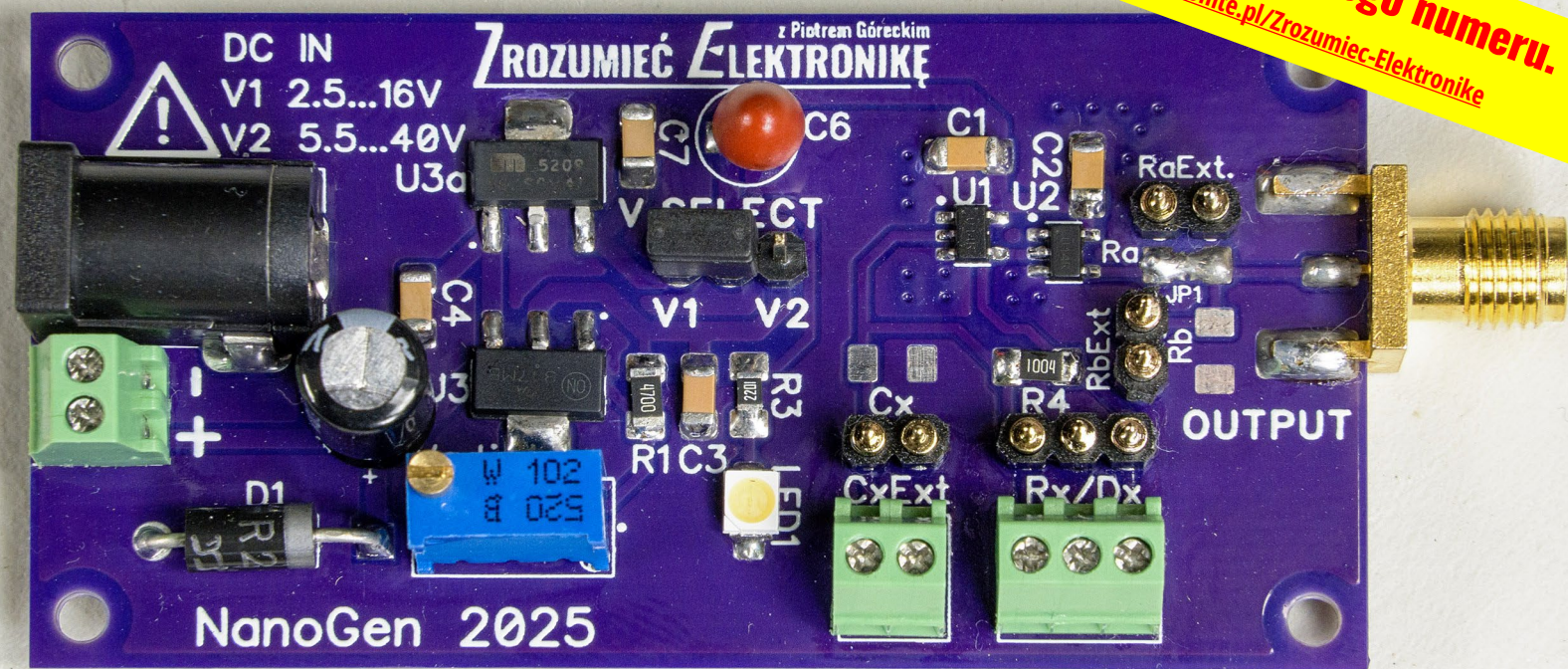
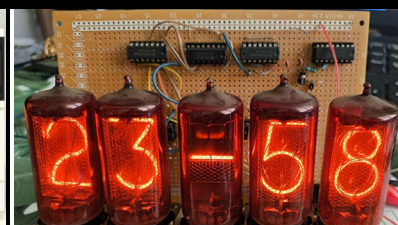
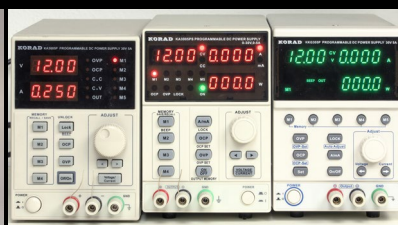
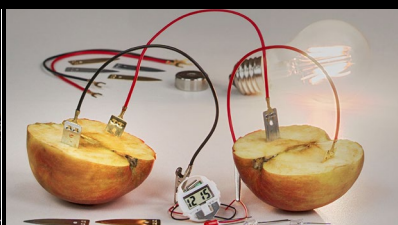
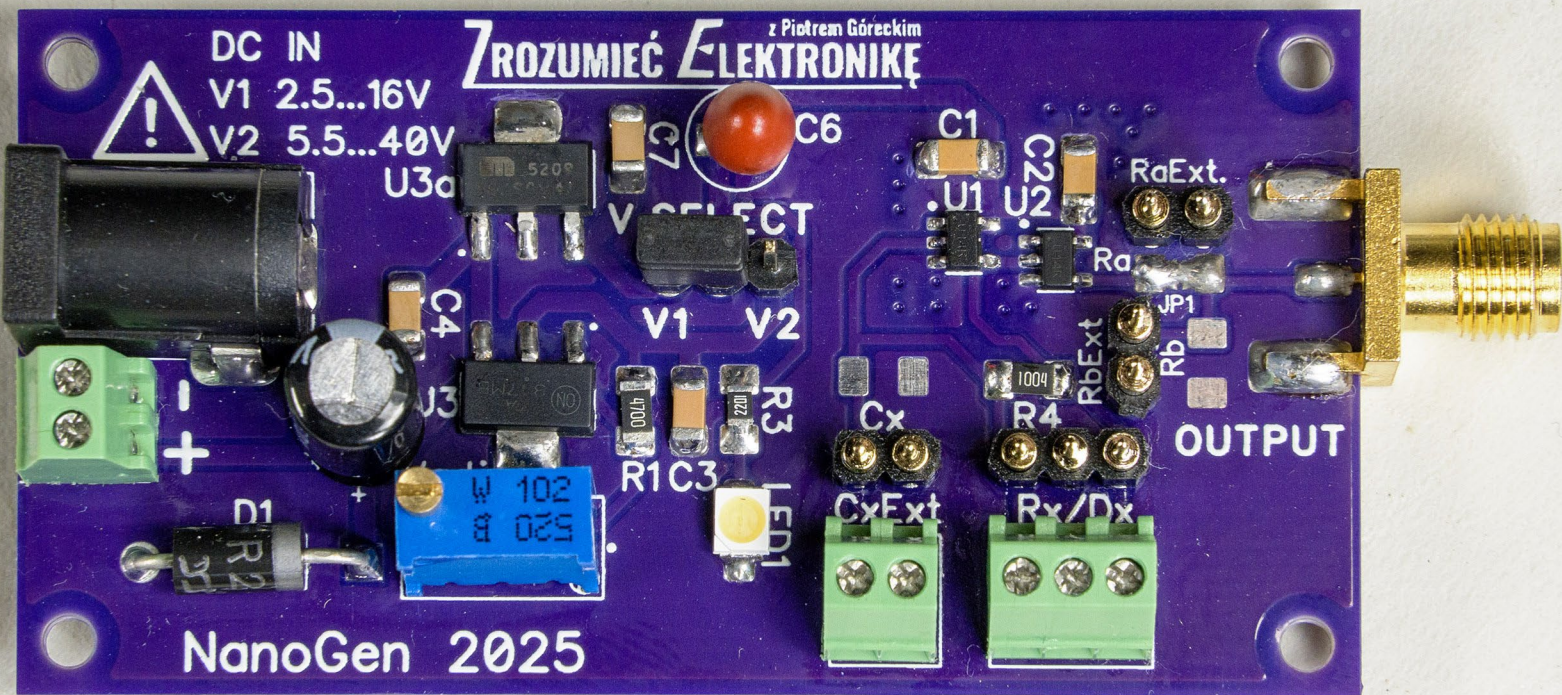




