

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>

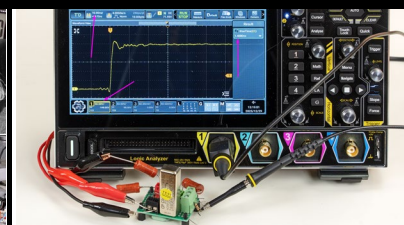
5/2026 Maj (41)

piotr-gorecki.pl



Rezystancja to nie (tylko) rezystor!

- Wysokonapięciowy generator impulsów • Przełączniki ze stykami zwilżanymi rtęcią
- Aktywne obciążenie 40 A • Generator liczb losowych • Dopasowanie szumowe – a co to właściwie jest?
 - Realne parametry rezystorów dużej mocy • Jak działa smartfon? 4G – telefonia, której nie ma!
- Lampowe konfiguracje SRPP i mu follower • Fascynujące przemiany energii: Nietypowa fotowoltaika



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



Rezystancja to nie (tylko) rezystor!

W poprzednim artykule serii zaczęliśmy omawianie elementów elektronicznych poczynając od rezystora. Ale w cyklu o fundamentach elektroniki przewodowej nie może zabraknąć odpowiedniej podbudowy. Dlatego w tym artykule omawiam ogromnie ważną kwestię właściwego rozumienia sensu rezystancji.

Rezystancja „uogólniona”

W poprzednich artykułach przedstawiając niewzruszone fundamenty elektroniki omawiałem pojęcie napięcia, prądu oraz oporności – rezystancji. W poprzednim artykule serii wyjaśniłem, dlaczego rezystor jest „przeklętym elementem elektronicznym”, którego należy w miarę możliwości unikać oraz dlaczego, wbrew powszechnej opinii, rezystory wcale nie są najpopularniejszymi elementami elektronicznymi. A w poniższym artykule spróbujemy „odkręcić” niektóre błędne poglądy na temat rezystancji i rezystorów. Otóż większości z nas oporność, czyli rezystancja bardzo mocno kojarzy się z opornikiem, czyli rezystorem (przykłady na **fotografii tytułowej**).

Rezystancja „materiałowa”

W zasadzie słusznie, jednak dziś **wiele osób zdecydowanie zbyt mocno kojarzy rezystancję z rezystorem i z prawem Ohma**. Znam ludzi, którzy mówią, że do układu trzeba wlutować rezystancję, a chodzi im o wlutowanie rezystora. A **rezystancja to zdecydowanie nie jest to samo, co rezystor!** To prawda, że każdy opornik, czyli rezystor, ma jakąś rezystancję, wyrażoną w omach. Ale to jest tylko wierzchołek góry lodowej. Mówiąc o oporności (rezystancji) opornika (rezystora) mamy na uwadze coś, co można nazwać **rezystancją materiałową**. A rezystancja to pojęcie zdecydowanie szersze. Rezystancja to nie tylko rezystor! I są bardzo różne „odmiany” rezystancji!

☰ Rezystancja [edytuj wstęp]

🌐 85 języków ▾

Rezystancja, **opór (elektryczny, czynny)**, **oporność**^[1] (**czynna**) (z łac. *resistere* — sprzeciwiać się, stawiać opór^[2]) – wielkość charakteryzująca relację między **napięciem** a **natężeniem prądu elektrycznego** w obwodach **prądu stałego**. W obwodach **prądu przemiennego** rezystancją nazywa się **część rzeczywistą impedancji zespolonej**^[3].

Zwyczajowo rezystancję oznacza się symbolem *R*. Jednostką rezystancji w **układzie SI** jest **om**, którego symbolem jest Ω .

Rezystancja zdefiniowana jest wzorem: $R = \frac{U}{I}$, gdzie: *R* – opór przewodnika elektrycznego, *U* – **napięcie** między końcami **przewodnika**, *I* – **natężenie prądu elektrycznego**^[4].

Rysunek 1

Prawidłowo, bardzo dobrze zdefiniowana jest ona w polskiej Wikipedii – **rysunek 1: Rezystancja, opór, oporność to wielkość charakteryzująca relację między napięciem i prądem elektrycznym**. Pisałem o tym **w pierwszym artykule tej serii**. Nieprzypadkowo też wielokrotnie podkreślałem, że naprawdę **warto rozumieć oporność (rezystancję) jako „okoliczności przekazywania energii”**.

Najprościej biorąc: zawsze, gdy mamy do czynienia z napięciem i prądem w obwodzie, możemy mówić o oporności, rozumianej jako stosunek napięcia i prądu. W najprostszym przypadku o rezystancji. O rezystancji jako stosunku, relacji między napięciem i prądem zawsze. Obecność napięcia i prądu nie zawsze świadczy o obecności opornika (rezystora), czyli o elementu elektronicznego.

Oporność to zawsze stosunek wartości napięcia i prądu ***U / I***, ale jest mnóstwo różnych rodzajów oporności, nie tylko rezystancja przy prądzie stałym.

I tu dochodzimy do kilku ważnych i obszernych zagadnień, które szczegółowo omówię w tym i w następnych artykułach, a na razie tylko zasygnalizuję dwie kwestie, w tym sprawę rozróżnienia „rezystancji materiałowej” od „rezystancji ogólnej”.

Rezystancja „uogólniona”

Podkreślam, że nawet wielu dobrych elektroników zdecydowanie zbyt mocno kojarzy oporność, rezystancję z rezystorem, czyli z elementem elektronicznym. A to błąd, utrudniający zrozumienie trudniejszych zagadnień.

Zapamiętaj raz na zawsze: **jeżeli w jakimś obwodzie jednocześnie mamy do czynienia i z napięciem *U*, i z prądem *I*, to zawsze związany jest z tym przepływ lub przemiana energii w tempie określonym przez wzór $P = U \times I$** .

I co najważniejsze w tym artykule: wtedy **zawsze możemy też mówić o jakiejś oporności, w najprostszym przypadku o rezystancji *R*, wyrażonej wzorem $R = U / I$** . Szeroko pojęta oporność to „okoliczności przekazywania energii”. I teraz mówimy o oporności czy rezystancji „uogólnionej”, która dotyczy

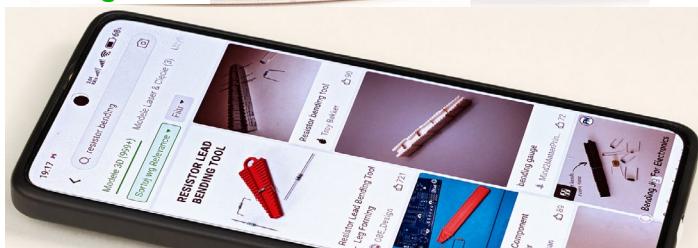
Prostym przykładem jest bateria z dołączoną żarówką. Na pewno bateria nie jest obciążona rezystorem, ale niewątpliwie jest obciążona rezystancją – (nieliniową) rezystancją żarówki.

Dużo lepszym przykładem jest bateria pracująca w kalkulatorze – **fotografia 2**. Kalkulator na pewno nie jest rezystorem i najprawdopodobniej nie zawiera ani jednego rezystora. Jednak zasilająca go bateria daje jakieś napięcie *U* i płynie jakiś prąd *I*, więc bateria ta niewątpliwie jest obciążona jakąś **rezystancją $R = U / I$** . W tym przypadku, podczas pracy mikroprocesora i wyświetlacza cała energia elektryczna z baterii zostaje finalnie zamieniona na ciepło (czyli podobnie jak w rezystorze).

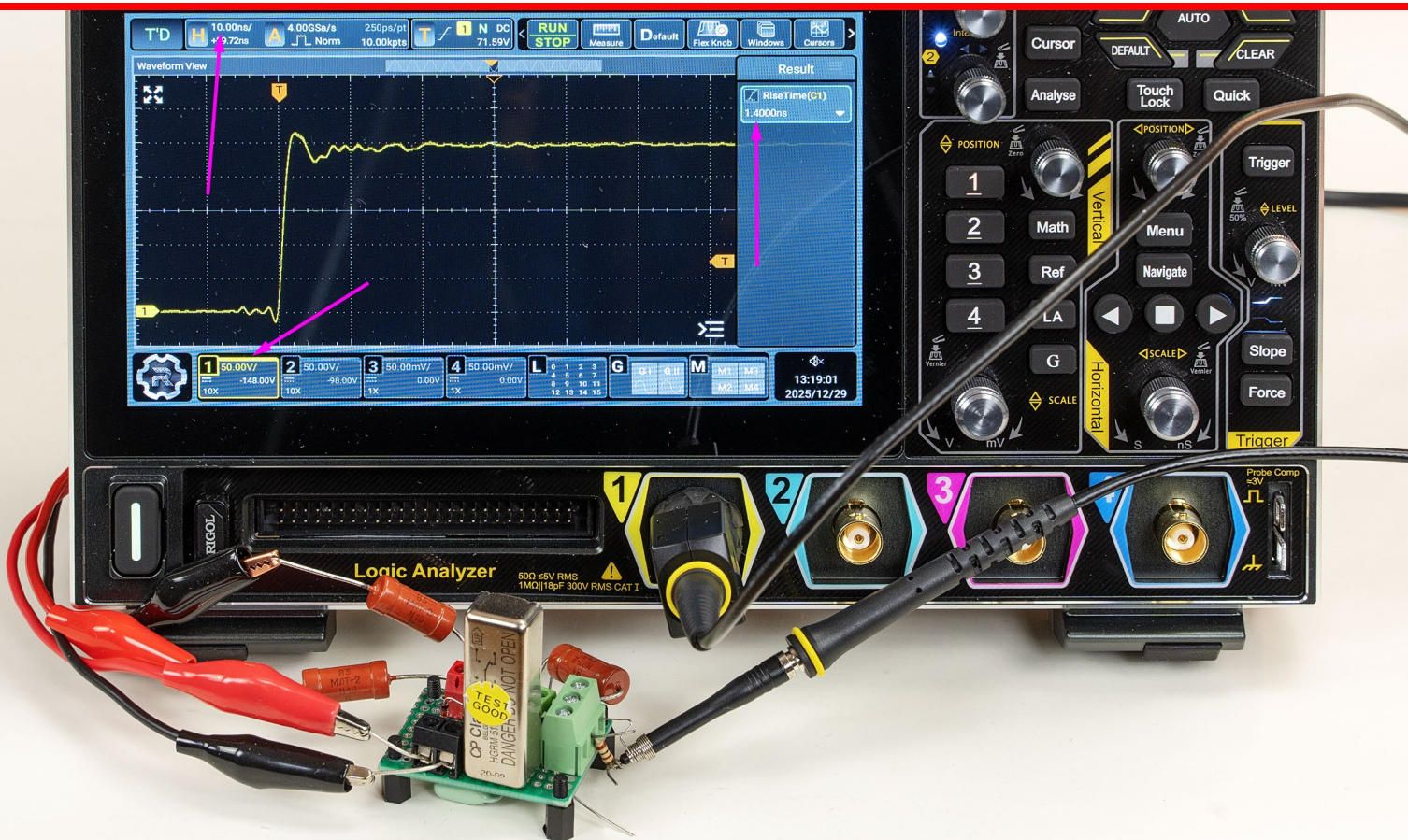
Jeszcze innym przykładem jest smartfon – **fotografia 3**. Gdy z niego korzystamy, energia elektryczna zawarta w akumulatorze jest zamieniana na różne inne rodzaje energii. Częściowo na energię światła, którym świeci ekran (podświetlany LCD lub świecący OLED). Częściowo na energię mechaniczną: akustyczną energię dźwięku z głośnika i energię wibracji.



