scalony US1. Kondensatory C5 i C6 filtrują zasilanie dla niego i na płytce znalazły się blisko jego obudowy.

Tym, co generuje pulsację, jest obwód składający się z rezystora R1, kondensatora C1 i bramki US1A. Bramka tego typu jest funktorem NAND, ale



po zwarciu jej wejść staje się zwykłym negatorem – zmienia wartość logiczną, która jest na jej wejściu, na przeciwną. Dodatkowo jej wejścia są opatrzone tzw. przerzutnikami Schmitta, co zapewnia histerezę. Przełączenie stanu wyjścia następuje wtedy, kiedy napięcie na wejściu przekroczy określony próg podczas narastania albo inny, kiedy ono opada. Pomiędzy tymi progami bramka nie reaguje.

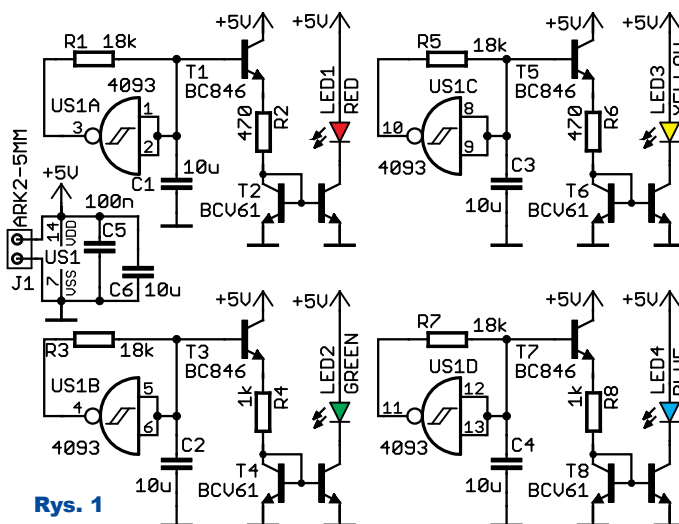
Rezystor R1 powoli ładuje kondensator C1 w czasie, kiedy wyjście bramki przyjmuje wysoki stan logiczny (równy niemal napięciu zasilania) – czyli potencjał jej wejścia jest niski. Sytuacja zmienia się po naładowaniu C1 do wystarczającego poziomu, gdyż wtedy

wyjście bramki przełącza się, przyjmuje stan niski (zblizony do 0V) i R1 zaczyna rozładowywać C1. To będzie trwało tyle czasu, aż C1 nie rozładuje się, osiągając napięcie równe dolnemu progowi przerzutu bramki. I cykl się powtórzy.

Zazwyczaj w tego typu generatorach pobieramy z wyjścia bramki sygnał prostokątny o określonej częstotliwości. Ale w tym zastosowaniu wyjście bramki nas nie interesuje, bowiem jego potencjał zmienia się skokowo. Ciekawszymi jest quasi-trójkątny przebieg napięcia na zaciskach C1. Przedrostek „quasi” nie został postawiony przypadkowo, ponieważ przeładowywanie kondensatora przez rezystor ma przebieg wykładniczy. Na oscylogramach widzimy to jako

łukowate zaokrąglenie zboczny tego sygnału trójkątnego.

Napięcie z górnej okładki tego kondensatora jest „powtarzane” przez wtórnik napięciowy na tranzystorze T1. Jego baza pobiera niewielki prąd, więc w równie niewielkim stopniu obciąża generator. Ponieważ ten tranzystor jest cały czas w stanie przewodzenia, jego napięcie baza-emiter jest z grubsza stałe i wynosi około 0,7V. Lecz potencjał bazy ulega zmianie, gdyż C1 cyklicznie



Rys. 1

ładuje się i rozładowuje, zatem potencjał emitera również ulega zmianom – odzwierciedla zmiany potencjału bazy, tyle że są one pomniejszone o napięcie baza-emiter, czyli o wspomniane 0,7V. Ten układ pracy tranzystora nazywamy wtórnikiem napięciowym.

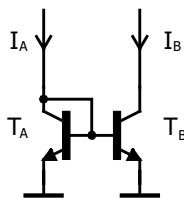
Obciążeniem emitera T1 jest rezystor R2, który dolnym zaciskiem jest podłączony do wejścia elementu T2, będącego w istocie układem dwóch tranzystorów tworzących tzw. lustro prądowe. Lewy tranzystor w T2 ma zwarty kolektor z bazą, bowiem istotne jest dla nas tylko jego złącze baza-emiter. Prąd wpływający do tego złącza wywołuje na nim spadek napięcia. Im większe jest natężenie tego prądu, tym większa jest wartość owego napięcia, choć zależność ta jest silnie nieliniowa. Możemy przyjąć, że wartość ta wynosi 0,7V i zmienia się nieznacznie dla małych prądów.

Jak zatem działa ten obwód? Skoro dolne wyprowadzenie R2 jest podłączone do złącza baza-emiter, na którym będzie się odkładało napięcie o niemal niezmienniej wartości, a górne wyprowadzenie tego samego rezystora ma cyklicznie zmieniany potencjał (dzięki wtórnikowi na tranzystorze T1), to napięcie na zaciskach tego rezystora również ulega zmianom. Można sobie to wyobrazić jako „rozciąganie” i „ściskanie” sprężyny, która dolnym końcem jest zaczepiona do podłoża, a ruszamy jej równym końcem.

Co nam daje zmiana napięcia na zaciskach R2? Tutaj wystarczy odwołać się do prawa Ohma: zmienia się płynący przez niego prąd. Jest to zmiana liniowa – zwiększenie napięcia między wyprowadzeniami rezystora daje proporcjonalny wzrost natężenia prądu, jaki przez niego płynie. Mamy bardzo prosty, a jednocześnie bardzo skuteczny, przetwornik napięcia na prąd.

Ten prąd wypływa z R2 (dostarcza go emiter tranzystora T1, pobierając go kolektorem ze źródła zasilania) w kierunku wspomnianego już wejścia lustra prądowego. Napięcie, jakie odkłada się między bazą i emitrem lewego tranzystora T2, wymusza identyczne napięcie między bazą i emitrem prawego tranzystora, który jest z kolei wyjściem lustra. Prawy tranzystor lustra jest tak sterowany, że próbuje swoim kolektorem „wessać”

prąd o natężeniu niemal identycznym jak ten, który wpływa do wejścia lustra. Mówiąc inaczej, lustro prądowe „odbija” prąd z rezystora R2 na swoje wyjście –



Rys. 2

rysunek 2. Wymuszamy prąd I_A , zaś prąd I_B jest pobierany z innego fragmentu obwodu i stanowi wyjście lustra. W sytuacji idealnej $I_A = I_B$.

A co jest podłączone do wyjścia? Dioda LED1. Bez jakiegokolwiek dodatkowych rezystorów ograniczających natężenie płynącego przez nią prądu, ponieważ jest on dokładnie kontrolowany przez lustro prądowe. Prawy tranzystor T2 pobiera ze źródła zasilania prąd o zmieniającym się natężeniu, zaś prąd ten „przy okazji” przepływa przez LED1. Dlatego dioda ta pulsuje.

Na początku wspomniałem, że układ ma cztery „prawie” identyczne bloki. Dwa z nich różnią się wartością rezystora „konwertującego” zmiany napięcia na zmiany prądu wchodzącego do lustra. Diody świecące na niebiesko i zielono świecą znacznie intensywniej (nasze oczy są na nie bardziej wyczulone) niż czerwone i żółte. Dlatego natężenie prądu dwóch wymienionych najpierw diod jest mniejsze (bo rezystory mają większą rezystancję) niż pozostałych, aby nie przyciemniały ich świecenia.

Jak wygląda przebieg prądu na zaciskach takiej diody? To nie tak łatwo zmierzyć, ale możemy poznać napięcie na zaciskach rezystora, który konwertuje napięcie na prąd. **Rysunek 3** pokazuje oscylogram napięcia na zaciskach R8 wraz z automatycznym pomiarem podstawowych wielkości. Możemy z niego odczytać, że minimalna wartość tego napięcia wynosi 860mV, a maksymalna 1,22V. Ponieważ jego rezystancja wynosi 1kΩ, przeliczenie napięcia na prąd jest proste: 1V odpowiada 1mA. Zatem prąd diody LED4 zmienia się w przedziale od

0,86mA do 1,22mA, więc jej świecenie nie jest zbyt intensywne. Częstotliwość tych zmian to około 10Hz. W rezultacie mamy szybko pulsującą diodę, która nigdy nie gaśnie, ponieważ stale płynie przez nią prąd. Kształt tego przebiegu jest taki, jak wcześniej omówiony przebieg napięcia na zaciskach C1, czyli trójkątny o zaokrąglonych zboczach.

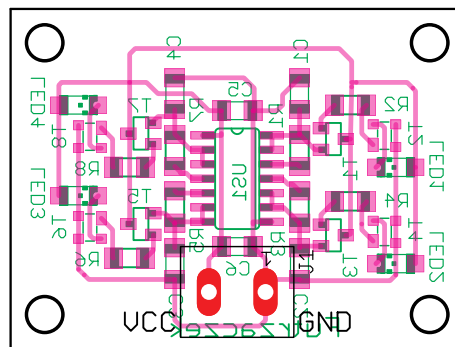
Jako lustro prądowe zostały wybrane układy typu BCV61, które są dostępne w różnych kategoriach wzmocnienia prądowego – A, B lub C. Im większe wzmocnienie wybieramy, tym bardziej będzie zbliżony prąd wyjściowy do wejściowego. Najwyższe ma grupa C, najniższe A, dlatego najlepsze byłyby elementy o oznaczeniach BCV61C. Ale ten układ jest na tyle mało wymagający, że nawet BCV61A będą działały poprawnie, ponieważ „ubytek” prądu będzie niezauważalnie niski.

Montaż i uruchomienie

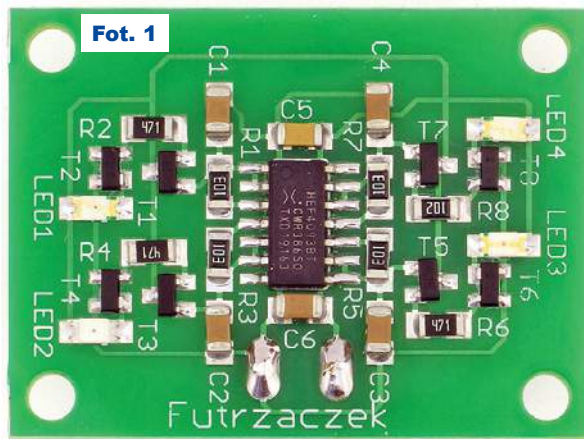
Układ prototypowy został zmontowany na jednostronnej płytce drukowanej o wymiarach 30×40mm. Wzór jej ścieżek i schemat montażowy przedstawia **rysunek 4**. W odległości 3mm od krawędzi płytki znajdują się otwory montażowe o średnicy 3,2mm każdy.

Ciąg dalszy na stronie 61

Rys. 4 – skala 150%



Rys. 3



Z potrzeby chwili...

Zasilacz do laptopa w roli ładowarki akumulatora

W cyklu „Z potrzeby chwili...” przedstawiamy opisy układów, urządzeń i instalacji elektronicznych, które powstały szybko dla zaspokojenia konkretnych potrzeb i te potrzeby zaspokoili. Szybki proces powstawania zwykle oznacza, że urządzenie nie jest do końca dopracowane i że w przyszłości może być lub będzie ulepszone, co też może zostać opisane w EdW. Zachęcamy do nadsyłania tego rodzaju materiałów do publikacji.

Pewnego razu spotkała mnie przykra niespodzianka w postaci rozładowanego akumulatora w samochodzie. Przyczyną rozładowania było prawdopodobnie błędnie działające urządzenie do automatycznego wzywania pomocy, które przez parę dni postoju zdążyło wyczerpać akumulator. Rozładowanie było na tyle poważne, że rozrusznik przy próbie zapłonu ani drgnął, a elektronika wariowała.

Niestety, nie miałem pod ręką prostownika, a odpalenia „na kable” z pewnych względów nie brałem pod uwagę.

Postanowiłem więc, że na szybko wykonam ładowarkę, używając komponentów, które aktualnie miałem pod ręką. Jako źródło prądu wybrałem zasilacz do laptopa o wydajności około 4A

i napięciu 19,5V. Do tego garstka dyskretnych elementów elektronicznych.