NanoGen – Generator nanosekundowy

- Lampowe konfiguracje z dwiema triodami • Wspólnie projektujemy: Wykorzystanie dawnych wyświetlaczy
- Mikroprocesory i mikrokontrolery – interfejs I2C • Fundamenty elektroniki – analogie hydrauliczne
- Fascynujące Przemiany Energii Proste ogniwa galwaniczne • Zasilacze KORAD, czyli nie tylko KA3005D
- GNSS RTK, czyli nawigacja o dokładności centymetrowej • Opowiadanie o „nieznaczących” drucikach





NanoGen – Generator nanosekundowy

Opisywany poniżej, prościutki generator znajdzie różne zastosowania, ponieważ może wytwarzać przebieg prostokątny o dowolnej częstotliwości aż do 380 MHz. Może wytwarzać impulsy prostokątne o bardzo stromych zboczach, o dowolnej długości i częstotliwości powtarzania, nie krótsze niż 1,25 nanosekundy.

Rozważania projektowe
Dlaczego akurat 74AUCxx?

Opis układu generatora

W artykule przedstawiony jest generator przebiegu prostokątnego – prosty, ale o zaskakujących parametrach. Pozwala uzyskać „prostokąt” o dowolnej częstotliwości, w zakresie od drobnych ułamków herca aż do imponujących 380 megaherców! Może wytwarzać impulsy o dowolnych parametrach; niezależnie ustawiany jest czas impulsu i czas przerwy, czyli częstotliwość i współczynnik wypełnienia. Co bardzo ważne, zbocza wytwarzanego przebiegu prostokątnego są bardzo strome: czas narastania i opadania jest krótszy niż pół nanosekundy. A w czasie 1 nanosekundy światło przebiega drogę około jednej trzeciej metra, a sygnał elektryczny w kablu pokonuje odległość około 20 centymetrów.

