

Uwaga – to skrócona zapowiedź – artykuły wstępnie planowane do następnego numeru.
Zapowiedź pełną mogą pobrać tylko Patroni z progów ≥ 20 zł: <https://patronite.pl/Zrozumiec-Elektronike>

4/2026 Kwiecień (40)

piotr-gorecki.pl

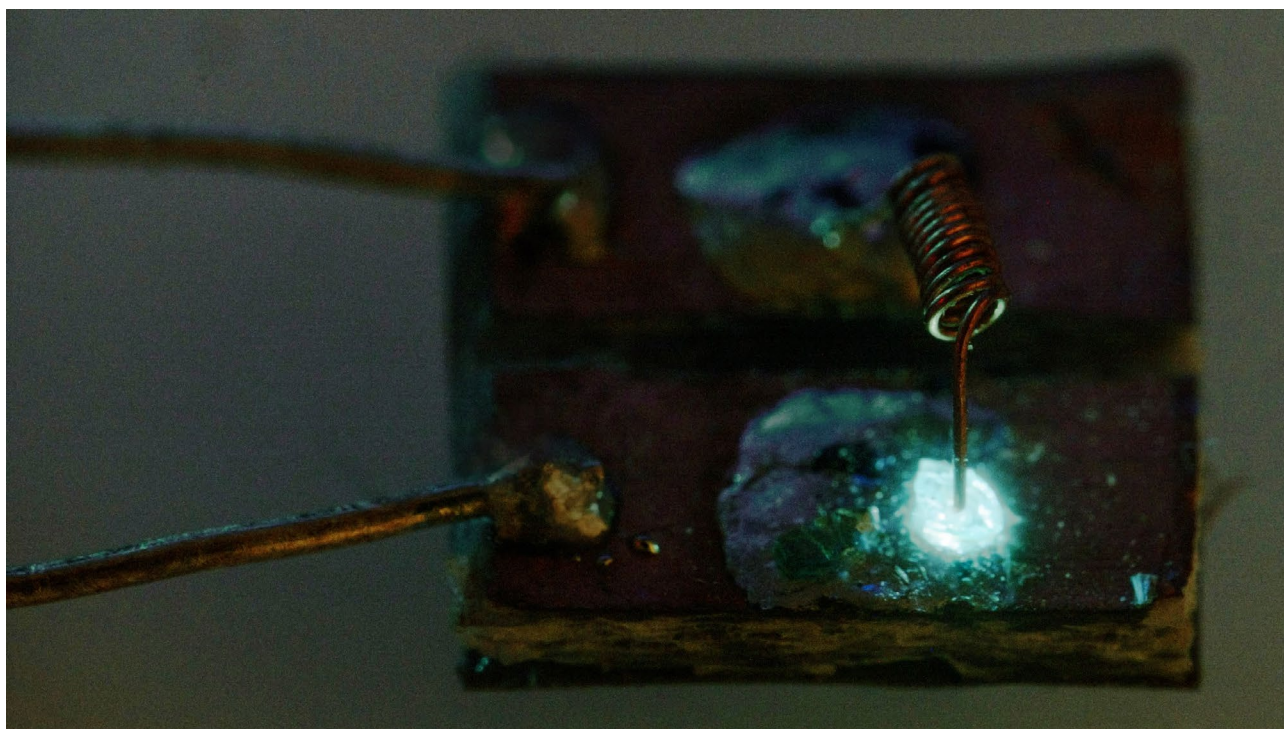


Zbuduj diodę LED!

- Przyrząd do parowania tranzystorów bipolarnych • Moduły klawiatur dotykowych
- Reanimacja zasilacza lampowego IZS-5/71 • Fascynujące przemiany energii: nietypowa fotowoltaika
- Rezystor – element przeklęty, ale nadal niezbędny • Wspólnie projektujemy: Montaż prototypów
- GNSS RTK, testy terenowe nawigacji • Xplained Mini – moduł z ATMEGA328 od Atmel



Inicjatywa Zrozumieć Elektronikę realizowana jest dzięki wsparciu Patronów i Mecenasów poprzez Patronite.pl



Zbuduj diodę LED!

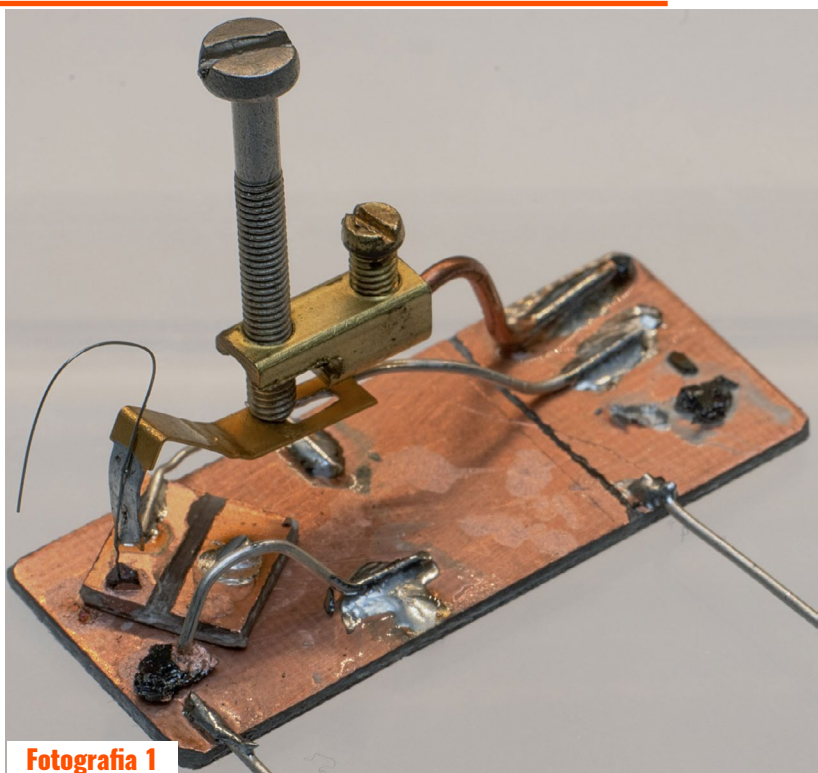
Artykuł jest wstępem do obszernego cyklu przedstawiającego diody LED. Dziś z kilku powodów nie buduje się we własnym zakresie diod LED, ale samodzielne zbudowanie diody LED, nawet mocno niedoskonałej, daje ogromną satysfakcję oraz doskonale wprowadza w temat tych, ogromnie ważnych dziś, elementów.

Odrobina historii i podstawy
Oczekiwania i realne wyniki
Materiały na diodę LED

Realizacja ostrzowych diod LED
Zbuduj i podziel się doświadczeniem!

Powyższa fotografia tytułowa pokazuje najprawdziwszą półprzewodnikową diodę LED podczas pracy. Zbudowana jest na centymetrowym kawałeczku płytki drukowanej (miedziowanego laminatu), a półprzewodnikowym materiałem czynnym jest kryształ węgla krzemu o wzorze chemicznym SiC. Jest to ostrzowa dioda LED.

Fotografia 1 przedstawia „stanowisko badawcze” do eksperymentów z podobnymi diodami LED, budowanymi z wykorzystaniem kryształów odpowiedniego półprzewodnika. Tutaj duże znaczenie ma możliwość zmiany nacisku drucika na powierzchnię kryształu, co realizowane jest za pomocą śrubki M3 i kawałka łączki elektrycznej. Oczywiście tego rodzaju „stanowisko badawcze” można zrealizować inaczej, w każdym razie budowa własnej diody LED wcale nie jest trudna!

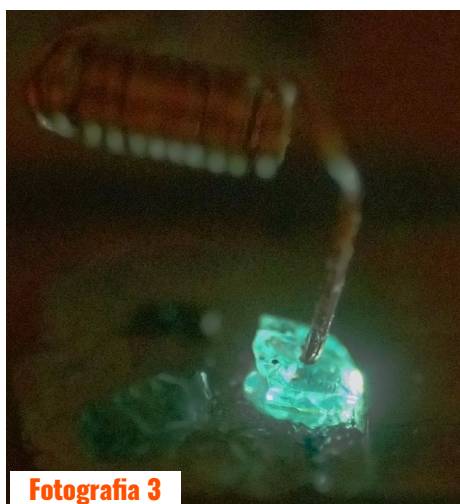
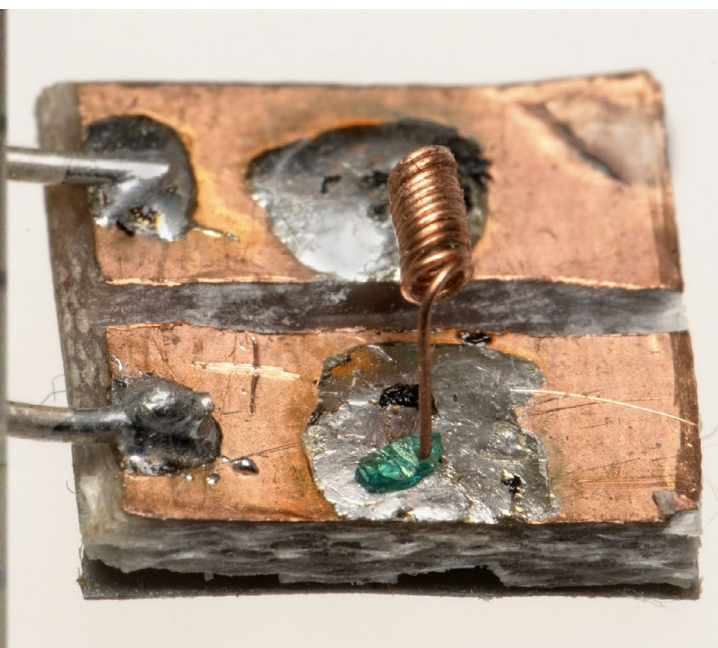


Fotografia 1

Na **fotografii 2** widać szczegóły budowy diody z fotografii tytułowej. **Fotografia 3** przedstawia zbliżenie świecącego kryształu. **Fotografia 4** pokazuje wtopione w cynę (stop lutowniczy) kryształki SiC w stanowisku z fotografii 1. Natomiast na **fotografii 5** widać jak jeden z tych kryształów świeci.



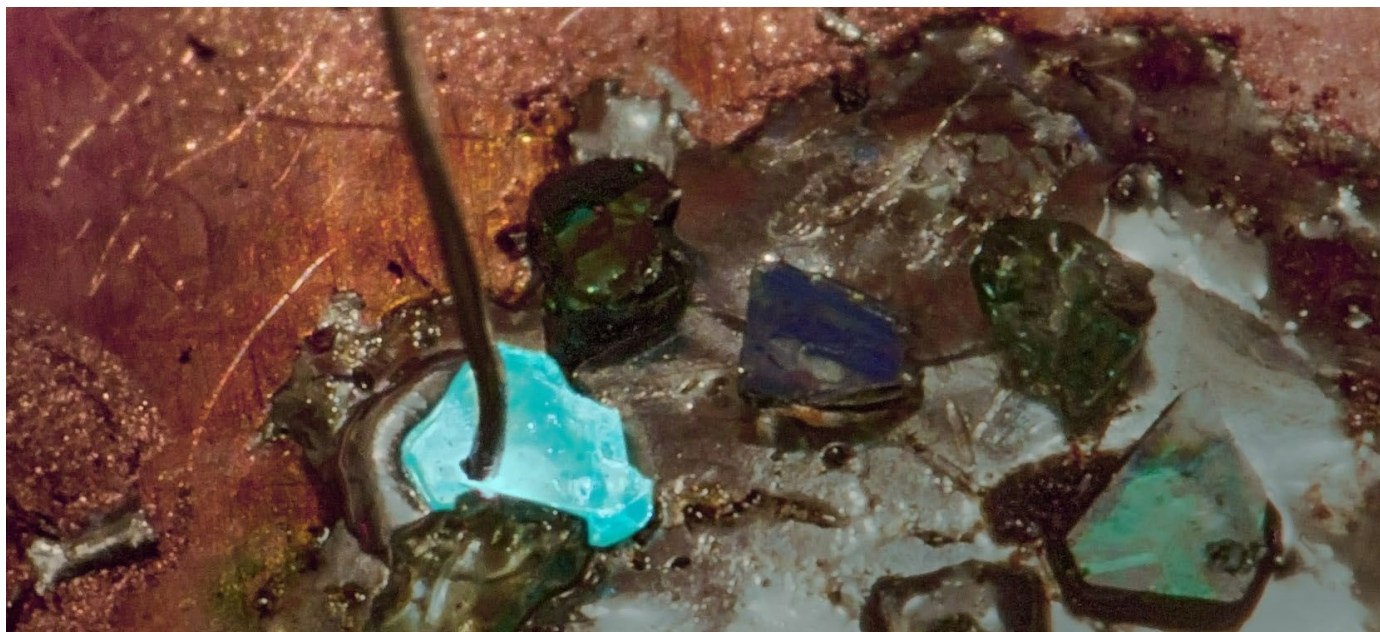
Fotografia 2



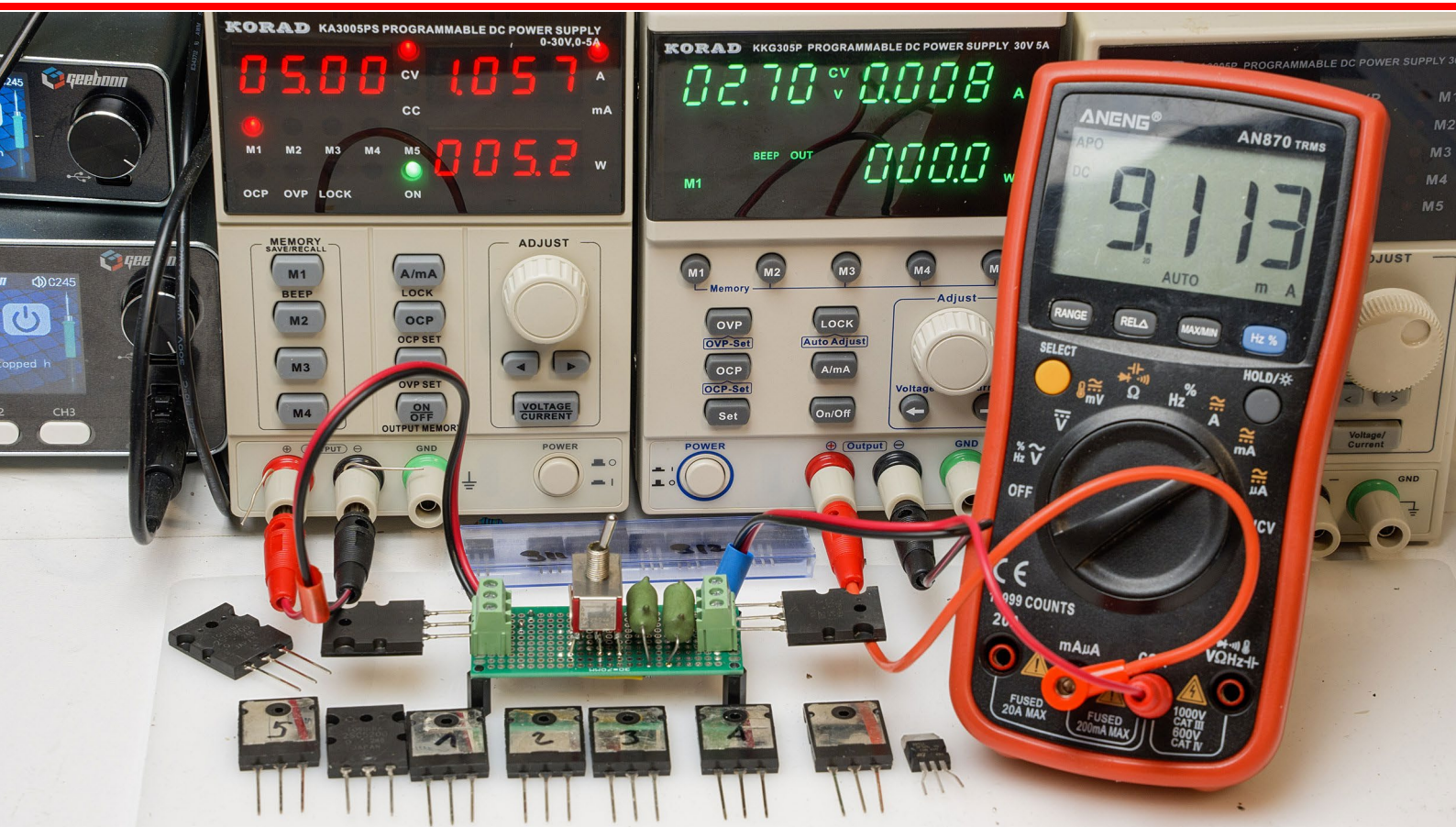
Fotografia 3



Fotografia 4



**Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.
W pełnej wersji ten artykuł oczywiście ma więcej stron.**



Przyrząd do parowania tranzystorów bipolarnych

W artykule opisany jest bardzo prosty, użyteczny, sprawdzony w praktyce przyrząd, pozwalający w szybki i łatwy sposób sprawdzać wzmacnienie tranzystorów bipolarnych oraz dobierać je w pary, także pary tranzystorów komplementarnych o jednakowym wzmacnieniu, potrzebne np. do wzmacniaczy mocy audio.

Rozważania projektowe

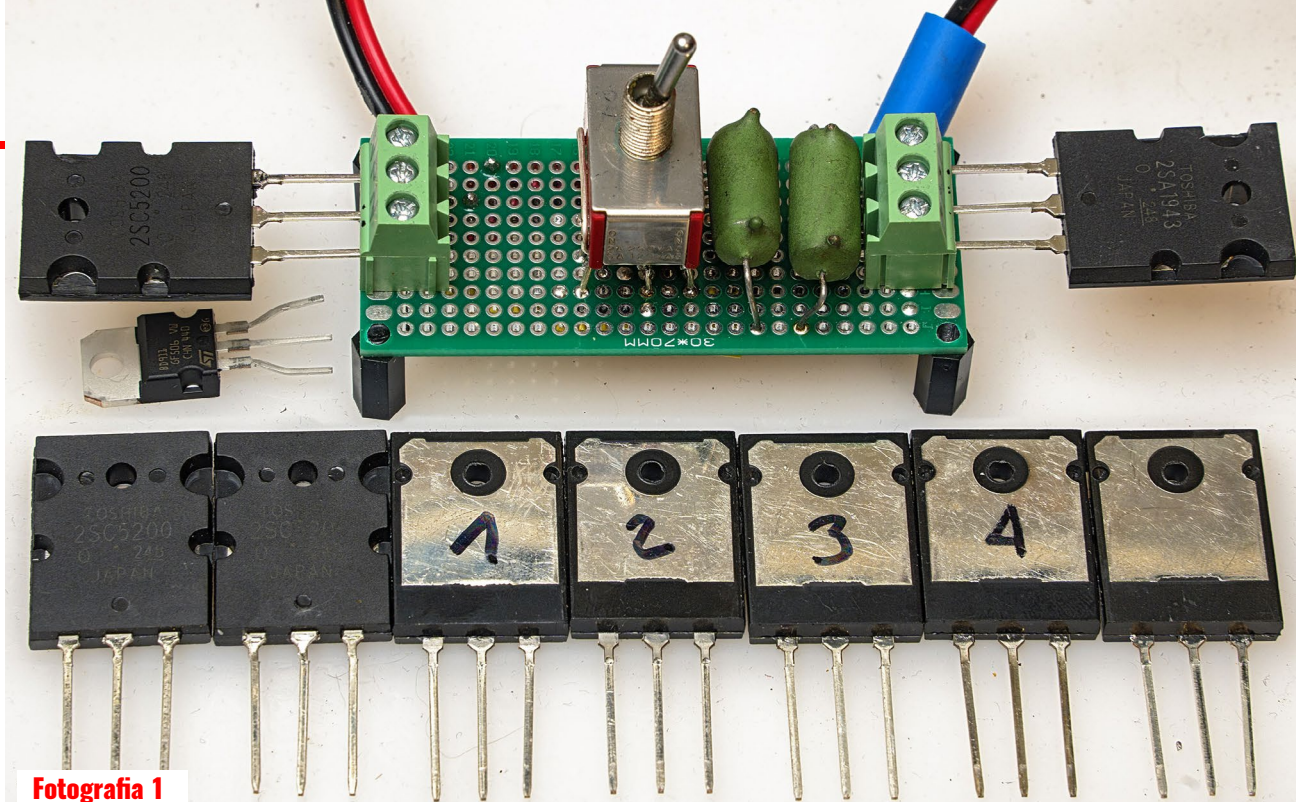
Szczegóły konstrukcyjne i możliwości zmian

Pomiary

Pomiar MOSFET-ów?