W zasadzie należałoby przeprowadzić przeróbkę zasilacza, by obniżyć napięcie i zastosować ogranicznik prądowy. Jednak nie jest to łatwe, a ja „z potrzeby chwili” wykrzystałem sposób bodaj najprostszy, choć wcale nie najlepszy: dodałem na wyjściu liniowy ogranicznik.

Pojemność akumulatora wynosiła 60Ah, czyli naładowanie do pełna zajęłoby ponad 15h. Na szczęście okazało się, że do odpalenia samochodu wystarczy tylko naładować częściowo – wystarczyło ładowanie przez około 8h prądem 2A.

Oprócz funkcji awaryjnej ładowarki akumulatorów układ spełnia się jako zasilacz w elektronicznej praktyce.

Opis układu

Akumulatory są obciążeniem o bardzo małej rezystancji dynamicznej, więc bezpośrednie podłączenie zasilacza od laptopa, którego napięcie wyjściowe, zależnie od modelu, wynosi około 19V, spowodowałoby zadziałanie zabezpieczenia nadprądowego lub przy jego braku przeciążenie i w konsekwencji zniszczenie zasilacza. Przystosowanie zasilacza polega na dodaniu obwodu ogranicznika prądu.

W pewnych warunkach można po prostu użyć rezystancji szeregowej, na której odłoży się nadwyżka

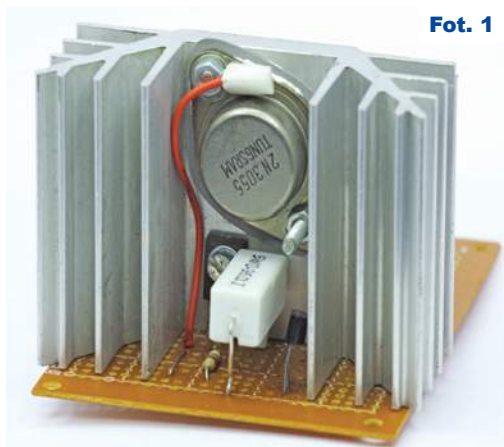
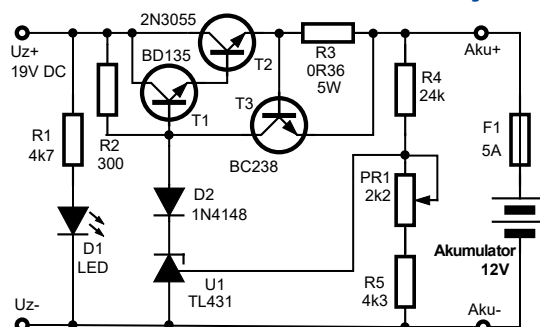
napięcia, przez co prąd zostanie ograniczony. Znany jest od lat sposób z żarówką, której rezystancja nie jest stała i rośnie w miarę wzrostu prądu.

Jednak rozwiązanie z szeregowym oporem jest mało efektywne, bo moc zasilacza nie jest w pełni wykorzystywana, a czas ładowania się wydłuża. Ponadto, co bardzo groźne, gdy akumulator przyjmie już odpowiednio duży ładunek, napięcie na nim nadal rośnie, a wartość powyżej 15..16V staje się niebezpieczna dla samego akumulatora i otoczenia, ponieważ rozpoczyna się proces gazowania.

Układ, który wykonałem z potrzeby chwili, jest pozbawiony tych wad, ponieważ nie dopuszcza do zwiększenia się napięcia na zaciskach akumulatora oraz prądu pobieranego z zasilacza.

Schemat urządzenia jest na **rysunku 1**. Układ można krótko scharakteryzować jako

Rys. 1



Fot. 1

liniowy regulator napięcia z ogranicznikiem prądu wyjściowego.

Tranzystor T2 pełni funkcję zaworu sterującego przepływem prądu do obciążenia. Razem z tranzystorem T1 tworzą układ Darlingtona, którego baza polaryzowana jest przez rezystor R2. Jako T1 zastosowałem niegdyś popularny, „pancerny” 2N3055. Zakładając, że rozładowany akumulator miał napięcie 9V, moc oddłożona na tranzystorze wykonawczym wynosi $(19V-9V) \cdot 2A = 20W$. Istnieje więc spory zapas mocy, nawet przy zastosowaniu niewielkiego radiatora.

Funkcję stabilizatora napięcia pełni układ TL431. Steruje on układem T1/T2, podbierając prąd z rezystora R2 dostarczającego prąd do bazy T1. Prąd wpływający do końcówki „C” jest tym większy, im większa jest różnica między napięciem na końcówce Adj U1 a wartością wewn. źródła napięcia referencyjnego 2,5V.

Dzięki układowi Darlingtona prąd polaryzujący bazę T1 może mieć niewielką wartość, rzędu mA. Prąd rezystora R2 wynosi około $(19,5V-12V-2,1V)/300\Omega = 18mA$, co daje straty mocy na U1 około 0,25W i jest wartością dopuszczalną.

Elementy R3 i T3 tworzą obwód ogranicznika prądu. Prąd z emitera T2, płynąc przez rezystor R3, powoduje pewien spadek napięcia. Gdy napięcie to zaczyna być bliskie 0,7V, otwiera się tranzystor T3, obniżając napięcie między bazą T1 a wyjściem układu, wskutek czego maleje prąd ładujący akumulator.

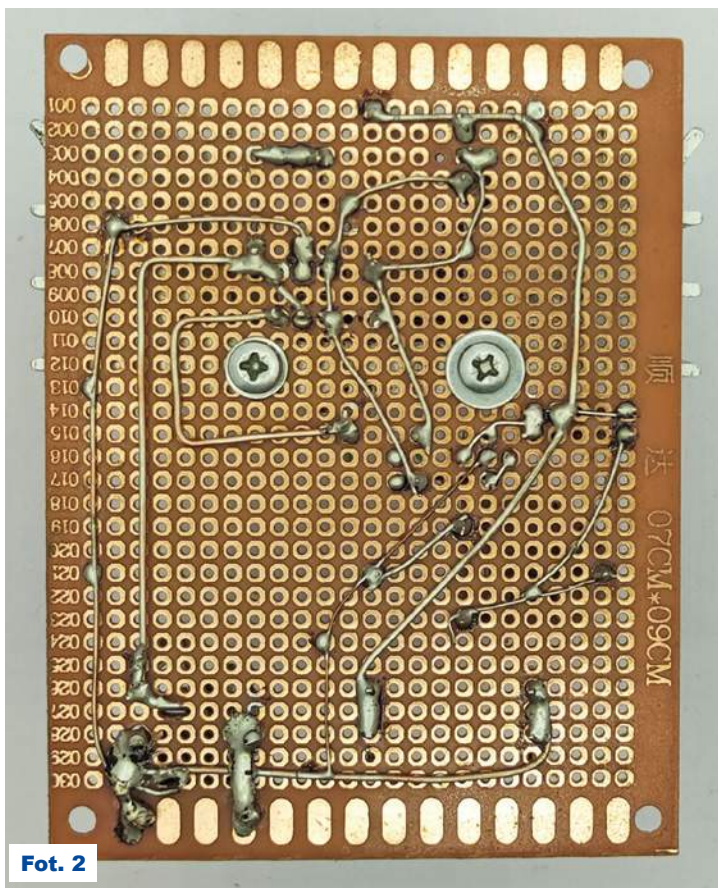
Rezystor R3 w trakcie pracy ładowarki nagrzewa się. Przy prądach rzędu 5A musi on oddawać do otoczenia moc $5A \cdot 0,7V = 3,5W$. Ja użyłem rezystora o rezystancji $0,36\Omega$ i mocy 5W; taka wartość rezystancji powoduje, że prąd zostaje ograniczony do około 2A, natomiast straty mocy wynoszą około 1,4W.

Końcówka Adj U1 podłączona jest do dzielnika napięciowego R4(PR1+R5). Co ważne, napięcie zasilające ten dzielnik pochodzi z za ogranicznika prądowego, żeby uniknąć błędu spowodowanego spadkiem napięcia na R3. Wartości rezystorów i potencjometru dobrano tak, aby możliwa była regulacja napięcia wyjściowego w zakresie od 12 do 16V. Dioda D1 zabezpiecza obwód przed przypadkowym, odwrotnym włączeniem akumulatora do zacisków wyjściowych. Układ U1 zawiera wewnątrz diodę włączoną zaporowo, która mogłaby spowodować przepływ prądu drogą od masy do emitera T1.

Obawy może budzić brak zabezpieczenia przed prądem wstecznym z akumulatora, gdy zasilacz nie jest podłączony. Gdy napięcie wyjściowe jest ustawione na wartość powyżej 12,8V (napięcie na sprawnym i naładowanym akumulatorze) prąd wsteczny nie popłynie.

Montaż i uruchomienie

Układ zmontowałem na uniwersalnej płytce, razem z radiatorem. Szczegóły spodu płytki widoczne są na **fotografii**



Fot. 2

Wykaz elementów

R1		4,7k Ω
R2		300
R3		0 Ω 36/5W
R4		24k Ω
R5		4,3k Ω
PR1		2,2k Ω
D1		LED
D2		1N4148
T1		BD135
T2		2N3055
T3		BC238
U1		TL431
F1		bezp. 5A
Złącza płaskie do druku 2x		
Gniazdo zasilania HP/Dell		
Radiator		

2. Gniazdo zasilania kompatybilne jest z wtykami zasilaczy HP oraz Dell. Wyposażone jest ono w diodę LED sygnalizującą podłączenie. Niektóre zasilacze wymagają „oszukiwania” środkowego pinu poprzez podanie napięcia o odpowiedniej wartości np. z dzielnika rezystorowego. Ja użyłem zasilacza HP, który tego nie wymagał.

Zaciski wyjściowe to stosowane w motoryzacji złącza płaskie, charakteryzujące się niską rezystancją połączeń. Sprawdzają się tutaj lepiej niż złącza śrubowe.

Fotografie pokazują szczegóły montażu z wierzchu i spodu płytki. Układu nie umieszczałem w obudowie, jedynie

na podkładce izolacyjnej. Brak obudowy to nawet zaleta, ze względu na lepsze chłodzenie, szczególnie że ładowanie powinno przebiegać na zewnątrz lub w pomieszczeniach dobrze wentylowanych. Ładowanie przeprowadziłem, ustawiając wcześniej napięcie wyjściowe bliskie 15V z użyciem woltomierza, bez podłączonego akumulatora. Dla bezpieczeństwa zastosowałem bezpiecznik 5A tuż przy dodatnim słupku akumulatora.

Piotr Wójtowicz
piotr.wojtowicz@elportal.pl

Odkrywamy schematy

część 6

Klasyczny forward z jednorozprężającym kluczem

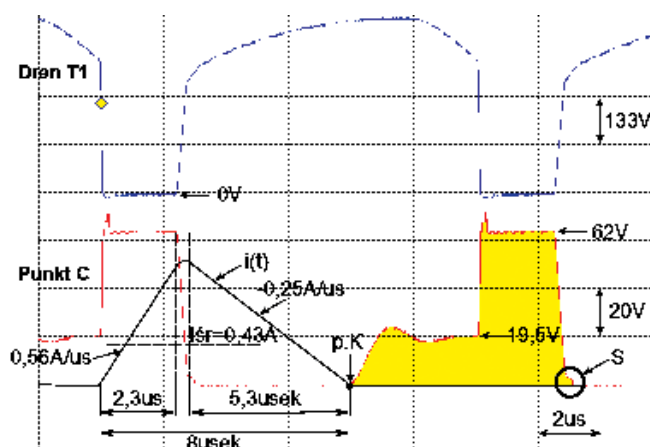
Poprzednią część cyklu zakończyliśmy pytaniem i teraz wypada dać odpowiedź. Chodziło o to, które z czynników: napięcie wejściowe, wyjściowe i obciążenie zasilacza, mają wpływ i jaki na pracę przetwornicy i na przebiegi obserwowane w kluczowych punktach układu: dren klucza, punkty A i B obwodu odzysku energii z resetowania rdzenia oraz po stronie wtórnej w węzle diod i indukcyjności gromadzącej energię.