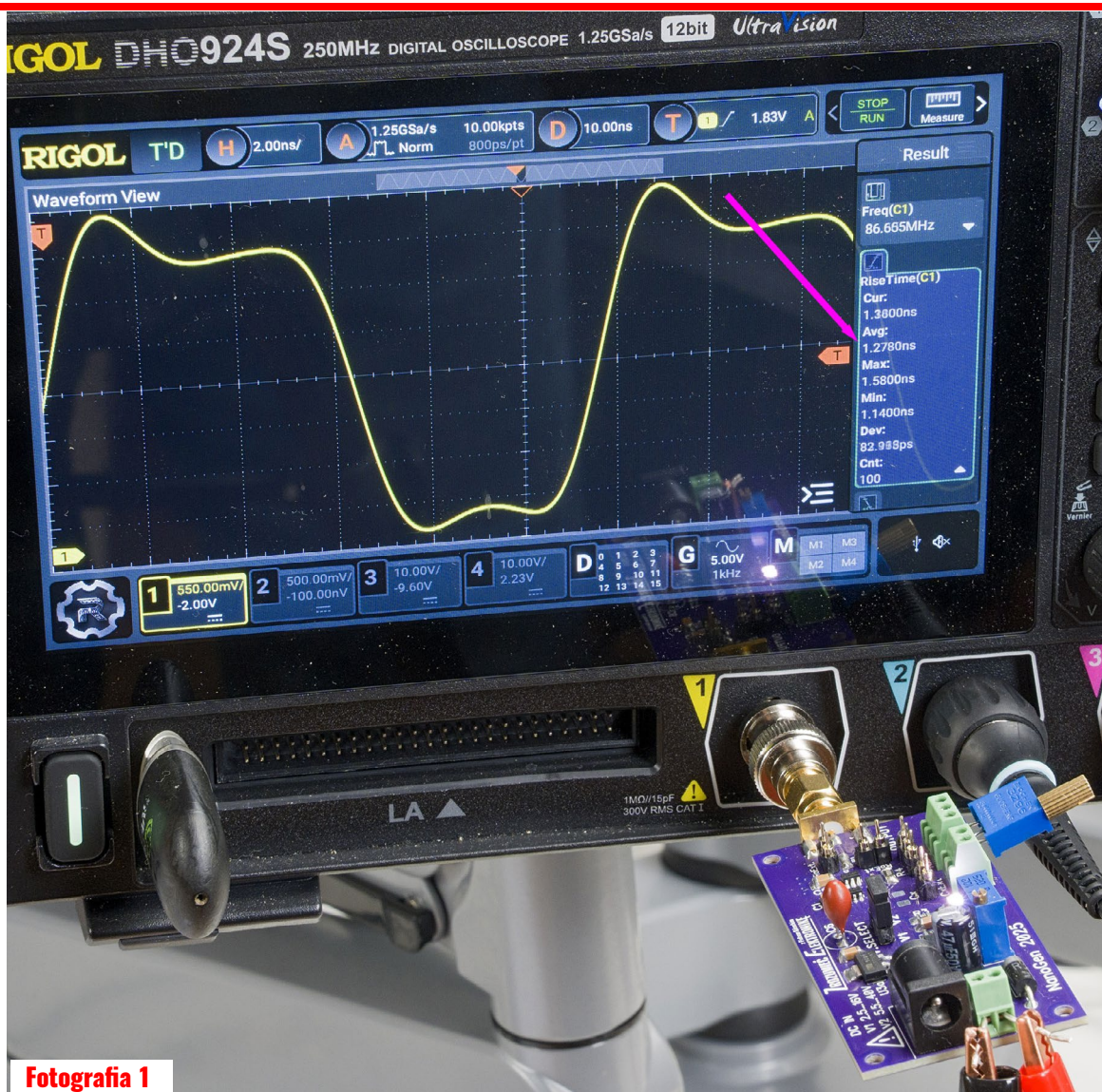
W czasie zmiany stanu logicznego sygnału, sygnał w kablu przebiegnie mniej niż 10 centymetrów.

Do czego potrzebny jest taki generator?

Generatory przebiegu prostokątnego o bardzo stromych zboczach potrzebne są do różnych celów. W licznych zastosowaniach częstotliwość i wypełnienie przebiegu nie są istotne, a najważniejsza jest właśnie stromość zboczy. Impulsy o stromych zboczach pozwalają przeprowadzać rozmaite pomiary i eksperymenty związane z dopasowaniem falowym. Strome impulsy pozwalają przeprowadzać kompensację częstotliwościową rozmaitych obwodów i układów. Pozwalają też w prosty sposób określać pasmo przenoszenia przyrządów, w tym oscyloskopów.

Idealny przebieg prostokątny jest złożeniem nieskończonej liczby sinusoidalnych składowych (nieparzystych). Właśnie te harmoniczne pozwalają w prosty sposób sprawdzać pasmo przenoszenia i jego liniowość w takich urządzeniach jak wzmacniacze, tłumiki, sondy. Co istotne, pozwalają też sprawdzić pasmo przenoszenia oscyloskopów. Otóż powszechnie wykorzystuje się prosty wzór $BW = 0,35 / tr$, gdzie BW to szerokość pasma, a tr , to czas narastania. Dla ścisłości należy dodać, że dotyczy on prostych obwodów 1 rzędu, czyli na przykład popularnych obwodów – filtrów RC z jednym

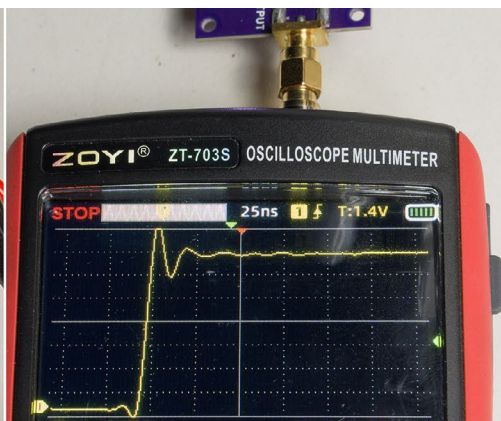
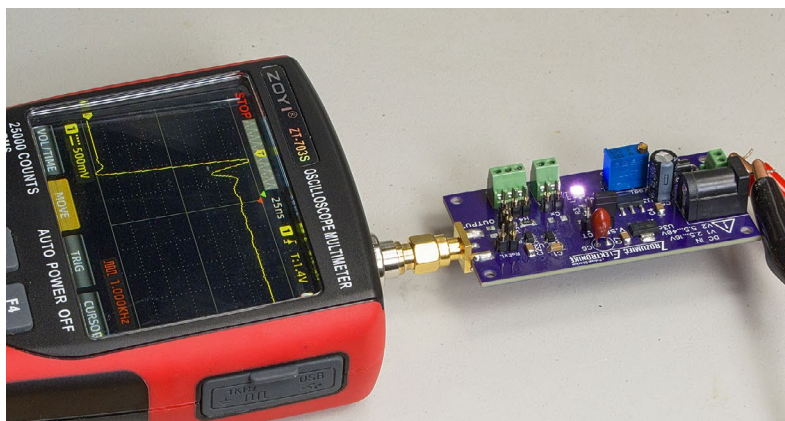
rezystorem i jednym kondensatorem. Jeżeli chodzi o układy mające charakter filtrów wyższych rzędów, zależność między pasmem przenoszenia i czasem narastania może być nieco inna i współczynnik we wzorze zamiast 0,35 może mieć wartość 0,4, a nawet więcej. Warto o tym wiedzieć, ale najczęściej wykorzystuje się właśnie wzór $BW = 0,35 / tr$.