Zacząłem się od tego, że Leszek, jeden z moich przyjaciół, wspomniał, że przydałby się charakterograf. Podchwyciłem temat, bo od dawna planuję budowę sterowanego zdalnie zasilacza czterokwadrantowego (tak!), który mógłby być podstawą budowy dobrego charakterografu. Zaczęliśmy próby budowy kilku wersji zasilacza czterokwadrantowego, co opiszę w oddzielnym artykule. Okazało się jednak, że Leszek wspominał o charakterografie w kontekście sprawdzania tranzystorów mocy, a konkretnie chodziło o dobieranie par bipolarnych tranzystorów mocy do wzmacniaczy mocy audio, budowanych z elementów dyskretnych. Porozmawialiśmy o realnych potrzebach i możliwościach.

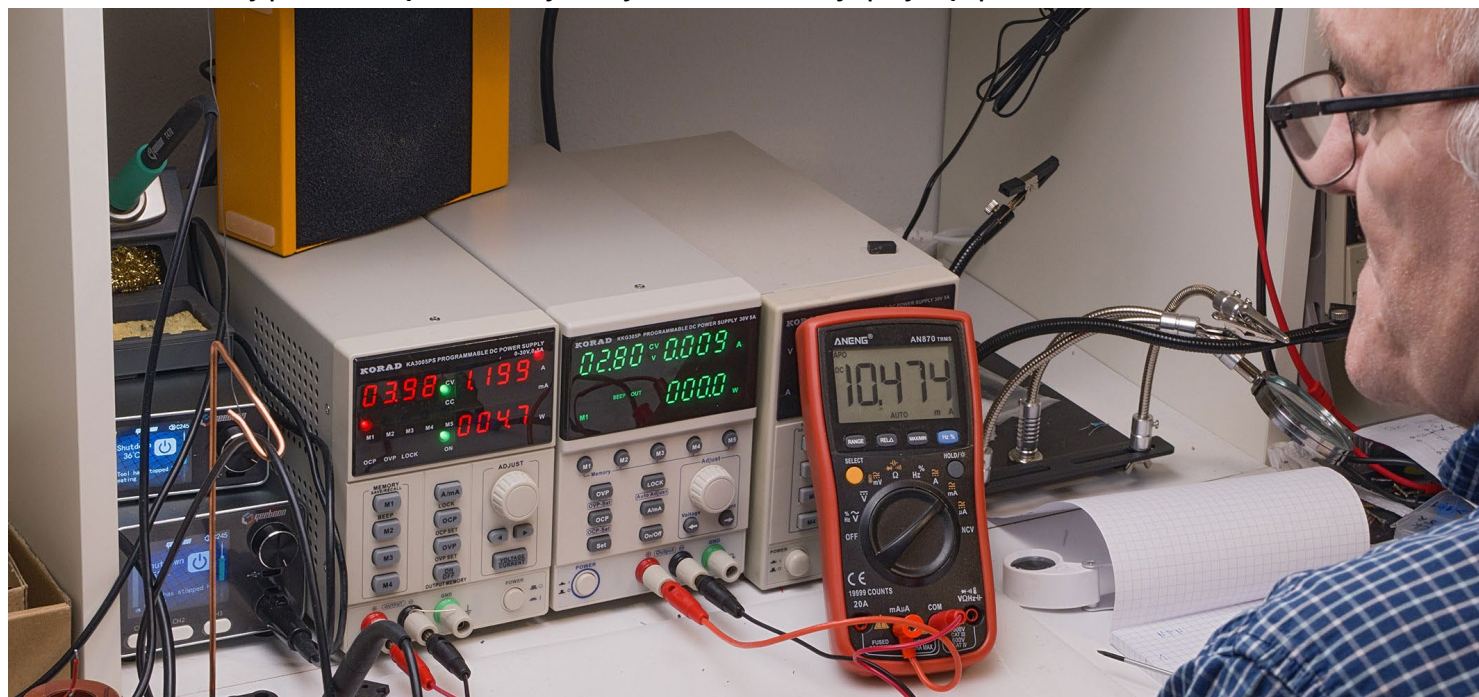
Okazało się, że wystarczy coś nieporównanie prostszego od charakterografu, do czego nie trzeba dedykowanej płytki drukowanej. Podstawą są karty katalogowe, gdzie podaje się wartość wzmacnienia prądowego, zmierzoną w konkretnych warunkach pracy, przy danym napięciu kolektora U_{CE} (kilka woltów) i przy podanym prądzie kolektora I_C , zwykle od jednego do kilku amperów. Podstawowa idea była taka, żeby nie wykorzystywać żadnych dodatkowych elementów elektronicznych, a jedynie dwa regulowane zasilacze laboratoryjne. Jeden zasilacz dostarczałby prąd bazy, ustawiony za pomocą ogranicznika prądu. Drugi dostarczałby prąd do obwodu kolektora przy takim napięciu, jak podaje karta katalogowa.



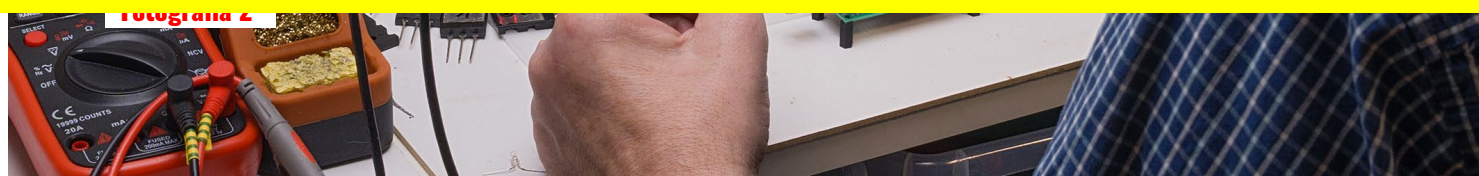
Fotografia 1

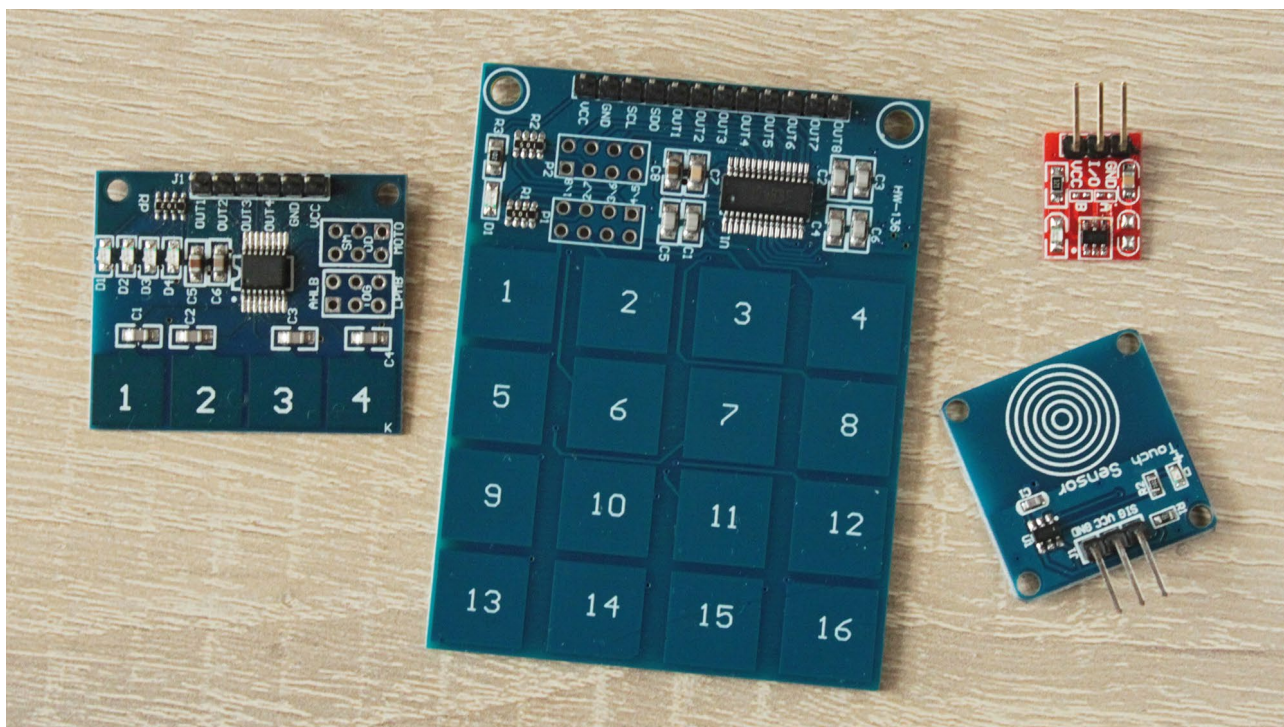
Jeszcze tego samego dnia narysowałem schemat śmiesznie prostego urządzenia do testowania komplementarnych tranzystorów bipolarnych z wykorzystaniem dwóch regulowanych zasilaczy. W pewien wtorkowy wieczór wzięliśmy się we dwóch do roboty. Leszek montował dla siebie właśnie ten prosty przyrząd do sprawdzania i parowania tranzystorów, a ja zająłem się „szybkim” generatorem impulsów o dużej amplitudzie do 500 V. Impulsów, których zbocza byłyby bardzo krótkie, rzędu nanosekund, potrzebnych do testowania wysokoomowych dzielników i sond. Efekty przedstawię w oddzielnym artykule.

W ciągu wieczora powstały dwa proste przyrządy, w tym tester tranzystorów, którego model przedstawiony jest na **fotografii 1** oraz na **fotografii tytułowej**. Fotografii te pokazują przykład wykorzystania do sprawdzania potężnych komplementarnych tranzystorów 2SA1943 oraz 2SC2500. Dwupozycyjny (czteroobwodowy) przełącznik pozwala wybrać albo tranzystor NPN, albo PNP, które dołączone są do dwóch zacisków ARK3. W praktyce, oprócz dwóch zasilaczy, wykorzystany został też miliamperomierz do dokładnego pomiaru prądu bazy. **Fotografia 2** prezentuje przyrząd podczas testów.



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE. W pełnej wersji ten artykuł oczywiście ma więcej stron.





Moduły klawiatur dotykowych

Elastyczność zastosowań szeroko rozumianych ekranów zintegrowanych z panelami dotykowymi wzbudziła zainteresowanie rozwiązaniami wykorzystania „przycisków” dotykowych w mniej zaawansowanych konstrukcjach.

Klawiatura z jednym przyciskiem
Klawiatura o czterech przyciskach
Moduł o 16 przyciskach

Testy modułów
Zastosowanie modułów

Przyciski mechaniczne, choć genialne w swojej prostocie, mają kilka wad. Jedną z nich jest konieczność eliminacji niekorzystnego zjawiska określanego jako dzwonienie styków. Można z nim walczyć na dwa sposoby: sprzętowo lub w ramach oprogramowania mikrokontrolera. Patrząc na finalną konstrukcję nie bez znaczenia jest konieczność odpowiednich rozwiązań mechanicznych, jak otwory w płycie czołowej by udostępnić przyciski użytkownikowi. Każdy przycisk musi mieć jakiś kapturek, a z tym wiąże się kolejny szczegół: trwały opis przycisku.

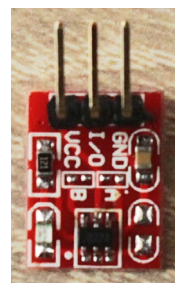
A gdyby tak zastosować sensory dotykowe? Wiele problemów technologicznych może rozwiązać się automatycznie. Obecnie dostępne na rynku sen-

sory mają wejścia o charakterze pojemnościowym, co oznacza, że sensor nie musi mieć kontaktu elektrycznego z palcem użytkownika, toteż można go umieścić pod nalepką z wydrukowanym symbolem przycisku.

Aby sprawnie posługiwać się tymi modułami warto zapoznać się z ich budową i funkcjonalnością.

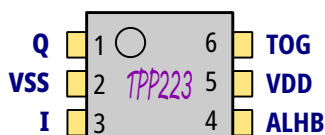
Klawiatura z jednym przyciskiem

Na **fotografii 1** widać jednoprzyciskowy moduł, w którym wykorzystany jest układ scalony TPP223.



Fotografia 1

Według dostępnej dokumentacji do użytego układu scalonego, występuje on

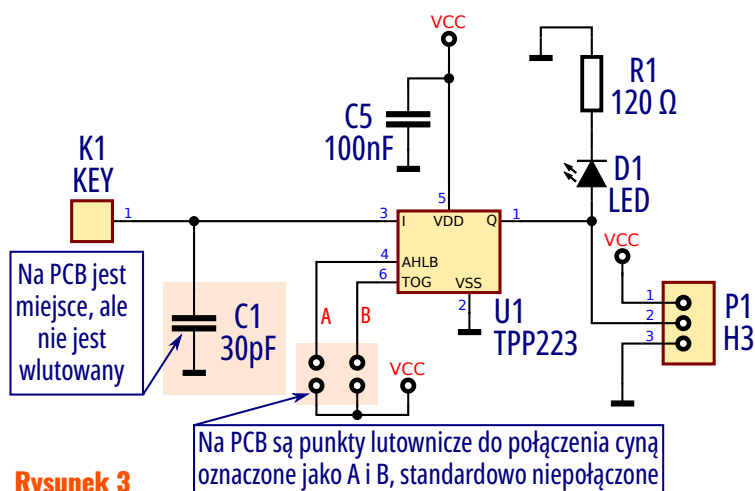


Rysunek 2

w kilku odmianach obudowy. Wariant w obudowie 16-pinowej jest bardziej funkcjonalny, jednak w prezentowanych modułach występuje w obudowie SOT23-6 (obudowa 6-pinowa, **rysunek 2**). Znaczenie poszczególnych pinów układu jest następujące:

- Q – wyjście typu CMOS, sygnalizujące stan wykrycia dotknięcia sensora,
- I – wejście do przyłączenia sensora dotyku,
- TOG – określa funkcjonalność sensora, możliwe opcje to: tryb przełączania (dla TOG=1) oraz tryb bezpośredni (dla TOG=0); wejście ma wbudowany do masy rezystor, toteż przy braku jakiegokolwiekysterowania tego pinu, układ rozpoznaje to jako zero logiczne,
- AHLB – wejście określające stan aktywny na wyjściu Q: sygnalizacja dotknięcia sensora stanem wysokim (dla AHLB=0) lub stanem niskim (AHLB=1); wejście ma wbudowany rezystor do masy, toteż przy braku jakiegokolwiekysterowania tego pinu, układ rozpoznaje to jako zero logiczne,
- VSS – pin zasilania (masa),
- VDD – pin zasilania (2 V...5,5 V).

Schemat modułu zbudowanego na bazie tego układu scalonego pokazuje **rysunek 3**. Istnieją tu możliwości konfiguracyjne (poprzez zalanie cyną pól lutowniczych, które „fabrycznie” nie są połączone), czyli zmiana stanu wejść TOG oraz AHLB. Biorąc pod uwagę ustawienia domyślne, moduł sygnalizuje naciśnięcie przycisku włączeniem diody LED i całość pracuje w trybie bezpośrednim (przykładowo dioda świeci, dopóki trzymany jest przycisk). Istnieje możliwość ustalenia pracy do trybu przełącznika, jedno naciśnięcie sensora włącza, a kolejne wyłącza (pracuje jak dwójka licząca). Możliwe do uzyskania tryby pokazuje tabelka na **rysunku 4**. Ja dolutowałem sobie do modułu delikatne przewody do listy pi-



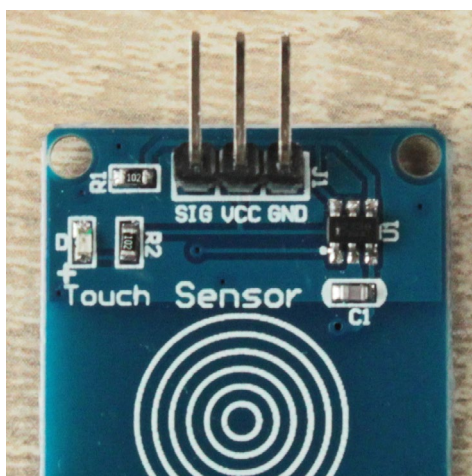
Rysunek 3

TOG	AHLB	Znaczenie
0	0	Tryb bezpośredni, aktywacja sensora sygnalizowana stanem wysokim
0	1	Tryb bezpośredni, aktywacja sensora sygnalizowana stanem niskim
1	0	Tryb przełącznika, aktywacja sensora sygnalizowana stanem wysokim
1	1	Tryb przełącznika, aktywacja sensora sygnalizowana stanem niskim

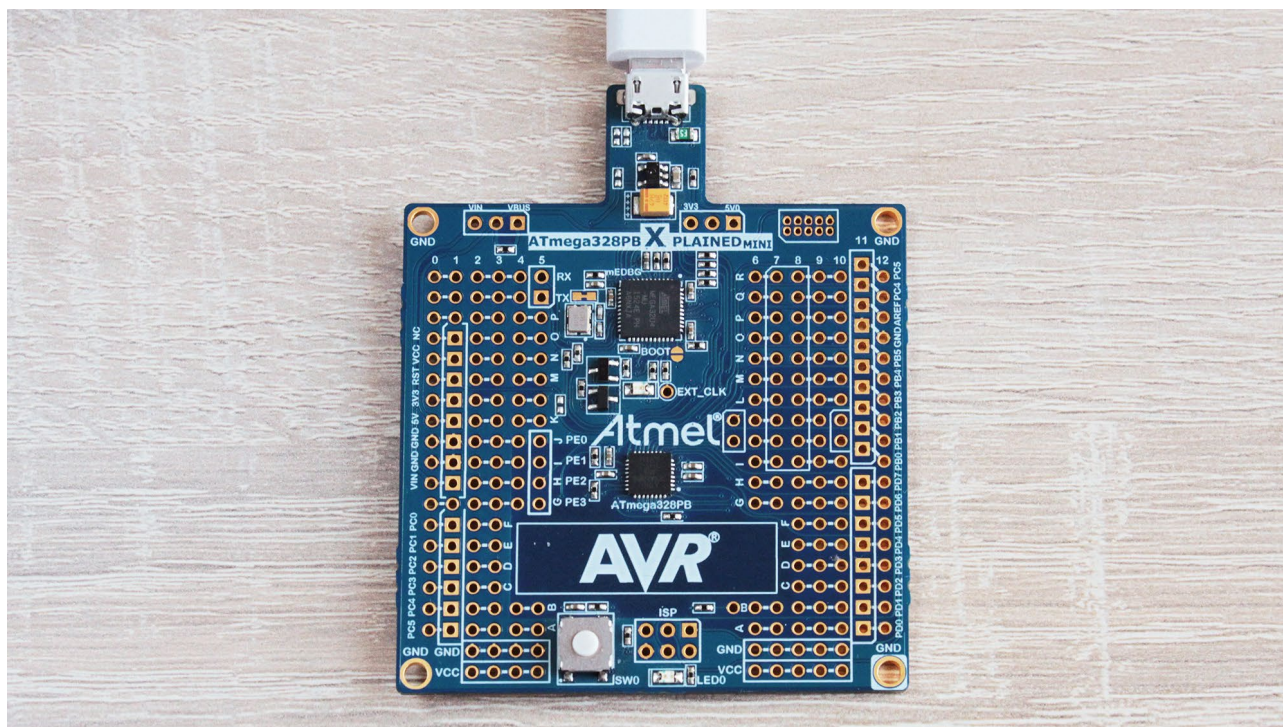
Rysunek 4

zaobserwowałem istotną cechę układu TPP223, dla której później znalazłem potwierdzenie w dokumentacji. Stan ustawień na wejściach TOG oraz AHLB jest rejestrowany przez układ w chwili włączenia zasilania, toteż jakakolwiek zmiana w trakcie działania wymaga wyłączenia na chwilę zasilania, by układ „zatribił” nowe ustawienia w chwili pojawienia się zasilania.