Fotografia 2



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



Wysokonapięciowy generator impulsów

Artykuł opisuje zaskakująco prosty generator impulsów o dużej amplitudzie, do 500 V, a nawet jeszcze więcej, i o bardzo stromych zboczach. Takie generatory silnych impulsów o stromych zboczach bywają potrzebne do różnych celów, między innymi do sprawdzania i kompensacji dzielników i sond pomiarowych.

Rozważania projektowe
Opis układu

Szczegóły dotyczące pomiarów

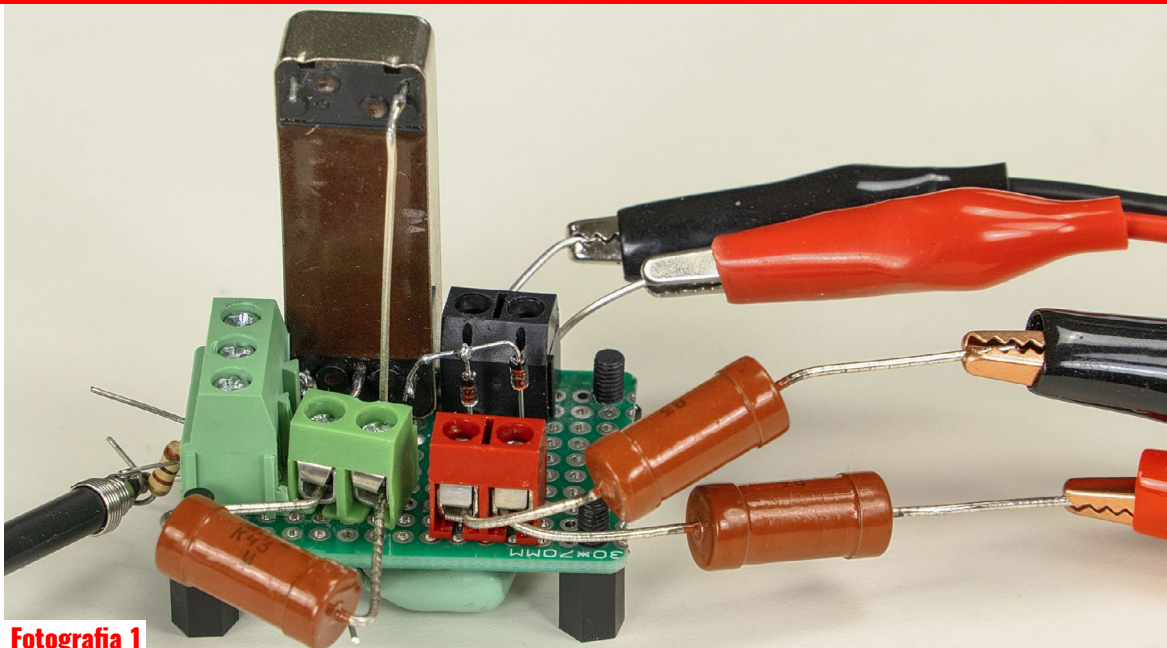
Budowę generatora impulsów o dużej amplitudzie i stromych zboczach planowałem od dawna. Ostatnio pojawiły się okoliczności, które pozwoliły doprowadzić zamierzenie do końca. Jedną to fakt, że kupiłem i doprowadziłem do pełnej sprawności zasilacz IZ55-71, którego napięcie wyjściowe można skokowo regulować w zakresie 0,1 V do 500,0 V. Drugą to artykuł, dotyczący wysokoomowych dzielników napięcia stałego. Otóż aż się prosi, żeby „pociągnąć” temat dzielników i rozszerzyć go także na napięcia zmienne. Wiadomo jednak, że dzielniki, czyli tłumiki napięć zmiennych muszą być skompensowane, by współczynnik tłumienia był jednokowy w jak najszerszym paśmie częstotliwości.

Problem w tym, że jeśli mają to być dzielniki 100000:1 lub nawet 10000:1, to do ich testowania potrzebne są napięcia zmienne o dużej amplitudzie, nawet rzędu kilkuset woltów. Podstawowym przebiegiem jest sinusoida, ale jest ogromny problem, żeby zbudować przestrajany generator sinusoidalny o tak dużym napięciu wyjściowym, i jednocześnie o dużej częstotliwości i liniowości. W warunkach amatorskich jest to praktycznie niemożliwe. Jedyną szansą jest wykorzystanie impulsów prostokątnych o dużej amplitudzie i jak najbardziej stromych zboczach. Realizacja wysokonapięciowego generatora impulsowego też nie jest łatwa, ale realna nawet w warunkach amatorskich.

Fotografia 1

przedstawia mocno nietypowy wysoko-napięciowy generator impulsów. Zbudowany jest na kawałku pytki uniwersalnej, zawiera dosłownie kilka elementów, a jeżeli chodzi o półprzewodniki, to widać tylko dwie diody 1N4148, które impulsów nie wytwarzają.

Na fotografii tytułowej widać rosnące zboczne impulsu.



Fotografia 1

Warto zauważyć, że skala napięcia to 50 woltów na działkę, czyli że widzimy impuls 250-woltowy! Z kolei skala czasu to 10 nanosekund na działkę i oscyloskop pokazuje czas narastania 1,4 nanosekundy!

Impulsy o tak zadziwiających parametrach wytwarza niepozorny stary przełącznik, pokazany na **fotografii 2**. Nie jest to jednak zwyczajny przełącznik kontaktowy, tylko specyficzny **przełącznik ze stykami zwilżanymi rtęcią**. Stąd literki Hg w oznaczeniu i ostrzeżenie, żeby nie rozbierać obudowy. Stąd też strzałka z napisem UP (do góry) i właśnie dlatego przełącznik w moim generatore nietypowo zamontowany jest w pionie.

Bardzo pozytywnym zaskoczeniem jest zmierzona stromość zboczy – jest wręcz rewelacyjna. Na fotografii tytułowej widać zmierzony czas narastania – to 1,4 ns. Ale jak widać na **fotografii 3**, po rozciągnięciu skali czasu (2 ns/dz), zmierzony czas narastania 250-woltowego impulsu takiego prymitywnego generatora wynosi 1,28 ns.

Niewiarygodne, bo zgodnie z popularnym wzorem $f = 0,35 / t_r$, pasmo przekracza 270 megaherców, a takie parametry impulsu pozwalają sprawdzić i przeprowadzić kompensację częstotliwościową dzielników napięcia – sond w paśmie grubo ponad 100 megaherców! Nawet



Fotografia 2



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





Przełączniki ze stykami zwilżanymi rtęcią

Przełączniki są coraz częściej zastępowane odpowiednikami półprzewodnikowymi. Ponadto zabronione jest wprowadzanie na rynek urządzeń powszechnego użytku, zawierających znaczące ilości rtęci. Dlatego więc obszerny artykuł poświęcam kontaktom rtęciowym, które coraz trudniej znaleźć na rynku wtórnym?

Styki kontaktronowe suche i mokre

Przechylne wyłączniki rtęciowe

Klasyczne kontaktrony rtęciowe

Rtęciowe kontaktrony „wszechpozycyjne”

Stare wersje z podstawkami

Przełączniki rtęciowe – aspekty techniczne

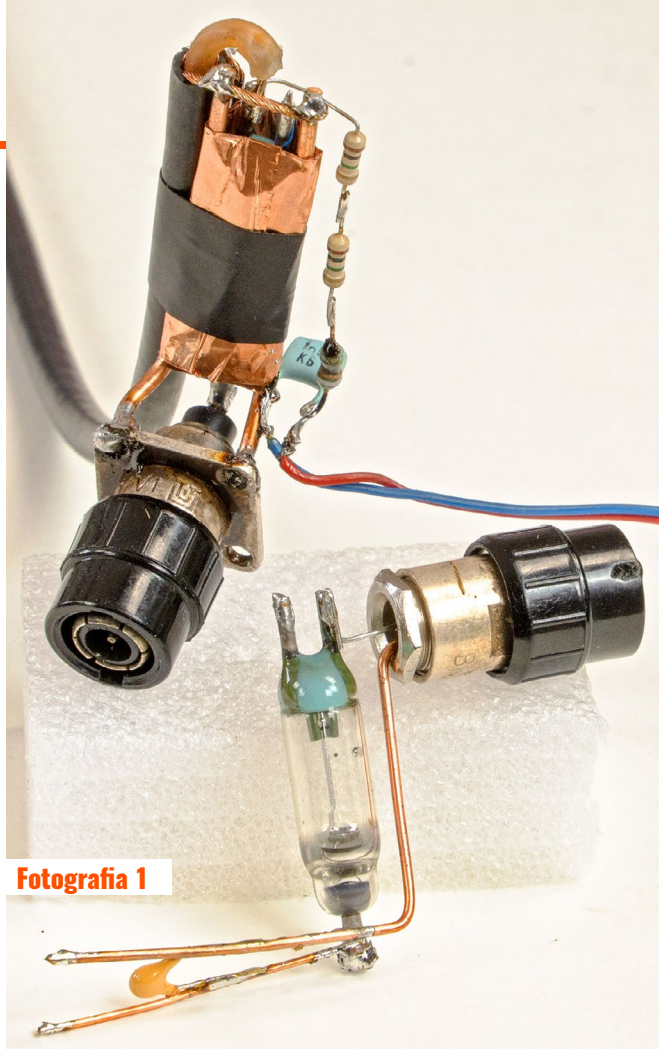
Przełączniki rtęciowe – inne aspekty

Najprościej, technicznie biorąc, **kontaktrony rtęciowe są wielokrotnie lepsze od „zwykłych” kontaktronów, a także od wielu innych przełączników.**

Nieco szersza odpowiedź, dlaczego warto wiedzieć o kontaktronach rtęciowych jest taka: **mogą one przełączać w obwodach o większym napięciu (500...1000 V) i mocy, mają rewelacyjnie dużą trwałość – setki, a nawet tysiące razy lepszą od klasycznych przełączników oraz wykazują znakomitą stałość rezystancji styku w czasie użytkowania.** Styki nie ulegają wypalaniu, a **co ważne – podczas przełączania nie występują drgania styków.** Przełączniki rtęciowe mogą przełączać małe i bardzo małe sygnały. Z uwagi na ogromną trwałość, mogą

niezawodnie pracować tam, gdzie następuje bardzo częste przełączanie. Ponieważ nie ma w nich drgań styków, od dawna **są wykorzystywane do wytwarzania impulsów o bardzo stromych zboczach.**

Problem w tym, że takie kontaktrony zawierają rtęć, której nie wolno wykorzystywać w popularnym sprzęcie. Obecnie większość wytwórców zaprzestała ich produkcji. Niemniej nadal są dostępne nie tylko na aukcjach, ale też w sklepach, które wyprzedają dawne zapasy (są też nadal produkowane, np. w Chinach). Dobra wiadomość jest taka, że w chwili pisania tego artykułu w dwóch znanych polskich sklepach internetowych można kupić kilka odmian kontaktronów rtęciowych, i to w atrakcyjnych cenach.