Otóż interesujące jest, że najmniejszy wpływ ma obciążenie zasilacza. Dopóki w indukcyjności L_{out2} prąd płynie ciągle (CCM), to teoretycznie wielkości prądu obciążenia nie powinniśmy zauważyć na żadnym z cytowanych oscylogramów. Jedynie bez obciążenia zasilacza wszystkie przebiegi mocno się deformują, co jest efektem pracy w trybie przewodności nieciągłej DCM. Na osobną część opracowania zasługiwałoby wyjaśnienie wszystkich konsekwencji, wad i zalet trybów CCM i DCM. Na **rysunku 38a i b** zamieszczamy przebiegi zdjęte podczas pracy zasilacza komputerowego z bardzo małym obciążeniem.

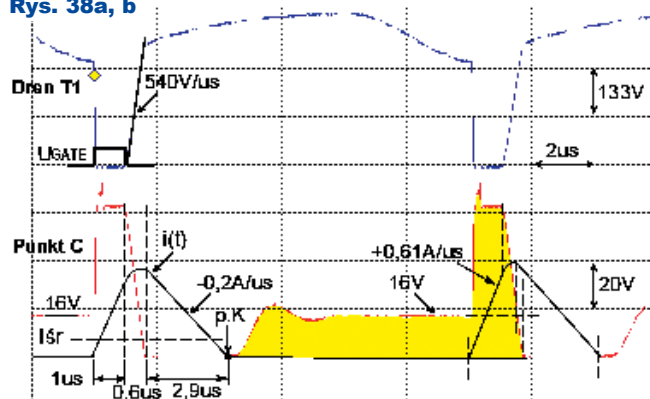
Napięcie wejściowe przetwornicy ma zdecydowany wpływ na współczynnik wypełnienia kluczowania: wydłuża/skraca się aktywny czas włączenia klucza. Proporcjonalnie do tego obserwujemy szersze i niższe (węższe i wyższe) impulsy po stronie wtórnej. Tak naprawdę tu jest punkt wyjścia. Praca całego zasilacza musi się dostosować tak, aby wartość średnia, czyli składowa stała, w tym miejscu miała określoną zadaną wartość, wymaganą przez obwód ujemnego sprzężenia zwrotnego. Poza tym wielkość napięcia U_{we} nie ma wpływu na pracę obwodu demagnetyzacji. Przebiegi w punktach A i B, poza sze-

rokością, zachowują charakterystyczne poziomy napięcie. To jest potwierdzeniem, że najistotniejszym czynnikiem jest tu wielkość iloczynu woltosekund przyłożonych do indukcyjności magnetyzacji L_m . A te wielkości (wysokość-napięcie i szerokość-czas) zmieniają się w odwrotnej proporcji, gdy zmienia się U_{we} .

Zdecydowany wpływ na wszystkie przebiegi ma natomiast wartość napięcia wyjściowego U_{wy} . Zwiększając napięcie U_{wy} , obserwujemy poszerzanie impulsów na drenie klucza i na wyjściu, podobnie jak obserwowaliśmy podczas obniżania U_{we} . W przetwornicy *forward* istotny jest bowiem stosunek tych napięć odniesiony do przekładni transformatora. Podnosząc U_{wy} , nawet przy stałej czerpanej z wyjścia mocy, obserwujemy także zdecydowany wzrost napięcia szczytowego na drenie tranzystora MOSFET-a. Także obserwacja punktów A i B ujawnia, że proporcjonalnie rośnie energia z obwodu demagnetyzacji, natomiast jej zwrot jest w dalszym ciągu trudno zauważalny na przebiegach-oscylogramach. Podczas regulacji U_{wy} zmienia się iloczyn woltosekund, a za tym idzie, strumień



Rys. 38a, b



magnetyczny, wielkość prądu i energii w L_m , przeciwnie niż w sytuacji regulacji U_{we} i I_{wy} . Aby to pokazać, zdjęto trzy oscylogramy z zasilacza ustawionego na: $U_{wy} = 7,5V$ z obciążeniem $5,3A$ – **rysunek 39a**, $U_{wy} = 12,2V/7A$ – **rysunek 39b** i $U_{wy} = 19V/9A$ – **rysunek 39c**.

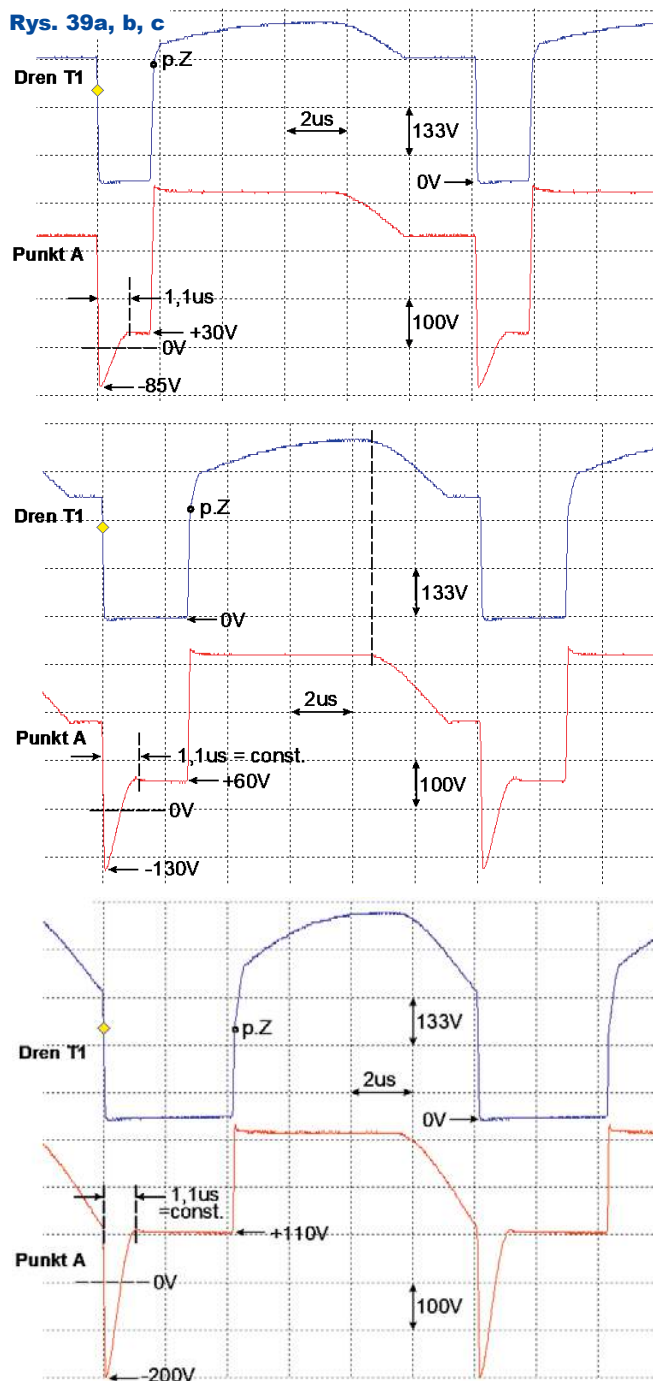
Na **rysunku 39a** ujemny szczyt napięcia w punkcie A to tylko $-83V$, a po przeładowaniu kondensatora w kierunku

dotadnim to tylko +30V. Na rysunku 39b to wartości odpowiednio -128V i +60V, a na rysunku 39c: -200V i +110V. Warto zauważyć, że czas przeładowania kondensatora C1 jest stały. Nie zależy on od ilości przeładowywanej energii. Jest on równy połowie okresu oscylacji własnych obwodu L1-C1, tu ok. 1 mikrosekundy. Najistotniejszym punktem obserwacji jest to, jak przesuwają się punkty Z na przebiegu obracającym napięcie na drenie klucza. Na kolejnych rysunkach punkt ten obniża się i o to właśnie chodzi! To punkt, w którym zaczyna przewodzić dioda D1 i zaczyna się zwrot energii. Na górnym przebiegu jest to widoczne jedynie jako spowolnienie zbocza napięcia narastającego na drenie klucza, a gdzie indziej tego ważnego procesu praktycznie nie widać. Na rysunkach 39a, b i c widać natomiast, co warunkuje takie, a nie inne, położenie punktu Z. Na dolnym przebiegu, w punkcie A schematu widać, że tuż po wyłączeniu klucza napięcie dąży do napięcia zasilania Vcc 320V. Wysokość stromego zbocza po wyłączeniu klucza na górnym i dolnym przebiegu jest zawsze jednakowa. Odcinek ten jest równy różnicy Vcc i odzyskanego napięcia na C1 po przeładowaniu obwodem komutacji. Zatem punkt Z napięcia na drenie leży o tyle poniżej napięcia zasilania, jakie napięcie i energię dało się odzyskać na kondensatorze C1. To jest wniosek najważniejszy i widać też, że straty w procesie odzysku energii są duże. Na oscylogramach z rysunku 39 warto jeszcze zwrócić uwagę, jak zmienia się czas włączenia klucza i napięcie na kluczu wraz ze zmianą wartości Uwy. A także jak zmienia się w czasie, na osi czasu, punkt szczytowego napięcia na mosfecie kluczującym. Także na to, jakie są te relacje czasowe względem częstotliwości kluczowania, która jest tu wartością stałą, co wynika z budowy sterownika.

Na tym kończymy odpowiedź na pytanie zadane w poprzedniej części artykułu bieżącego cyklu. Ale zanim przejdziemy do przeróbek zasilacza, przyjrzymy się jeszcze bliżej rysunkom 38a i 38b. Tylko te dwa dotyczą pracy zasilacza z bardzo małym obciążeniem, a można z nich wysnuć wnioski, których nie widać na innych oscylogramach. Tu badaniom poddano zasilacz HEC-385AD, którego fragment schematu pokazany był na rysunku 36 w poprzedniej części. Uwy ustawiono

na wartość 19,5V i obciążono je prądem zaledwie 300mA – przebiegi na rysunku 38a. Na rysunku 38b napięcie wyjściowe ustawione było na 16V, bez obciążenia. Jedynym obciążeniem pozostał rezystor o wartości 180Ω istniejący oryginalnie w zasilaczu. To przy 16Voltach znikome obciążenie ok. 1,5 wata. Na obu rysunkach wyraźnie widać moment, w którym kończy się przepływ prądu w indukcyjności L_{out12}. To oznacza przewodność nieciągłą DCM w tej indukcyjności. Od tego momentu napięcie w węzle katod diod D4a, b chce wrócić do poziomu, jakie istnieje na kondensatorze wyjściowym. Czyli do wartości Uwy. Nie ma już bowiem żadnej przyczyny, aby utrzymywać tu inną wartość; przed punktem K było to 0V, a dokładniej ok. -1V. Ze względu na istniejące tu pojemności, jakieś rozproszone, pasywniczne, nie dzieje się to skokowo, lecz oscylacyjnie. Widać to, mimo dużego zmniejszania tłumienia obwodu LC. L to L_{out12}, indukcyjność gromadząca energię, której wartość to 75uH. Z przebiegu można się doliczyć, że wszystkie występujące tam pojemności muszą mieć w sumie około dwóch nanofaradów. Na tym samym rysunku widać też inną cechę, która może początkowo dziwić. Mimo dużej nastawionej wartości Uwy, czas włączenia klucza jest stosunkowo krótki. I co za tym idzie, mała jest wartość wypełnienia (Duty Cycle). A powiedziano przecież, że w przetwornicy przepustowej parametr ten jest wprost proporcjonalny do stosunku Uwy/Uwe. A Uwe nie zmieniano! W warunkach pomiarowych rysunku 38 była to wartość ok. 320V. Wytlumaczeniem tej pozornej niezgodności jest fakt, że teraz, w trybie DCM kalkulacja jest inna! Niezmienną

Rys. 39a, b, c



wartość musi mieć powierzchnia pod dolnym przebiegiem zaznaczona na obu rysunkach 38a i b. Tyle, ile wkładu wniesie zafalshowanie spowodowane trybem DCM, o tyle musi się skrócić czas aktywny włączenia klucza. I ujemne sprzężenie zwrotne zadba, że tak właśnie będzie.

Na przebiegach widzimy napięcie, trudniej jest zmierzyć przebiegi prądów. Na rysunku 38 zaznaczono, jak powinien wyglądać przebieg prądu w indukcyjności L_{out12}. Od punktu K do momentu następnego cyklu włączenia klucza prąd powinien być zerowy.