Istnieje inny wariant modułu z tym samym układem scalonym (**fotografia 5**), w którym istotnymi elementami są inne rozłożenia sygnałów na trzypinowym złączu (jeżeli zrobimy kabelek do modułu z fotografii 1, to nie będzie pasował do nowego) oraz brak możliwości zmian konfiguracyjnych (piny TOG oraz AHLB są na stałe przy-



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE. W pełnej wersji ten artykuł oczywiście ma więcej stron.



Xplained Mini – moduł z ATMEGA328 od Atmel

Współczesna technologia w dużej mierze nie jest przyjazna hobbystom ze względu na obudowy układów scalonych. Jednak w „drugiej ręce” oferuje gotowe rozwiązania w postaci modułów. Takim przykładem jest prezentowany moduł Xplained Mini.

Różnice w stosunku do Arduino Uno Uruchomienie modułu

Firma Atmel zaoferowała użytkownikom moduł z mikrokontrolerem ATMEGA328PB (takim samym, jaki znajduje się w Arduino Nano), określony jako *Xplained Mini*. Ma on na swoim pokładzie oprócz mikrokontrolera istotny element jakim jest programator. Aby załadować kod programu do pamięci Flash mikrokontrolera nie jest konieczny dodatkowy programator (choć również istnieje możliwość zaprogramowania go w standardowy sposób, poprzez 6-pinowe złącze KANDA). Przydatnym może się okazać również istnienie jednego przycisku – do dowolnego wykorzystania przez użytkownika. Częstotliwość sygnału taktującego dla mikrokontrolera również wynosi 16 MHz. W dużej mierze można

Przykładowy program Programowanie mikrokontrolera

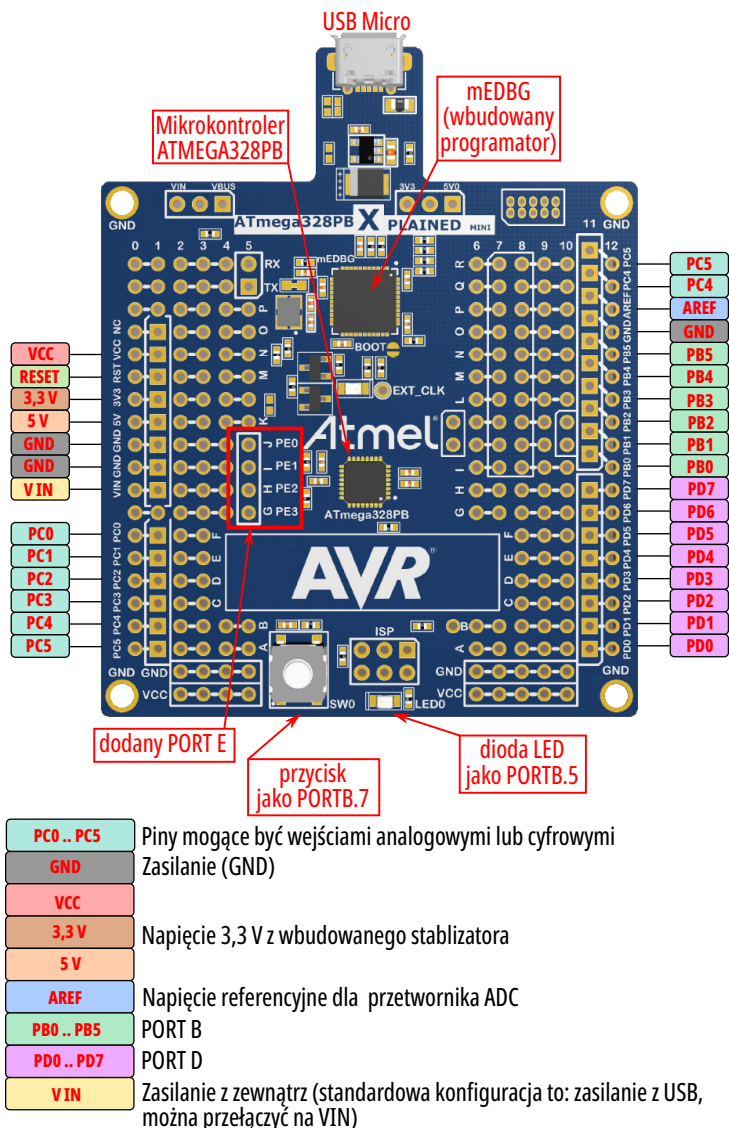
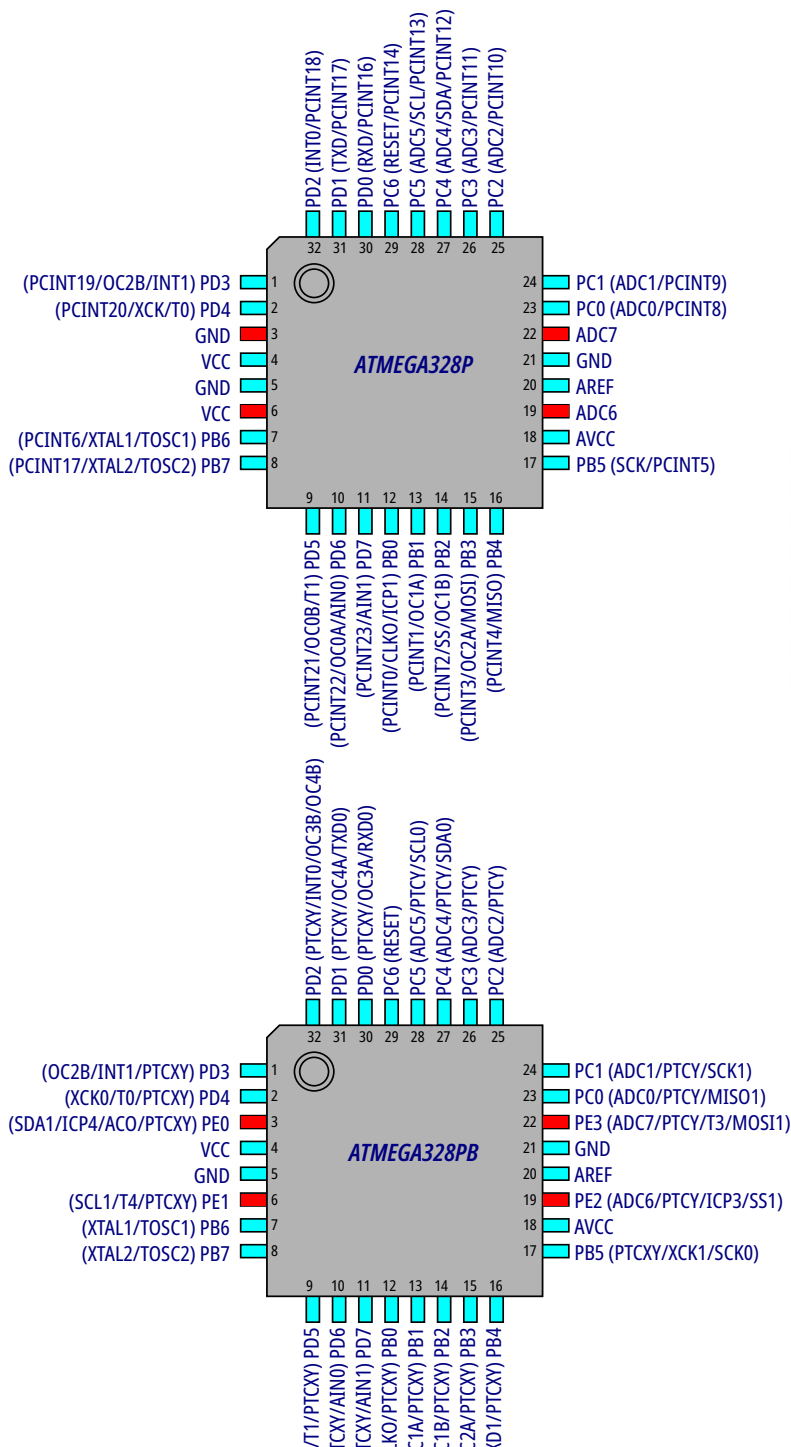
uznać ten moduł za znacząco rozszerzające rozwiązanie w stosunku do Arduino Uno, gdyż jest tu zastosowany nowszy mikrokontroler o większych zasobach sprzętowych. Porównanie do Arduino Nano jest niemiarodajne, gdyż Nano nie ma na swoim pokładzie programatora.

Różnice w stosunku do Arduino Uno

Pierwszą istotną różnicą w *Xplained Mini* w stosunku do Arduino Uno jest zastosowanie mikrokontrolera ATMEGA238PB (Xplained) w stosunku do ATMEGA328P (Uno). Oprócz większej liczby wbudowanych układów licznikowo-zegarowych, większej liczby kontrolerów transmisji szeregowej,

E – ELEMENTY I MODUŁY

w ATMEGA328PB zawarty jest kontroler dotykowy (z pojemnościowymi czujnikami, można bez dodatkowych układów utworzyć kilka pól reagujących na dotyk, Xplained jest w to wyposażony). Moduł Arduino Uno wykorzystuje mikrokontroler w obudowie DIP28, więc ma zredukowaną liczbę wejść dla przetwornika ADC, Arduino Nano posiłkuje się układem w obudowie QFP32. Większa liczba wypro-



Rysunek 2

wadzeń daje dodatkowe wejścia do przetwornika ADC, jednak rozwiązanie jest nietypowe jak na standardy występujące w AVR: te wejścia mogą jedynie pełnić funkcje wejść analogowych.

Mikrokontroler z *Xplained Mini* występuje wyłącznie w obudowie przeznaczonej do montażu powierzchniowego o 32 pinach. Jednak w stosunku do starszego wariantu (ATMEGA328P), tutaj nie zachodzi pełna kompatybilność wyprowadzeń. W miejsce nietypowych wejść analogowych oraz kosztem wyprowadzeń do zasilania dodany jest czterobitowy PORT E, jako wielofunkcyjne rozwiązanie (wejścia/wyjścia cyfrowe oraz wejścia analogowe do przetwornika ADC). Istotne różnice pokazuje **rysunek 1**.

Uruchomienie modułu

Moduł przeznaczony jest do używania w środo-

Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.
W pełnej wersji ten artykuł oczywiście ma więcej stron.



Reanimacja zasilacza lampowego IZS-5/71

W artykule opisane są perypetie związane z naprawą kupionego okazjnie, starego, ale bardzo solidnego polskiego zasilacza, dającego napięcie w zakresie 0,1 V...500 V i prąd do 1 ampera, wyprodukowanego mniej więcej 50 lat temu. Zasilacza, którego główny regulator tworzy sześć lamp elektronowych.

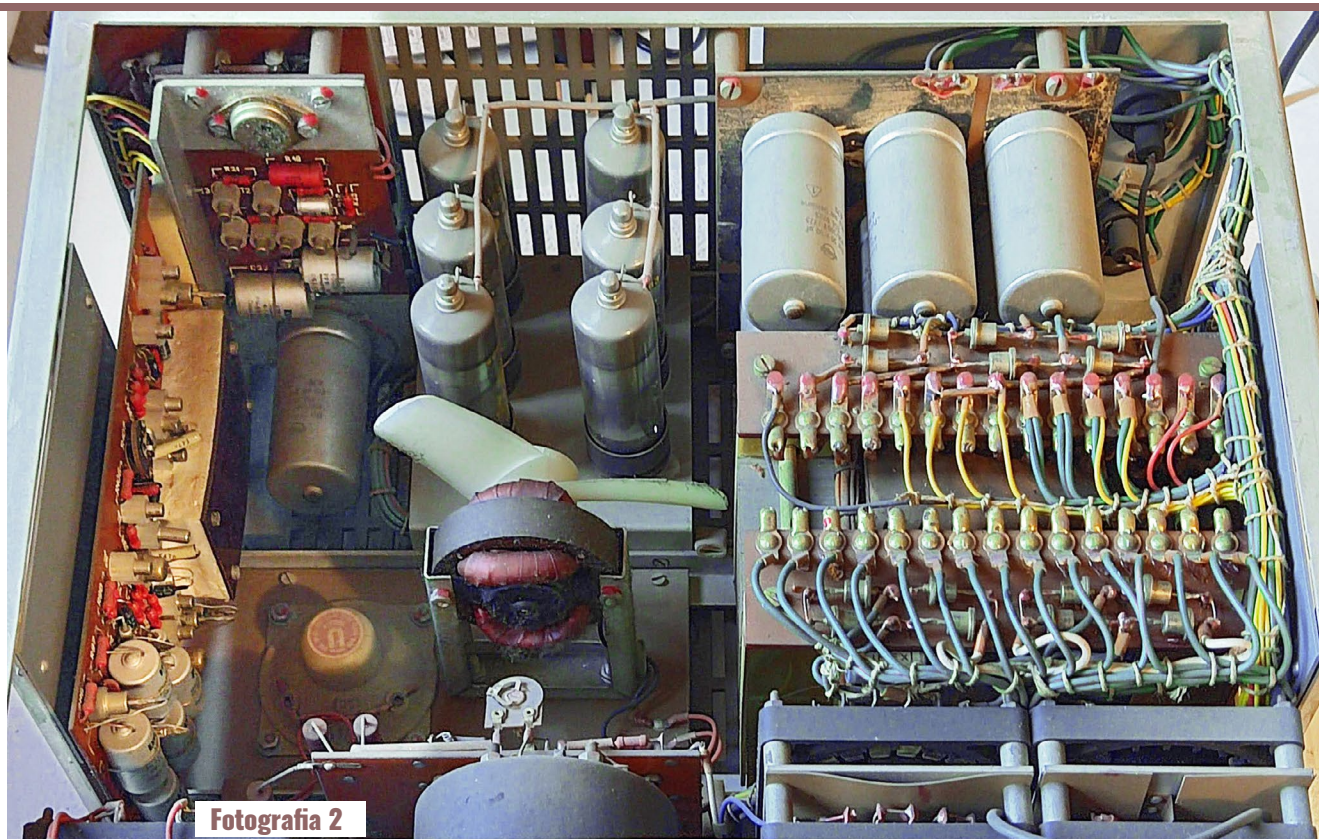
Jak pisałem wcześniej, za pośrednictwem jednego z portali aukcyjnych drogą kupna nabyłem stary zasilacz IZS-5/71, pokazany na fotografii tytułowej. Jest bardzo atrakcyjny m.in. dlatego, że napięcie wyjściowe można skokowo regulować co 0,1 V w zakresie 0,1...500,0 V. Zasilacz ma wbudowany prymitywny ogranicznik prądu (10 mA, 100 mA, 1 A). Konstrukcja z początku lat 70. zawiera dostępne wtedy w kraju elementy. Wtedy nie było tranzystorów, które mogłyby wytrzymać napięcie ponad 500 woltów i prąd 1 ampera. Dlatego głównym regulatorem jest sześć połączonych równolegle lamp – pentod EL36.



Fotografia 1

Nabyłem go z nadzieją, ale i z niepewnością. Wstępne oględziny wskazały, że prawdopodobnie nie był otwierany od nowości, bo na boku była oryginalna plomba – **fotografia 1**. To była i dobra, i zła wiadomość. Dobra, że nikt w nim „nie grzebał”, a zła, że od nowości pracują tam te same lampy, które być może są mocno zużyte. A zakup sześciu lamp PL36 byłby wydatkiem dużo większym niż nabycie zasilacza.

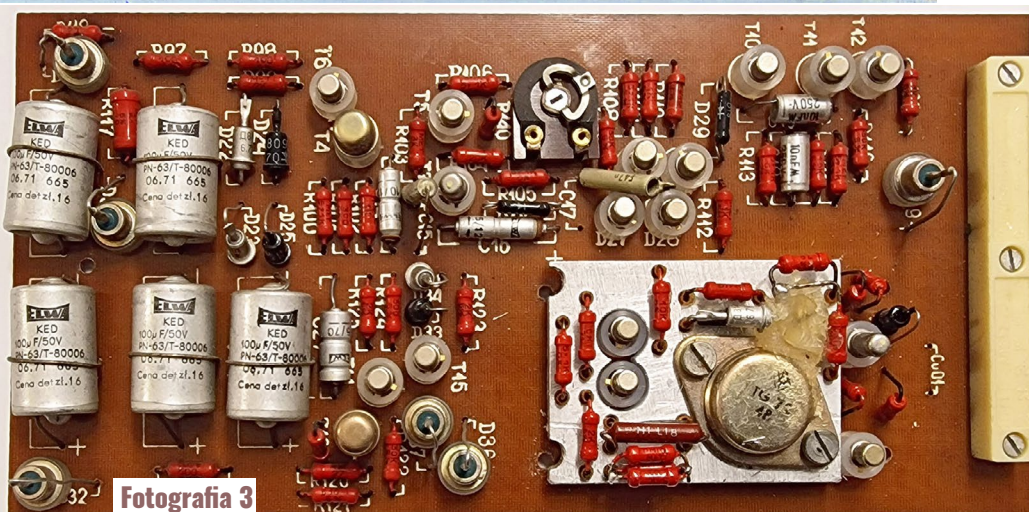
Po pierwszym włączeniu zasilacz zadziałał, a przełączniki regulowały napięcie wyjściowe. To był dobry znak na początek.



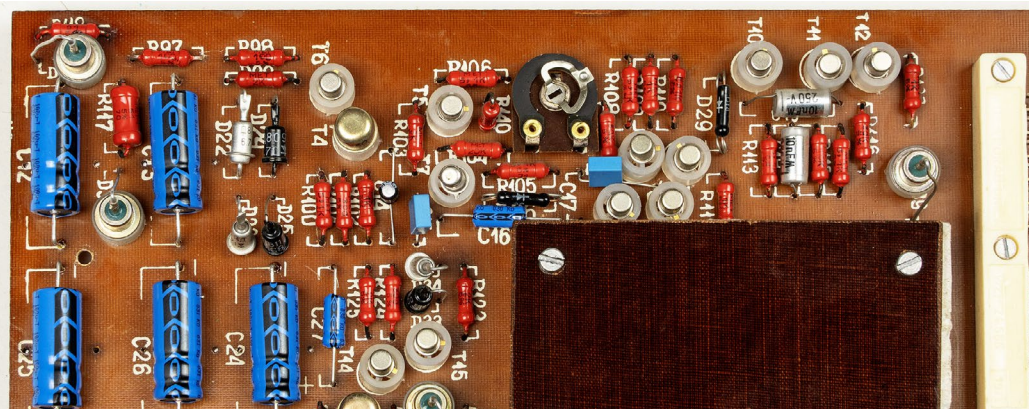
Fotografia 2

Do remontu i testów zgłosili się Leszek i Arek. Wspólnie podziwialiśmy przedstawiony na **fotografii 2** widok, jaki ukazał się po otwarciu obudowy. Imponuje potężny transformator, a niepokój budzą lampy oraz wszystkie „elektrolity”, które mają około pół wieku.