Fotografia 1

Przełączniki ze stykami zwilżanymi rtęcią są pokazane na fotografii tytułowej. Na **fotografii 1** widoczne są układy ze stykami rtęciowymi, służące do wytwarzania króciutkich, nanosekundowych impulsów, których zbocza mają czasy narastania poniżej pół nanosekundy. Wykorzystane są rurki kontaktronowe, które muszą „pracować w pionie”. Sterowanie odbywa się „ręcznie”, przez zbliżenie magnesu trwałego.

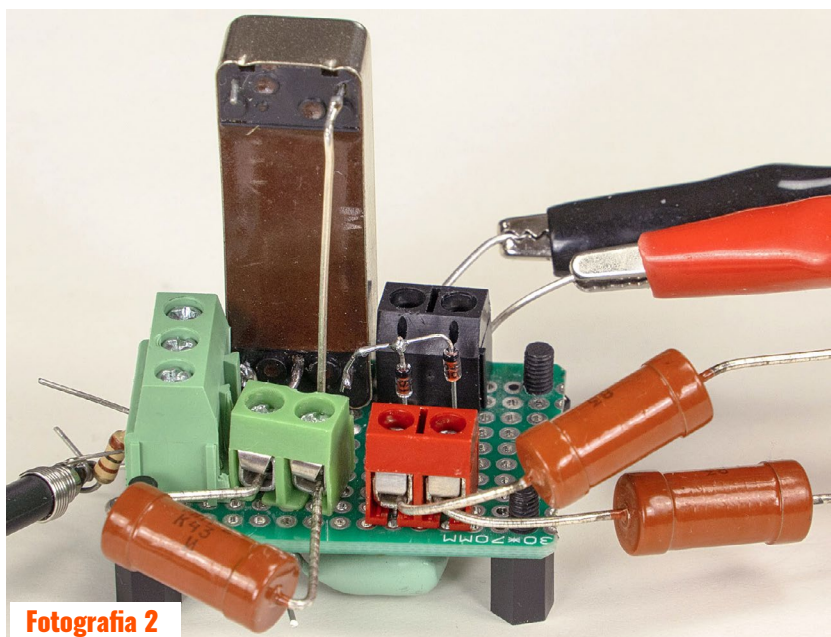
Fotografia 2 przedstawia elektromechaniczny generator impulsów o amplitudzie do 500 woltów i o bardzo stromych zboczach, który jest projektem okładowym tego numeru ZE.

Natomiast na **fotografii 3** pokazany jest uniwersalny tester przełączników, który będzie opisany w jednym z najbliższych numerów ZE.

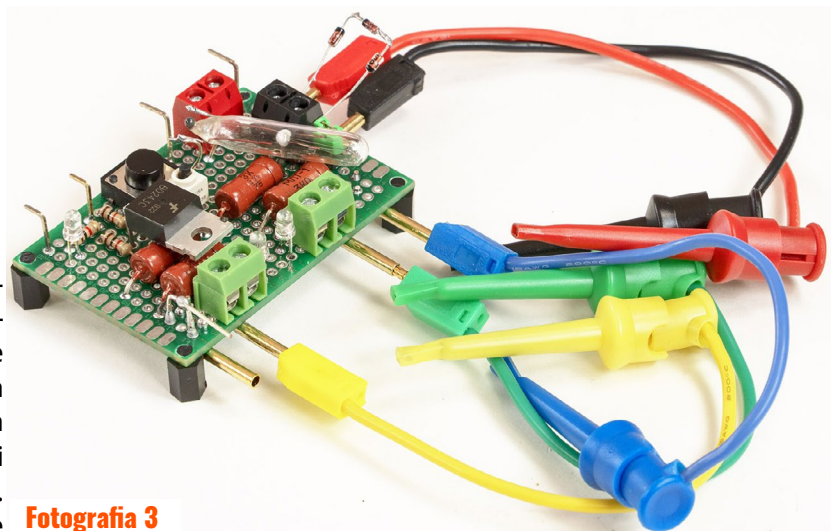
Styki kontaktronowe suche i mokre

Fotografia tytułowa prezentuje rtęciowe przełączniki kontaktronowe, które po angielsku nazywane są **mercury relays**, **mercury wetted relays**, częściej **wet reed relays**. *Wet*, czyli *mokre*, w przeciwieństwie do znacznie popularniejszych *suchych* (*dry reed relays*).

W języku polskim nie mówimy o przełącznikach trzcinowych ani o łącznikach trzcinowych (*reed*), tylko właśnie o **kontaktronach**. Ich cechą charakte-



Fotografia 2



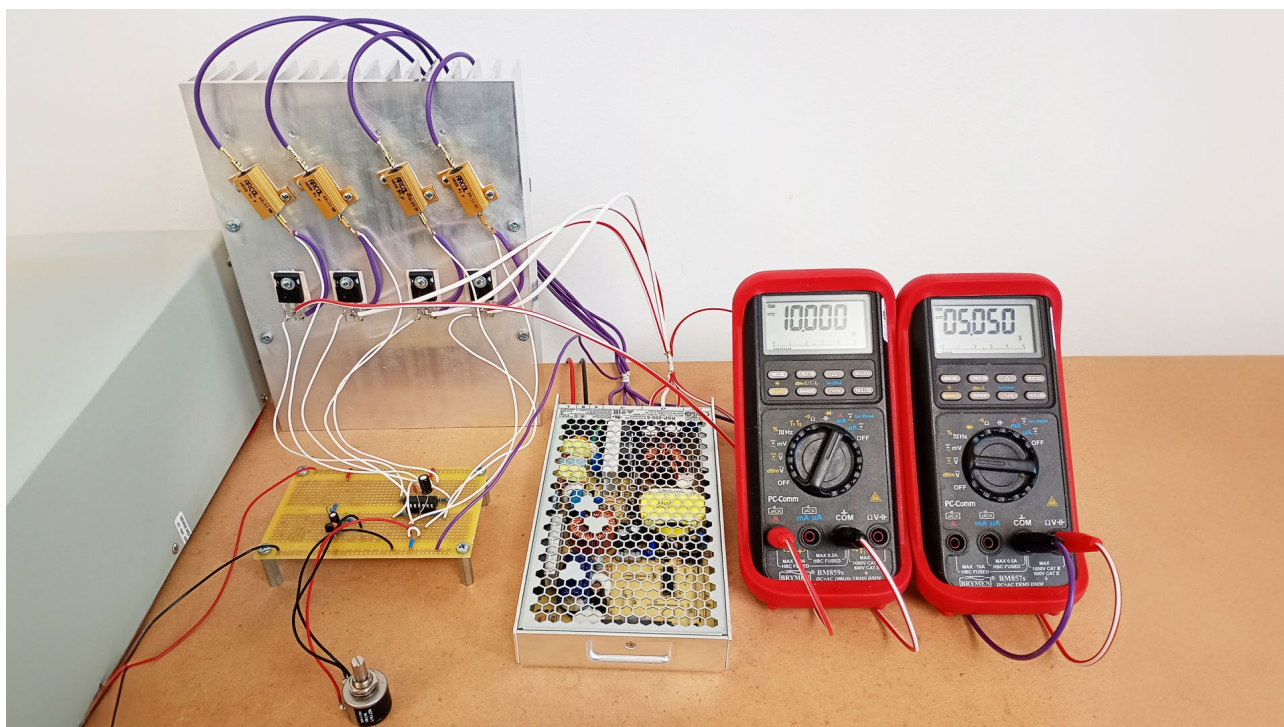
Fotografia 3

Na **fotografii 4** widać popularne rurki kontaktronowe klasyczne, czyli suche (*dry*). Wewnątrz są metalowe elastyczne pręciki z materiału magnetycznego. W obecności pola magnetycznego styk zmienia stan. Najpopularniejsze są styki typu A, czyli zwierne, mniej popularne są styki typu C – przełączne. Rzadko spotykane są styki B – rozwierane przez pole magnetyczne, niekiedy z użyciem wbudowanego dodatkowego magnesu trwałego.

Przewodzące pręciki zawarte w rurkach kontaktronowych muszą być wykonane z materiału ferromagnetycznego (zwykle jest to żelazonikiel), który ma niekorzystne niektóre istotne tu właściwości.



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



Aktywne obciążenie 40 A

W artykule zaprezentowano prosty układ sztucznego obciążenia przeznaczonego do sprawdzenia stałonapięciowego zasilacza o mocy wyjściowej 200 W (5 V / 40 A). Pomogło ono w zweryfikowaniu deklarowanych przez producenta parametrów oraz zdjęciu jego charakterystyk.

Wstępne przymiarki
MOSFET-y równolegle
Wersja finalna

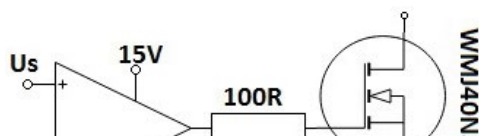
Opis układu
Charakterystyki zasilacza

Aktywne obciążenie to urządzenie warsztatowe przydatne w każdym warsztacie elektronika. Jego zadaniem jest obciążanie badanego układu zadany prądem. Zaletą obciążenia aktywnego w stosunku do obciążenia rezystorowego jest możliwość płynnej regulacji prądu oraz jego niezależność od napięcia testowanego układu. Jest niezbędne podczas testów zasilaczy firmowych oraz układów własnej konstrukcji. Dzięki niemu można sprawdzić zabezpieczenia prądowe oraz termiczne. Przyda się również do badania przetwornic, baterii oraz akumulatorów.