Przez czas włączenia klucza – liniowo narastający, a od momentu wyłączenia klucza do punktu K, czyli utraty ciągłości prądu – liniowo opadający. Stromość w każdym przypadku jest proporcjonalna do napięcia i odwrotnie proporcjonalna do wartości indukcyjności. Zatem zbocze narastające powinno mieć nachylenie $(62\text{V}-19,5\text{V})/75\mu\text{H}$, co daje wartość $0,56\text{A}/\text{mikrosekundę}$. Odpowiednio zbocze opadające $(19,5-1\text{V})/75\mu\text{H}=0,25\text{A}/\text{us}$.

Czy to wnosi jakieś istotne informacje? Znajac nachylenie i czas, doliczymy się wartości szczytowej, maksymalnej prądu. Zgodnie z rysunkiem 38a będzie to $0,56\text{A}/\text{us} \times 2,3\mu\text{s} = 1,28\text{A}$. Licząc ze zbocza opadającego, powinniśmy uzyskać tę samą wartość. Tu wyjdzie $1,32\text{A}$. Choć ta niewielka rozbieżność może być efektem błędów pomiaru, jest to istotny czynnik, który także widać na dolnym przebiegu rysunku 38a. Zbocze opadające przebiegu napięcia w węźle C, jak i narastające na drenie jest stosunkowo łagodne, a o przyczynach tego faktu powiemy dalej. Fragment przebiegu prądu w tym odcinku czasu nie jest liniowy, a powinien być kawałkiem paraboli. W każdym razie jeszcze ma rosnąć. Zafałszowanie spowodowane tym fragmentem jest tym większe, im dłużej trwa powrót ze stanu wysokiego do zera woltów w węźle C. I wyraźnie widać, że na rysunku 38b jest zdecydowanie większe niż na rysunku 38a. Pomijając ten niuans, liczymy dalej opierając się na rysunku 38a. Mając wartość szczytową, czas narastania i opadania prądu oraz pełny okres powtarzalności przebiegów, nietrudno wyliczyć wartość średnią. Chcąc uwzględnić odcinek paraboli, rachunki tylko niewiele się komplikują. Wyjdzie około $0,43\text{A}$. Gdy ten prąd przemnożymy przez $19,5\text{V}$, to powinniśmy uzyskać moc dostarczaną do wyjścia. Wychodzi $8,4\text{W}$. W warunkach pomiarowych rysunku 38a zasilacz był obciążony niewielką żarówką 24-woltową, która przy $19,5\text{V}$ pobierała prąd $0,3\text{A}$. To jest 6W . Do tego należy uwzględnić rezystor 180Ω istniejący w zasilaczu. Przy napięciu $19,5\text{V}$ wydzielił się na nim ok. 2W . Zgodność jest zatem zaskakująco duża.

Jak się dobrze przyjrzeć rysunkowi 38a, widać tu jeszcze jedną istotną informację. Szczegół oznaczony S to zwrot energii. Ta mała kłopotliwość, która wygląda na jakieś zakłócenie, to

właśnie zwrot energii przekazanej przez obwód demagnetyzacji do wyjścia.

Porównanie rysunków 38a i 38b wnosi dodatkowe informacje i ujawnia dalsze szczegóły pracy analizowanego zasilacza, przy czym można się tu posługiwać schematem nr 21. Napięcie wyjściowe nieznacznie obniżono do 16V , a żarówkę obciążającą usunięto. Czas aktywny włączenia klucza jeszcze się skrócił, choć nieproporcjonalnie w stosunku do obniżenia napięcia. Punkt K utraty ciągłości prądu w L_{out} przesunął się w lewą stronę. Czas powrotu napięcia w węźle C do wartości U_{wy} wygląda identycznie, choć zmieniła się wartość napięcia. Przeliczając wg danych odczytanych z oscylogramów na rysunku 38b, dostaniemy czasy narastania i opadania prądu odpowiednio: $(62\text{V}-16\text{V})/75\mu\text{H} = 0,6\text{A}/\text{us}$ i $(16\text{V}-1\text{V})/75\mu\text{H} = -0,2\text{A}/\text{us}$. Dalsze przeliczenie wartości maksymalnej i średniej da wartości $0,6\text{A}$ i $0,12\text{A}$. Przemnożone przez 16V daje ok. 2W , co znów z dużą dokładnością pasuje do rezystora obciążającego 180Ω . Ale z tego rysunku odczytamy informacje jeszcze ciekawsze. Dlaczego tutaj zbocze narastające na górnym przebiegu i opadające na dolnym jeszcze złagodziło?

Tu w całej okazałości widać cechy pracy przetwornicy przepustowej! Przebieg na uzwojeniu wtórnym odzwierciedla to, co jest na pierwotnym. Napięcia i prądy transformują się jak w normalnym transformatorze. Impuls napięcia na bramce MOSFET-a jest krótki i wobec braku trzeciego kanału w oscyloskopie dorysowano go na rysunku 38b. Ale aktywna część przebiegu na drenie jest szersza, aniżeli impuls na bramce. Z dolnego przebiegu widać, że w czasie przejściowym, zanim napięcie na drenie osiągnie poziom U_{zas} , prąd po stronie wtórnej jest bliski szczytu, a w tym przypadku to tylko $0,6\text{A}$. Po stronie wtórnej prąd ten płynie przez diodę D4a, a więc płynie przez uzwojenie wtórne. I jak przystało na normalny transformator, transformuje się do uzwojenia pierwotnego zgodnie z przekładnią. Z relacji napięć odczytamy, że przekładnia wynosi tu ok. $5:1$. Prąd $0,6\text{A}$ da więc swoje odbicie w uzwojeniu pierwotnym 120mA . Jakie zbocze da taki prąd na kondensatorze 220pF ? Przeliczenie da wartość ok. $540\text{V}/\text{us}$, co się dość dobrze zgadza ze stromością zbocza, odczytaną z górnego

przebiegu rysunku 38b. Te zależności ładnie widać, gdy obserwujemy pracę zasilacza z bardzo małym obciążeniem. W warunkach znamionowych prądy po obu stronach są dużo większe i zbocza wyglądają na idealnie strome.

A dlaczego w powyższej kalkulacji wzięliśmy wartość pojemności 220pF ? To pojemność kondensatora C3 (patrz rysunek 36). Kondensator C1 jest tu 20-krotnie większy i w tych warunkach pomiarowych utrzymuje napięcie bliskie zeru. Możemy go zatem zupełnie pominąć, a C3 widzimy jako włączony równolegle do obwodu dren-źródło MOSFET-a, tak jak jest np. na schemacie nr 19. Zatem na nim występuje zbocze narastające. A co ze zboczem opadającym w fazie włączenia klucza? Jest szybkie, krótkie, a niedobra wiadomość jest taka, że energia z C3 rozładowuje się w tranzystorze! Nie jest to jednak błąd czy niedopatrzenie konstruktor-skie. Warto tę energię poświęcić, aby zmniejszyć straty dynamiczne w kluczu poprzez zmniejszenie nakładania się prądu i napięcia w fazach przełączania. Przeliczenia ujawniają, że nie poświęcamy dużo. Energia na $C=220\text{pF}$ przy napięciu 320V , to ok. $11\mu\text{J}$. Przy częstotliwości pracy, jaką tu mamy, 82kHz daje moc na poziomie $0,8\text{W}$.

Przy okazji może budzić się pytanie i wątpliwość, czy na pewno równoważne są układy z kondensatorem C3 włączonym względem masy i obwodu plusa zasilania? W pierwszym przypadku wydaje się logiczne, że energia zgromadzona w pojemności rozładowuje się w tranzystorze. Ale w przypadku drugim włączenie klucza przeciętnie ładuje kondensator, a nie rozładowuje go? I tu powstaje ładne zadanie z elektrotechniki: udowodnić, że podczas ładowania kondensatora z napięcia stałego przez jakąkolwiek rezystancję, małą czy dużą, zawsze na straty pójdzie połowa energii! To znaczy, że niezależnie od wartości R, w rezystancji wydzieli się tyle samo energii, ile zgromadzi się na kondensatorze. To był przerywnik, kontynuujemy temat naszego zasilacza.

Do celów dydaktycznych autor zmierzył wartości wszystkich elementów, w tym indukcyjności po kłopotliwym wylutowaniu z płyty zasilacza i przeliczył, czy wszystko się zgadza. To lepiej obrazuje pracę układu, niż analiza gołego schematu. Ale te same rachunki można powtórzyć dla znalezienia dowolnej niewiadomej opierając się na pomiarach

oscilloskopowych zasilacza. Obmiary przeprowadzone w tej części dotyczyły zasilacza HEC-385AD. Zaczęliśmy od tego, że tu obwód odzysku energii magnetyzacji jest niby taki sam, jak na schemacie z rysunku 33, ale odniesienie obwodu komutacji względem masy, a nie źródła tranzystora, nie jest najszybsze, żeby nie ocenić tego jako konstruktorski błąd. Poza tym inne wartości mają tu elementy tego obwodu. Wartość pojemności $C3=220\text{pF}$ wygląda na dobrą rozsądną. Obwód komutacji ma tu większą pojemność

kondensatora, a mniejszą indukcyjność cewki. Częstotliwość własna obwodu jest zbliżona i w każdych warunkach półokres tej częstotliwości mieści się w czasie włączenia klucza ($f \approx 0,5\text{MHz}$ i $T/2 \approx 1\mu\text{s}$). Ale czy to jest obojętne? Czy ważny jest tylko iloczyn LC? Tu obwód ma niemal dwukrotnie mniejszą impedancję charakterystyczną, ok. 69Ω , co skutkuje większym prądem komutacji. Na tym zakończymy rozważania teoretyczne przybliżające, ale mimo wszystko dalekie od pełnego, zrozumienie pracy zasilacza

komputerowego w konfiguracji forward z jedn tranzystorowym kluczem. W następnej części zajmemy się już przebudówką tego typu zasilacza PC.



Karol Świerc
rtv@silnet.pl

Ciąg dalszy ze strony 15

Natomiast z maksymalną prędkością transmisji, w megabitach na sekundę (Mbps), sprawa jest jeszcze bardziej skomplikowana. Maksymalna prędkość wyznaczona jest przez liczbę (stopę) błędów – przekłamań, dopuszczalną w danym systemie. Tabela świadczy, że określana w klasach szerokość pasma przenoszenia nie zawsze przekłada się wprost proporcjonalnie na maksymalną prędkość transmisji. Ta bowiem zależy też od liczby wykorzystywanych par – skrętek. Dziś w naszych popularnych kablach ethernetowych wykorzystujemy tylko dwie z czterech par. W szybszych

systemach Ethernetu gigabitowego wykorzystuje się wszystkie cztery pary do jednoczesnego przesyłania danych.

Istotne jest też, jaką odmianę modulacji wykorzystano w danym standardzie. Otóż w kablu nie są przesyłane „surowe bity”, tylko są one w pewien sposób zakodowane. To jest kolejny szeroki i niełatwy temat, którego nie sposób tu omówić choćby w skrócie.

W podsumowaniu trzeba mocno podkreślić, że w grę wchodzi nie tylko szerokość pasma, zależna od szczegółów budowy kabla i jego długości. Dlatego dopiero wszystkie wspomniane wcześniej czynniki decydują, jaka jest maksymalna prędkość przesyłania informacji

w danym kablu, a właściwie w danym systemie – standardzie.

W grę wchodzi długość kabla: otóż zagwarantowana dla danej kategorii prędkość transmisji dotyczy określonej w normie długości – odległości. W praktyce oznacza to, że w krótszym kablu prędkość transmisji może być większa. Dla praktyka znaczenie ma pokrewny wniosek: jeżeli kabel jest krótki, to wymaganą prędkość transmisji uda się uzyskać z kablem niższej kategorii. Dlatego na niewielkie odległości komputerowy sygnał ethernetowy (np. z kamery IP) udaje się przesłać przez kable niskiej kategorii: telefoniczne, a nawet alarmowe.