Generalną zasadą jest, że w tak starych układach przede wszystkim trzeba sprawdzić właśnie wszystkie kondensatory elektrolityczne. Wiadomo, że stare „elektrolity” potrafią wysychać, co może zmniejszyć ich pojemność niemal do zera. Pozytywnym zaskoczeniem było to, że wszystkie największe „elektrolity” produkcji niemieckiej (DDR) mają pojemność bardzo bliską nominalnej. Tych mniejszych, widocznych na płycie drukowanej z lewej strony oraz na **fotografii 3**, nie sprawdzaliśmy, bo Leszek zaoferował się,



Fotografia 3



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE. W pełnej wersji ten artykuł oczywiście ma więcej stron.



Fascynujące przemiany energii: nietypowa fotowoltaika (1)

Poniższy artykuł także jest częścią nietypowej instrukcji obsługi do zestawu *Fascynujące przemiany energii*, dostępnego w sklepie Botland, służącego do przeprowadzenia mnóstwa interesujących eksperymentów. Zaczynamy badać przemiany energii świetlnej na elektryczną (i na odwrót).

Fotoelementy krzemowe

Ogniwa fotoelektryczne z innych materiałów

Diody LED jako ogniwa fotowoltaiczne

Co się może nie udać?

Główne cele ćwiczeń

Ten artykuł jest rozszerzeniem i uzupełnieniem mojego filmu B103, dotyczącego przemian energii świetlnej. Ściśle biorąc, energia świetlna to energia fal elektromagnetycznych, czyli okresowych zmian, drgań specyficznie połączonych pól elektrycznego i magnetycznego (czykolwiek te pola są).

W zasadzie dla elektroników najbardziej istotne jest promieniowanie elektromagnetyczne (EM) radiowe i mikrofalowe, czyli fale elektromagnetyczne o częstotliwościach rzędu megaherców, najwyżej gigaherców. Jednak na razie będziemy się zajmować promieniowaniem świetlnym – głównie światłem widzialnym, a to są fale elektromagnetyczne (EM) o częstotliwościach rzędu setek teraherców, czyli setek tysięcy gigaherców. Pomijamy różne groźne

rodzaje promieniowania jak promieniowanie rentgenowskie i promieniowanie gamma, które też są falami EM, tylko o jeszcze większych częstotliwościach. W prezentowanych przeze mnie eksperymentach wykorzystuję światło widzialne, ale „leciutko zahaczam” o ultrafiolet oraz o podczerwień, ale na razie nie o mikrofały ani o fale radiowe. Badamy różne przemiany energii i najłatwiej je badać na przykładzie widzialnego promieniowania świetlnego.

Wracamy do prostych konkretnych: światło słoneczne niesie duże ilości energii, czego doświadczamy latem na plaży i nie tylko tam. Zdecydowana większość energii światła słonecznego zamienia się na ciepło, ale potrafimy część tej energii zamienić na energię elektryczną za pomocą paneli fotowoltaicznych.

Jesteśmy przyzwyczajeni, że wokół nas są montowane instalacje zawierające charakterystyczne, duże panele fotowoltaiczne, takie jak widać na **fotografii 1**. Mają one specyficzne właściwości i dla większości z nas okazują się elementami tajemniczymi. W następnych filmach i artykułach zajmiemy się klasycznymi panelami fotowoltaicznymi, może nie aż tak dużymi, ale na razie poznajmy bardzo interesujące podstawy.

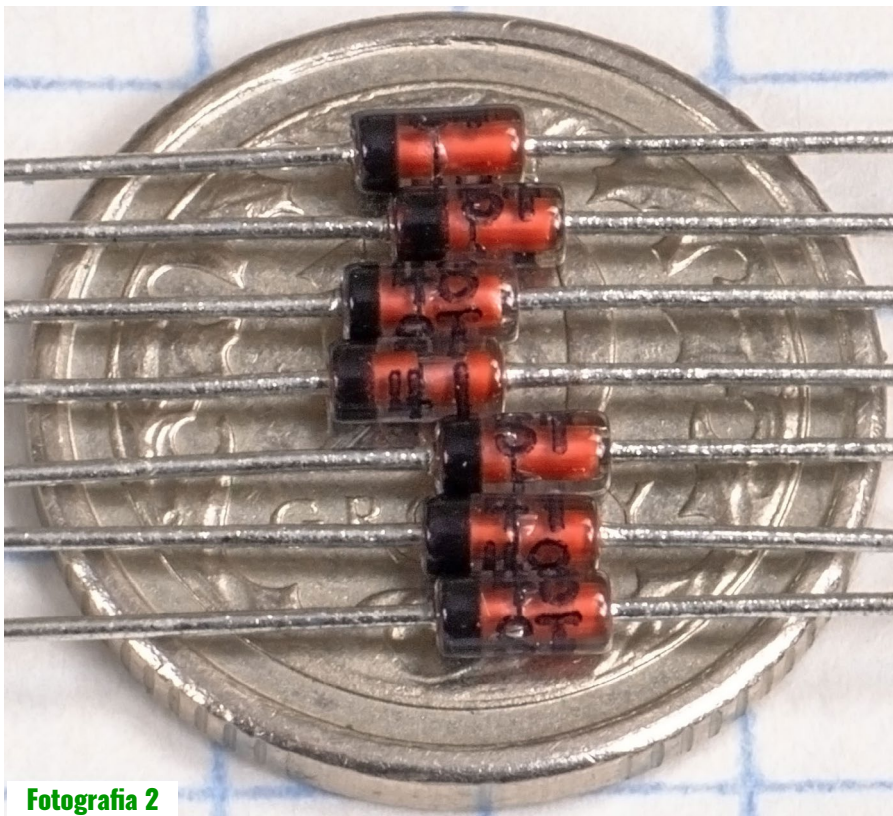
Niewątpliwie panel fotoelektryczny, fotowoltaiczny, w skrócie często oznaczany PV (photovoltaic) zamienia energię światła słonecznego na energię elektryczną w jakimś tempie (z jakąś mocą), określonym przez napięcie i prąd. Zwykle nie zastanawiamy się nad szczegółami jego pracy, bo to nie dziwi!

Natomiast dla wielu osób zaskoczeniem jest fakt, że funkcję niedoskonałych paneli, a raczej ogniw fotowoltaicznych, mogą pełnić diody prostownicze, a także tranzystory. W filmie pokazałem, że silnie oświetlona, najpopularniejsza na świecie mała krzemowa dioda prostownicza 1N4148 (**fotografia 2**) wytwarza napięcie rzędu kilkudziesięciu miliwoltów. Trudno byłoby zmierzyć znikomo mały prąd i moc, ale faktem jest, że jest to też panel fotowoltaiczny, co prawda śmiesznie słaby, ale jednak zamieniający energię światła na energię elektryczną.

Z taką przemianą energii lepiej radzą sobie inne elementy. Nieprzypadkowo poświęciłem sporo czasu i wysiłku, żeby „rozpruć” obudowy kilku tranzystorów. Na **fotografii 3** widać kilka krajowych starych tranzystorów w metalowych obudowach. Wnętrze jednego z nich pokazuje dokładniej **fotografia 4**. Tutaj światło dociera do krzemowej struktury dużo lepiej, niż w przypadku diody 1N4148, dlatego taki tranzystor, a właściwie jedno z jego złączy, czyli jedna dioda, okazuje się lepszym ogniwem – panelem PV.



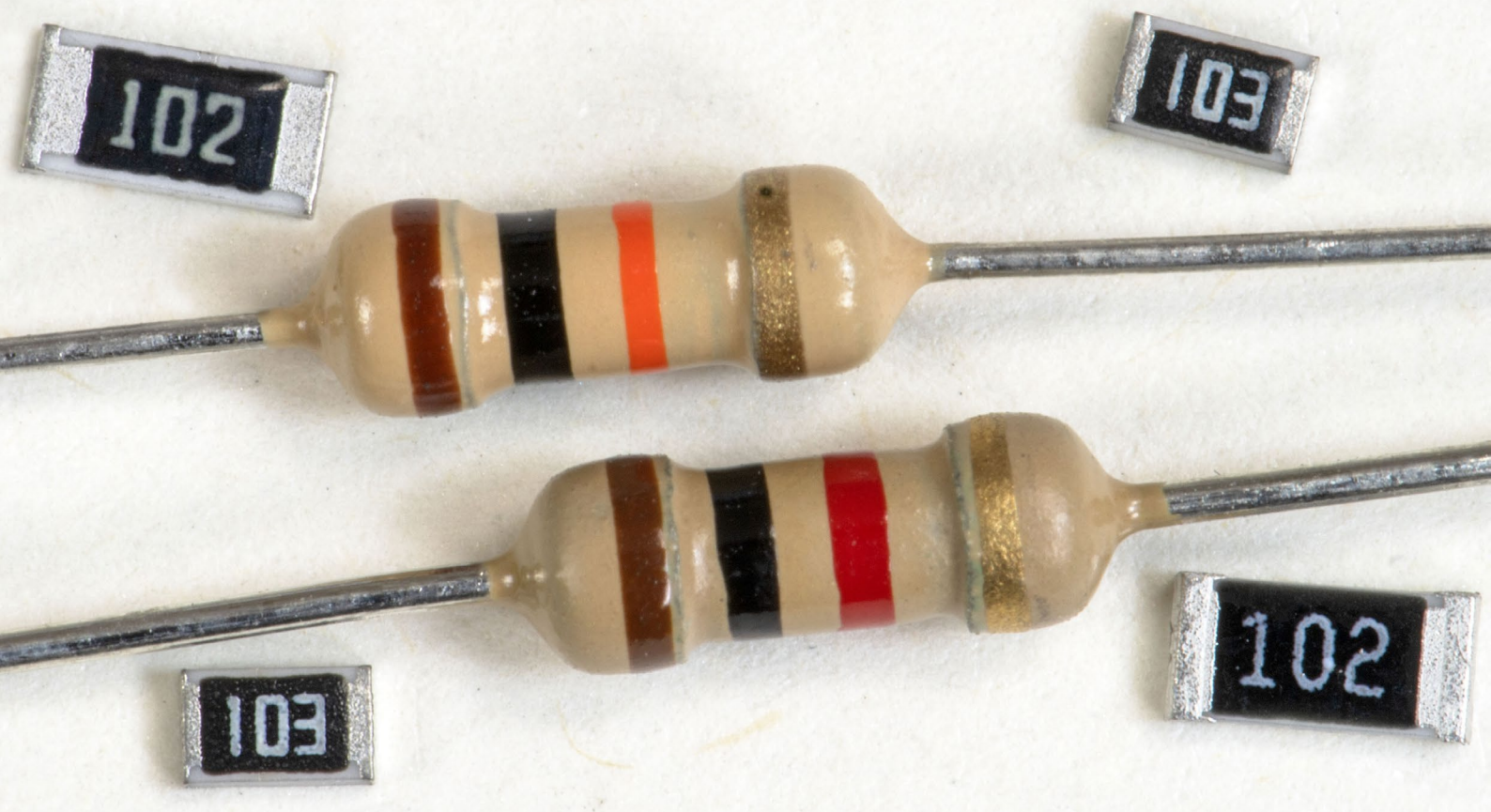
Fotografia 1



Fotografia 2



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE. W pełnej wersji ten artykuł oczywiście ma więcej stron.



Rezystor – element przeklęty, ale nadal niezbędny

W poprzednich artykułach tej serii przedstawiłem elementarne informacje o fundamentach elektroniki od strony matematyczno-fizycznej. W tym artykule zaczynamy omawianie elementów elektronicznych, ale będziemy też wracać do fundamentalnych zależności i praw. Nieprzypadkowo zaczynamy od rezystora.

[Dlaczego rezystor jest przeklętym elementem?](#)

[Czy rezystor to najpopularniejszy element?](#)

[Połączenie szeregowe i równoległe rezystorów](#)

[Rezystorowe dzielniki napięcia](#)

[Czy istnieją dzielniki prądu?](#)

[Schemat zastępczy dzielnika napięcia](#)

[Rezystancja widziana od strony...](#)

W tym artykule uzasadnię, dlaczego rezystor to przeklęty element elektroniczny, którego należy unikać, o ile to tylko możliwe. Często jest to do zrobienia, ale niekiedy okazuje się praktycznie niemożliwe lub po prostu nieekonomiczne, dlatego niestety nadal musimy wykorzystywać te przeklęte elementy.

Fotografia tytułowa pokazuje przykładowe rezystory: dawne i współczesne. Każdy rezystor ma dwa podstawowe parametry: rezystancję oraz obciążalność (potocznie nazywaną mocą). Parametry rezystorów szerzej omówione są w artykule [E103](#), który zawiera najważniejsze wiadomości o parametrach rezystorów. A teraz omówię coś bardzo ważnego.

Dlaczego rezystor jest przeklętym elementem?

Energia jest najważniejsza i uzasadnienie, dlaczego rezystor jest przeklętym elementem elektronicznym ma ścisły związek z energią. Dlatego zacznę od podstawowego, oczywistego wzoru: $P = U \times I$. Jeżeli w jakimś obwodzie jednocześnie występuje napięcie U i płynie prąd I , to niewątpliwie mamy do czynienia z przesyłaniem lub przemianą energii w większym lub mniejszym tempie, które nazywamy mocą, oznaczmy literką P i wyrażamy w watach [W].

Jeżeli na przykład z baterii pobierany jest prąd, to magazynowana w niej energia chemiczna, przekształcona do postaci związanej z napięciem U i prądem I ,

jest „wysyłana na zewnątrz” – do jakiegoś obciążenia, którym może być na przykład rezystor lub dioda LED.

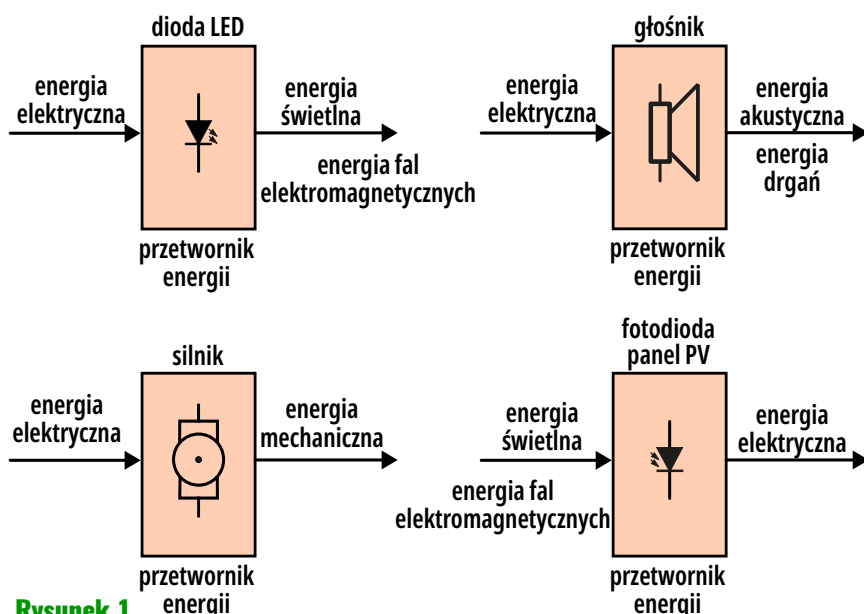
I jeżeli na tym obciążeniu, na diodzie LED lub rezystorze występuje jakieś napięcie U i płynie przez niego prąd I , to w tych elementach mamy do czynienia z przemianą formy energii. Element (obciążenie) zamienia doprowadzaną energię elektryczną (związaną z napięciem i prądem stałym) na inną formę energii.

I tu zbliżamy się do ogromnie ważnego punktu tego artykułu: **zadaniem diody LED jest zamiana energii elektrycznej „stałoprądowej” na energię świetlną czyli najprościej biorąc – na światło widzialne**, co w pewnym uproszczeniu pokazane jest na **rysunku 1**. Analogicznie **zadaniem głośnika czy brzęczyka jest zamiana energii elektrycznej na energię akustyczną**, najprościej mówiąc – na dźwięk. Tak samo **zadaniem silnika elektrycznego jest zamiana energii elektrycznej na energię (pracę) mechaniczną**.

To oczywiście, co jest celem tej zamiany. Na razie pomijamy fakt, że nie potrafimy takich przemian realizować w sposób idealny – tylko część energii elektrycznej pobieranej z baterii zostaje zamieniona na energię świetlną, akustyczną czy mechaniczną. Reszta (większość) nie tylko się marnuje, ale staje się niepożądanym odpadem, swego rodzaju śmieciami, a konkretnie – zamienia się na ciepło – na energię cieplną. Te zagadnienia omawiam stopniowo w oddzielnym cyklu „Fascynujące przemiany energii”. A teraz analizujemy działanie rezystora i pytamy o przemianę energii w rezystorze.

Najogólniej biorąc dioda LED, głośnik czy silnik zamieniają energię elektryczną na inny pożyteczny i potrzebny nam rodzaj energii. Niestety tylko częściowo, a niepożądanym efektem ubocznym jest zamiana energii elektrycznej na energię cieplną, którą można nazwać „najmniej szlachetną”. No cóż, niestety nie umiemy takich przemian energii przeprowadzać w sposób doskonały czy bliski doskonałości, a jedynie z ograniczoną sprawnością czy skutecznością. W uproszczeniu pokazuje to **rysunek 2**. Ale jednak otrzymujemy pożądaną formę energii. Inaczej jest z rezystorem.