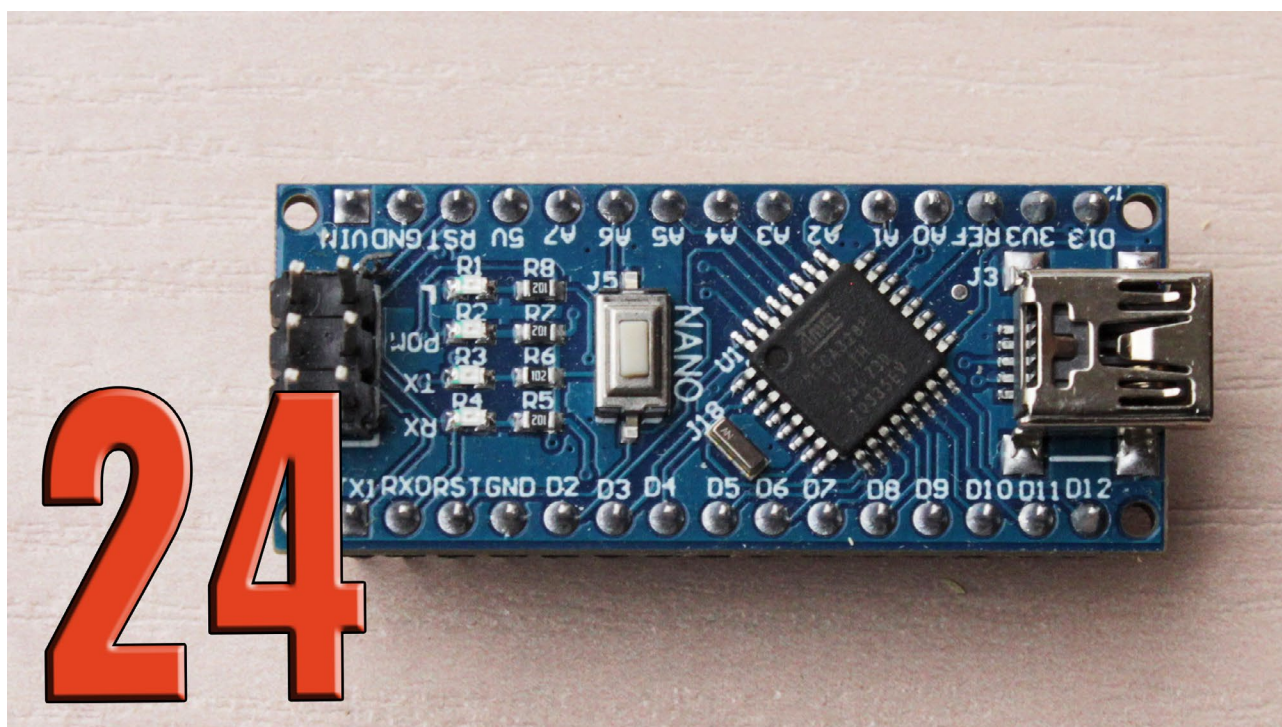
Wstępne przymiarki

trach wyjścia 5 V / 40 A, co oznacza, że może oddać 200 W mocy. Ponieważ mamy do czynienia z zasilaczem stałonapięciowym, do jego wyjścia musimy podłączyć regulowane źródło prądowe będące w stanie rozproszyć wspomnianą moc.

Pierwotnie obciążenie postanowiłem wykonać na źródle prądowym zbudowanym na jednym tranzystorze MOSFET WMJ40N50D1 o parametrach: $I_D = 40 \text{ A}$, $P = 416 \text{ W}$, $R_{thjc} = 0,3^\circ\text{C/W}$, $T_j = 150^\circ\text{C}$, według schematu przedstawionego na **rysunku 1**.



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



Mikroprocesorowa ośła łączka, część 24

Rosnące zapotrzebowanie na elastyczność rozwiązań w programowaniu komputerów doprowadziło do powstania nowego języka określanego jako C++. W gruncie rzeczy jest to nowy i zarazem stary język programowania, gdyż wywodzi się z klasycznego języka C.

[Pojęcie klasy](#)

[Obiekt](#)

[Przykład klasy](#)

[Implementacja metod klasy](#)

[Utworzenie obiektu](#)

[Przykład programu](#)

Zadaniem prowadzonego cyklu artykułów poświęconych programowaniu w języku C, jest przedstawienie jego składni, opisanie posługiwania się podstawowym narzędziem jakim jest środowisko Atmel Studio oraz konstruowanie algorytmów. Ta część cyklu stanowi zamknięcie dotychczasowych rozważań i przejście niejako na „wyższy poziom”. Z tym wiąże się nowe możliwości, gdyż język C++ udostępnia elementy, które nie istnieją w klasycznym C. Z innej perspektywy można to potraktować jako wstęp do programowania w środowisku Arduino (jest to jedynie środowisko narzędziowe, podobnie jak Atmel Studio), gdyż Arduino posługuje się językiem C++.

Nowy język programowania określany jako C++ jest ulepszonym wariantem C. Te udoskonalenia są

„kompatybilne wstecz”, czyli wszystkie zapisy z klasycznego C są w pełni rozumiane przez kompilator C++, program napisany w języku C może być przetwarzany przez kompilator C++. I tu może się zrodzić pytanie: czym te języki się różnią?

Język C++ wprowadza nowe pojęcie klasy oraz obiektu. Opisywanie nowych elementów „na sucho” jest mało interesujące, toteż nowe, istotne elementy języka C++ będą ilustrowane obsługą 8-cyfrowego wyświetlacza z wykorzystaniem układu MAX7219. Taki wyświetlacz wraz z jego obsługą był szczegółowo opisany w artykule *Mikroprocesorowa ośła łączka, część 9* [ZE 2/2025]. Ten wybór stwarza możliwość porównania wariantu dla klasycznego C oraz C++.

Pojęcie klasy

Klasa to specyficzny typ danych łączący dane oraz funkcje na nich operujące (te informacje obecnie mogą wyglądać tajemniczo, ale w kolejnych krokach powinny stopniowo się wyjaśnić). Szukając analogii w języku C, to klasa jest podobna do struktury: składa się z pól i na tym podobieństwo się kończy. Podobnie jak każdy typ strukturalny (*typedef struct*) wymaga własnej definicji, tu również programista musi zdefiniować klasę. Ogólną postać klasy pokazuje **rysunek 1**. Wystąpiły tu cztery nowe słowa kluczowe: *class*, *public*, *protected* i *private* (nie mogą być używane w innym znaczeniu niż w definicji klasy). Słowo *class* rozpoczyna definicję klasy (tak jak *typedef struct*). W obrębie pól klasy występują trzy rodzaje dostępu do nich, jako *public* (składowe klasy dostępne dla wszystkich), *private* (składowe klasy są dostępne jedynie dla funkcji należących do klasy) oraz *protected* (składowe klasy są dostępne dla samej klasy oraz jej klas pochodnych, dziedziczących). Dziedziczenie w C++ to mechanizm programowania obiektowego pozwalający klasie pochodnej (dziecko) przejmować pola i metody klasy bazowej (rodzic). Kolejność specyfikacji *public*, *private*, *protected* jest dowolna i może wystąpić wielokrotnie.

Występujące na rysunku 1 elementy klasy „*definicja elementów klasy w kategorii*” to pola (dane składowe, zmienne) przechowujące stan obiektu oraz metody (funkcje składowe) określające jego zachowanie.

W każdej definicji klasy musi wystąpić jedna metoda (funkcja wchodząca w skład klasy) o takiej samej nazwie jak sama klasa. Jest ona określana jako konstruktor obiektu, czyli funkcja, której zadaniem jest zainicjowanie jej stanu początkowego. W stosunku do każdej innej metody ta funkcja nie zwraca żadnego wyniku (nawet nie ma słowa kluczowego *void* na oznaczenie, że funkcja nie zwraca wyniku). Obok konstruktora w definicji klasy może wystąpić (nie jest obowiązkowy) destruktor jako funkcja, której zadaniem jest „posprzątać” po działaniu klasy, destruktor nie ma parametrów wywołania. Podobnie jak konstruktor, ta metoda również ma tę samą nazwę co klasa i dla odróżnienia od konstruktora jej nazwa jest poprzedzona znakiem tyldy „~”. Konstruktor jak i destruktor musi być dostępny „z zewnątrz”, czyli umieszczony w sekcji *public*.

Obiekt

Definicja klasy jest jedynie zgłoszeniem identyfikatora określającego odpowiednią strukturę da-

```
class <nazwa klasy>
{
    public:
        <definicja elementów klasy w kategorii public>
        ...
        <definicja elementów klasy w kategorii public>
    protected:
        <definicja elementów klasy w kategorii protected>
        ...
        <definicja elementów klasy w kategorii protected>
    private:
        <definicja elementów klasy w kategorii private>
        ...
        <definicja elementów klasy w kategorii private>
};
```

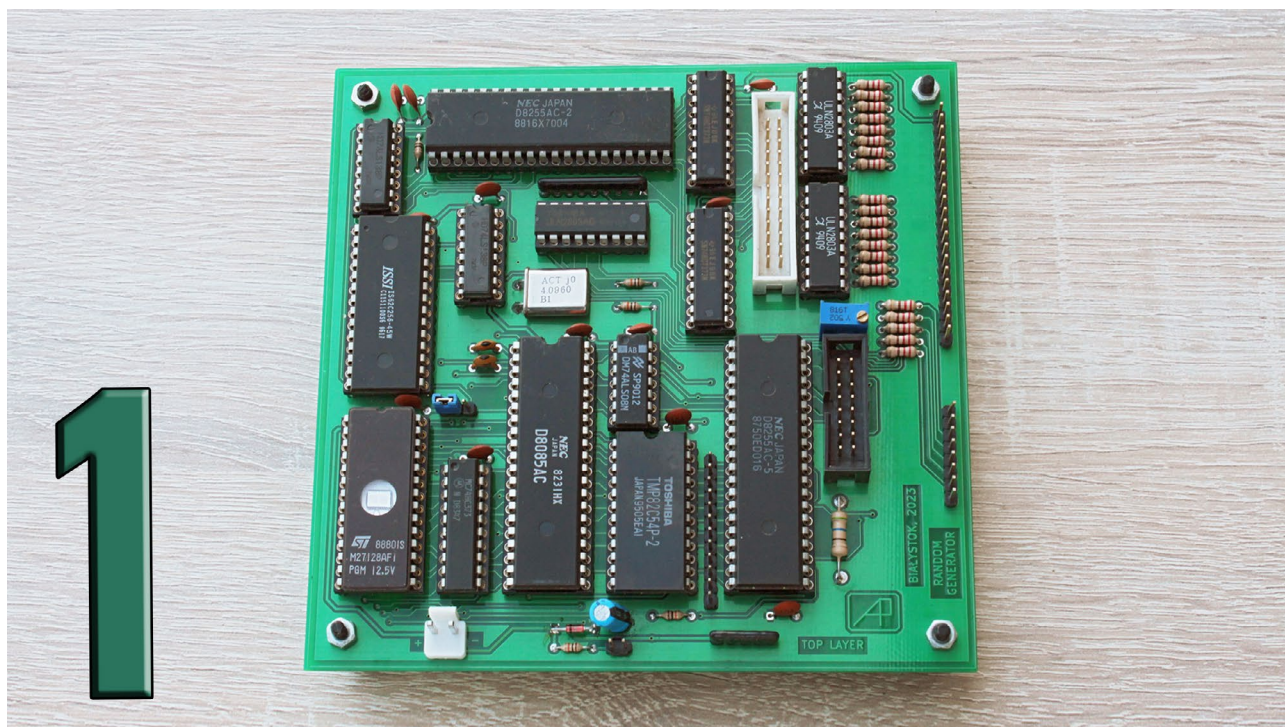
Rysunek 1

utworzyć w przestrzeni adresowej pamięci obszar odpowiadający danemu typowi. Każde utworzenie takiej zmiennej stanowi indywidualne, niezależne „ciało” zmiennej. Identycznym zasadom podlegają obiekty: jest to obszar w pamięci (jak zmienna) mający reprezentację odpowiadającą klasie (klasa jako typ zmiennej). W stosunku do zwykłych zmiennych, które istnieją z chwilą ich utworzenia (mają swoje miejsce w przestrzeni pamięci operacyjnej i ewentualnie są zerowane przed uruchomieniem funkcji *main*), obiekty należy utworzyć. Nie jest to jakaś skomplikowana procedura, do utworzenia obiektu używa się jego konstruktora. Obiekt w sensie zajęcia miejsca w pamięci istnieje, a operację utworzenia należy rozumieć jako wykonanie wszelkiego rodzaju działań inicjujących stan początkowy (niekoniecznie obszar w pamięci związany z obiektem musi zostać wyzerowany, jak w przypadku zwykłych zmiennych) i w sumie sam programista umieszczając w funkcji konstruktora odpowiednie instrukcje decyduje, jakie działania związane są z utworzeniem obiektu.