Ciąg dalszy ze strony 55

Zdecydowana większość zamontowanych na płytce elementów jest w obudowach przystosowanych do montażu powierzchniowego (SMD) i to od nich zalecam rozpocząć lutowanie. W przypadku diod LED warto zwrócić uwagę na prawidłową biegunowość, ponieważ stosowne oznaczenie zazwyczaj jest umieszczone na spodzie diody. Jeszcze jednym elementem, wymagającym większej uwagi podczas montażu, jest układ scalony U1. Warto przylutować go na samym początku, ponieważ jest na środku płytki i pozostałe elementy mogą utrudniać późniejszy dostęp do jego nóżek. Na samym końcu polecam wlutować łącznik J1. Całkowicie zmontowany

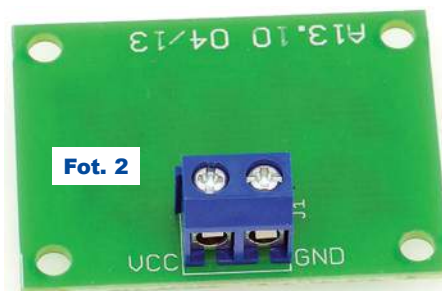
Wykaz elementów

R1,R3,R5,R7	18k Ω SMD1206
R2,R6	470 Ω SMD1206
R4,R8	1k Ω SMD1206
C1-C4,C6	10 $\mu\text{F}/16\text{V}$ SMD1206
C5	100nF SMD1206
LED1	czerwona SMD1206 np. LED SMD R1
LED2	zielona SMD1206 np. LED SMD G1
LED3	żółta SMD1206 np. LED SMD Y1
LED4	niebieska SMD1206 np. LED SMD B1
T1,T3,T5,T7	BC846 lub podobne
T2,T4,T6,T8	BCV61 np. BCV61C (opis w tekście)
U1	CD4093 S014
J1	ARK2 5mm

Komplet podzespołów z płytką jest dostępny w Sklepie AVT jako zestaw AVT3301

układ jest widoczny na **fotografii 1** oraz na **fotografii 2**.

Jeżeli cały montaż przebiegł pomyślnie, nie trzeba dodatkowych czynności uruchomieniowych. Układ



jest przystosowany do zasilania napięciem 5...6V i pobiera około 15mA. Dobrym źródłem zasilania dla niego może być ładowarka USB albo cztery baterie AA lub AAA połączone szeregowo. Po włączeniu naszym oczom powinno ukazać się światło emitowane przez szybko pulsujące diody LED.

Michał Kurzela
michal.kurzela@ep.com.pl

Stały konkurs: Co to jest?

Zadanie CoTo2202

Zadanie konkursowe brzmi:
Co przedstawia zamieszczona obok fotografia?

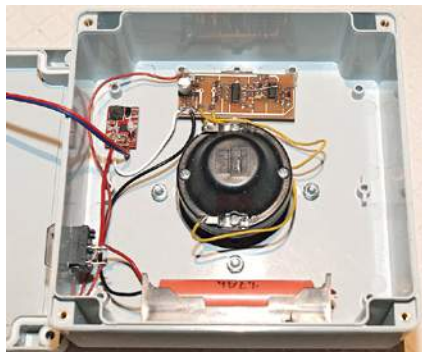
Prosimy o krótkie odpowiedzi.

E-maile z odpowiedziami należy przysyłać w ciągu miesiąca od ukazania się numeru, na adres:

konkursy@elportal.pl,
nie zapominając o podaniu adresu niezbędnego do wysyłki upominku.

W tytule e-maila należy podać nazwę konkursu, numer zadania i własne nazwisko, np. *CoTo2202Kowalski*.

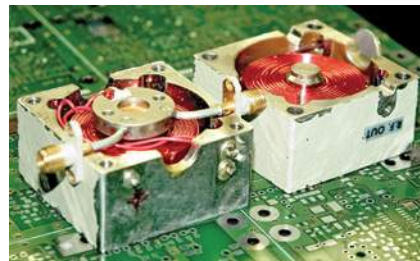
Wśród autorów prawidłowych odpowiedzi rozlosowane zostaną 3 kity AVT.



Rozwiązanie zadania CoTo2111

Fotografia pochodzi z artykułu „Uniwersalny sterownik filtra YIG”, którego autorami są Ireneusz Szulski i Rafał Orodziński.

Artykuł ukazał się w EdW 12/2020, na stronie 18.



Za prawidłowe odpowiedzi upominki w postaci kitów AVT otrzymują:

Zbigniew Jarzyna – Czechowice-Dziedzice,
Jarosław Węgliński – Warszawa,
Duliński Andrzej – Kraków.

W najbliższych numerach EdW planujemy

EdW 3/2022

Generator nanosekundowy

Generatory są ogromnie ważnym wyposażeniem warsztatu każdego elektronika. Rozwój generatorów DDS nie zmienił faktu, że jednym z najważniejszych przebiegów testowych jest fala prostokątna, mająca jak największą stromość zboczy.



EdW 4/2022

Ładowarka superkondensatorów

Naładowanie kondensatora wydaje się zadaniem trywalnym, niegodnym uwagi. Zadanie przestaje być trywialne, jeżeli ładowany kondensator ma pojemność rzędu tysiąca faradów lub więcej. Projekt przedstawia zaskakująco prosty sposób rozwiązania poważnego problemu.



EdW 5/2022

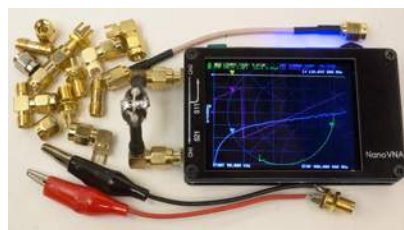
Zgrzewarka superkondensatorowa

Planowana do publikacji Ładowarka superkondensatorów ma służyć między innymi do współpracy ze Zgrzewarką superkondensatorową. Układ elektroniczny zgrzewarki jest dziecinnie prosty, ale działanie i efekty są imponujące.



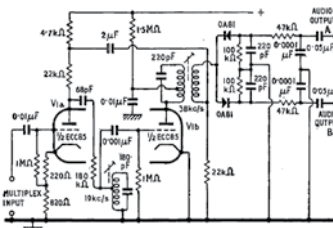
W kolejce na publikację czekają m.in.:

Różne projekty i artykuły edukacyjne przygotowywane przez Piotra Góreckiego przedstawiane są na stronie: <https://bit.ly/3aj0ixL> osiągalnej także za pomocą QR-kodu: gdzie możesz zadecydować o kolejności ich publikacji.



Pomiary pH wody i gleby

Elektronika stopniowo wkracza w nowe dziedziny. Między innymi pozwala mierzyć jakże ważny w wielu dziedzinach współczynnik pH.



Lampowe dekodery stereo?

Czy w epoce lampowej można było zrealizować dekodery stereo? Czy to jest w ogóle możliwe bez zastosowania tranzystorów oraz bez układów scalonych?

UWAGA!!

**Przysyłając rozwiązanie dowolnego konkursu,
NIE ZAPOMINAJCIE o podaniu w e-mailu pełnych danych adresowych.
Ich brak uniemożliwia wysłanie, a więc także przyznanie Czytelnikowi nagrody/upominku.**

Mamy do dyspozycji 4 rezystory o wartościach:
1Ω, 2Ω, 3Ω, 4Ω.

Wartości tych rezystorów należy wpisać do diagramu tak, aby każda wartość występowała dokładnie raz w danej kolumnie i rzędzie. Wartości oddzielone pustym polem tworzą grupy. Suma wartości (połączenie szeregowo rezystorów) w danej grupie musi odpowiadać liczbie podanej poza diagramem.

Jako rozwiązanie należy podać wartości rezystancji i puste miejsca w zaznaczonym obszarze.

Dla ułatwienia poniżej zamieszczamy przykładowy, uzupełniony diagram.

			3		3		4	2	4
			5	1	3	9	1	3	1
			2	9	4	1	5	5	5
1	3	6	x	1	x	3	x	2	4
3	7		3	x	1	2	4	x	x
9	1		x	3	2	4	x	x	1
6	4		4	2	x	x	1	3	x
8	2		1	4	3	x	x	x	2
10			x	x	x	1	2	4	3
2	4	4	2	x	4	x	3	1	x

					1					5	
					7	3	7	4	6	9	1
					3	6	3	6	4	1	4
	1	7	2								
4	2	1	3								
	3	1	6								
		7	3								
		1	9								
			10								
	2	3	5								

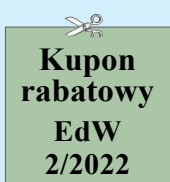
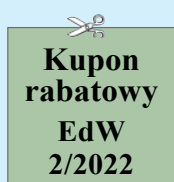
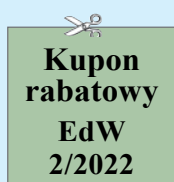
Autorem krzyżówki jest **Sławomir Kabat** z Niepołomic.

**Autor w nagrodę otrzymuje
6-miesięczną e-prenumeratę EdW.**

AVT stosuje system rabatów dla wszystkich wiernych Czytelników EdW, dokonujących zakupów w sieci handlowej AVT drogą sprzedaży wysyłkowej. Naklejenie na kartonik zamówienia trzech kuponów wyciętych z trzech kolejnych najnowszych wydań EdW uprawnia do: **10% zniżki** na zakup kitów AVT, TSM, Vellemana, **10% zniżki** na książki w ramach Księgarni Wysyłkowej AVT. **Już zakup na sumę 139 zł pozwala zaoszczędzić kwotę równą cenie jednego numeru EdW.**

Uwaga!

Zniżki dotyczą wyłącznie zamówień osób prywatnych.



Rozwiązaniem zagadki z EdW 10/2021 jest hasło: **REZYSTOR**

Upominki w postaci kitów AVT otrzymują:

Marek Kowal – Wojcieszków,

Krzysztof Mirosław – Puławy.

Damian Ząbczyk – Nowa Osuchowa.

Rozwiązania z tego numeru (tylko hasło) należy nadsyłać w ciągu 45 dni od ukazania się tego numeru EdW.

E-maile z rozwiązaniami powinny w tytule zawierać nazwę konkursu, numer zadania i nazwisko Czytelnika, np.

Krzyżówka2202Kowalski.

Listy powinny być opatrzone podobnym dopiskiem.

Uwaga! Przysyłając rozwiązanie krzyżówki, nie zapominajcie o podaniu e-mailu pełnego adresu. Jego brak uniemożliwia wysłanie, a więc także przyznanie Czytelnikowi upominku.

Natomiast przysyłając propozycję zagadki napiszcie: **Krzyżówka – propozycja** (żeby nie myliło się z rozwiązaniami). Wraz z propozycją nowej krzyżówki należy przysłać oświadczenie, że krzyżówka jest oryginalnym dziełem podpisanego i że nie była nigdzie publikowana. Redakcja nie ingeruje w treść merytoryczną (precyzję sformułowań) haseł krzyżówki.

EdW 2/2022 – lista osób nagrodzonych:

Emil C. Poznań
Circuit Chaos. Kraków
Andrzej Duliński. Kraków
Zygmunt Flisak. Opole
Mariusz Hejto. Łowczówek
Zbigniew Jarzyna. Czechowice-Dziedzice
Michał Jurkiewicz. Gdańsk
Sławomir Kabat. Niepołomice
Łukasz Kamiński. Jemiołowo

Jacek Konieczny. Poznań
Marek Kowal. Wojcieszów
Krzysztof Mirosław. Puławy
Andrzej Nowicki. Warszawa
Rafał Orodziński. Białystok
Bartłomiej Radzik. Warszawa
Paweł Sobótka. Szydłowiec
Michał Stach. Kamionka Mała
Tomasz Sukiennik. Kraków

Tadeusz Suszał. Warszawa
Tomasz Szombara. Jastrzębie-Zdrój
Karol Świerc. Ruda Śląska
Grzegorz Turek. Miedzyrzecz
Jarosław Węgliński. Warszawa
Teodor Woźniak. Łódź
Damian Ząbczyk. Nowa Osuchowa
Krzysztof Zych. Kraków

Uwaga! Jeśli do końca lutego poczta nie dostarczy osobie z powyższej listy przesyłki z nagrodą, prosimy zgłosić ten fakt redakcji (22 783 00 20, ewa.dudzik@elportal.pl)

Zajrzyj do interesujących materiałów „Świat Radio” 1–2/2022

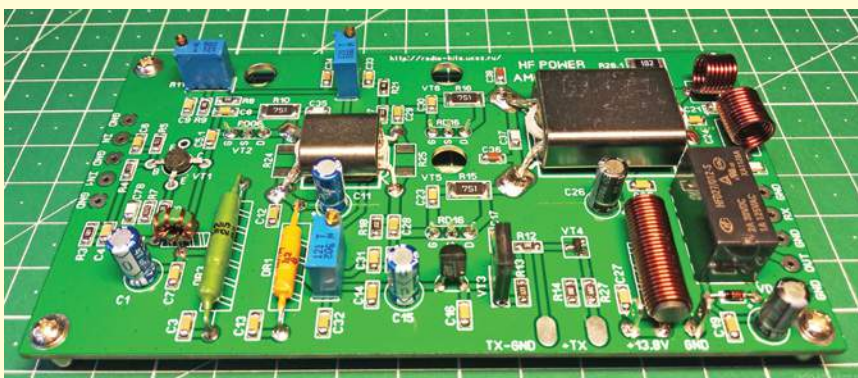


Prosty minitransceiver DSB/80m

Budowa urządzeń jednowstęgowych wymaga pewnego doświadczenia. Początkującym radioamatorom proponujemy wykonanie najpierw transceivera DSB/QRP np. na bazie AVT 174, który jest prostszy w uruchomieniu niż SSB, a zastosowany układ UL1242 (TBA120S) wciąż jest dostępny na rynku.