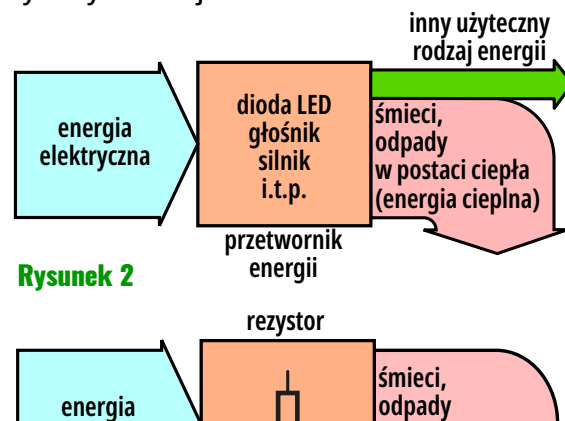
I teraz najważniejsze: **w przeciwieństwie do innych pożytecznych elementów, rezystor jest swego rodzaju pasożytem, który zamienia, a wręcz marnuje całą dostarczaną energię elektryczną na**



Rysunek 1

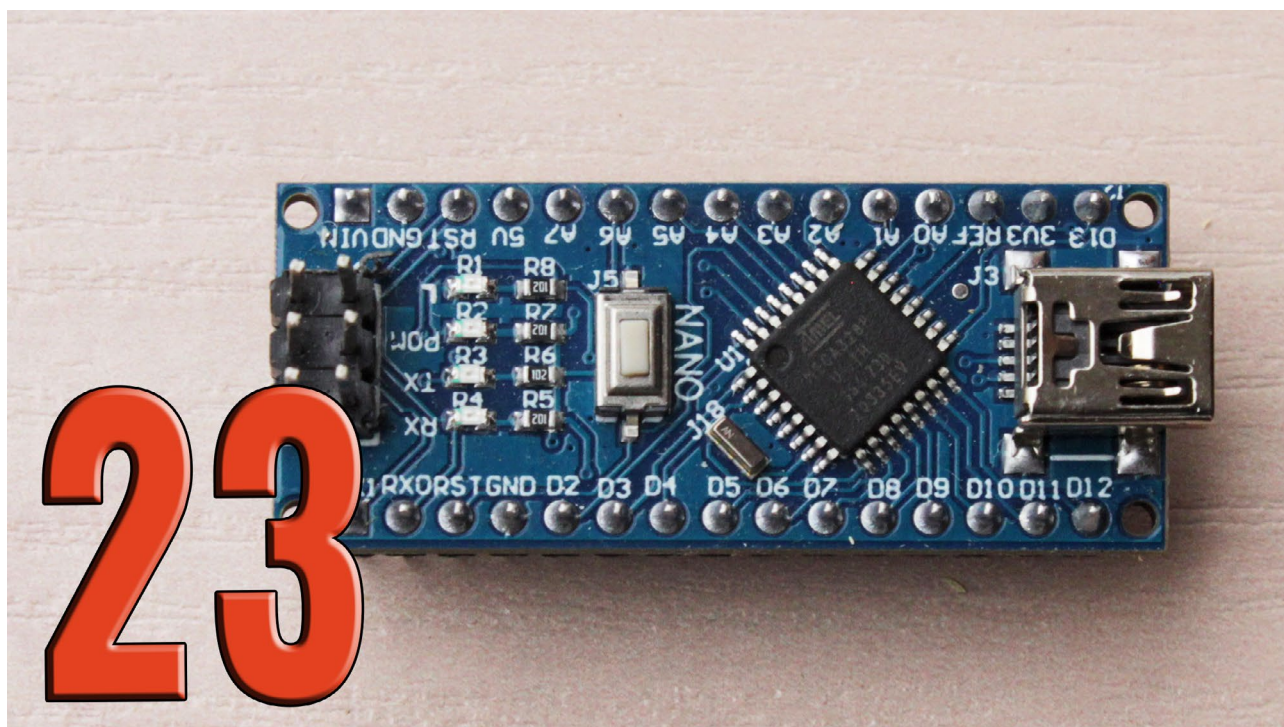
w elektronice energia cieplna (ciepło) jest przekleństwem. Ja w innych, trudniejszych artykułach i filmach omawiam rodzaje i przemiany różnych form energii. A tutaj przypominam w wielkim skrócie: **jeżeli w układzie elektronicznym wydziela się ciepło, energia cieplna, to następuje wzrost temperatury. A wysoka temperatura jest śmiertelnym wrogiem elementów i układów elektronicznych**. Właśnie wydzielanie ciepła jest przeszkodą w miniaturyzacji elementów elektronicznych, bo **wydzielane w układach ciepło trzeba rozprzyszczyć do otoczenia, żeby nie dopuścić do nadmiernego wzrostu temperatury**. Urządzenia elektroniczne nie mogą być gorące po pierwsze z uwagi na możliwość poparzenia użytkownika, a po drugie z uwagi na ryzyko uszkodzenia.

Wydzielanie ciepła i wzrost temperatury to jeden aspekt poważnego problemu. **Drugi ważny aspekt to marnowanie cennej energii baterii i akumulatorów**. Dawniej mówiło się, że rezystor służy do ograniczania prądu w obwodzie. Często tak bywa – niestety. Przykładem jest taśma LED.



Rysunek 2

Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE. W pełnej wersji ten artykuł oczywiście ma więcej stron.



Mikroprocesorowa ośła łączka, część 23

Każdy system mikroprocesorowy wymaga komunikacji z użytkownikiem. Poza podstawowymi możliwościami, jakie daje wykorzystanie klawiatury oraz wyświetlaczy, duże znaczenie ma transmisja szeregową.

USART w ATMEGA328

Prędkość transmisji szeregowej

Obsługa przerw od USART

Zastosowanie bufora cyklicznego

Program maszyny do pisania

Realizacja fizyczna

Transmisja szeregową jest jednym z podstawowych „mediów” pozwalających na przesyłanie nawet sporych objętościowo, szeroko rozumianych danych w obu kierunkach: do mikrokontrolera oraz z mikrokontrolera. Obecnie produkowane układy standardowo są wyposażone w tego typu kontrolery a często zdarza się, że jest ich więcej niż jeden (zależy to od modelu mikrokontrolera). Na przykładzie *Arduino Nano* typowo wykorzystywany był ATMEGA328 (takie były pierwsze, obecnie montowane są ATMEGA328PB), który ma do dyspozycji jeden kontroler transmisji szeregowej, jak pokazuje **rysunek 1** (na następnej stronie). Tego typu podzespoły

często są identyfikowane skrótem UART (ang. Universal Asynchronous Receiver and Transmitter), czyli kontroler szeregowej transmisji asynchronicznej (nadający oraz odbierający dane) lub USART (ang. Universal Synchronous and Asynchronous Receiver and Transmitter) – kontroler szeregowej transmisji synchronicznej i asynchronicznej (również umożliwiający nadawanie oraz odbiór danych). Występuje tu pewna różnica między tymi dwoma rodzajami kontrolerów obsługi transmisji szeregowej (jako synchroniczna/asynchroniczna). W pierwszym przypadku przesyłanie danych jest zsynchronizowane dodatkowym sygnałem zegarowym i jest

stosowane głównie w systemach telekomunikacji do przesyłania większych bloków danych oraz budowania sieci (przykładowo protokoły wyższych warstw, jak SDLC czy HDLC bazujące na transmisji synchronicznej). W zastosowaniach amatorskich praktycznie nie jest używane (nie będziemy się tym zajmować). Wariant asynchroniczny to powszechnie stosowane rozwiązanie do przesyłania danych między systemami mikroprocesorowymi.

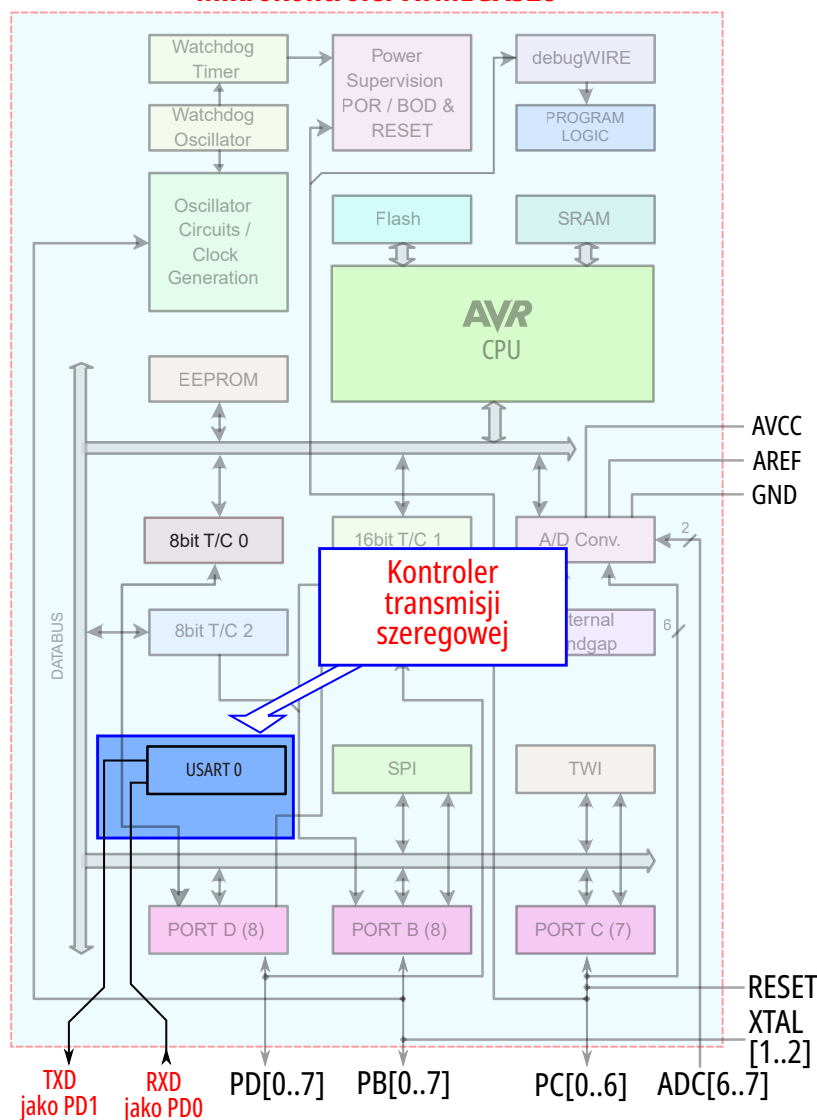
Z transmisją szeregową wiązą się dodatkowo dwa określenia: tryb transmisji oraz prędkość transmisji. Zadaniem układu USART, zawartego w mikrokontrolerze, jest zamiana danych równoległych (mikrokontroler wpisuje do odpowiednich rejestrów USART dane równoległe jako 8-bitów) na dane szeregowe, które są wysyłane jeden po drugim z kolejnych pozycji bitowych na odpowiednim pinie mikrokontrolera. Z tym wiąże się tempo wysyłania tych danych, określane jako prędkość transmisji szeregowej (w literaturze jest określane jako *baud rate*). Z kolei tryb transmisji definiuje takie informacje jak liczba transmitowanych bitów (nie zawsze musi to być po osiem bitów) oraz dodatkowe szczegóły związane z weryfikacją poprawności transmisji. Taką elementarną kontrolą jest dodawanie do każdego wysłanego bajtu danych bitu parzystości. Oczywiście można przysyłać dane bez bitów kontroli. Więcej informacji na ten temat jest w artykule *Mikroprocesory i mikrokontrolery – interfejs szeregowy RS232*, ZE 04/2024.

Naturalnym jest, by obie strony (nadająca oraz odbierająca dane) robiły to z tą samą prędkością oraz stosowały ten sam tryb transmisji. W prostych mikrokontrolerach najczęściej wykorzystywany jest tryb 8N1 (8 bitów danych, bez kontroli parzystości i z jednym bitem stopu) oraz prędkość 9600 bps (ang. bits per second – liczba bitów na sekundę). Prędkości transmisji są znormalizowane, a 9600 bps jest jedną z nich.

USART w ATMEGA328

Mikrokontroler ATMEGA328 zawiera w swojej

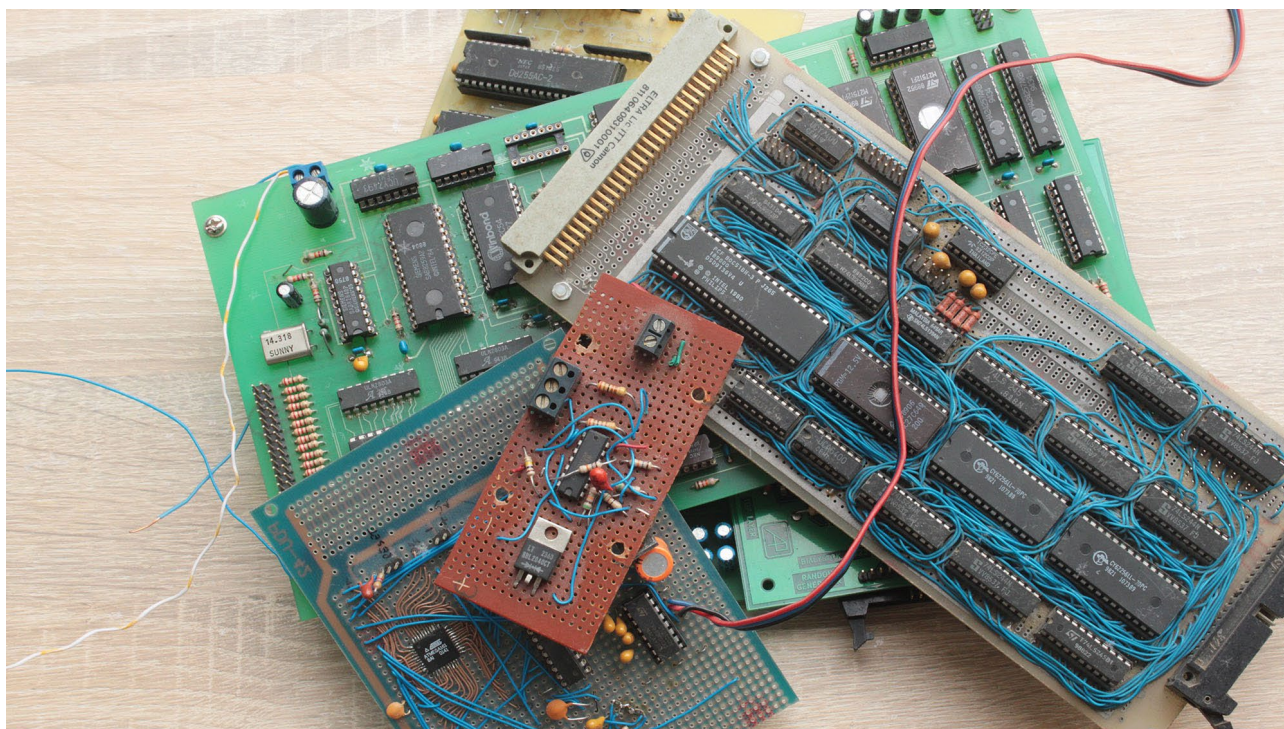
Mikrokontroler ATMEGA328



Rysunek 1

na tryb transmisji oraz definiujących jej prędkość. Najczęściej stosowanym trybem transmisji jest 8N1. Zapewne wynika to z tego, że pierwsze mikrokontrolery nie obsługiwały dodawania (przy nadawaniu) i weryfikacji (przy odbieraniu) bitów kontroli parzystości i te działania musiał realizować programista, umieszczając wyliczony bit parzystości (lub nieparzystości) w odpowiednich polach rejestrów przy nadawaniu. Te wszystkie szczegóły określa się ustawiając odpowiednie bity w rejestrach konfiguracyjnych mikrokontrolera AVR. Jego odpowiednie bity są ustawiane na początku programu. Tu można wykorzystać tę cechę mikrokontrolerów AVR, że po operacji zerowania (naciśnięciu przycisku reset lub akcji automatycznego resetu po włączeniu zasilania) bity są zerowane, więc wystarczy „zająć się” bitami,

Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE. W pełnej wersji ten artykuł oczywiście ma więcej stron.



Wspólnie projektujemy: Montaż prototypów

Kwestia budowy prototypów to wręcz niewyczerpany temat. Tę tematykę można rozpatrywać w dwóch kategoriach: badawczo-rozwojowych oraz „budowlanych”.

„Druciany” prototyp czy fabryczne PCB Własne PCB

Prawdą jest, że wiedza postrzegana jedynie w kategoriach „teoretycznych” nie ma żadnego zastosowania i właściwie niczemu nie służy. Można stworzyć „na papierze” dowolną konstrukcję, jednak zawsze pozostaje pytanie, czy będzie to działać zgodnie z naszym zamysłem. Istnieje tylko jedna droga, która jest w stanie to zweryfikować: należy zbudować prototyp. Tu pojawia się dosyć istotny problem technologiczny: jak to zrobić? Większość moich zainteresowań skupia się wokół układów cyfrowych oraz mikroprocesorowych. Te mają taką specyfikę, że pomimo koncepcyjnej prostoty są często dosyć rozbudowane, gdyż występuje wiele układów, które należy ze sobą połączyć.

„Druciany” prototyp czy fabryczne PCB

Budowę prototypów można widzieć z dwóch perspektyw: finalnego urządzenia lub jakiegoś jego

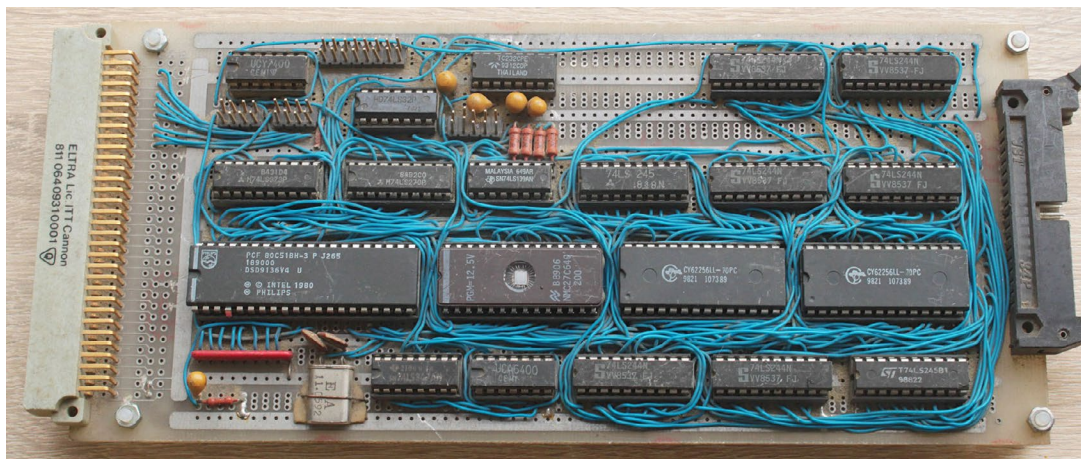
Współczesne wymagania

fragmentu. Przy tworzeniu konstrukcji zawsze występuje wiele niewiadomych. Ograniczając się do układów cyfrowych, takich szczegółów jest całe mnóstwo. Przykładowo biorąc pod uwagę klasyczne liczniki cyfrowe, realizują one przede wszystkim dwie operacje: wyzerowanie licznika oraz zliczanie impulsów. Wyzerowanie to kwestia podania na odpowiednie wejście stanu logicznego (są dwie możliwości: stan wysoki lub niski). Z kolei zliczanie odbywa się w wyniku wystąpienia zbocza sygnału (również są dwie możliwości: zbocze narastające lub opadające). Łatwo stworzyć układ, w którym nie działa licznik, bo przykładowo jest na stałe zerowany.

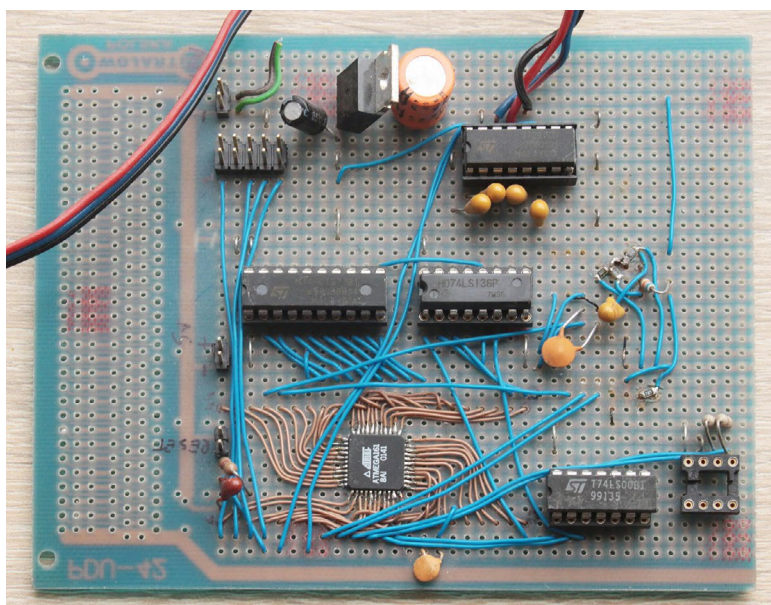
W takich przypadkach warto „przećwiczyć” odpowiednie fragmenty całości i przekonać się, że działają zgodnie z założeniami. To wymaga zbudowania realnego układu. Na przestrzeni kilkudziesięciu lat bardzo zmieniały się możliwości technologiczne.