Przykład klasy

Do zilustrowania powyższych opisów posłuży obsługa 8-cyfrowego wyświetlacza LED z układem MAX7219. Warto to zrealizować jako niezależną parę plików (plik nagłówkowy mający rozszerzenie *.h* oraz plik implementacji, który tym razem ma rozszerzenie *.cpp*). W tej koncepcji definicja klasy musi być dostępna dla każdego modułu programu

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



Generator liczb losowych

Generator liczb losowych zrealizowany na bazie mikroprocesora pochodzącego z „zamierzchłych” czasów. Niekóre z tych układów odcisnęły piętno na historii rozwoju techniki mikroprocesorowej. Takim przykładem był mikroprocesor z oferty Intel, o symbolu 8085.

Algorytmiczne rozwiązanie generatora liczb losowych
Realizacja sprzętowa

Płytką PCB

Obecna oferta mikrokontrolerów jest bardzo szeroka. Ich różnorodne wyposażenie w układy peryferyjne zintegrowane w obrębie jednej struktury półprzewodnikowej, pozwala na realizację nawet złożonych projektów. Jednak zanim do tego doszło, konstruktorzy musieli wcześniej sami borykać się z problemami, o których obecnie wielu nawet nie wie, że istniały. Tytułowy mikroprocesor 8085 (**fotografia 1**) miał swoją premierę w 1976 roku i był znacząco ulepszoną wersją bardzo popularnego poprzednika 8080.

Potrzebowałem jakiegoś rozwiązania sprzętowego pozwalającego na diagnostykę i ocenę zachowania się układów cyfrowych w warunkach sterowania chaotycznego. Taką możliwość zapewnił mi generator liczb losowych. Zagadnienie nie jest aż tak wymagające, by wykorzystać nowoczesne mikrokontrolery, mając w szufladzie układy, których

świećność minęła już wiele lat temu. Z drugiej strony, zapotrzebowanie na zasoby (jako liczbę pinów dostępnych do użycia oraz wielkość pamięci operacyjnej) jest większe, niż możliwości wielu popularnych obecnych mikrokontrolerów. By sprostać tym potrzebom można uzupełnić podstawowy mikrokontroler o dodatkowe układy, jednak nie wszystkie mają możliwość prostej rozbudowy (szczególnie o dodatkową pamięć operacyjną). Nie bez znaczenia jest również możliwość wykorzystania układów,



Fotografia 1

które od dawna zalegają na dnie szuflady – wykorzystania ich zamiast oddania do recyklingu złota. Dodatkowo jest to swoista podróż w głąb historii do pionierskich czasów, gdzie wszystko zależało od nas samych (wielkość pamięci, liczba portów oraz funkcjonalność dodatkowych kontrolerów do realizacji specyficznych operacji).

Algorytmiczne rozwiązanie generatora liczb losowych

Najbardziej znanym sposobem generowania liczb pseudolosowych jest metoda opracowana przez Lehmera, zwana LCG (ang. LCG – Linear Congruential Generator – liniowy generator kongruentny). Polega ona na obliczaniu kolejnych liczb pseudolosowych: zgodnie z prostą formułą, bazując na poprzedniej liczbie losowej (wynik obliczeń zawsze obcięty do 16-bitów):

$$x_{n+1} = a \cdot x_n + b$$

Oczywiście algorytm wymaga pierwszej liczby losowej, określanej jako ziarno (ang. seed). Do celów testowych jako ziarno można wybrać dowolną liczbę 16-bitową, w rzeczywistej realizacji jest to wartość wygenerowana przez mikroprocesor w następujący sposób: mikroprocesor w każdym przerwaniu inkrementuje pewną 16-bitową zmienną. Użycie przycisku z klawiatury uruchamiającego generowanie liczb, bierze jako ziarno tę inkrementowaną liczbę (użytkownik uruchamiając przyciskiem generowanie liczb działa wystarczająco losowo, a mikroprocesor bierze od chwili włączenia zasilania ciągle inkrementowaną liczbę w obsłudze przerwania od upływu czasu). Zanim przystąpiłem do realizacji projektu sprawdziłem w komputerze działanie algorytmu. Napisałem prosty program (**listing 1**), który generował 65536 liczb (tyle jest różnych liczb 16-bitowych) i zliczał ich liczbę wystąpień.

```
#include <stdio.h>
#include <stdlib.h>
#define TabSize 0x10000
#define MultConst 65
#define AddConst 17
FILE * OutFile ;
unsigned short Status [ TabSize ] ;
int main(int argc, char *argv[])
{
    unsigned long Loop ;
    unsigned short RandomV ;
    /*-----*/
    OutFile = fopen ( „random.txt” , „w” ) ;
    for ( Loop = 0 ; Loop < TabSize ; Loop ++ )
        Status [ Loop ] = 0 ;
    RandomV = 1234 ;
    fprintf ( OutFile , „Badania losowosci: stala multiplikatywna=%d, stala addytywna=%d\n” ,
        MultConst , AddConst ) ;
    for ( Loop = 0 ; Loop < TabSize ; Loop ++ )
    {
        RandomV = ( ( MultConst * RandomV ) + AddConst ) ;
        fprintf ( OutFile , „Liczba losowa numer %d=%d\n” , Loop , RandomV ) ;
        Status [ RandomV ] ++ ;
    } /* for */ ;
    for ( Loop = 0 ; Loop < TabSize ; Loop ++ )
    {
        fprintf ( OutFile , „Krotnosc liczby %d=%d\n” , Loop , Status [ Loop ] ) ;
    } /* for */ ;
    fprintf ( OutFile , „Tropienie statystyki\n” ) ;
    for ( Loop = 0 ; Loop < TabSize ; Loop ++ )
    {
        if ( Status [ Loop ] == 0 )
            fprintf ( OutFile , „blad liczba %d nie wygenerowala\n” , Loop ) ;
        if ( Status [ Loop ] > 1 )

```

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



Fascynujące przemiany energii: Nietypowa fotowoltaika (2)

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

Kwantowy, czyli granulowany świat
Bariera energetyczna w półprzewodnikach
„Snop czarnego światła”...

Dlaczego diody LED nie są dobrymi fotoogniwami?
Czy dioda krzemowa może być diodą LED?
Białe diody LED

Ten artykuł też jest rozszerzeniem i uzupełnieniem moich filmów o numerach [B103](#) oraz [B104](#), dotyczących przemian energii świetlnej. Z podręczników dowiadujemy się, że energia świetlna to energia fal elektromagnetycznych, czyli okresowych zmian, drgań specyficznie połączonych pól: elektrycznego i magnetycznego (czykolwiek te pola są). Nie jest to jednak cała prawda o energii świetlnej.

Kwantowy, czyli granulowany świat

„Od zawsze” elektryczność i magnetyzm były zjawiskami bardzo tajemniczymi. Dawni naukowcy próbowali je tłumaczyć na podstawie różnych wyobrażeń opartych na intuicji. Takie intuicyjne wyobrażenie, oparte na XVIII-wiecznej teorii fluidów mówi, że

prąd elektryczny płynie w przewodach podobnie jak woda płynie w rurach. To wyjaśnienie było proste i przekonywało także ówczesnych naukowców.

Jednak już w połowie XIX wieku za sprawą takich geniuszy jak Faraday, Maxwell i Heaviside okazało się, że elektryczność i magnetyzm są ze sobą nierozłącznie związane i że zjawiska elektryczne mają charakter falowy. To było trudne do zrozumienia nie tylko dla ich współczesnych. Do dziś mnóstwo osób nie jest w stanie pogodzić się z faktem, że to nie prąd elektryczny przenosi energię, lecz energię przenosi współdziałanie pola elektrycznego i pola magnetycznego.

Rozważania Maxwella, oparte na koncepcjach Faradaya zaowocowały odkryciem i praktycznym

wykorzystaniem fal radiowych, które tak samo jak światło, są okresowymi, specyficznymi zmianami pól: elektrycznego i magnetycznego. A podstawowa różnica to częstotliwość zmian, drgań tych pól.

Maxwell, Heaviside i ich następcy wyjaśnili bardzo wiele, ale już na przełomie wieków XIX i XX było jasne, że nie są to wyjaśnienia pełne.

To długa historia. Między innymi w eksperymentach wychodziło na to, że natężenie światła to nie wszystko. Okazywało się mianowicie, że natężenie światła mogło być duże, a najprościej mówiąc, czujnik nie reagował. Wszystko wskazywało na to, że, upraszczając, oprócz „ilości światła”, w grę wchodziło coś innego, coś w rodzaju „siły światła”.

Światło niewątpliwie jest falą elektromagnetyczną, ale w wielu sytuacjach zachowuje się tak, jakby było strumieniem swego rodzaju granulek, strumieniem cząstek. I do dziś wielu osobom nie mieści się w głowie, że światło jest jednocześnie i falą, i cząstką, zwaną fotonem, co w fizyce nazywa się *dualizmem korpuskularno-falowym*.

Okazało się, że ta „siła światła”, a raczej energia pojedynczych granulek – kwantów światła, fotonów – zależy od częstotliwości fali elektromagnetycznej, czyli od długości fali, od barwy światła.

Elektronik zapisałyby to w postaci wzoru z użyciem literki **f** do oznaczenia częstotliwości. W fizyce zamiast literki **f** wykorzystuje się grecką literkę ni (**v**), a w grę wchodzi też maleńka stała fizyczna, zwana stałą Plancka. Dlatego wzór na energię pojedynczej „granulki światła”, fotonu zwykle zapisuje się w postaci $E = h\nu$, gdzie **h** to stała Plancka, a **v** – częstotliwość.

To jest artykuł o przemianach energii, o energii światła. Możemy słusznie mówić, że światło – fale elektromagnetyczne o ogromnych częstotliwościach rzędu setek teraherców – niosą energię. Tak, ale światło jest też jednocześnie strumieniem „granulek” – fotonów. Dlatego opisując energetyczne właściwości światła należy podać nie tylko, najprościej mówiąc, „ilość światła”, czyli „liczbę granulek”, ale też energię poszczególnych „granulek”.

Czym większa częstotliwość (mniejsza długość fali), tym energia pojedynczego fotonu jest większa. Częstotliwość fali światła czerwonego (o długości fali około 700 nanometrów) jest mniejsza, niż niebieskiego (o częstotliwości fali około 400 nanometrów), więc **kwant światła czerwonego ma mniejszą**



innymi z tymi pokazanymi na **fotografii 1**, udowodniły, że wszystkie diody LED są jednocześnie ogniwa-
mi fotowoltaicznymi. Wszystkie reagują na światło białe, które jest mieszaniną światła o różnych barwach. Natomiast niejednakowo reagują na światło o różnych barwach. Okazało się też, że niebieskie diody LED pracujące jako fotoogniwa nie reagują na „niskoenergetyczne” światło czerwone, a tym bardziej na jeszcze „mniej energetyczne” promieniowanie podczerwone o długości fali około 940 nanometrów. Reagują na „wysokoenergetyczne” światło niebieskie i tym bardziej ultrafioletowe (o długości fali światła mniej więcej 300 nanometrów).