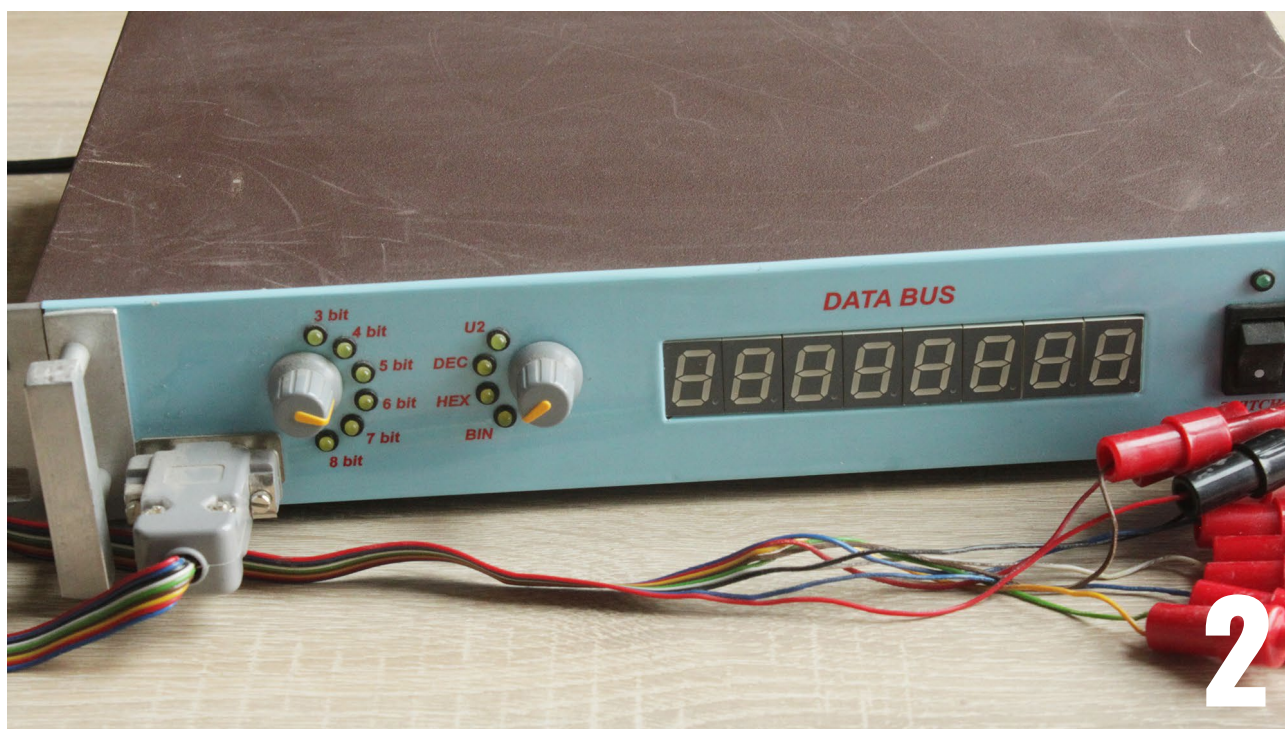
Fotografia 1

Zgodnie z tym wzorem i zgodnie z **rysunkiem 1**, mój 250-megahercowy oscyloskop DH0924S ma pasmo przenoszenia 274 MHz.

Z **fotografii 2** można oszacować, że ręczny, 50-megahercowy multimetrosocyloskop ZT-703S ma czas narastania około 7 nanosekund, czyli zgodnie z tą zależnością pasmo około 50 MHz.



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



Ośmiobitowy monitor szyny danych, część 2

Oto kolejna część artykułu dotyczącego budowy urządzenia, które pozwala pokazać na wyświetlaczu w różnej formie informację o deklarowanej liczbie bitów (nie więcej niż 8). Część ta jest poświęcona uruchomieniu oraz przedstawia propozycję rozwiązania obudowy urządzenia.