Na początku mojej przygody z elektroniką operowałem elementami, które mają niewiele wyprowadzeń (najwięcej miał tranzystor). Ta cecha implikowała określone rozwiązania. Używałem płytek PCW, w których rozgrzanym gwoździem robiłem otwory (w latach 70. mieszkalem niedaleko budowanego bloku, gdzie takie płytki były standardowym wyposażeniem w mieszkaniach i wiele takich ścinków wałało się na budowie). Wkładałem tam elementy i łączyłem odizolowanym drucikiem po drugiej stronie (nawet nie była potrzebna lutownica, gdyż wyprowadzenia owijałem drucikiem). Później nadeszły czasy własnych płytek drukowanych, gdzie lakierem do paznokci malowane były ścieżki a całość trawiona w kwasie azotowym. W tej technologii uzyskanie nawet standardowego rastru wyprowadzeń (2,54 mm) jest problematyczne. Następny postęp technologiczny to pojawienie się uniwersalnych płytek prototypowych. **Fotografia 1** pokazuje realizację pewnego urządzenia z mikrokontrolerem, które na dnie szuflady zachowało się do obecnych czasów. Ta technologia jest pracochłonna i łatwo o pomyłkę, ale z drugiej strony ma bardzo pozytywną cechę, jaką jest możliwość modyfikacji i „naprawy” własnych błędów. Znaczącą jej wadą jest ograniczenie do układów w obudowach typu DIP. Pojawienie się układów w obudowach mało przyjaznych hobbystom, również odpowiednim wysiłkiem dawało się pokonać. **Fotografia 2** jest przykładem, że tego typu ograniczenia są do rozwiązania.

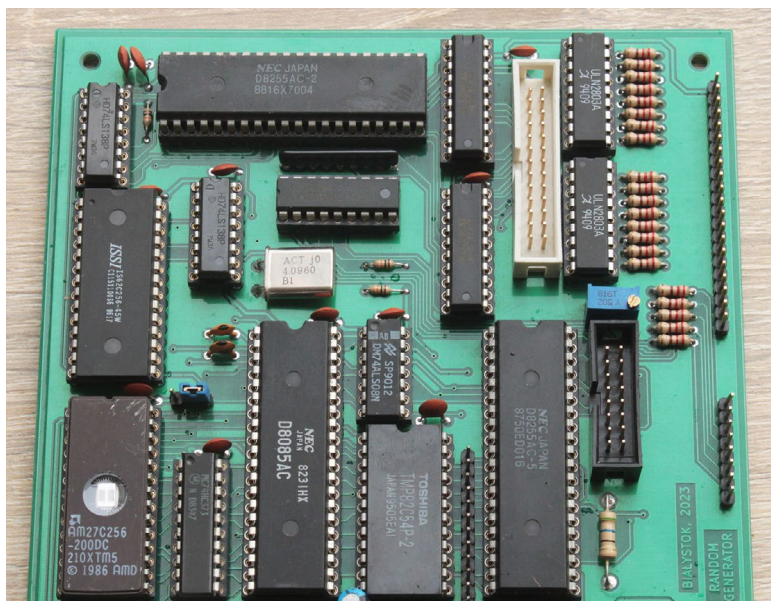
Współczesne układy stają się bardziej złożone, a to pociąga za sobą konieczność większej znajomości ich cech i warunków współdziałania z innymi komponentami. Nawet w obrębie układów cyfrowych, gdzie mogłoby się wydawać, że nie powinny wystąpić problemy, rzeczywistość przynosi niespodzianki. Tu za przykład może posłużyć urządzenie z wykorzystaniem mikroprocesora. Z punktu widzenia konstrukcyjnego nie wystąpił żaden problem (urządzenie jest w trakcie budowy i jak dotychczas nie objawił się żaden problem hardware’owy dotyczący połączeń między elementa-



Fotografia 1



Fotografia 2



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.
W pełnej wersji ten artykuł oczywiście ma więcej stron.



Elektroenergetyka – przesył i dystrybucja, część 2

Dziś zajmiemy się drugą częścią systemu przesyłu i dystrybucji energii elektrycznej – napięciem średnim i niskim. To właśnie ta część systemu jest tym, co widzimy na każdej osiedlu, ulicy czy przy każdym domu.

Konstrukcja linii niskiego i średniego napięcia
Słupy linii systemu dystrybucyjnego

Przewody
Izolacja

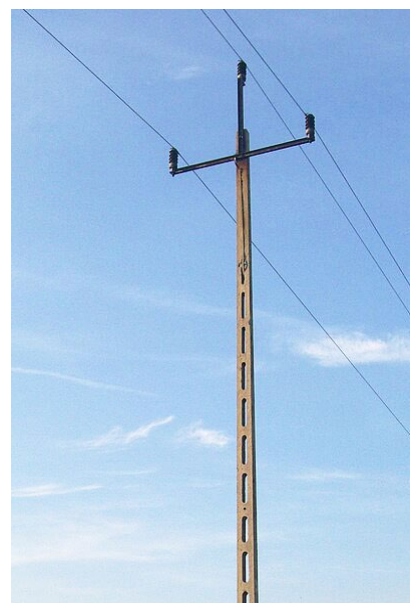
Konstrukcja linii niskiego i średniego napięcia

Linie średniego i niskiego napięcia to w większości linie napowietrzne (szczególnie na terenach wiejskich), zaś w miastach to praktycznie linie kablowe. Nowe odcinki linii, w miarę możliwości, szczególnie w terenie zabudowy mieszkaniowej, budowane są jako linie kablowe. Głównym uzasadnieniem jest brak uszkodzeń sieci elektroenergetycznej w wyniku zjawisk atmosferycznych, takich jak złamane drzewa i inne podobne.

Główną różnicą pomiędzy liniami SN i nN a liniami WN i NN jest brak przewodów odgromowych. Do-

Słupy linii systemu dystrybucyjnego

W odróżnieniu od słupów linii WN i NN, słupy linii SN i nN to zazwyczaj konstrukcje betonowe – konstrukcje kratowe stosowane są wyjątkowo, zazwyczaj w przypadku konieczności wykorzystania słupa o znacznej wysokości. Słupy betonowe to tzw. żerdzie. Dawniej powszechnie



**Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.
W pełnej wersji ten artykuł oczywiście ma więcej stron.**





GNSS RTK, testy terenowe nawigacji

W poprzednim numerze ZE opisywałem moje zmagania związane z uruchomieniem bazującego na Raspberry Pi odbiornika GNSS RTK, czyli nawigacji o dokładności centymetrowej. Wyniki pierwszych prób okazały się na tyle obiecujące, że postanowiłem zweryfikować działanie tego urządzenia w terenie.

Źródło zasilania

Dostęp do sieci

Interfejs użytkownika

Obudowa

W poszukiwaniu precyzyjnych map

QGIS

Testy z mapami uzbrojenia terenu

Testy z podstawą geodezyjną

Podsumowanie

Choć nawigacja poprawnie działała na moim biurku, to przeprowadzenie kolejnych testów poza domem wymagało wprowadzenia pewnych zmian, aby instalacja stała się mobilna: należało zadbać o niezależne zasilanie, łącze internetowe oraz interfejs użytkownika. Poniżej pokrótce opiszę rozwiązania, na jakie się zdecydowałem.

Źródło zasilania

Ponieważ Raspberry Pi zasilane jest poprzez gniazdo USB, oczywistym rozwiązaniem jest użycie powerbanku. Zmierzony pobór prądu mojego Raspberry Pi 3 wraz z podłączonym odbiornikiem

GNSS i wyświetlaczem LCD wynosi około 0,8 A. Pozwala to oczekiwać, że powerbank o pojemności 20 000 mAh zapewni urządzeniu wiele godzin nieprzerwanej pracy.

Dostęp do sieci

W domu Raspberry Pi mogło korzystać z routera Wi-Fi, jednak w terenie należało zapewnić inne łącze internetowe. Zdecydowałem się na „hotspot osobisty”, który można dziś aktywować w niemal każdym smartfonie. W zasięgu sieci 5G szybkość transmisji okazała się w pełni wystarczająca do przesyłania poprawek RTK do modułu GNSS.

Wydaje się, że ze względu na niskie zapotrzebowanie na przepustowość danych, całość powinna działać również przy połączeniu LTE bądź 3G, jednak nie miałem jeszcze okazji sprawdzić tego w praktyce.

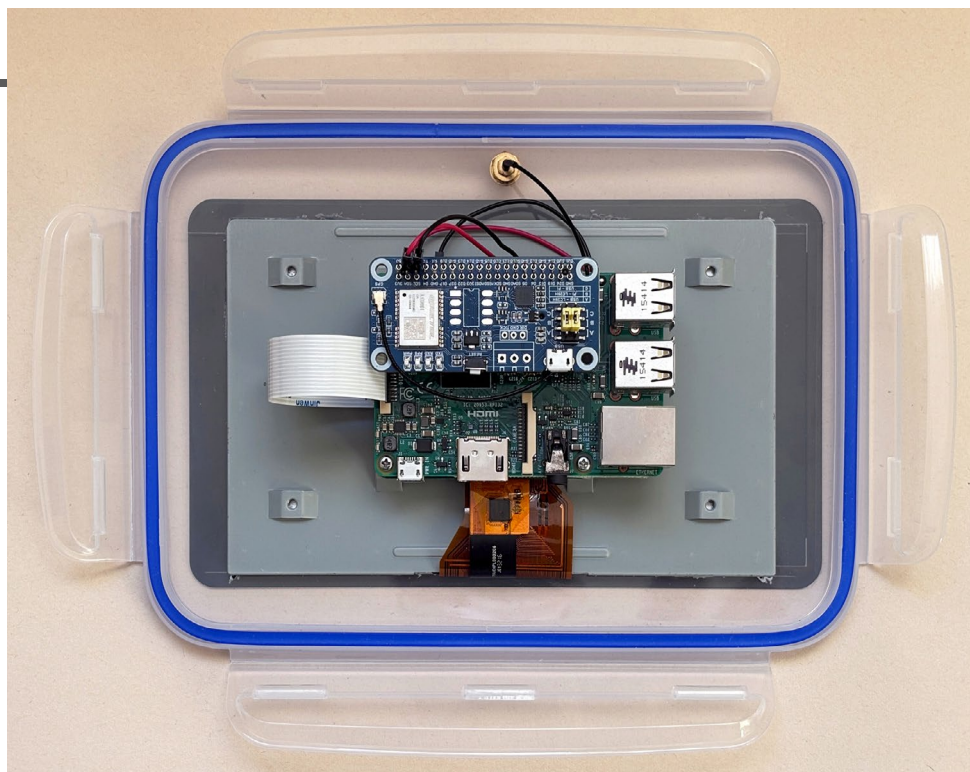
W tym miejscu warto zauważyć, że mój pracujący pod kontrolą systemu iOS smartfon nie umożliwia podglądu adresu IP klienta korzystającego z udostępnianego łącza. Stanowi to pewne utrudnienie, ponieważ aby połączyć się z Raspberry Pi w ramach lokalnej sieci hotspota, powinienem znać jego adres IP. Ten zaś jest przydzielany dynamicznie, a zatem może ulegać zmianom. Problem ten staje się istotny w sytuacji, gdy klient nie ma podłączonego fizycznego ekranu.

Interfejs użytkownika

Pracując z Raspberry Pi w domu, zazwyczaj korzystam ze zdalnego pulpitu, co pozwala mi na używanie monitora, myszy i klawiatury podłączonych do komputera stacjonarnego. W terenie można oczywiście posilkować się laptopem lub odpowiednią aplikacją na smartfonie.

Jednak tym razem, ze względu na wspomniane wcześniej wątpliwości związane z dynamicznym adresem IP, zdecydowałem się na podłączenie fizycznego ekranu. Wykorzystałem posiadany od lat dedykowany dla Raspberry Pi wyświetlacz dotykowy 7" Touchscreen Display, którego wcześniej nie miałem jeszcze okazji przetestować. Montaż polegał na przykręceniu minikomputera czterema śrubami do tylnej części obudowy wyświetlacza. Za połączenie elektryczne odpowiada specjalna taśma przesyłająca obraz oraz dane dotyku, a także dwa przewody wpinane w piny GPIO, które służą do zasilania. System automatycznie wykrywa

Fotografia 1



ekran oraz funkcje dotykowe, dzięki czemu konfiguracja nie jest wymagana.

Obudowa

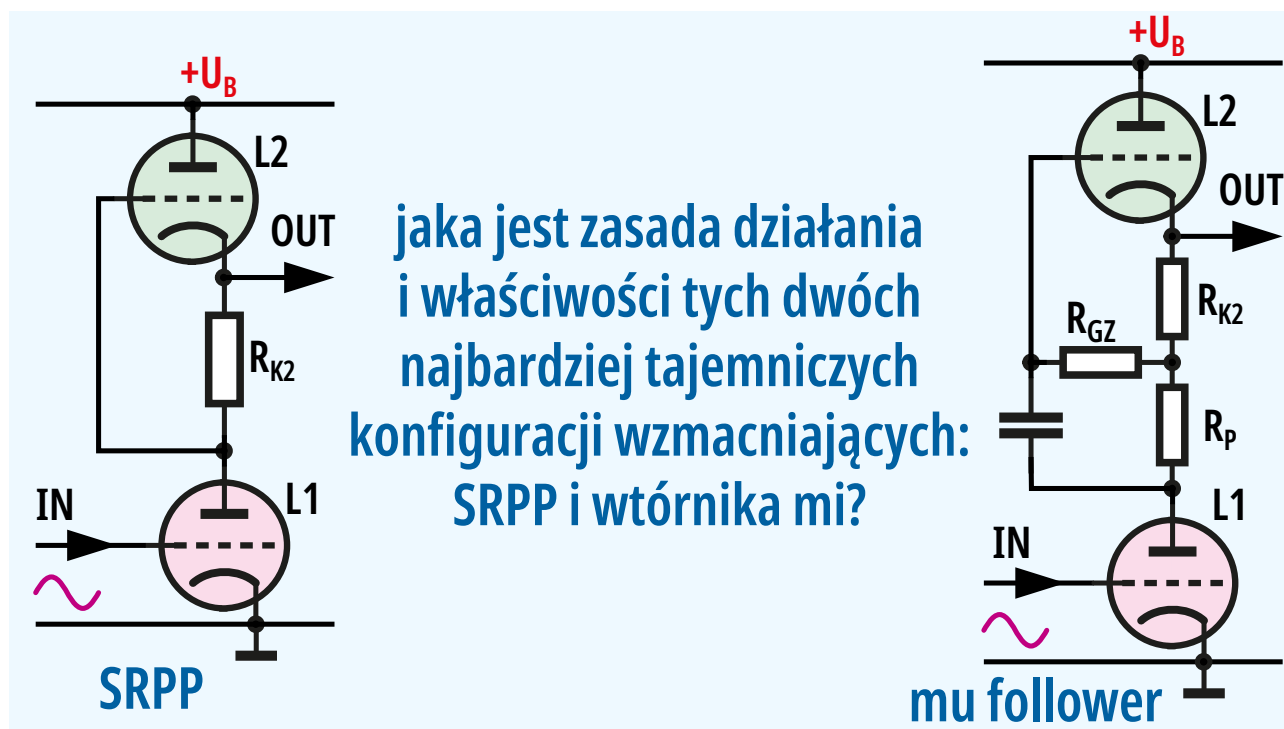
Moja nawigacja była już niemal gotowa do wyprawy w teren, jednak całość należało jeszcze zabezpieczyć przed uszkodzeniem. Zdecydowałem, że na czas testów wszystko umieszczę w plastikowym pudełku, widocznym na **fotografii tytułowej**. Pod względem gabarytów idealny okazał się wykonany z przezroczystego polipropylenu pojemnik spożywczy o pojemności 2,2 litra.

W pokrywie wyciąłem prostokątny otwór o wymiarach nieco mniejszych niż metalowa ramka wyświetlacza. Dzięki temu, po wciśnięciu ekranu w przygotowane miejsce, nie było potrzeby stosowania dodatkowych mocowań. Powyżej ekranu wywierciłem otwór, w którym przykręciłem gniazdo antenowe zamontowane na końcu krótkiego przewodu, wpinanego bezpośrednio do modułu GNSS. Wnętrze pokrywy z zamontowanymi podzespołami szczegółowo prezentuje **fotografia 1**.

Moduł odbiornika został wpięty w dedykowane złącze za pośrednictwem dostarczonej przez producenta przejściówki z długimi goldpinami, które umożliwiają podłączenie kolejnych podzespołów. W moim przypadku wpiąłem w nie jedynie dwa przewody dostarczające napięcie 5 V do

**Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.
W pełnej wersji ten artykuł oczywiście ma więcej stron.**

Fotografia 2



Lampowe konfiguracje SRPP i mu follower (1)

W cyklu o lampach elektronowych chcemy zająć się dwiema konfiguracjami, które są najmniej rozumiane i najbardziej tajemnicze. Niewiele osób ma jasne wyobrażenie o ich działaniu i właściwościach. Zanim je omówię, muszę koniecznie przedstawić obszerne informacje wstępne i wyjaśnić liczne nieporozumienia.

Nie ma komplementarnych lamp!

Co to jest push-pull?

W tym artykule zaczynamy omawianie popularnych, a słabo rozumianych konfiguracji znanych jako SRPP oraz mu-follower. Wcześniej omawiałem głównie konfiguracje dwutriodowe, będące prostym połączeniem układów pojedynczych, które po kolei omawiałem we wcześniejszych artykułach. Te wcześniej omawiane konfiguracje wcale nie są tajemnicze, tylko trzeba dobrze zrozumieć podstawowe, elementarne układy pracy z jedną triodą. A potem już można realizować konfiguracje dwutriodowe niemal na zasadzie „każdy z każdym”. Negatywnym, niepraktycznym wyjątkiem jest tam konfiguracja, gdzie na wejściu jest wtórnik (układ ze wspólną anodą), który współpracuje ze wzmacniaczem ze wspólną katodą.

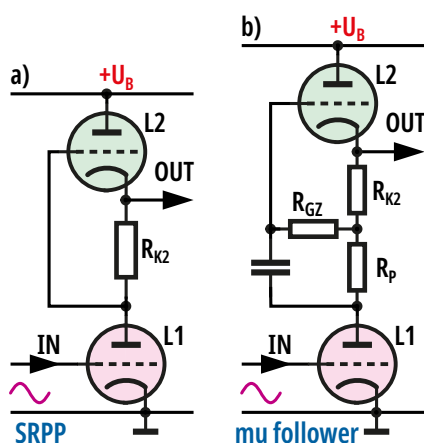
Spośród omówionych dotychczas konfiguracji dwulampowych tylko wtórnik White'a nie jest prostym, bezpośrednim złożeniem elementarnych układów, bo występuje tam dodatkowa interakcja, coś w rodzaju sprzężenia zwrotnego.

Co ogromnie ważne, także w konfiguracjach SRPP i mu-follower występuje specyficzna interakcja między lampami składowymi. I właśnie ta interakcja jest nieprzeniknioną tajemnicą nie tylko dla początkujących, ale też dla mnóstwa osób, które uważają się za znawców techniki lampowej. Wiele osób ma fałszywe, z gruntu błędne wyobrażenia o działaniu i właściwościach dwutriodowych stopni wzmacniających, zwanych SRPP i mu-follower.

Właśnie dlatego konieczne jest szerokie naświetlenie tła i wyprostowanie błędnych wyobrażeń, czego nie da zrobić się w jednym artykule. Dlatego konfiguracjom tym poświęcone są aż cztery obszernie artykuły.

Na **rysunku 1a** pokazana jest w pewnym uproszczeniu konfiguracja nazywana **SRPP**, co powszechnie rozszyfrowywane jest jako **Series Regulated Push-Pull**, czyli jako szeregowy układ przeciwsoalny, bo „push-pull” znaczy dosłownie „pchaj i ciągnij”, co po polsku nazywamy pracą przeciwsoalną. W Internecie można znaleźć lepsze, gorsze oraz zupełnie błędne analizy działania takiego wzmacniacza. Mało tego!