Natężenie światła czerwonego (liczba granulek), może być duże, ale niebieska fotodiody nie reaguje nawet przy bardzo silnym świetle czerwonym. Dlatego, że poszczególne granulki mają za małą energię żeby „przeskoczyć próg”.

Natomiast czerwone diody LED w roli fotoogniw reagowały na światło widzialne o dowolnych barwach. Najprościej biorąc dlatego, że „granulki” światła o innych barwach niosą większą energię, niż czerwone.

Bariera energetyczna w półprzewodnikach

A co znaczy „przeskoczyć próg”? Tu dochodzimy do ogromnie ważnych zagadnień związanych ze wszystkimi elementami półprzewodnikowymi. Ich działania absolutnie nie da się wytłumaczyć ani za pomocą prymitywnych analogii wodnych, ani nawet za pomocą genialnej koncepcji fal elektromagnetycznych Maxwella.

Te XIX-wieczne, uproszczone, skądinąd bardzo pożyteczne wyobrażenia, w przypadku półprzewodników okazują się bezużyteczne, bo elektromagnetyzm jest dużo bardziej skomplikowany i nadal nie do końca poznany i opisany. Dlatego, żeby zrozumieć właściwości półprzewodników, czymkolwiek one są,

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



Realne parametry rezystorów dużej mocy

Wbrew pozorom, prawidłowe wykorzystanie rezystorów dużej mocy nie jest tak łatwe, jak się wydaje. Informacje podawane przez sprzedawców bardziej szkodzą niż wyjaśniają. W artykule przedstawione są informacje dotyczące praktycznych aspektów wykorzystania rezystorów mocy w aluminiowych obudowach.

Parametry „handlowe” oraz katalogowe
Jaką moc naprawdę ma rezystor 100-watowy?
Od teorii do najprostszej praktyki

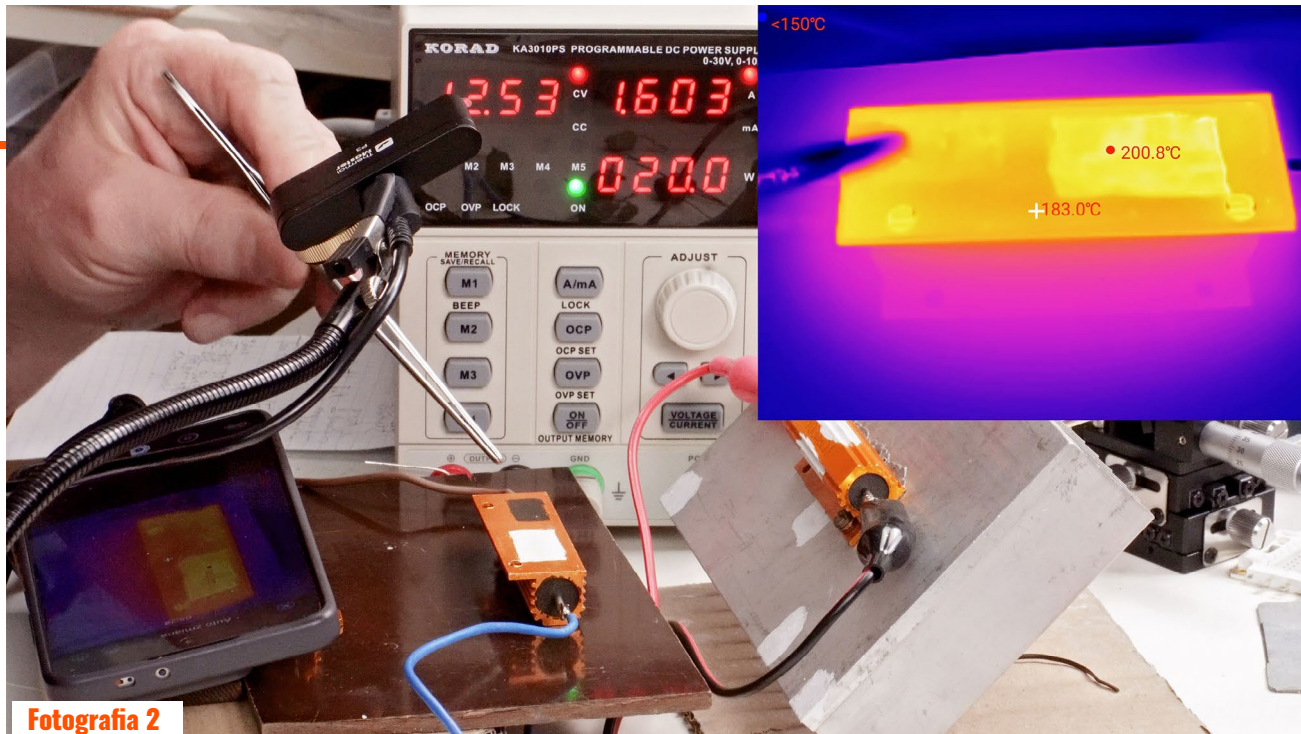
Dokładniejsza analiza kart katalogowych
Parametry termiczne
Jak w praktyce określać obciążalność?

Artykuł jest uzupełnieniem [filmu oznaczonego E028](#). Z tego filmu pochodzi też powyższa fotografia tytułowa. Pokazałem tam kilka interesujących eksperymentów, w tym przypadki samowylutowania rezystorów. Ale głównie zajmowałem się rezystorami, które potocznie nazywane są aluminiowymi. Ich obudowa jest naprawdę porządna – to aluminiowy odlew z żebrami pełniącymi funkcję radiatora, wyposażony w nóżki pozwalające przykręcić go do jakiejś podstawy. Jak pokazuje **fotografia 1**, na obudowie podana jest moc strat, czyli obciążalność. Rezystor 100-watowy ma korpusu o długości 60 mm. Rozmiary i wykonanie obudowy budzą zaufanie, ale w filmie pokazałem,

że już przy mocy strat około 20 W, czyli przy 1/5 mocy nominalnej, temperatura obudowy osiąga +200°C. A woda wrze w temperaturze +100 stopni.



Fotografia 1



Fotografia 2

Trzeba też pamiętać, że nadal bardzo często używany klasyczny lut ołowiowy (Sn63/Pb37, a także Sn60/Pb40) mięknie i topi się w temperaturze $+183^{\circ}\text{C}$...

Starsi elektronicy znają informacje o przypadkach samowylutowania z płytek drukowanych niektórych drutowych rezystorów mocy. Przykład pokazałem w filmie. Końcówki elementów uzyskiwały wtedy temperaturę ponad 183 stopni, ale to wcale nie dowodziło, że przekroczone były katalogowe parametry rezystora. To było i jest niezrozumiałe dla wielu elektroników, którzy nie wniknęli w praktyczne aspekty dotyczące katalogowych parametrów takich rezystorów. Nadal jest problem z prawidłowym rozumieniem realnego sensu parametrów rezystorów dużej mocy, w szczególności właśnie rezystorów w aluminiowych obudowach. Ich „radiatoropodobny” wygląd przekonuje mnóstwo osób, że taki rezystor może niejako z natury rozproszyć podaną w katalogu moc nominalną. A tak nie jest, co pokazałem w filmie.

Czy są to przykłady oszustwa? Czy producent i sprzedawca oszukują?

Nie! Testowany w filmie rezystor rzeczywiście jest 100-watowy, tylko pracuje bez niezbędnego dodatkowego radiatora. **Fotografia 2** przedstawia pomiar tego rezystora bez radiatora. Zasilacz pokazuje pobór mocy 20,0 W, a kamera termowizyjna świadczy, że obudowa rezystora ma temperaturę 200 stopni.

Fotografia 3 przedstawia pomiar takiego rezystora, ale przymocowanego do bardzo solidnego radiatora. Tym razem zasilacz wskazuje, że moc tracona w rezystorze w postaci ciepła wynosi 118 watów, ale temperatura obudowy rezystora jest niższa – wynosi około 185 stopni.

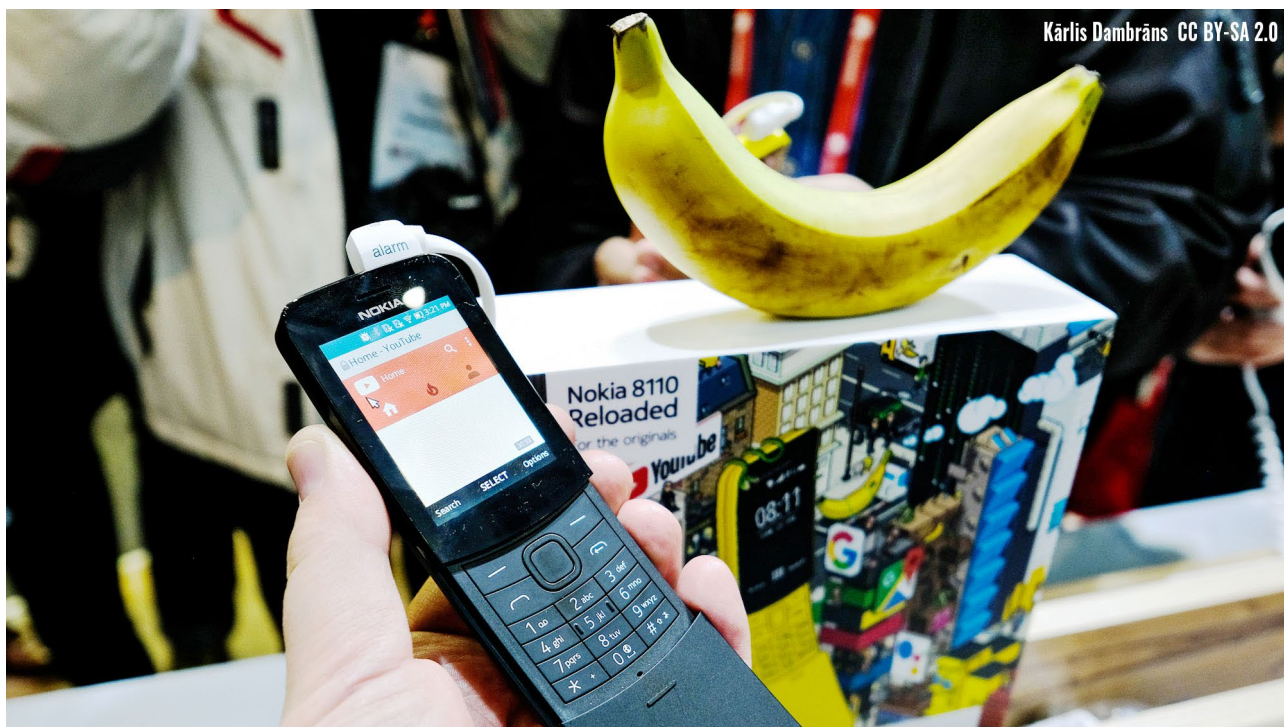
Co prawda mniej niż w sytuacji z poprzedniej fotografii, ale i tak około 160 stopni ponad temperaturę otoczenia. Czy to jest dopuszczalne? Czy grozi szybkim uszkodzeniem? O tym za chwilę.