Wzmacniacz mocy UT2FW

Opisany w ŚR 1–2/22 wzmacniacz mocy na fale krótkie jest skonstruowany na tranzystorach RD06HVF1 oraz 2xRD16HHF1 i został przystosowany do transceiverów QRP w celu zwiększenia mocy maksymalnie do 50W. Konstrukcja jest oparta na opracowaniu UT2FW, a kit urządzenia przygotowany przez UR3IQH, łącznie z filtrem LPF, jest dostępny w sieci.



Errare Humanum Est

EdW 1/2022

Str. 5, w połowie wstępniaka jest: RS484, powinno być: **RS485**.

Str. 27, pierwsza kolumna, 9 wiersz od dołu – jest: –18V na stopień Celsjusza, powinno być: **–18mV**.

Uwagi nadesłał Marian Gabrowski z Polkowic

Str. 17. W felietonie „Wynalazki i wynalazcy” błędnie zostało opisane włączanie silników trójfazowych dużej mocy. Powinno być: W silnikach dużej mocy stosowane jest połączenie gwiazda–trójkąt. W momencie rozruchu silnik jest przełączany w gwiazdę, dzięki czemu uzwojenia silnika zasilane są niższym napięciem i silnik powoli wchodzi na obroty. Gdy silnik osiągnie około 0.7–0.8 obrotów nominalnych należy przełączyć go na trójkąt i dalej pracuje w normalnych warunkach. Autor felietonu dziękuje Panu D. Adamskiemu za przekazaną uwagę.

R E K L A M A

AVT 1980 Czasowy włącznik zbliżeniowy

Czas załączenia jest regulowany w zakresie od ok. 10 sekund do 5 minut. Urządzenie jest przystosowane do zasilania napięciem 8...12 V DC.



Znajdź nas na 



Porady warsztatowe

Kilka przydatnych wskazówek

Myjka ultradźwiękowa

Zdaniem autora jest to urządzenie potrzebne elektronikowi, tak samo jak lutownica. Głównym parametrem myjki jest moc ultradźwięków. Myjki ultradźwiękowe można kupić zarówno w dyskontach, jak i na portalach aukcyjnych. Z myjek dyskontowych najlepiej sprawdziła się autorowi myjka kupiona... w Lidlu – **fotografia 1**. Lepsze są myjki profesjonalne o mocy ultradźwięków ponad 100W. Zbudowane są one całkowicie z metalu, ale są znacznie droższe niż te, które możemy kupić w dyskontach. Wsadzenie palca do myjki o dużej mocy ultradźwięków i dotknięcie do jej ścianki, gdy jest ona wypełniona płynem, powoduje, że odczuwamy silne ciepło, a martwy naskórek odrywa się od powierzchni skóry. W myjce, dzięki silnym drganiom cząsteczek rozpuszczalnika eliminowany jest wpływ napięcia powierzchniowego, co powoduje, że wnikają one w najdrobniejsze szczeliny i usuwają zabrudzenia. Podczas mycia wzrasta temperatura środka myjącego, co zwiększa rozpuszczalność wielu substancji. Za pomocą myjki ultradźwiękowej można usunąć resztki topnika np. spod układu SMD, czego nie możemy zrobić w żaden inny sposób.

Aby proces mycia był skuteczny, myjka powinna być wypełniona cieczą zgodnie z zaleceniami producenta. W przypadku dużej ilości zabrudzeń autor ogranicza objętość rozpuszczalnika, co zwiększa moc ultradźwięków przypadającą na jednostkę objętości a tym samym poprawia skuteczność czyszczenia. Środki, jakimi czyścimy płytki drukowane, zależą od rodzaju substancji stosowanych podczas jej lutowania. Pamiętajmy zasadę „podobne rozpuszcza się w podobnym”, a więc substancje polarne w rozpuszczalnikach polarnych, zaś niepolarne w niepolarnych. Topniki do lutowania stopów niklu, zawierające kwas fosforowy, rozpuszczają się dobrze w wodzie. Autor do wstępnego czyszczenia takich płytek stosuje np. denaturat lub

izopropanol o zawartości wody około 20 procent. Dodatek do wody kilku kropli detergentu obniża napięcie powierzchniowe i ułatwia czyszczenie elementów. Pozostałości wody z płytki usuwa za pomocą alkoholu izopropylowego o jak najmniejszej zawartości wody. Izopropanol bardzo dobrze pochłania wodę. Autor odradza stosowanie denaturatu do mycia końcowego ze względu na zawartość w nim wody. Autor spotkał się z denaturatami, które zawierały do 30% wody.

Kalafonię z płytki drukowanej i typowe topniki można usunąć za pomocą alkoholu izopropylowego. Do mycia wstępnego takiej płytki kilkuprocentowy dodatek wody absolutnie nie szkodzi, jednak do mycia końcowego należy stosować izopropanol o jak najmniejszej zawartości wody. Obecnie izopropanol ma cenę niższą niż denaturat. W początkowym okresie pandemii cena izopropanolu bardzo wzrosła, gdyż ma on działanie odkażające silniejsze od alkoholu etylowego.

Autor stosuje dwufazową procedurę czyszczenia płytki drukowanej. Czas czyszczenia zależy od mocy ultradźwięków, jaką ma myjka. Pierwsza faza, alkohol izopropylowy z dodatkiem około 5–7 procent wody. W czasie tego czyszczenia następuje przerwa, podczas której „grubsze” zabrudzenia usuwane są za pomocą miękkiej szczoteczki do zębów. Następnie proces czyszczenia płytki drukowanej kontynuowany jest w myjce ultradźwiękowej. Gdy płytka jest już w miarę czysta, płytka czyszczona jest jeszcze raz szczoteczką do zębów i przepłukiwana za pomocą alkoholu z myjki. Pozostały alkohol zlewany jest do szczelnego pojemnika (butelka po mleku) przez lejek z watą, na której osadzają się grubsze zanieczyszczenia.

Dalej następuje drugi etap czyszczenia, podczas którego płytka ponownie jest myta w alkoholu izopropylowym,



Fot. 1

tym razem o jak najmniejszej zawartości wody. Gdybyśmy pozostawili płytkę do wyschnięcia po pierwszym etapie, to pozostałyby na niej brzydkie ślady resztek topnika, rozłożone na całej płytce drukowanej. Alkohol użyty kilkakrotnie do ostatniego czyszczenia autor wykorzystuje podczas pierwszego etapu czyszczenia. Czyszczenie za pomocą myjki należy prowadzić w dobrze wentylowanym pomieszczeniu, a podczas czyszczenia używać rękawiczek jednorazowych (można ich używać wielokrotnie) w celu ochrony skóry przed alkoholem izopropylowym, który ma silne działanie wysuszające na skórę. W przypadku stosowania myjek dyskontowych i izopropanolu prawdopodobnie możemy pożegnać się z gwarancją producenta, gdyż pod jego wpływem schodzi np. część farby stosowanej w wykończeniu obudowy myjki.

Ponieważ podczas mycia część alkoholu, acetonu zawsze odparowuje, należy stosować pokrywkę, która redukuje straty rozpuszczalnika. W przypadku mycia acetonem autor stosuje prowizoryczną przykrywkę z pokrywki pojemnika do lodów. Aceton rozpuszcza oryginalną pokrywkę plastikową myjki.

Na zakończenie artykułu autor chce podziękować Waldkowi 3Z6AEF za uwagi do tego tekstu.

Rafał Orodziński
sq4avs@gmail.com

Moje pasje – własnej roboty magnetofon szpulowy

część 1

Moje zainteresowanie magnetofonem zaczęło się kilka lat później niż radiem. Gdy jednak połąknałem bakcyla magnetofonowego, to obydwie dziedziny uprawiałem z jednakową pasją.

W czasie nauki w szkole zawodowej miałem okazję obejrzeć przystawkę magnetofonową Viola. Gdy dowiedziałem się, że można za jej pomocą nagrywać z mikrofonu lub bezpośrednio z radia, postanowiłem poznać temat bliżej. Bardzo mnie ciekawiło, jak przebiega nagrywanie dźwięku na taśmę, czy taśma przechodzi przez szczelinę głowicy, czy tylko się przed nią przesuwają. Nie mając na ten temat żadnej literatury, niektóre sprawy mogłem tylko sobie wyobrazić, polegając na niedokładnych, zasłyszanych opisach.

Po ukończeniu Zasadniczej Szkoły Zawodowej w Kole zostałem przez dyrekcję skierowany do Technikum Przemysłowo-Pedagogicznego w Pabianicach. W tym mieście działał Radioklub LPŻ dość dobrze wyposażony i skupiający grupę miłośników radia. Szybko się do niego zapisałem. Mimo że radiotechnika była moją wielką pasją, to coraz bardziej wciągała mnie również technika magnetofonowa. W radioklubie nawiązałem kontakt z innymi pasjonatami, z którymi wymieniałem się zdobytą wiedzą oraz skromną literaturą. Ta dotycząca magnetofonów dopiero zaczynała się trochę rozkręcać. W latach 50. w „Radioamatorze” zaczęły się ukazywać interesujące artykuły na temat budowy i działania magnetofonu. W roku 1956 ukazała się książka „Magnetofon taśmowy”, R. Girulskiego i J. Różyckiego. Była to pierwsza przystępnie opracowana pozycja dotycząca działania i budowy magnetofonu. Autorzy książki publikowali również artykuły w „Radioamatorze” i w latach 50. i 60. należeli do ścisłej czołówki znawców dziedziny magnetofonowej.

Młodym Czytelnikom pewnie niełatwo zrozumieć, że tak trudno było dosłownie o wszystko. Z perspektywy

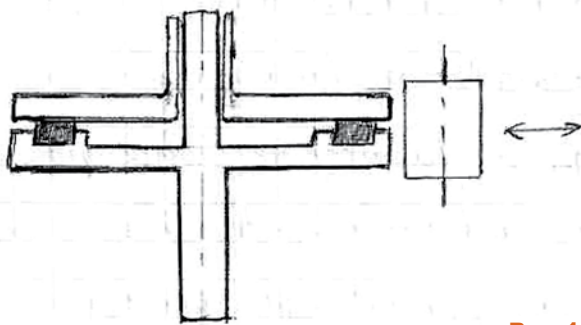
czasu dostrzegam jednak także dobre strony tych niedostatków. Ta z trudem zdobywana wiedza, ciągle poszukiwanie potrzebnych części i głowienie się nad rozwiązywaniem problemów dopingowały do doskonalenia swej wiedzy i umiejętności. Niedobór części był również jednym z bodźców do opracowania i wdrażania własnych pomysłów.

Z tamtych czasów zachowałem wielki szacunek do książek i do ludzi nauki, którzy swoją wiedzą dzielą się z innymi. Dziś do takich osób należą Piotr Górecki, a także Karol Świerc.

W drugiej połowie 1959 roku, gdy zdobyłem podstawową wiedzę o budowie i działaniu magnetofonu, zacząłem czynić przygotowania do zbudowania własnego. Rok później zostałem zatrudniony jako tokarz w Pabianickiej Fabryce Urządzeń Mechanicznych i dzięki temu mogłem samodzielnie wykonać również część mechaniczną.

Budowę zacząłem od przygotowania dokumentacji technicznej oraz dostosowania schematu do posiadanych części. Niestety, nie mogę zilustrować mojej opowieści żadnymi fotografiami, ponieważ podczas budowy magnetofonu nie zadbałem o dokumentację fotograficzną, a później wykonanie zdjęć było niemożliwe z powodów, o których napiszę dalej. Dlatego pozostaje opis słowny, mam nadzieję dość dokładny. Na początek przedstawię wersję pierwszą.