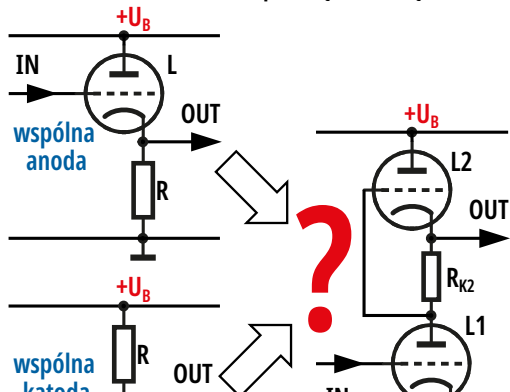
Otóż oprócz tajemniczego układu z rysunku 1a, w literaturze spotykamy jeszcze bardziej tajemniczą



Rysunek 1

I tu **podstawowe pytanie: czy w konfiguracjach z rysunku 1 górna lampa pracuje jako wtórnik, czy-li układ ze wspólną anodą?** Jak Ty uważasz?

To jest naprawdę najważniejsze pytanie! Bez prawidłowego zrozumienia działania górnej lampy nie można zrozumieć działania całości. Wiele osób jest przekonanych, że SRPP z rysunku 1a to połączenie układu ze wspólną katodą (dolna lampa) z układem



ze wspólną anodą – wtórnikiem (górna lampa), według **rysunku 2**.

Kwestię, czy górna lampa L2 to wzmacniacz, czy wtórnik, czy może jeszcze coś innego,

Do wyjaśnienia jest też **druga ważna kwestia: dlaczego taką konfigurację nazywa się układem przeciwsoalnym (push-pull)?** Jak w takim układzie dopatrzeć się pracy w trybie „pchaj i ciągnij”?

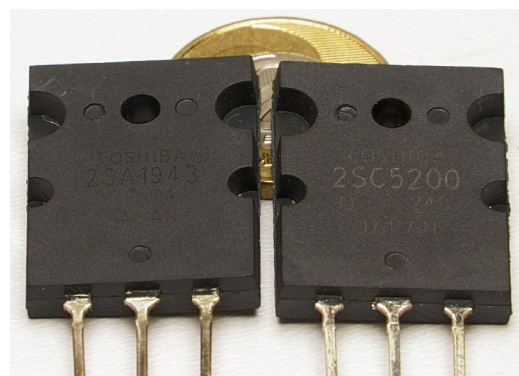
Jeżeli jednak wersja z rysunku 1a jest układem przeciwsoalnym, to co z tą „przeciwsoalnością” w podobnej wersji z rysunku 1b? **Jeżeli wersja z rysunku 1b też jest układem przeciwsoalnym, to dlaczego jest nazywana wtórnikiem?** I dlaczego w nazwie występuje litera mi (μ), która oznacza współczynnik amplifikacji, czyli maksymalne teoretyczne wzmocnienie napięciowe triody? Przecież wiadomo, że wtórnik ma wzmocnienie bliskie jedności, a współczynnik amplifikacji triod wynosi zwykle kilkadziesiąt? Skąd takie zamieszanie?

W Internecie jest trochę sensownych, ale niestety wielokrotnie więcej mało sensownych oraz błędnych informacji na temat SRPP i *mu follower*. Jest też sporo osób, które zajmują się lampami, ale nie mają głębszej wiedzy i „wciskają kit” swoim rozmówcom, głównie klientom, którym chcą za duże pieniądze sprzedać rzekomo rewelacyjne wzmacniacze z takimi właśnie tajemniczymi stopniami.

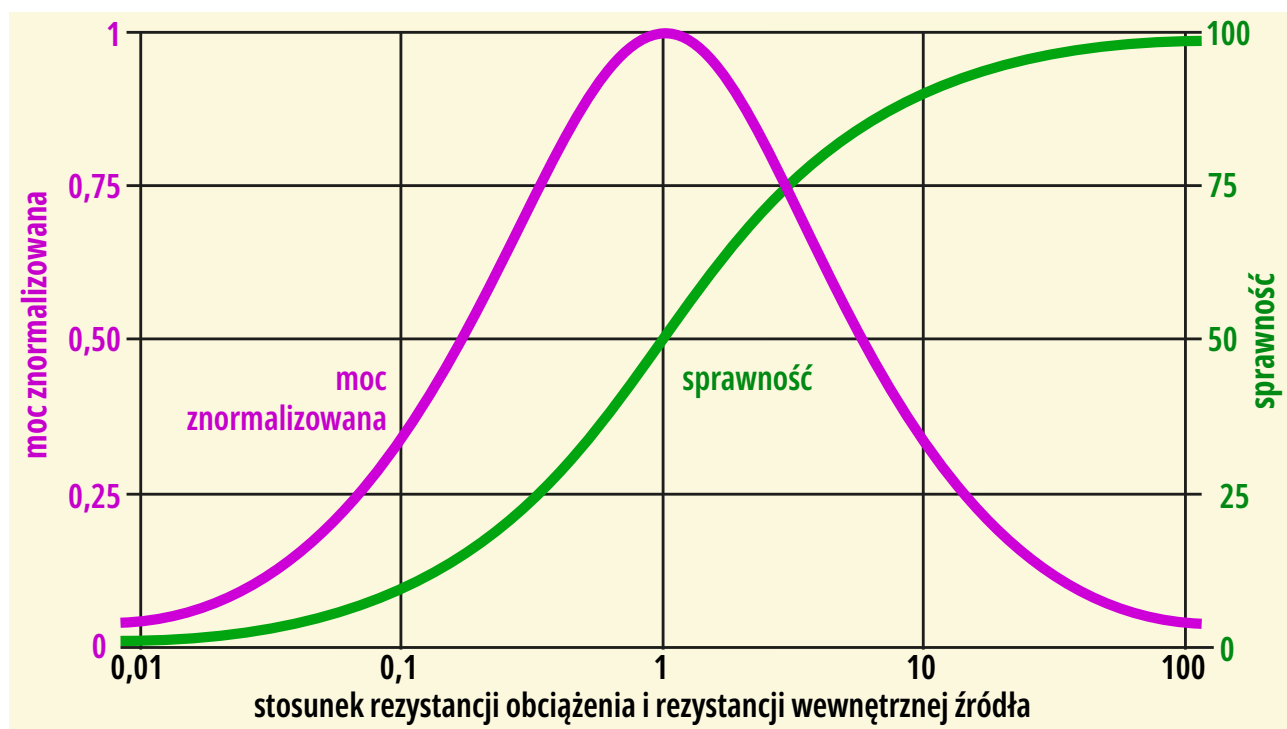
Ja nie jestem ekspertem od lamp, ale na elektronice trochę się znam. W kolejnych artykułach detalicznie wyjaśnię działanie konfiguracji SRPP i *mu follower*, ale ich analiza jest dość skomplikowana. Dlatego konieczne muszę wcześniej podać czy też przypomnieć pewne informacje podstawowe. Informacje dotyczące obciążenia dynamicznego, informacje o wzmacniaczach przeciwsoalnych oraz o komplementarnych elementach wzmacniających. Z początku będzie się wydawać, że omawiam zupełnie niepowiązane zagadnienia. Potem jednak te informacje ułożą się w harmonijną całość. Zaczniemy od zaskakującego tematu „komplementarnych lamp”.

Nie ma komplementarnych lamp!

Dla współczesnego elektronika jest „oczywista oczywistością”, że mamy do dyspozycji zarówno tranzystory typu NPN, jak i komplementarne, czyli dopełniające



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE. W pełnej wersji ten artykuł oczywiście ma więcej stron.



Dopasowanie energetyczne. Czy w ogóle jest potrzebne?

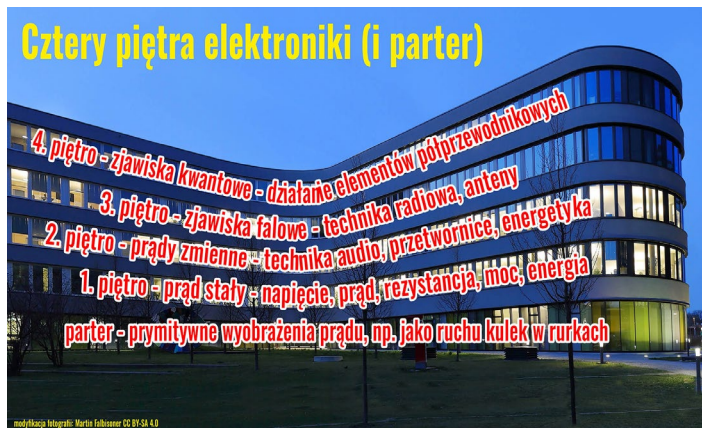
W poprzednim artykule serii omówiłem teoretyczny aspekt szerokiej kwestii dopasowania energetycznego, czyli dopasowania mocy. Z dopasowaniem energetycznym najczęściej mamy do czynienia w układach wysokiej częstotliwości. W poniższym artykule przedstawię zarys tego obszernego zagadnienia.

Dopasowanie energetyczne w układach w.cz.?
Dopasowanie energetyczne i „wydajność źródła”
Dopasowanie energetyczne anten

Generatory, wzmacniacze wysokiej częstotliwości
Reaktancja, impedancja i moc bierna

W poprzednim artykule pokazałem dobitne przykłady akumulatora i wzmacniacza mocy audio, gdzie absolutnie nie chcemy mieć dopasowania energetycznego. Tam, żeby nie zmniejszać napięcia, dużo korzystniejsza jest praca przy oporności obciążenia R_L dużo większej, niż rezystancja wewnętrzna źródła energii R_W .

Z dopasowaniem energetycznym często, a właściwie powszechnie mamy do czynienia w układach wysokiej częstotliwości. Dlaczego? Otóż zagadnienie ma kilka nietrywialnych aspektów, które przedstawię, a raczej tylko zasygnalizuję teraz w sposób jak najprostszy. Dopiero potem zajmę się dopasowaniem szumowym.



Dopasowanie energetyczne w układach w.cz.?

Wielu osobom wydaje się, że w układach wysokiej częstotliwości najważniejszą kwestią jest unikanie odbić i że to bezwzględnie wymaga dopasowania falowego, czyli równości rezystancji falowej kabla i rezystancji dopasowujących z jego dwóch stron.

Wielu osobom wydaje się, że w układach w.cz. omawiane teraz dopasowanie energetyczne jest niejako efektem ubocznym, wynikającym tylko z konieczności dopasowania falowego. Efektem ubocznym, ale nieuniknionym. Nieuniknionym, ale niepożądanym, bo jak pokazałem w poprzednim artykule serii, przy dopasowaniu energetycznym marnujemy połowę mocy w rezystancji wewnętrznej źródła.

Takie są częste wyobrażenia. Otóż są to wyobrażenia niepełne, zbyt uproszczone i w pewnych sytuacjach nieprawdziwe.

Sytuacja w układach w.cz. jest podwójnie, a nawet potrójnie skomplikowana. Jak już wcześniej pisałem, **dopasowanie falowe jest wymagane tylko w przypadku wykorzystywania linii długich – kabli. Jeśli takich długich kabli nie ma, to nawet w układach w.cz. dopasowanie falowe nie jest wymagane.**

Tak! Dopasowanie falowe **nie jest niezbędne** jeśli przesyłamy sygnał na „małe” odległości, gdy możemy zaniedbać opóźnienia i odbicia. I tu ważne pytanie: **czy wtedy potrzebne lub pożądane jest dopasowanie energetyczne bez dopasowania falowego?**

Odpowiedź brzmi: To zależy.

W grę wchodzi bowiem dwa niełatwe do zrozumienia aspekty. Jeden dotyczy mocy maksymalnej – omawiam go już w następnym śródtytuł. Drugi aspekt dotyczy szumów i jego zarys przedstawię w następnym artykule tej serii.

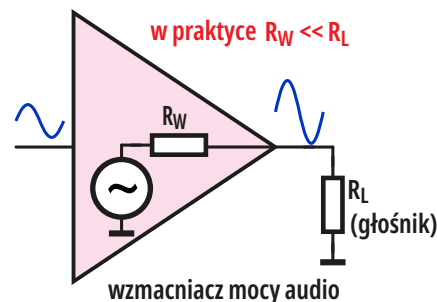
Dopasowanie energetyczne i „wydajność źródła”

W poprzednim artykule pisałem, że w elektronice bardzo często jak ognia unikamy pracy w warunkach dopasowania energetycznego. Przykładem jest akumulator, dla którego dłuższa praca w warunkach dopasowania energetycznego oznaczałaby katastrofę. Dlaczego więc w szkole i na studiach tak dużo mówi się o dopasowaniu energetycznym?

Otóż w układach wysokiej częstotliwości – radiowych – sytuacja jest zupełnie inna, niż w przypadku akumulatorów. Akumulatory dysponują dużą ilością energii, są źródłami napięciowymi, a nie prądowymi i prawie zawsze pracują przy rezystancji obciążenia wielokrotnie większej, niż ich rezystancja wewnętrz-

Z wyjątkiem krótkich chwil pracy rozrusznika samochodowego, wcale nie zależy nam na tym, żeby z akumulatora pobierać maksymalną moc. Wprost przeciwnie – z różnych względów wcale tego nie chcemy, także dlatego, żeby tego akumulatora zbyt szybko nie rozładować. Podobnie jest przy przebiegach zmiennych w systemach audio, przede wszystkim we wzmacniaczach mocy – **rysunek 1**. Zależy nam, by nie tłumić, **nie zmniejszać napięcia**.

Zupełnie inaczej jest w układach i obwodach radiowych. W szczególności w przypadku anten i wzmacniaczy wysokiej częstotliwości, ale to jest dla wielu trudne do zrozumienia. Zaczniemy w sposób chyba najprostszy – od anten odbiorczych.



Rysunek 1

Dopasowanie energetyczne anten

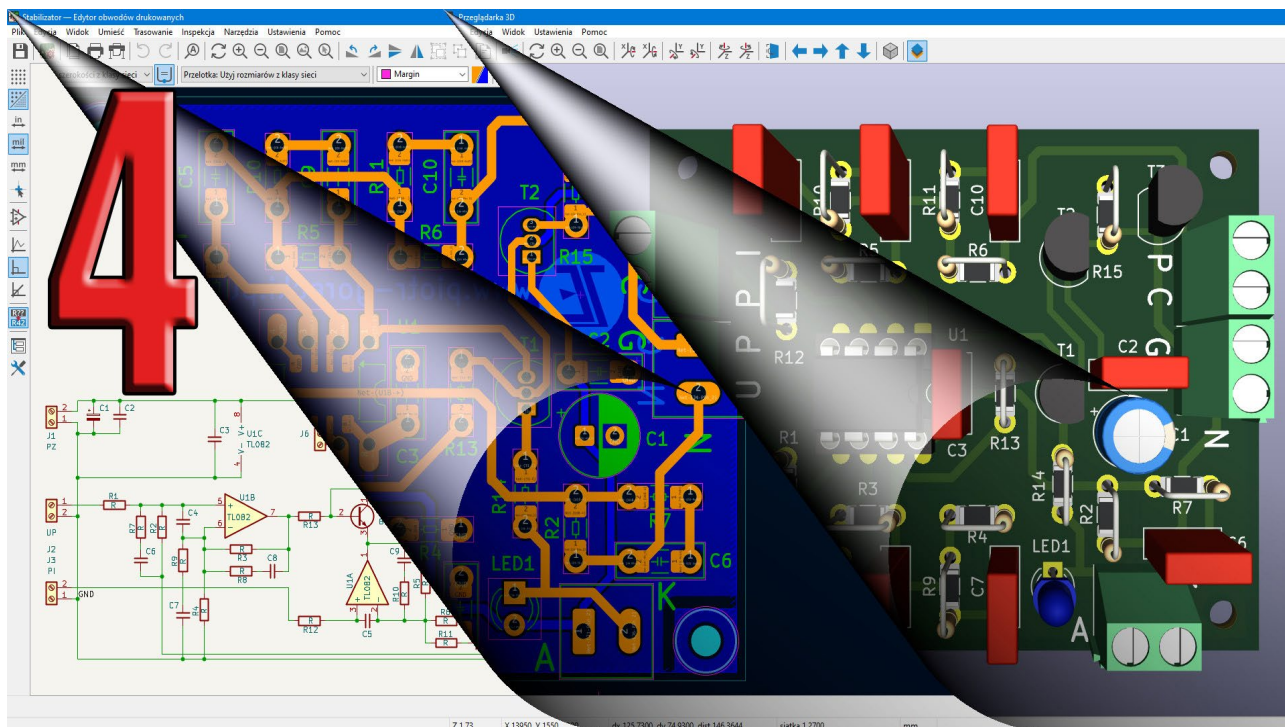
Fale elektromagnetyczne, czyli specyficznie połączone zmiany pola elektrycznego i pola magnetycznego, rozchodzące się w przestrzeni niosą energię. Antena odbiorcza niejako zbiera energię tych fal elektromagnetycznych.

Z punktu widzenia użytkownika, antena odbiorcza niewątpliwie staje się źródłem energii elektrycznej (prądu i napięcia), ale źródłem bardzo słabym, bo energia fal radiowych z dala od nadajnika jest znikoma. Najprościej biorąc, antena odbiorcza staje się „baterijką”, „baterijką napięcia i prądu zmiennego wysokiej częstotliwości” (dla uproszczenia pomijam wyobrażenie anteny jako źródła prądowego).

Antena odbiorcza potraktowana jako „baterijka” czy też jako źródło napięciowe ma oczywiście jakąś rezystancję wewnętrzną i indukuje się w niej jakieś niewielkie napięcie wysokiej częstotliwości. Jak wskazuje **rysunek 2**, antenę, podobnie jak baterię, możemy traktować jako źródło energii elektrycznej o jakimś napięciu U_{ANT} i o jakiejś rezystancji wewnętrznej R_W , przy czym ta rezystancja R_W nie zależy od siły fal radiowych, tylko od konstrukcji anteny i wynika z kilku czynników.



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE. W pełnej wersji ten artykuł oczywiście ma więcej stron.



Projektowanie płytek drukowanych za pomocą KiCad (4)

W dwóch pierwszych artykułach zapoznaliśmy się z pakietem KiCad oraz zrealizowaliśmy schemat. W trzecim zaczęliśmy najważniejszy krok: zapoznaliśmy się z programem „płytkowym” i jego ustawieniami, określiliśmy rozmiary płytki i rozmieściliśmy elementy. Czwarty artykuł pokazuje, jak dokończyć dzieła.

Trasowanie ścieżek
Wypełnienie płytki
Test DRC
Widok 3D

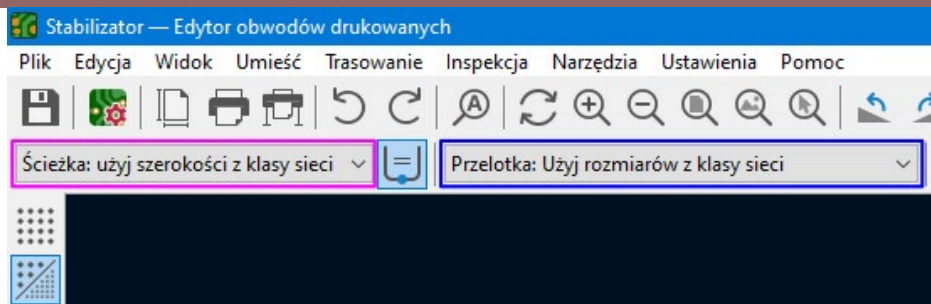
Dane produkcyjne
Przeglądarka plików Gerber
Podsumowanie

W tym artykule opisany jest sposób trasowania ścieżek. W przypadku płytek z jedną warstwą miedzi może to być zadanie trudne, wymagające zastosowania zwór z drutu. My projektujemy płytkę dwuwarstwową. W zasadzie ścieżki można byłoby równomiernie rozmieścić w obu warstwach miedzi, a program automatycznie wprowadzi tzw. przelotki, czyli połączenia między ścieżkami z obu stron płytki. W wielu układach z różnych powodów staramy się zrealizować obwód masy nie w postaci ścieżek, tylko w postaci płaszczyzny masy wypełniającej prawie całą powierzchnię płytki (zwykle dolna warstwa miedzi). Wtedy w tej warstwie poprowadzone zostaną tylko nieliczne, niezbędne ścieżki, a większość ścieżek prowadzona jest w tej drugiej warstwie (górnej).