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



Fotografia 3



Jak działa smartfon?

4G – telefonia, której nie ma!

To jest szósty artykuł serii zainspirowanej pytaniem mojej małżonki: jak działa smartfon? Niestety, nie da się tego wyjaśnić krótko. Ja do tej pory omówiłem radiotelefony oraz generacje od 0 G do 3 G. W tym artykule powinienem omówić czwartą generację telefonii komórkowej oznaczaną skrótem 4 G.

Telefonia 2 G, 3 G i Internet
Standardy telefonii komórkowej
Telefonia 4 G i LTE

Podsumowanie
Czym telefonia cyfrowa różni się od analogowej?
Nie tylko dla dociekliwych

Dla osoby „nietechnicznej” próba zrozumienia, jak działa smartfon, jak możliwe jest przesyłanie na dowolne odległości rozmów, tekstów, zdjęć i filmów, przypomina próbę wskoczenia do pędzącego ekspresu (**fotografia 1**). Odpowiedź musi być bardzo obszerna, a osobom „nietechnicznym” nie da się w prosty sposób wyjaśnić wszystkich zadziwiających szczegółów. Jednak można przystępnie wyjaśnić i przybliżyć, pokazać zarys, najważniejsze zasady, a także pewne ciekawostki. Można, tylko w tym celu trzeba wrócić do korzeni i pokazać, jak telefonia powstała i jak się zmieniała.

Z poprzedniego artykułu serii wiemy, że dużo lepszy system UMTS (3G) nie wyparł i nie zastąpił pocziwego GSM (2G). Telefonię 3G po cichu, bez rozgłosu wycofano i zamknięto w prawie wszystkich krajach. Do rozmów telefonicznych nadal wykorzystuje się mocno już dziś archaiczny system GSM, czyli 2G. Historia 3G była dziwna i podobnie jest z telefonią czwartej generacji (4G)! **Otóż można powiedzieć, że tak naprawdę nie było jej, nie ma i nie będzie!**



Fotografia 1

Telefonia 2G, 3G i Internet

Telefonię GSM i inne systemy drugiej generacji (2G) opracowywano w latach 80. a wykorzystano w praktyce od lat 90. Co ogromnie ważne, w latach 90. równolegle i zupełnie niezależnie zaczynał rozwijać się Internet, czyli pomalutku zaczęło się wykorzystanie globalnej sieci komputerowej.

Pierwotny standard GSM zapewniał i zapewnia przeprowadzanie rozmów telefonicznych, ale praktycznie nie miał związku z rodzącym się pomału Internetem. Owszem, GSM to system w pełni cyfrowy, czyli w telefonicznym standardzie GSM dźwięk przesyłany jest w formie cyfrowej i pozostaje trochę „zapasu”, żeby obok (albo nawet zamiast) sygnałów dźwiękowych przesyłać trochę informacji cyfrowych. Trochę, ale jak na dzisiejsze warunki – żałownie mało.

W telefonii 2G następuje utworzenie kanału cyfrowego do dwukierunkowego przesyłania „cyferek”, które w pierwszej kolejności mogą przekazywać „informacje głosowe”, ale nie tylko – „cyferki” mogą przekazywać dowolne, jakiejkolwiek informacje, na przykład zdjęcia, film, wideorozmowę, dokumenty.

Sieci telefonii komórkowej 2G według założeń miały zapewniać tylko rozmowy telefoniczne i przesyłanie SMS-ów – krótkich wiadomości tekstowych

Potem doszły MMS-y, czyli nie tylko tekst – literki, ale inne multimedia, także obrazki i dźwięk.

To jednak były, powiedzmy, „dodatki do telefonu”. A w tym czasie oddzielnie, niejako obok, intensywnie rozwijał się Internet. **Internet, czyli globalna sieć połączonych ze sobą komputerów służąca do wymiany dowolnych informacji cyfrowych.**

Internet na początku przewidziany był do przekazywania tekstów, dokumentów, obrazów – ogólnie jakichkolwiek informacji w postaci cyfrowej. Zaczął się w wojsku, następnie sieci internetowe wykorzystano na wyższych uczelniach, a potem dość szybko „trafił pod strzechy”.

Internet był oddzielną inicjatywą, zupełnie niezwiązaną z telefonią, działającą na zupełnie innych zasadach. Tak, pomimo pozornych podobieństw, na zdecydowanie innych zasadach.

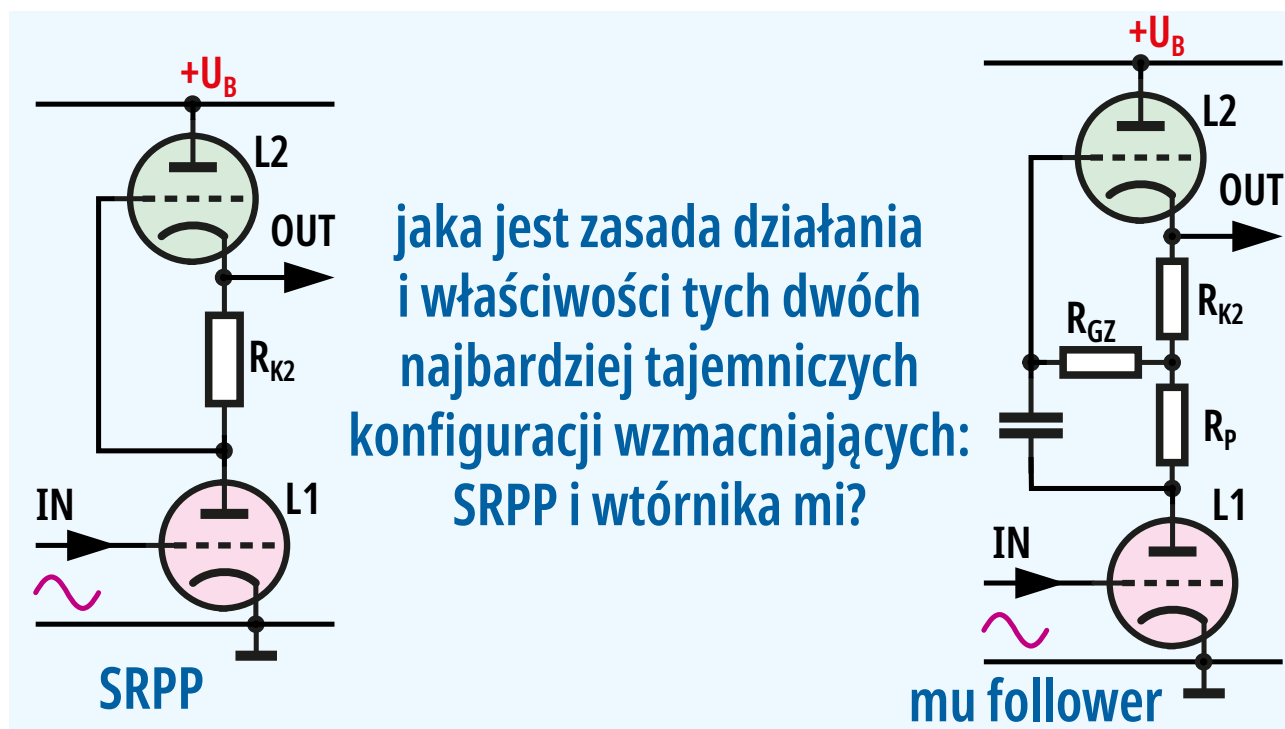
W telefonii (telekomunikacji) abonent ma swój numer telefonu. W Internecie użytkownik nie ma numeru telefonu, tylko zupełnie inny numer, tak zwany adres IP (czytaj: adres aj-pi), o którym zresztą większość użytkowników Internetu nic nie wie. Można najogólniej stwierdzić, że adres IP nie jest własnością użytkownika i znajomość samego adresu IP nie pozwala zidentyfikować użytkownika. Dziś zresztą ta sprawa jest niesamowicie poplątana.



Fotografia 2 Nokia 8210 4G fot. PJ: CC BT-SA 4.0

w Internecie (IPv4), i późniejsze próby wprowadzenia „dokładniejszego adresowania” (IPv6) zakładały, że adres „aj-pi” będzie identyfikował, może nie abonent, ale konkretny sprzęt, urządzenie. Niestety,

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



Lampowe konfiguracje SRPP i mu follower (2)

W cyklu o lampach elektronowych zajmujemy się dwiema konfiguracjami, które są najmniej rozumiane i najbardziej tajemnicze. Niewiele osób ma jasne wyobrażenie o ich działaniu i właściwościach. Dlatego koniecznie muszę przedstawić dalsze obszernie informacje wstępne i wyjaśnić liczne nieporozumienia.

Układy lampowe – różne punkty widzenia
Po czym poznajemy wtórnik (lampowy)?

Maksymalne wzmocnienie napięciowe lampy
Dynamiczne obciążenie anodowe

Zgodnie z obietnicą, mam możliwie przystępnie przedstawić pracę i właściwości konfiguracji znanych jako SRPP i mu follower. Podstawowe schematy obu tych konfiguracji pokazane są na **rysunku tytułowym**. Problem w tym, że my, świadomie czy nie, chcielibyśmy mieć „czyste i jasne sytuacje” oraz jednoznaczne definicje. A w elektronice nie zawsze tak jest, czego przekładem są właśnie te dwie konfiguracje.

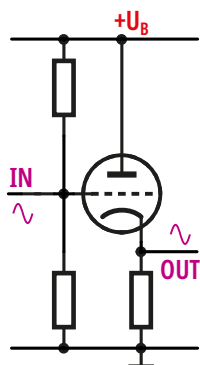
Układy lampowe – różne punkty widzenia

Omawiane konfiguracje układowo są nieskomplikowane, natomiast opisanie ich działania za pomocą dobrze rozumianych określeń nie jest ta-

kie proste. Do opisu ich działania można podejść na co najmniej dwa różne sposoby. Jeden związany jest ze wzmacniaczami przeciwsobnymi, co częściowo przedstawiłem w poprzednim artykule serii. Częściowo, dlatego do tego wątku jeszcze wrócimy. A w poniższym artykule przedstawię inne podejście.

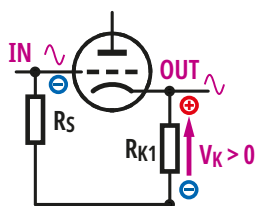
Zacznę o tego, że dość łatwo zdefiniować i opisać właściwości omawianych wcześniej podstawowych konfiguracji z jedną lampą. Mamy wtedy trzy podstawowe konfiguracje: ze wspólną katodą, ze wspólną siatką i ze wspólną anodą. Teraz interesuje nas ta ostatnia. Układ ze wspólną anodą, czyli wtórnik katodowy.