Otóż w założeniu miał to być prosty magnetofon oparty na opisach w książce oraz w „Radioamatorach”. Miał mieć jedną prędkość przesuwu taśmy – 19,05 cm/s, bez możliwości szybkiego przewijania. Nie podobało mi się takie rozwiązanie i postano-



Rys. 1

wilem zrobić magnetofon z szybkim przewijaniem taśmy w obie strony. Nie mogąc znaleźć rozwiązania sprzęgła, które by mi się podobało, wymyśliłem bardzo proste rozwiązanie. Postanowiłem wykonać obydwa talerzyki (dolny i górny) o tej samej średnicy a zaszprzęglenie uzyskiwać przez dociskanie do obydwu krawędzi gumowej obrotowej rolki (**rysunek 1**).

Ponieważ nie miałem pomysłu na sterowanie obydwoma sprzęgłami za pomocą jednego pokrętkła, każdą rolkę dociskową (sprzęgającą) wyposażyłem w odrębną dźwignię. Talerzyki górne według opisu powinny mieć po jednym bolcu zabezpieczającym szpulę przed swobodnym obracaniem się. Ja jednak wykonałem po trzy skrzydełka przylutowane do tulejki na wzór magnetofonu fabrycznego. Wałek, na którym miało być umieszczone koło zamachowe, wykonany jako całość z rolką przesuwu, toczyłem i szlifowałem w kłach tokarki z dokładnością 0,005–0,01 mm. Z tą samą dokładnością wykonałem też wałki do sprzęgieł oraz panewki z brązu. Końcową obróbkę koła zamachowego wykonałem na trzpieniu, aby jak najbardziej zminimalizować bicie promieniowe (ekscentryczność). Rolkę dociskową, rolki sprzęgające oraz kołki prowadzące taśmę wykonałem z dokładnością 0,02–0,05 mm. Koledzy z hartowni bardzo ładnie poczernili koło zamachowe i sprzęgła tak, że wyglądały jak

firmowe. Do napędu zastosowałem silnik podobny do adapterowego, lecz o większej mocy, stosowany w wentylatorach domowych. Silnik został umieszczony w ekranie z miękkiej, dobrze wyżarzonej blachy stalowej.

Cała część mechaniczna została umocowana na płycie stalowej o grubości 2mm i umieszczona w skrzynce drewnianej oklejonej dermą. Skrzynka miała z boku wycięty otwór na głośnik eliptyczny, który był zasłonięty plastikowym reflektorem od radia. Przy wykonywaniu sprzęgieł popełniłem poważny błąd, stosując łożyska kulkowe zamiast panewek. Po uruchomieniu okazało się, że łożyska te wytwarzały bardzo duży szum, wzmacniany jeszcze przez płytę nośną i skrzynkę. Ponieważ, jak mawiał Seneka „Errare humanum est”, sprzęgła osadziłem w panewkach i błąd został naprawiony.

Po dokonaniu tej poprawki mechanizm pracował bardzo cicho, przyszła więc kolej na część elektryczną: zasilacz, wzmacniacz zapis-odczyt, generator oraz głowice. Mimo że zasilacz był mniej więcej taki jak w przeciętnym radiu, to transformator musiałem wykonać sam. Trafa stosowane w odbiornikach radiowych miały zbyt duży prąd jałowy rzędu 120–150mA i wytwarzały duże pole rozproszenia. Nawinięty przeze mnie transformator miał prąd jałowy ok. 30mA i dodatkowe uzwojenie 6,3V do zasilania pierwszej lampy wzmacniacza. Uzwojenie to zostało spięte drutowym opornikiem nastawnym, którego ślizgacz był dołączony do punktu masy pierwszej lampy. Wzmacniacz wzorowałem na schemacie umieszczonym w książce, lecz dostosowałem go do posiadanych lamp. Jako pierwszą zastosowałem EF86, która jest najlepszą lampą do tego celu. Jako drugą EF21, jako głośnikową EL84, a generator zbudowałem na 6Π9. Do dziś pamiętam, jak z palca poleciała smułka dymu i zapiekło, gdy podczas regulacji prądu kasowania dotknąłem gorącego końca głowicy.

Wzmacniacz i generator zostały zmontowane w blaszanym pudełku, lampy zaekranowane i całość umiesz-

czona daleko od zasilacza i silnika. Po zmontowaniu całości i sprawdzeniu, czy dobrze działa, pozostało do wykonania głowice według opisu zamieszczonego w książce.

Materiał na rdzenie głowic pozyskałem ze znajdującego w radioklubie transformatora międzystopniowego, który miał rdzeń z dobrych blach o grubości 0,3mm. Najpierw narysowałem kształt połówki rdzenia, potem wyciąłem 40 prostokątów o wymiarach połówki rdzenia z ok. 1mm nadładkiem na ostateczną obróbkę. Po zaznaczeniu miejsc na otwory do znitowania, wywierciłem je wiertłem 1mm i każdy otworek dokładnie oczyściłem. Następnie znitowałem nitami miedzianymi 4 pakiety po 10 blaszek. Potem pakiety zostały obrobione pilniczkami według narysowanego kształtu. Po obróbce nity usunąłem, wszystkie blaszki dokładnie oczyściłem i z jednej strony pomalowałem cienką warstwą lakieru nitro. Po jego wyschnięciu poszczególne pakiety ponownie znitowałem, ale bardzo ostrożnie, aby nie uległy zdeformowaniu.

Najtrudniejszą i najbardziej pracochłonną czynnością było takie doszlifowanie wewnętrznych połówek rdzenia, aby po ich złożeniu szczelina między nimi była prawie niewidoczna. W tym celu zdobyłem tzw. kamień missisipi, o bardzo równej powierzchni i na nim szlifowałem połówki rdzeni. Po skończeniu szlifowania należało zrobić karkasy do cewek. Wyciąłem z bakelitu dwie kostki o odpowiednich wymiarach i z nich za pomocą wiertła i iglaków w ciągu kilku godzin wykonałem karkasy. Na karkas głowicy uniwersalnej nawinałem około 1800 zwojów ϕ 0,06mm CuE, a na karkas głowicy kasującej 150 zwojów ϕ 0,25. Szczelinę roboczą głowicy uniwersalnej uzyskałem przez włożenie między przednie strony połówek rdzenia paski folii aluminiowej o grubości 0,008mm. Szczelinę tylną uzyskałem przez włożenie między połówki rdzenia paski papieru o grubości 0,1mm. Szczelinę roboczą głowicy kasującej uzyskałem przez włożenia paski blachy mosięż-

nej o grubości 0,3mm. Szczeliny tylnej głowicy kasującej nie ma.

Tylne połówki rdzeni zostały włożone do cewek, a całość ściśnięta między mosiężnymi nakładkami śrubą M4. Przed skręceniem nakładek połówki rdzeni zostały wyrównane i dociśnięte. Po skręceniu, czoła głowic były jeszcze dokładnie szlifowane na marmurku (to taka drobnoziarnista ośleka, używana kiedyś do ostrzenia brzytwy). Po zamontowaniu głowicy na płycie nośnej, głowicę uniwersalną dokładnie zaekranowałem osłoną z miękkiej blachy stalowej i przystąpiłem do prób.

Po włączeniu usłyszałem w głośniku wyraźne buczenie – przydźwięk sieci. Między jeden koniec głowicy uniwersalnej a masę włączyłem cewkę kompensacyjną i ustawiając ją w różnych położeniach, znalazłem takie, w którym przydźwięk sieci był niezauważalny. Potem było ustawianie głowicy według taśmy wzorcowej, dobieranie właściwego prądu podkładu i kasowania oraz pierwsze nagranie.

Do dziś pamiętam pierwszą nagraną piosenkę, (a była to *Zapomniana piosenka*) i radość, gdy odtworzona brzmiała czysto i bez zniekształceń. Mój swojej roboty magnetofon wernie odtwarzał nagrane utwory i umożliwiał szybkie przewijanie taśmy w obie strony.

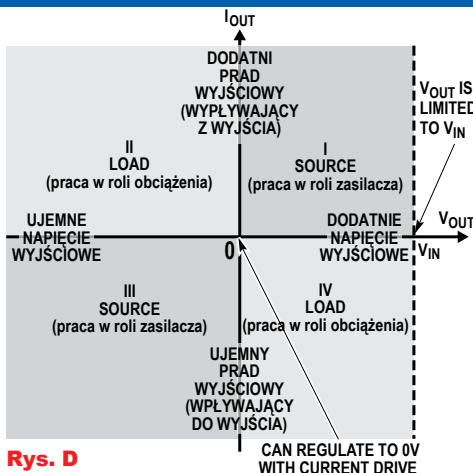
Ja jednak czułem pewien niedosyt. Byłem zafascynowany magnetofonem „Smaragd”, którego funkcje były sterowane przełącznikiem klawiszowym. Dlatego po miesiącu postanowiłem zrobić część mechaniczną prze-

robić na wzór „Smaragda”. O tym opowiem w drugiej części artykułu.

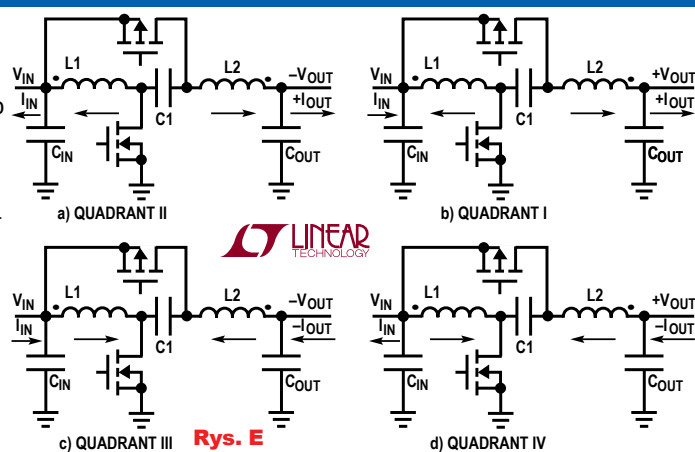


Jerzy Szymański
j.szymanski@wp.eu

Otóż samo nazwanie przetwornicy impulsowej czteroćwiartkową (lub czterokwadrantową – *four quadrant converter*) nie do końca określa jej funkcje i możliwości. Wygląda też na to, że nie wszyscy uczestnicy zadania prawidłowo zrozumieli jej działanie. Niedawno także w EdW mówiliśmy o przetwornicach dwukierunkowych – dwukierunkowych w sensie kierunku przekazywania energii. Z kolei o czterech ćwiartkach mówiliśmy w przypadku silników elektrycznych. Samo sformułowanie „cztery ćwiartki” w przypadku przetwornicy i zasilacza słusznie wskazuje na pracę przy napięciach i prądach zarówno dodatnich, jak i ujemnych.



Rys. D



Rys. E

Dlatego na początek trzeba wyjaśnić prosty, ale dość ważny szczegół, nie dla wszystkich oczywisty. Owszem, na wyjściu (V2) mogą występować napięcia oraz prądy zarówno dodatnie, jak i ujemne, ale na wejściu (V1) napięcia nie mogą być ujemne. Napięcie wejściowe V1 może być jedynie dodatnie (0...80V), przy czym prawidłowa praca jest gwarantowana dla napięć V1 od 4,5V wzwyż.