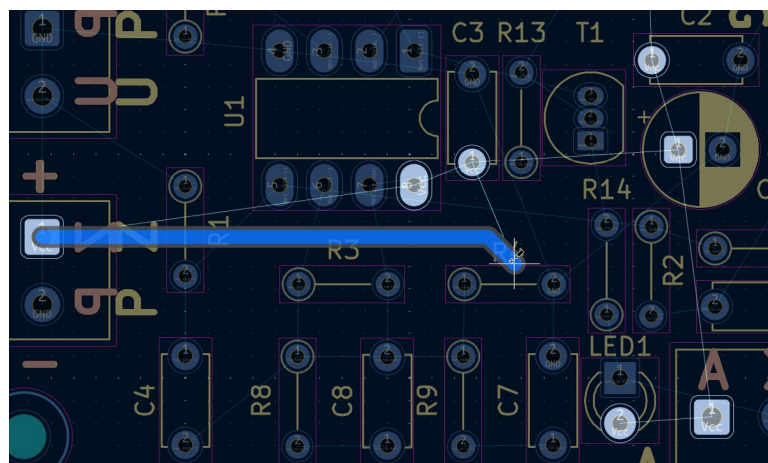
Trasowanie ścieżek

Przed rozpoczęciem trasowania upewniamy się, czy na pasku narzędzi z rozwijanej listy mamy wybrane w różowej ramce **Ścieżka: użyj szerokości z klasy sieci** oraz w niebieskiej ramce **Przelotka: Użyj rozmiarów z klasy sieci**, jak na **rysunku 1**. Pomiędzy tymi ramkami mamy ikonkę, które włączenie powoduje, że gdy rozpoczynamy rysowanie ścieżki od innej ścieżki, nowo rysowana ścieżka będzie miała szerokość ścieżki od której rozpoczęto jej rysowanie. Do trasowania ścieżek wybrałem siatkę 0,508 mm. Jako pierwszą wytrasujemy na warstwie dolnej B.Cu ścieżkę zasilania Vcc od złącza zasilania, do pola lutowniczego nr 8 układu U1, anody diody LED1, złącza śrubowego tej diody i kondensatorów C1 i C2.

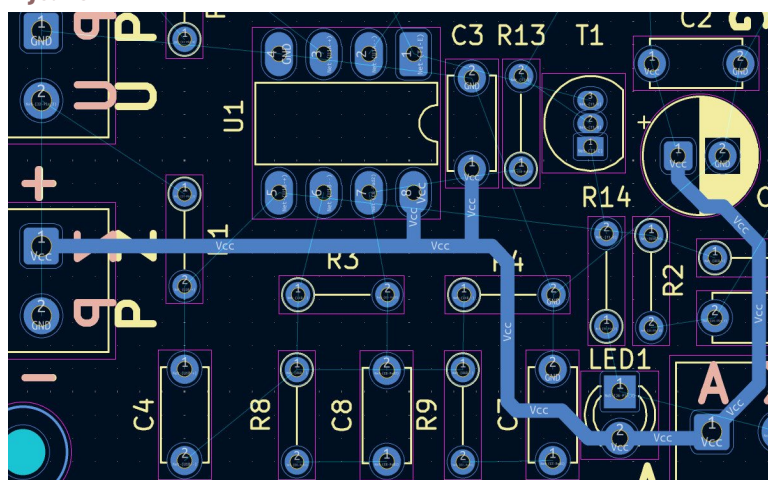
W tym celu klikamy ikonę trasowanie pojedynczej ścieżki, skrót klawiaturowy **X** i klikamy lewym klawiszem myszy na pole lutownicze Vcc złącza śrubowego. W tym momencie zostaną podświetlone pola lutownicze należące do tej sieci, a za kursorem myszki będzie prowadzona ścieżka. Na jej bokach widoczna jest szara obwódka – **rysunek 2**. W ten sposób wskazywany jest prześwit (odstęp) tej ścieżki. Ścieżkę tę prowadzimy w kierunku rezystora R4, a następnie kondensatora C7 pod tymi fotoprintami w kierunku anody LED1 i zacisku złącza śrubowego. Od złącza śrubowego prowadzimy ścieżkę dalej pod footprintami C6 i R7 do kondensatora C1. Na końcu do wytrasowanej ścieżki dołączamy pole lutownicze nr 8 układu U1. Trasowana ścieżka podąża za kursorem myszki załamując się pod kątem 45°. Gdyby zaszła konieczność ręcznego załamania ścieżki w określonym miejscu płytki, wówczas w tym miejscu należy kliknąć lewym klawiszem myszy. Podobnie w przypadku konieczności połączenia ścieżek należących do tej samej sieci, należy kliknąć lewym klawiszem myszy w miejscu ich łączenia. Wytrasowaną ścieżkę Vcc widzimy na **rysunku 3**. Gdy dwukrotnie klikniemy lewym klawiszem myszy na wytrasowanej ścieżce zobaczymy okno właściwości ścieżki z **rysunku 4**. Jak widać w czerwonej ramce, ścieżka ta ma szerokość 0,8 mm, tak jak to zdefiniowaliśmy w klasach sieci. Nie musimy więc pilnować szerokości ścieżek, ponieważ robi to za nas program. Przełączamy się na warstwę górną **F.Cu** i trasujemy kolejne ścieżki. Jeśli jakaś ścieżka nie układa się nam ładnie podczas trasowania przykładowo z punktu A do B, możemy ją trasować z punktu B do A. Drobna zmiana, a potrafi poprawić



Rysunek 1



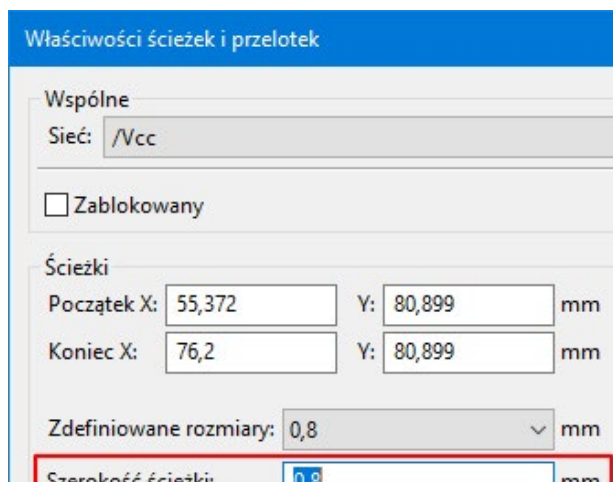
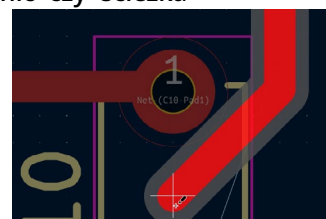
Rysunek 2



Rysunek 3

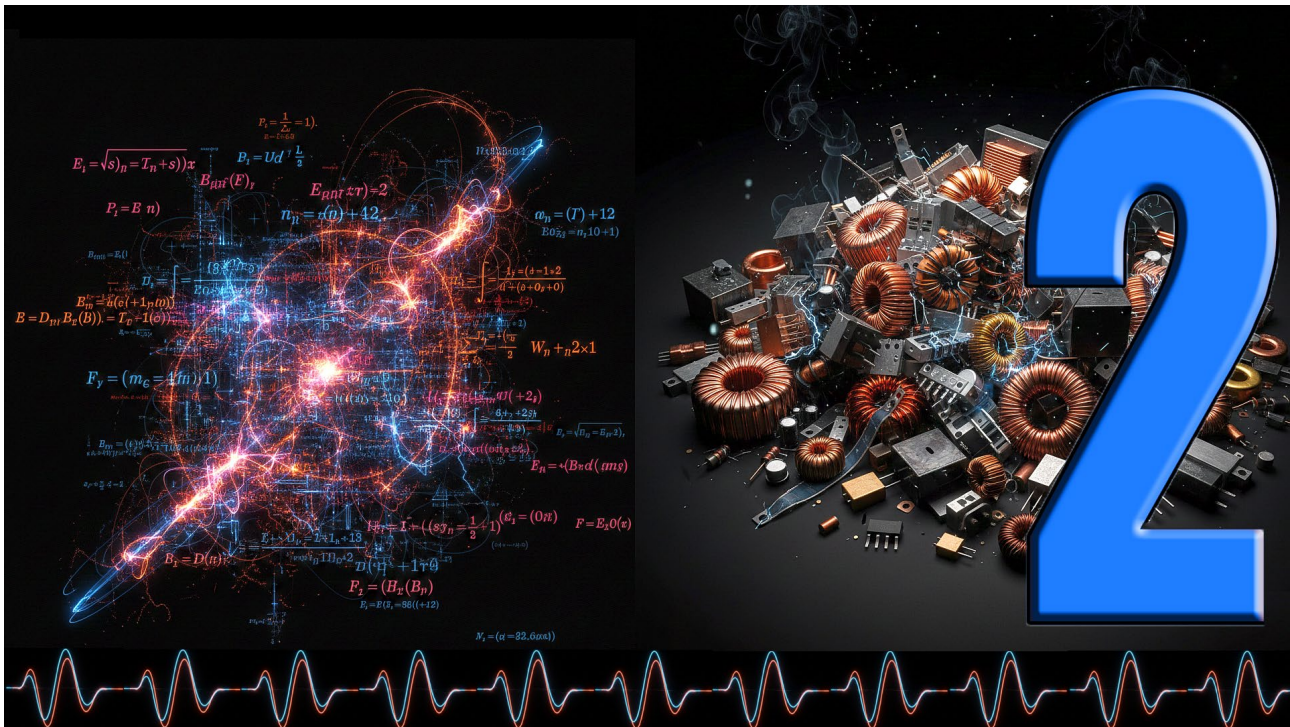
efekty trasowania. Spójrzmy jeszcze na **rysunek 5**, gdzie widzimy trasowaną ścieżkę i jej prześwit oznaczony szarym kolorem względem pola lutowniczego nr 1 kondensatora C10. Pole to również ma oznaczony swój prześwit przez okrąg w postaci cienkiej linii. W ten sposób łatwo możemy ocenić czy ścieżka zmieści się przy dużym zagęszczeniu elementów i ścieżek na płytce. Trasowanie rozpoczynamy od najkrótszych ścieżek, przechodząc do dłuższych.

Ważna uwaga! Nie trasujemy



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE. W pełnej wersji ten artykuł oczywiście ma więcej stron.

Rysunek 5



Opowiadanie o „nieznaczących” drucikach

Niniejszy minicykl stanowi niezbędne wprowadzenie, pozwalające ugruntować i odświeżyć wiedzę o polach oraz o indukcyjności. Zrozumienie tych fundamentów jest kluczowe, aby w pełni przyswoić treści zawarte w kolejnych częściach serii: „Moja konfrontacja z szybkimi sygnałami cyfrowymi”.

Natura świata potrafi zaskoczyć...
Dziwności nad dziwnościami, a to tylko drucik...

Kolejne niespodziewane zaskoczenie...
Elektrostatyka kontra drucik ogólnie

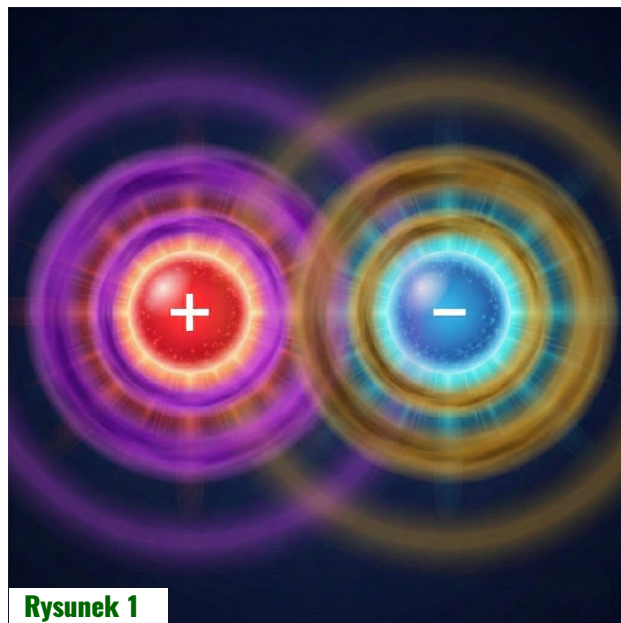
W poprzednim odcinku skupiliśmy się na omówieniu podstawowych pojęć oraz zjawisk towarzyszących przewodnikowi, przez który płynie prąd, ze szczególnym uwzględnieniem pola magnetycznego, które generuje się wokół niego. W tym wykładzie wejdziemy bardzo głęboko w samą strukturę miedzianego przewodnika, badając niezwykle zjawiska na styku fizyki klasycznej i mechaniki kwantowej. Scharakteryzujemy oddziaływanie elektromagnetyczne, którego pole jest nośnikiem energii w każdym obwodzie elektrycznym. Ponadto skupimy się na elektronach, bardzo dziwnych obiektach kwantowych, które nie dość, że są źródłem pola elektromagnetycznego w konkretnych okolicznościach,

to jeszcze odgrywają kluczową rolę w fizyce ciała stałego. Spróbujemy znaleźć odpowiednie analogie mechanizmów, odnosząc się do zjawisk w fizycznym realnym świecie, które dzieją się bezpośrednio pod maską zwykłego, niepozornego drucika w stanie pozornego spoczynku oraz w przypadku działania sił elektrostatycznych. Artykuł ten przygotowuje solidny grunt pod tematy związane z energią elektromagnetyczną w przypadku, gdy przez nasz tytułowy drucik płynie prąd elektryczny, jako „uporządkowany ruch elektronów”. Nie bez powodu zamknąłem tę definicję w cudzysłów, kolejny artykuł otworzy nam wrota do spojrzenia na to wszystko z zupełnie innej perspektywy...

Natura świata potrafi zaskoczyć...

Bezpośrednim źródłem pól elektromagnetycznych są cząstki obdarzone ładunkiem elektrycznym. *Lecz czym jest ładunek?* Tu znów muszę użyć fundamentalnego stwierdzenia, że jest to jedna z podstawowych cech materii, wrodzona jej pewna właściwość. Nie wiemy z czego składa się ładunek, ale wiemy, że obdarzone nim cząstki elementarne wytwarzają ogromne pola siłowe (**rysunek 1**) oraz mają zdolność reakcji na te pola. Ponadto ładunki mogą być dodatnie lub ujemne, a tutaj intuicja słusznie podpowiada, że mają zdolność do przyciągania lub odpychania się nawzajem. To tajemnicze pole siłowe, które ma realny wpływ na cząstki obdarzone ładunkiem, to jedno z fundamentalnych oddziaływań w naszym wszechświecie, czyli **oddziaływanie elektromagnetyczne**. Oddziaływanie to jest kluczowe dla istnienia naszego wszechświata – nie tylko trzyma wszystkie atomy w ryzach, ale też jest odpowiedzialne za wszelkie wiązania chemiczne. Jedną z jego oczywistych cech, z której nie zdajemy sobie do końca sprawy, jest zdolność „uszytniania materii”. To między innymi dzięki tej właściwości jestem w stanie pisać ten tekst, a moje palce nie przenikają przez klawisze klawiatury (**rysunek 2**). Co więcej ja ich nawet bezpośrednio nie dotykam – to elektrony w atomach moich palców z elektronami w atomach plastikowych klawiszy tak na siebie oddziałują – odpychają się, co daje złudzenie bezpośredniego dotyku...

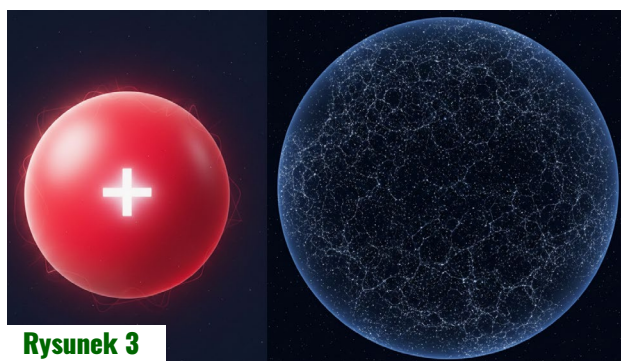
Oddziaływanie elektromagnetyczne to potężna, wręcz gigantyczna siła. Porównując to oddziaływanie do oddziaływania grawitacyjnego dla dwóch protonów, musimy zdać sobie sprawę, że różnica w siłach między nimi jest **kolosalna, aż 10^{36} razy większa na korzyść elektromagnetyzmu**. A jaka różnica sił będzie dla dwóch elektronów? Aby przekonać się jak potężny jest ładunek elektronu, ustawmy je obok siebie. Elektrony będą się odpychać, bo mamy do czynienia z jednoimiennymi ładunkami, ale działa na nie również grawitacja, ze względu na to, że posiadają masę, a siła ta zawsze dąży do przyciągania się obiektów. Obliczenia wskazują, że oddziaływanie elektromagnetyczne między nimi jest **około $4,17 \times 10^{42}$ razy silniejsze niż oddziaływanie grawitacyjne występujące między nimi!!!** Aby uzmysłowić sobie tę różnicę w siłach, wyobraźmy sobie średnicę protonu (siła grawitacji) i średnicę obserwowalnego wszechświata (siła elektromagnetyzmu). To zestawienie najlepiej obrazuje skalę zjawiska, z którym mamy do czynienia (**rysunek 3**). Nic dodać, nic ująć...



Rysunek 1



Rysunek 2



Rysunek 3

przez takie cząstki. W zależności od układu odniesienia i stanu ruchu, manifestują się one **jako pole elektryczne, magnetyczne lub ich nierozzerwalna jedność – pole elektromagnetyczne**. Odniesmy się do elektronów ze względu na to, że

**Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.
W pełnej wersji ten artykuł oczywiście ma więcej stron.**

ZROZUMIEĆ **ELEKTRONIKĘ** z Piotrem Góreckim

ZE 4/2026

piotr-gorecki.pl



Wydawca: Zrozumieć Elektronikę sp. z o.o. ul. Nadarzyn 23A 05-230 Kobyłka

Redaktor Naczelny: Piotr Górecki

e-mail: kontakt@piotr-gorecki.pl

Redakcja techniczna: Ewa Górecka-Dudzik (ewa@piotr-gorecki.pl)

**Stali współpracownicy: Andrzej Pawluczuk, Tadeusz Suszał, Karol Świerc,
Mateusz Ostrycharz, Paweł Pawłowicz, Rafał Kozik, Szymon Burian, Jacek Kosecki**

**Inicjatywa Zrozumieć Elektronikę realizowana jest
dzięki wsparciu Patronów i Mecenasów poprzez
konto autorskie Patronite: <https://patronite.pl/Zrozumiec-Elektronike>**

Uwaga! Ani autorzy artykułów, ani wydawca nie biorą odpowiedzialności za ewentualne szkody spowodowane wynikiem eksperymentów inspirowanych treścią czasopisma i strony internetowej.

Osoby, które chciałyby przeprowadzić eksperymenty związane z treścią artykułów powinny mieć odpowiednie kwalifikacje BHP dotyczące elektryczności oraz świadomość ryzyka.

Osoby niepełnoletnie i niedoświadczone mogą przeprowadzić takie działania jedynie pod opieką wykwalifikowanych opiekunów, np. nauczycieli.

Projekty przedstawiane w czasopiśmie mogą być wykorzystane jedynie do własnych potrzeb, a ich wykorzystanie do innych celów, zwłaszcza zarobkowych, wymaga zgody Autora.

Wszystkie materiały zamieszczane w czasopiśmie są własnością ich twórców, więc przedruk czy umieszczenie na stronach internetowych wymaga pisemnej zgody Autora.