Uruchomienie
Obudowa urządzenia

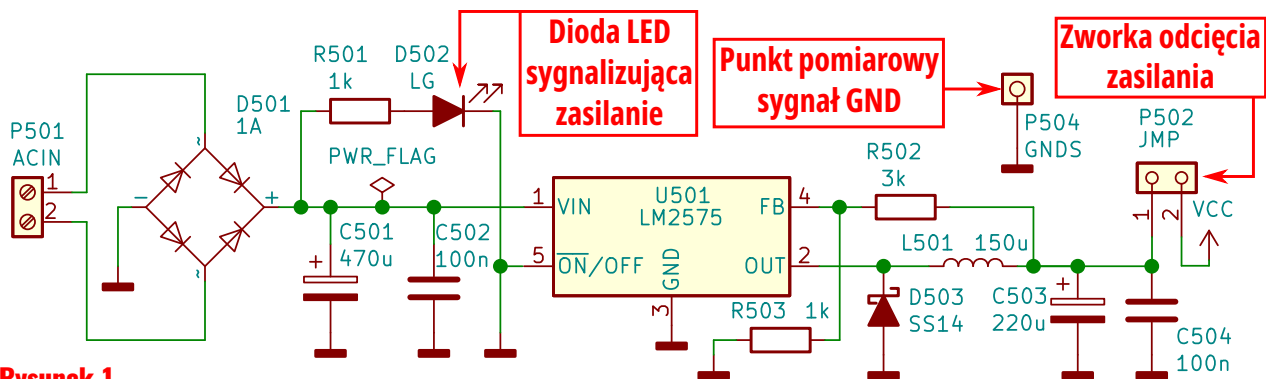
Prezentacja działania

Po zmontowaniu poszczególnych płytek wchodzących w skład monitora, nadchodzi czas na uruchomienie oraz integrację poszczególnych elementów w jedną całość. Do zasilania monitora przewidziany jest transformator o mocy kilku VA i napięciu wyjściowym w granicach 9...15 V. Wartość napięcia nie jest krytyczna, sam monitor zawiera impulsowy stabilizator napięcia, co daje wysoką sprawność w przetwarzaniu napięcia przemiennego do standardowo stosowanego w układach cyfrowych napięcia stałego +5 V.

Uruchomienie

Nigdy nie można wykluczyć błędów montażu, np. braku połączenia między elementami czy zim-

nego lutu. W niektórych przypadkach może to doprowadzić do przykrych niespodzianek i może się zdarzyć, że zasilacz monitora wypracuje niewłaściwe, zbyt wysokie napięcie zasilające. Takie elementarne zabezpieczenie jest przewidziane w konstrukcji. Występuje tu zworka (P502), której zadaniem jest odcięcie stabilizatora napięcia od całej reszty (patrz **rysunek 1**). Po przyłączeniu transformatora (złącze P501) należy zmierzyć napięcie między P504 (punkt odniesienia do pomiarów) a pin 1 w P502 (wyjście zasilacza), jak pokazuje **rysunek 2**. Przy poprawnym napięciu zasilającym +5 V (możne ono się różnić o kilkaset mV z uwagi na tolerancję rezystorów R502 i R503 – rysunek 1) można na stałe wlutować zworkę P502.



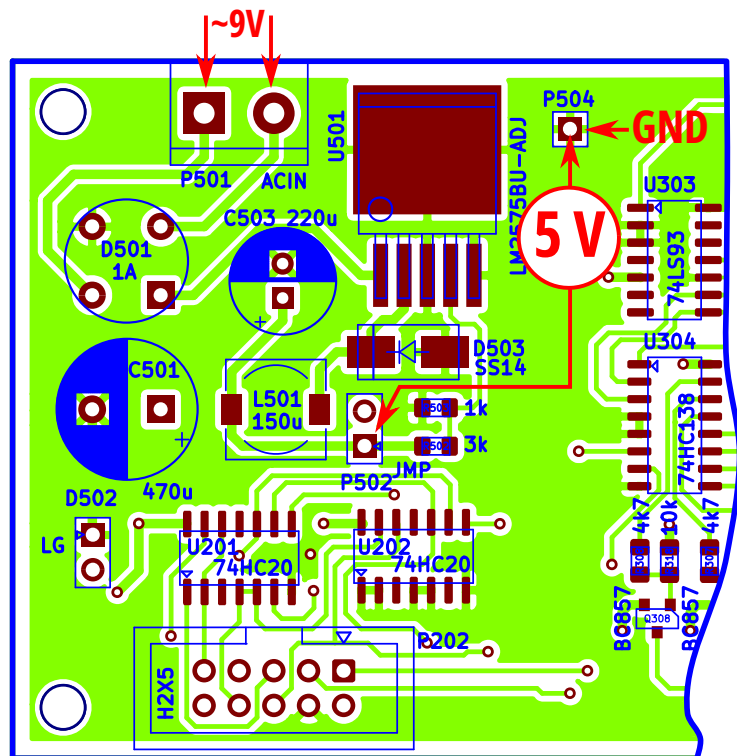
Rysunek 1

W przeciwnym wypadku koniecznością jest sprawdzenie połączeń w obrębie zasilacza (pomijam przypadek pomylnych elementów).