Podstawowy układ pokazany jest na **rysunku 1**. Ale już drobne zmiany mogą zdecydowanie zmienić właściwości. I właśnie dobrym przykładem jest wtórnik. Otóż powszechnie wiadomo, że dla uproszczenia układu i dla wygody, często w układach lampowych wykorzystujemy pomysł na autopolaryzację, czyli uzyskujemy potrzebne ujemne napięcie



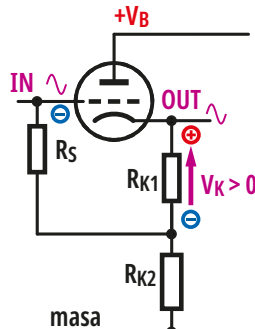
Rysunek 1

siatki względem katody za pomocą małego rezystora R_K w obwodzie katody, według idei **rysunku 2**. Autopolaryzację możemy też wykorzystać we wtórniku i wtedy zamiast kanonicznej wersji z rysunku 1 mamy układ jak na **rysunku 3**.



idea autopolaryzacji
Rysunek 2

Możemy wykorzystać, wykorzystujemy i cieszymy się sprytnym wykorzystaniem autopolaryzacji. Tak, tylko zwykle nie pamiętamy, że wersja wtórnika z autopolaryzacją, „po cichu”, niejako przy okazji wprowadza bootstrap – podciąganie z katody na siatkę, przez co jeszcze bardziej rośnie (i tak duża dynamiczna) rezystancja wejściowa takiego wtórnika. W praktyce wzrost dynamicznej rezystancji wejściowej we wtórniku nie ma praktycznego znaczenia, więc zazwyczaj w ogóle nie zwracamy na to uwagi.



Rysunek 3

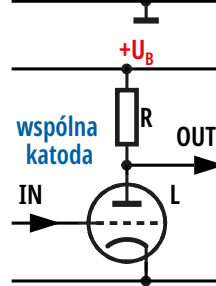
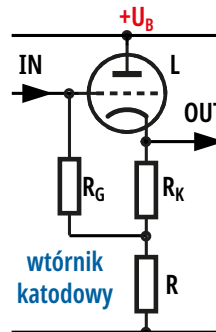
A ponadto sporo osób nie rozumie, co to jest bootstrap – podciąganie i jakie są jego konsekwencje. Dlatego tym bardziej niezrozumiałe okazuje się działanie i właściwości konfiguracji niezbyt słusznie nazywanych **SRPP** i **mu follower**. W nich nakłada się na siebie kilka czynników, w szczególności różne rodzaje sprzężenia, nie tyle jednak, dość dobrze rozumianego ujemnego sprzężenia zwrotnego, co właśnie bootstrapu, który można traktować jako „sprzężenie w przód”, a także jako dodatnie sprzężenie zwrotne.

Problem w tym, że my oczekujemy krótkiego, łatwego, intuicyjnego wyjaśnienia, a już choćby w przypadku bootstrapu z rysunku 3 takiego nie ma. Sytuacja jeszcze bardziej się komplikuje w przypadku omawianych układów „piętrowych”

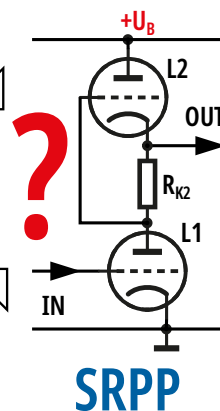
zrozumienia ich działania i specyfiki niezbędne są informacje z poprzednich artykułów, w szczególności z ostatniego, w którym głównie omawiałem wzmacniacze push-pull, ale też wspomniałem, że nie ma „komplementarnych lamp”. Te informacje będą nam jeszcze potrzebne, ale na razie muszę wprowadzić jeszcze inny wątek dotyczący wtórników.

Po czym poznajemy wtórnik (lampowy)?

Wiele osób uważa, że konfiguracje z rysunku tytułowego to specyficzne odmiany wtórnika, a w takim przekonaniu utwierdza nazwa *mu follower*, czyli *wtórnik mi*. Uważają, że jest to połączenie klasycznego wzmacniacza ze wspólną katodą oraz wtórnika według **rysunku 4**. Uważają, że oporność wyjściowa jest bardzo mała, a to właśnie dzięki



Rysunek 4



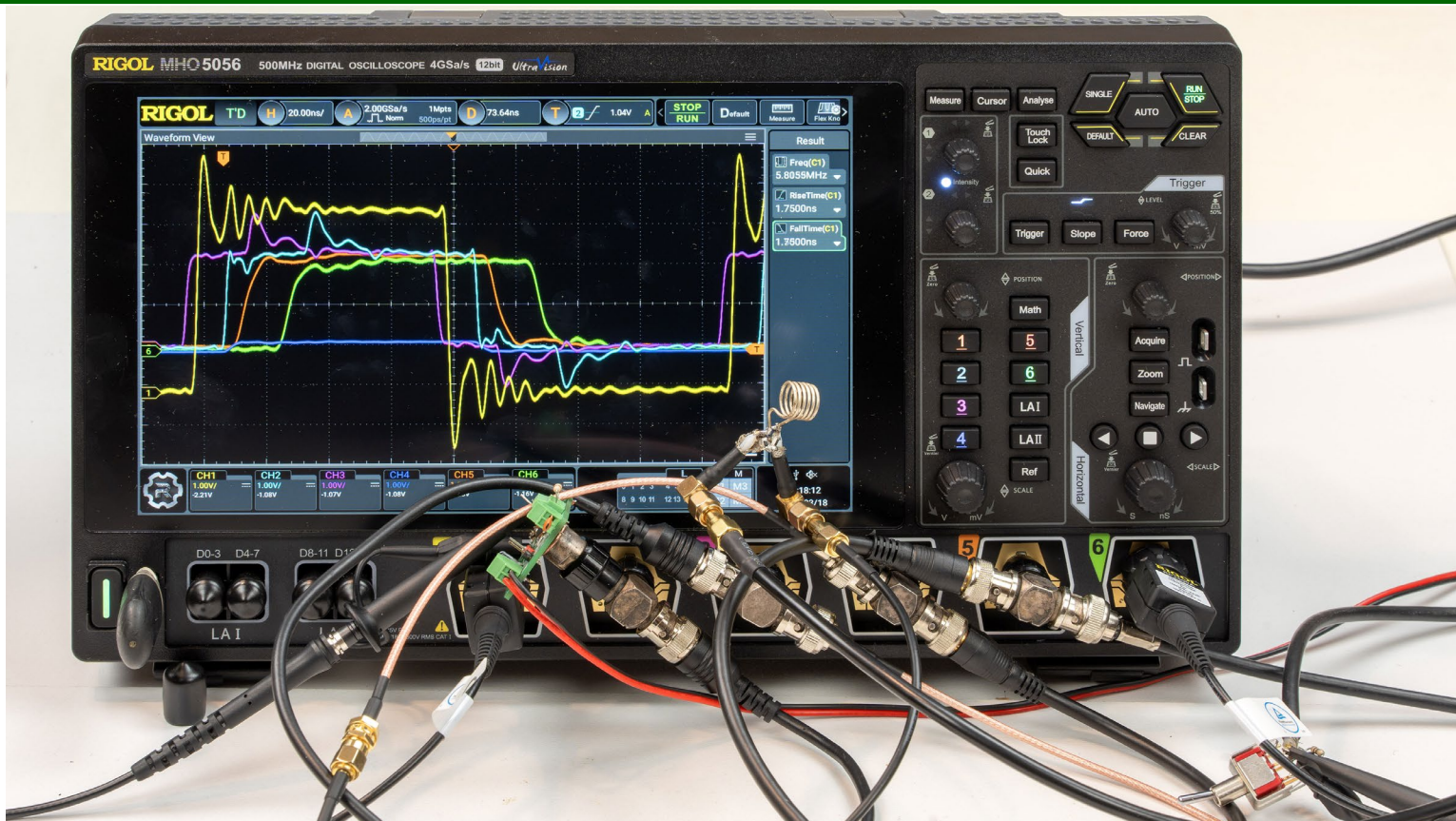
o b e c n o ś c i wtórnika.

Nie, nie, nie! Podobieństwo jest pozorne. Owszem, na dole niewątpliwie mamy układ ze wspólną katodą (L1). Ale na górze (L2) wcale nie mamy wtórnika! Mniej zorientowanych w błąd wprowadza fakt, że wyjściem sygnału jest katoda L2, bo oni jednoznacznie łączą to z wtórnikiem katodowym. Wniosek taki jest pochopny, a nawet błędny, i to błędny z kilku powodów.

To prawda, że w prawidłowo pracującej triodzie przebieg (niewielkiego zmiennego) napięcia na katodzie jest podobny do przebiegu napięcia na siatce. Przebiegi zmiennie na siatce i katodzie są przesunięte „w pionie”, bo generalnie napięcie siatki jest ujemne względem katody. Ale zgodność zmian napięcia stałego czy zmiennego na siatce i katodzie nie jest najważniejsza!

Wtórnik katodowy to konfiguracja, w której występuje bardzo silne ujemne sprzężenie zwrotne, przez co wzmocnienie jest bliskie, ale trochę mniejsze od jedności – równe jedności byłoby w idealnym

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



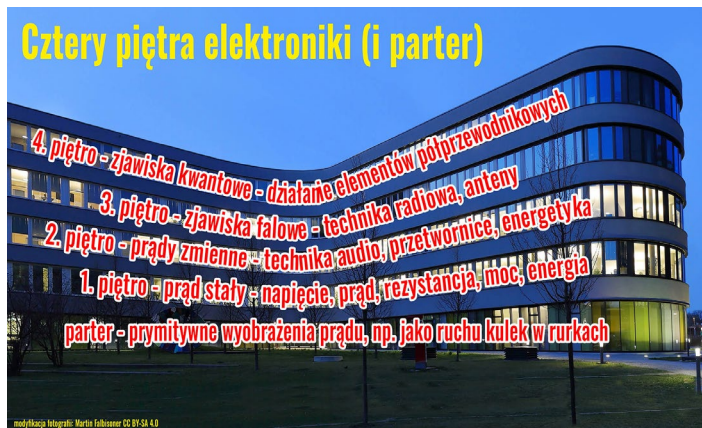
Dopasowanie szumowe – a co to właściwie jest?

Zgodnie z obietnicą, omawiając kolejne aspekty problemu dopasowania, zaczynamy zajmować się dopasowaniem szumowym. To wyjątkowo trudny temat i okazało się, że nawet zgrubne zasygnalizowanie go w jakiś prosty sposób wymaga dwóch artykułów, z których poniższy jest tylko wprowadzeniem.

Znów dopasowanie falowe i energetyczne
Co ważniejsze: napięcie czy moc?

Wzmocnienie i granice wzmacniania

Wcześniej, w artykule Podstawowe zjawiska falowe, omówiłem fundamenty ogromnie ważnego, a słabo rozumianego problemu zjawisk falowych i odbić, a w artykule serii Dopasowanie falowe – dlaczego i kiedy potrzebne? pokazałem jak można uniknąć odbić i wyjaśniłem, czym jest dopasowanie falowe. W poprzednim artykule tej serii Błogosławieństwo czy przekleństwo dopasowania? omówiłem dopasowanie energetyczne i pewne fałszywe wyobrażenia na jego temat, pokutujące w umysłach wielu osób. W poniższym artykule spróbuję nie tyle wyjaśnić, co wstępnie zasygnalizować i nieco przybliżyć naprawdę trudny problem dopasowania szumowego.



ZROZUMIEĆ **E**LEKTRONIKĘ z Piotrem Góreckim

ZE 5/2026

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.