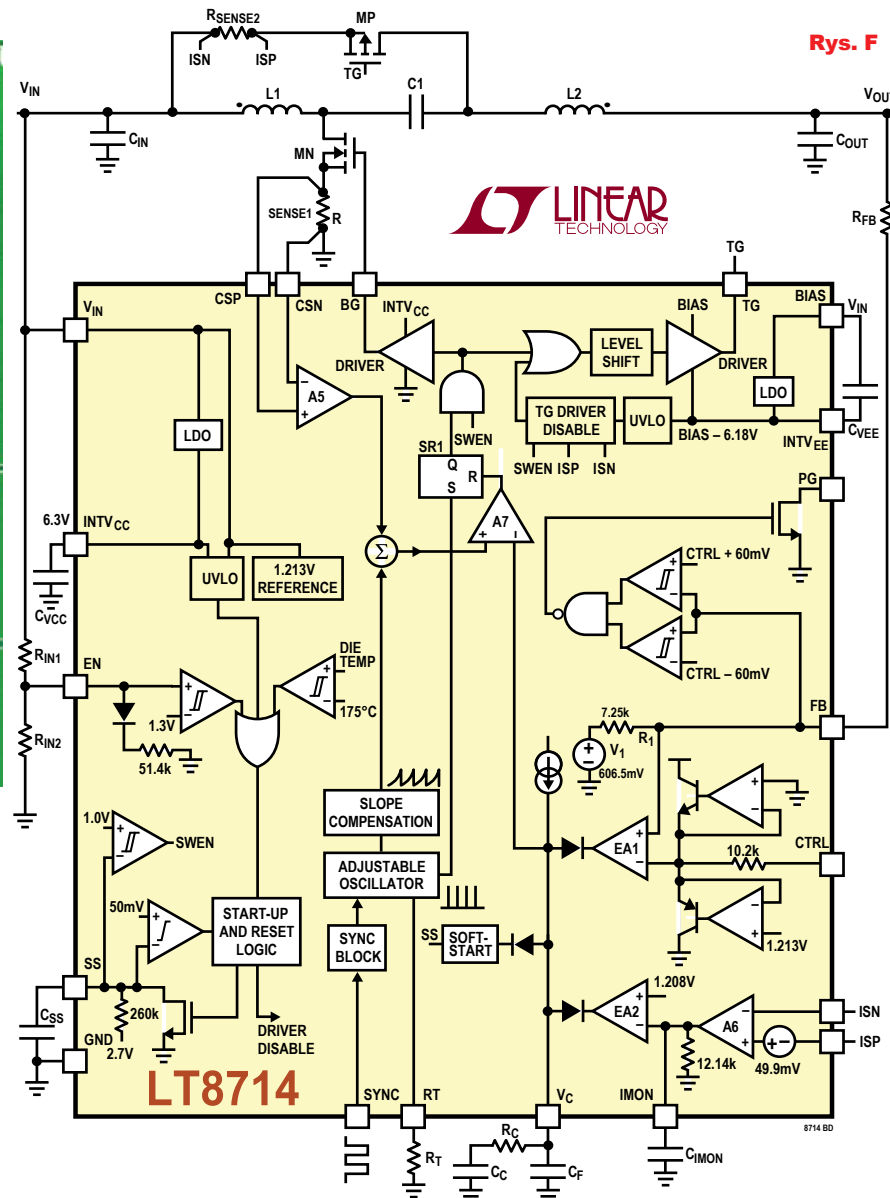
Napięcie V1 podczas pracy ma być dodatnie, ale przetwornica może przenosić energię w dwóch kierunkach, dlatego zależnie od trybu pracy prąd wejściowy może być dodatni (gdy przetwornica pobiera energię ze źródła napięcia V1) albo ujemny (gdy przetwornica „wycofuje” energię do źródła napięcia V1). Określenie: *czteroćwiartkowa* odnosi się tylko do wyjścia tej przetwornicy: do napięcia V2 i prądów tam płynących, jak ilustruje to **rysunek D**. Wyjście przetwornicy może być źródłem energii, źródłem zasilania, może dostarczać energii do obciążenia – podczas pracy w I i III ćwiartce. Wyjście przetwornicy może też być obciążeniem, zamiast dostarczać energię, wyjście przetwornicy może pobierać energię ze współpracującego obwodu, który nie jest wtedy obciążeniem, tylko nietypowo źródłem energii – to praca w ćwiartkach II i IV. Pochodzący z karty katalogowej, tylko nieco zmodyfikowany **rysunek E** pokazuje pracę układu w czterech ćwiartkach – kwadrantach. Zaznaczona jest biegunowość napięć wyjściowych, a strzałki pokazują podstawowy kierunek przepływu prądów (który dla cewek może być inny w poszczególnych fazach cyklu pracy).

Rysunek F pokazuje wewnętrzny schemat blokowy układu scalonego. Omawiana konstrukcja zawiera szereg interesujących rozwiązań. Najbardziej zainteresowani poszukają ich samodzielnie w karcie katalogowej.

Nagrody-upominki za zadanie **JakDziałal10** otrzymują:

Grzegorz Turek – Międzyrzecz,
Paweł Sobótko – Szydłowiec,
Michał Jurkiewicz – Gdańsk.

Wszyscy uczestnicy konkursu zostają dopisani do listy kandydatów na bezpłatne prenumeraty.



Rys. F



sklep.avt.pl

Produkty z oferty i wyroby AVT
można nabyć na kilka sposobów:

W sklepie internetowym:

sklep.avt.pl

W sklepie firmowym AVT:

Warszawa - Żerań
ul. Leszczynowa 11



Wypełniając poniższy formularz zamówienia

Formularz należy wysłać na adres:

AVT SPV Sp. z o. o.
03-197 Warszawa
ul. Leszczynowa 11



prześlij na adres:

AVT SPV Sp. z o.o.
03-197 Warszawa
ul. Leszczynowa 11

Miejsce na
kupon
rabatowy
EdW 12/2021

Miejsce na
kupon
rabatowy
EdW 1/2022

Miejsce na
kupon
rabatowy
EdW 2/2022

Tu wklej kupony z ostatnich 3 numerów EdW
a uzyskasz **zniżkę 10%** dla stałych czytelników.
(szczegóły na stronie 73)
Prenumeratzy nie muszą wklejać kuponów,
wystarczy, że podadzą nr prenumeraty!

ZAMÓWIENIE na artykuły z oferty AVT

Kity

Opis oznaczeń wersji kitów:

- [A] płytki drukowane PCB
- [UK] zaprogramowany układ
- [A+] płytki PCB i zaprogramowany układ
- [B] płytki PCB (lub płytki), UK (jeśli występuje) i komplet elementów elektronicznych wymienionych w dokumentacji zestawu.
- [C] zestaw zmontowany

Numer kitu AVT	A	A+	B	C	UK

Inne artykuły z oferty AVT

Kod – Nazwa	Ilość

Nadawca:
imię i nazwisko adres nadawcy prenumeraty

Adres:

☐ wysyłka pobraniowa kurierem: 19zł

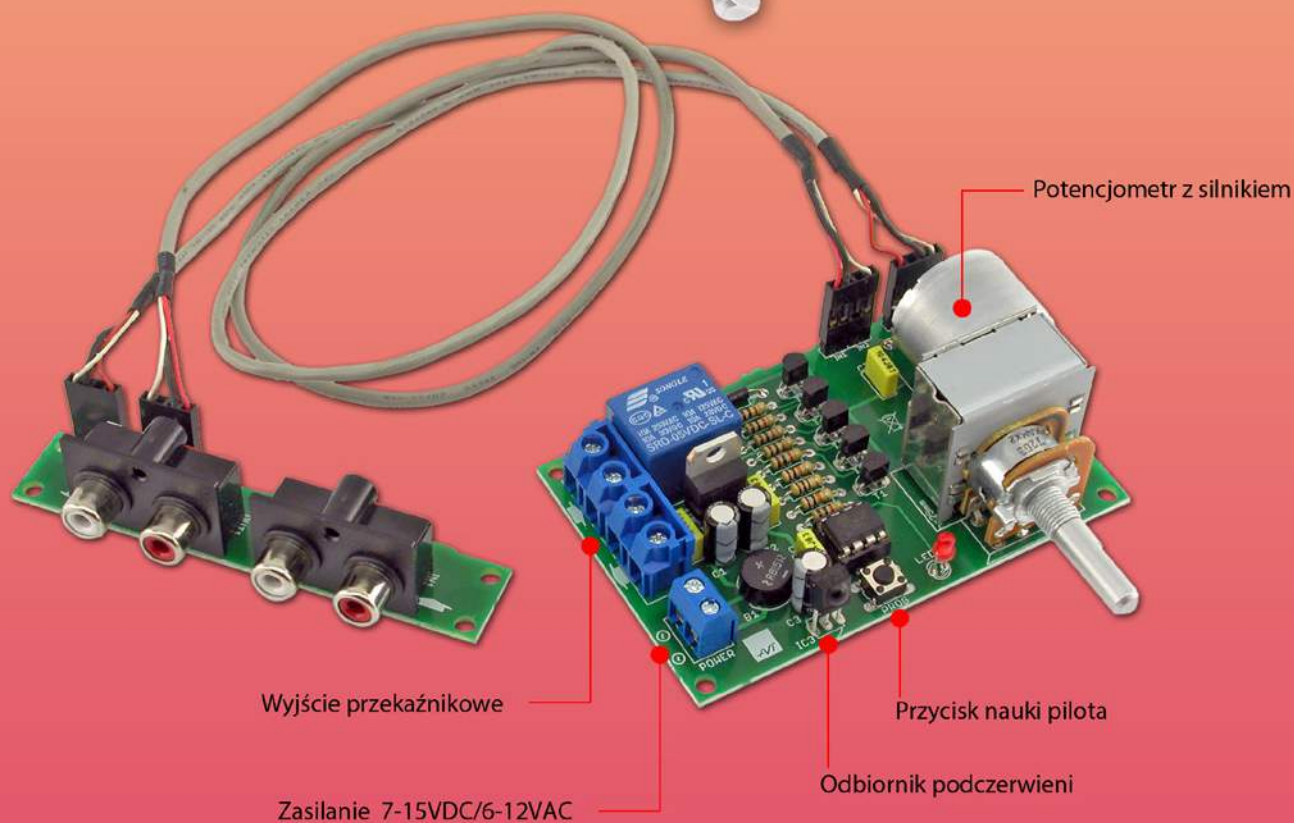
AVT 3222 Sterowany dowolnym pilotem potencjometr audio

Urządzenie doskonale nadaje się do każdego wzmacniacza audio wyposażonego w standardowy, „ręczny” potencjometr. Układ może być sterowany praktycznie dowolnym pilotem na podczerwień IR od sprzętu powszechnego użytku. Wymaga tylko przeprowadzenia prostej procedury zapamiętywania kodów pilota. Był testowany z kilkunastoma pilotami od różnych telewizorów, dekodów, DVD i sprzętu audio – z każdym działał prawidłowo. Dowolny przycisk pilota można przypisać do jednej określonej funkcji: włącz/wyłącz, obroty w lewo (ciszej) lub obroty w prawo (głośniej).

AVT3222 - KIT do samodzielnego zlutowania - 95zł

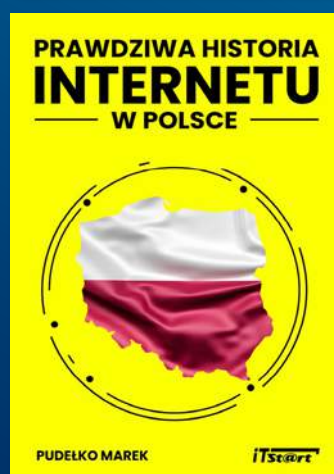
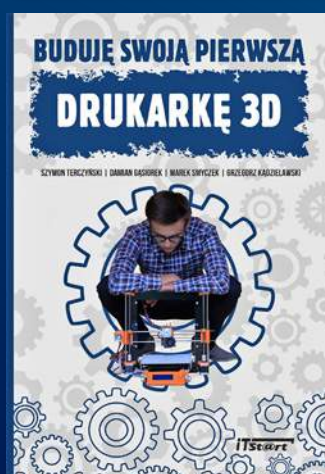
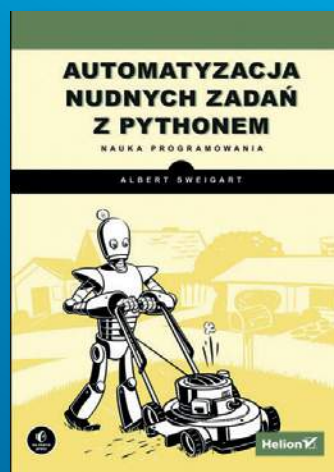
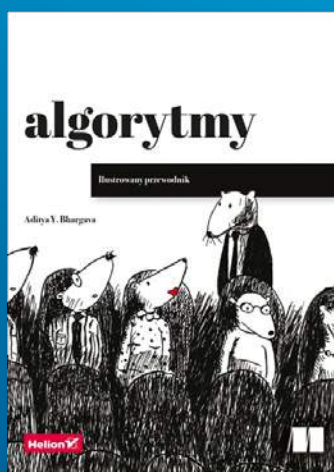


Zdjęcie przedstawia zlutowany KIT AVT3222



KSIĄŻKI W ULUBIONYM KIOSKU

Wybieraj spośród naszych nowości i bestsellerów!



Zobacz pełną ofertę książek na UlubionyKiosk.pl