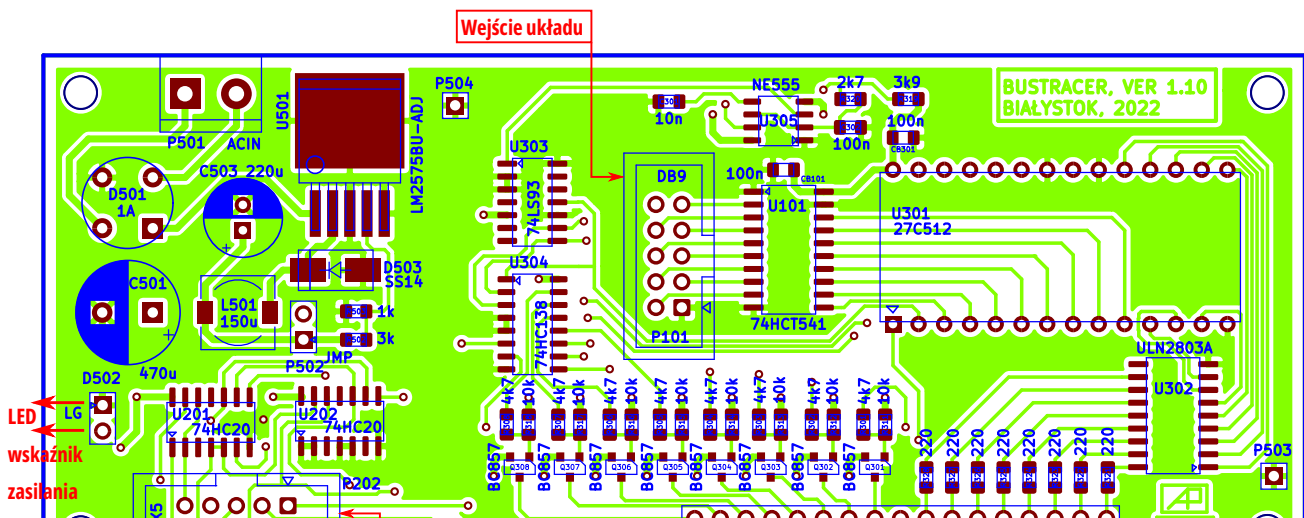
Znaczenie istotnych elementów (z punktu widzenia połączeń między płytkami) pokazuje **rysunek 3**.

Obudowa urządzenia

Ostatnim elementem prowadzącym do powstania diagnostycznego urządzenia jest jego obudowa. Ten mały elektroniczny element jest jednak równie istotny jak wszystkie komponenty elektroniczne tworzące całość. W praktyce amatorskiej często stosowane są jakieś obudowy plastikowe, gdyż ich obróbka mechaniczna nie wymaga zbyt dużego wyposażenia warsztatowego. Zaproponuję rozwiązanie obudowy metalowej, której wykonanie wbrew pozorom nie jest aż tak skomplikowane jak może się wydawać. Koncepcja takiego rozwiązania była opisana w artykule *Metalowe obudowy do urządzeń elektronicznych* (ZE, maj 2024).

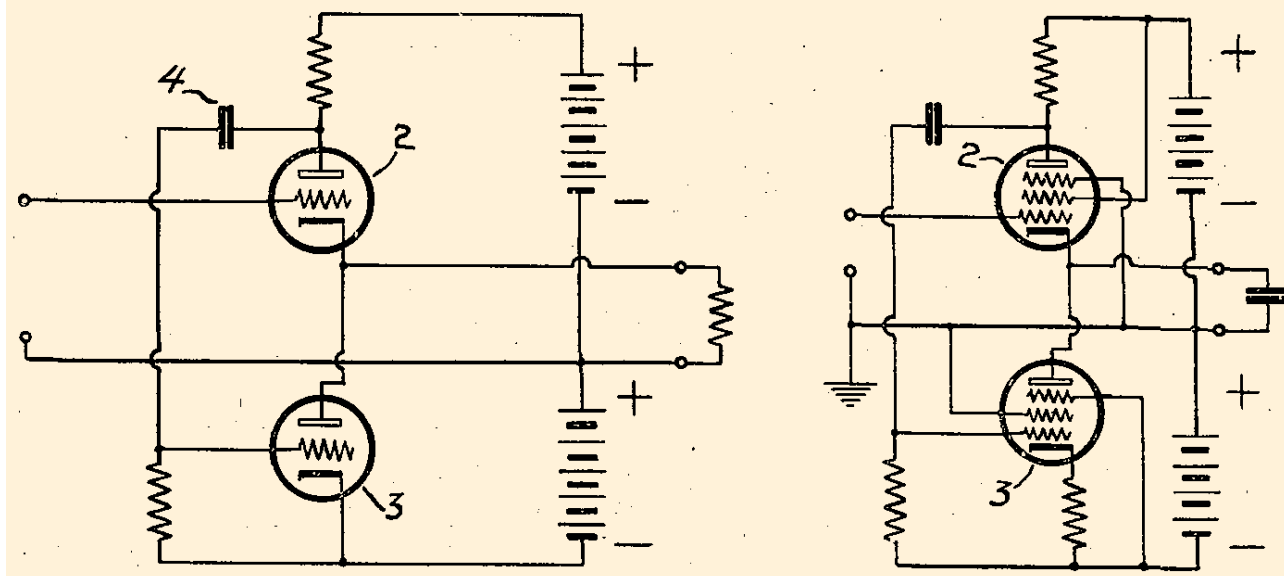


Rysunek 2



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

Sept. 19, 1944. E. L. C. WHITE 2,358,428
THERMIONIC VALVE AMPLIFIER CIRCUIT ARRANGEMENTS
Filed May 15, 1943



Lampowe konfiguracje z dwiema triodami (2)

Kontynuujemy omawianie i zgrubną analizę przedwzmacniaczy lampowych, zawierających dwie lampy triody, pracujące w różnych konfiguracjach, także jako obciążenie. Podstawą są triody, ale w artykule naświetlone jest też szersze tło, związane z pentodami a także elementami półprzewodnikowymi.

Wzmacniacz kaskadowy – kaskoda
Katodyna – co to tak naprawdę jest?
Inne koncepcje dwulampowe

Wtórnik z dynamicznym obciążeniem
Wtórnik White'a
PSRR i wtórnik Aikido

We wcześniejszych artykułach tej serii omówiłem podstawowe, proste układy pracy triod: ze wspólną katodą i ze wspólną anodą, a po części też ze wspólną siatką. W poprzednim artykule zacząłem nalizować układy, konfiguracje wzmacniające z dwiema triodami. Na razie omówiłem tam tylko najprostsze konfiguracje dwutriodowe i zachęciłem do samodzielnych eksperymentów. Analogicznie jest w poniższym artykule. Po pierwsze opisuję następne, ciekawsze konfiguracje układowe z dwiema triodami. Po drugie pokazuję, że w sumie to są zagadnienia łatwe, bo jest to po prostu zestawienie dwóch prostszych, „pojedynczych” konfiguracji. Po trzecie zachęcam do praktycznych eksperymentów, bo to najważniejsze.

Podkreślam, że konfiguracje z dwiema triodami tak naprawdę są prostym połączeniem układów pojedynczych, które omawiałem we wcześniejszych artykułach. Wcale nie są tajemnicze, tylko trzeba dobrze zrozumieć podstawowe konfiguracje z jedną triodą (może też wrócić do tych wcześniejszych artykułów). Nieprzypadkowo poświęciłem im wcześniej tak dużo uwagi. Zwracam uwagę, że analizujemy połączenia tych prostszych konfiguracji niemalże na zasadzie „każdy z każdym”. Niemalże wszystkich, bo w poprzednim artykule stwierdziłem, iż w układach audio nie bardzo ma sens konfiguracja, gdzie na wejściu jest wtórnik (układ ze wspólną anodą), który współpracuje ze wzmacniaczem ze wspólną katodą.

Wzmacniacz kaskodowy – kaskoda

Ma natomiast pewien sens połączenie układów ze wspólną katodą oraz ze wspólną siatką. To jest tak zwany układ kaskodowy. Zapamiętaj, że kaskoda (nie kaskada) to połączenie układu ze wspólną katodą z układem ze wspólną siatką, według idei z **rysunku 1**.

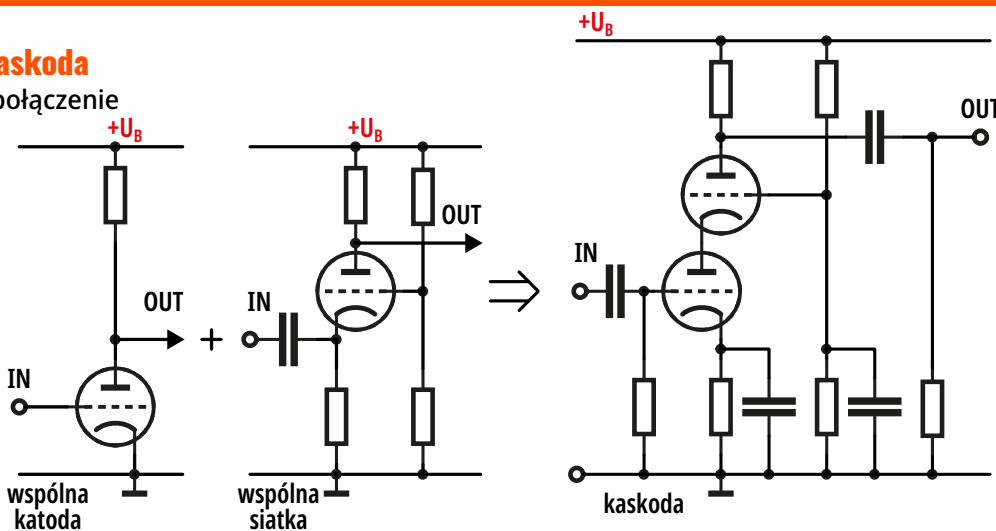
Ta konfiguracja jest (była) często stosowana w układach w.cz., bo pozwala na pracę lampy (triody) przy dużo większych częstotliwościach, likwidując tak zwany efekt Millera. Natomiast w urządzeniach audio nie ma ona istotnych zalet. Duża jest rezystancja wyjściowa, taka jak w układzie ze wspólną katodą. Wypadkowe wzmocnienie napięciowe może być trochę większe, niż jednej lampy, ale niewiele większe. Właściwości takiej konfiguracji przypominają właściwości lampy pentody.

Wzmacniacze kaskodowe z dwiema triodami są bardzo rzadko spotykane w urządzeniach audio, dlatego nie będziemy zagłębiać się w szczegóły.

Katodyna – co to tak naprawdę jest?

Z **katodyną** jest pewien problem. Otóż w źródłach polskich tak nazywa się połączenie układu wzmacniającego ze wspólną anodą (wtórnika) z układem ze wspólną siatką, według **rysunku 2**. W układach audio taka konfiguracja wzmacniająca nie ma istotnych zalet i jest rzadko spotykana.

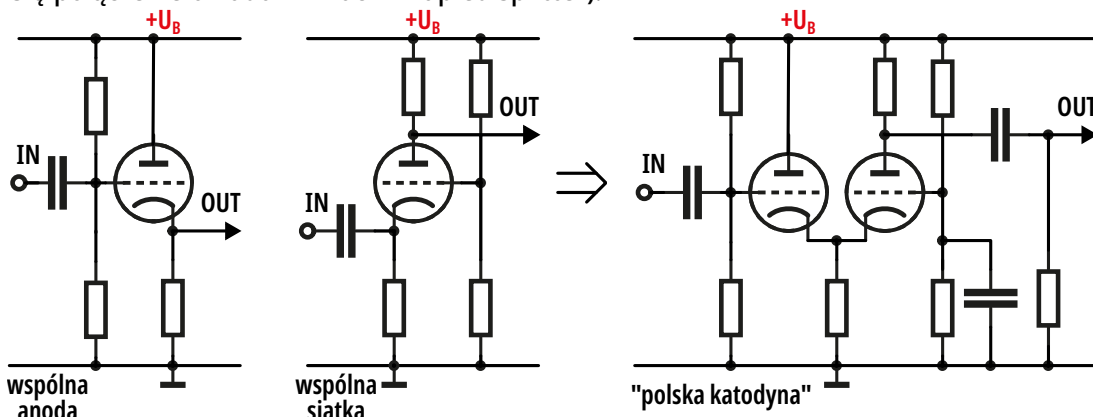
Natomiast w literaturze światowej katodyną nazywa się przede wszystkim układ z jedną lampą i dwoma wyjściami, według koncepcji z **rysunku 3**. Taka konfiguracja służy do uzyskania dwóch bliźniaczych sygnałów o jednakowej amplitudzie i przeciwnych fazach. W układach audio jest wykorzystywana w niektórych przeciwsobnych wzmacniaczach mocy (czyli push-pull) do wy-



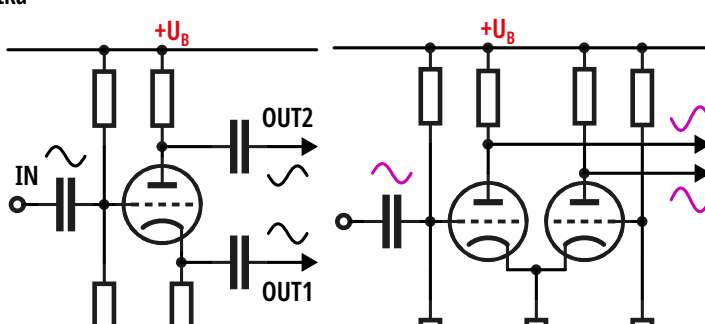
Rysunek 1

i katody powinny mieć identyczne wartości. W rzeczywistości rezystory te trzeba tak dobrać, żeby podczas pracy amplitudy sygnałów na obu wyjściach były jednakowe.

Do takiego samego celu, czyli do wytwarzania jednakowych sygnałów o przeciwnych fazach, służy też konfiguracja z dwiema lampami, według **rysunku 4**. Rysunek ten jest o podobny do rysunku 2, tylko w anodzie pierwszej lampy też jest umieszczony rezystor i dzięki temu układ ma dwa wyjścia, na których występują sygnały o przeciwnych fazach. W zagranicznej literaturze taki układ nie jest jednak nazywany katodyną, tylko najczęściej parą z długim ogonem (long tail pair – LTP) ewentualnie rozdzielaczem sprzężonym katodowo (cathode coupled splitter).



Rysunek 2



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



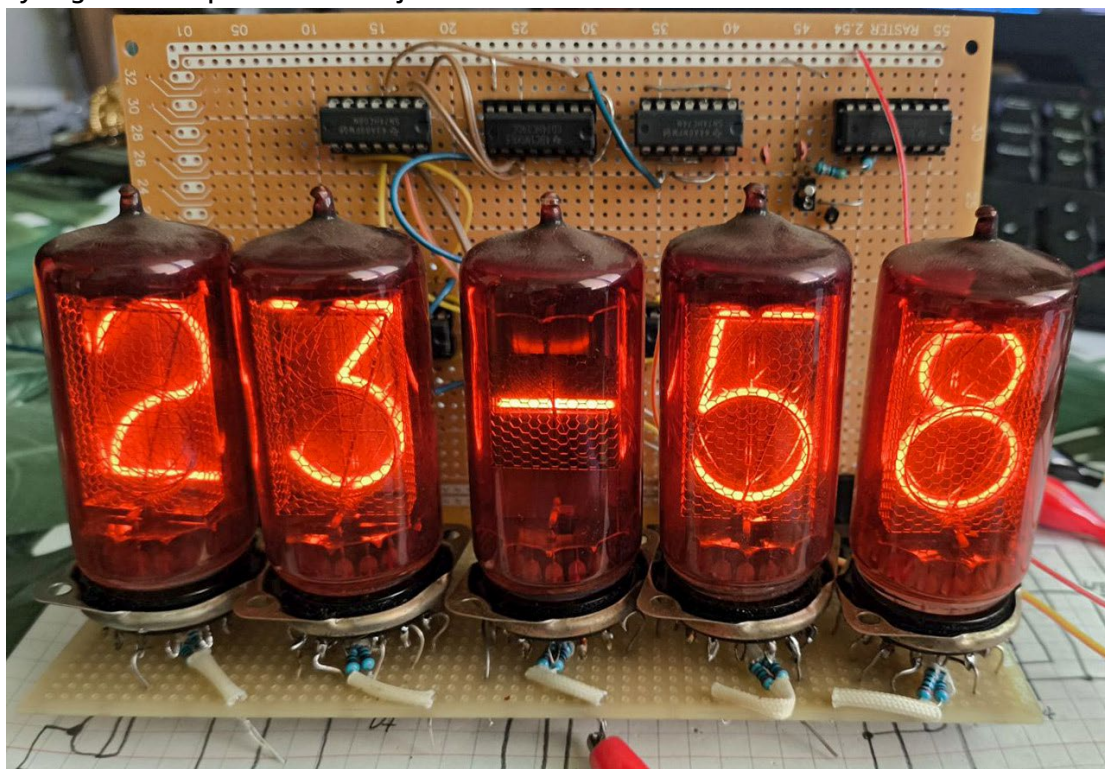
Wspólnie projektujemy: Wykorzystanie dawnych wyświetlaczy

Zadanie z cyklu „Wspólnie projektujemy”, oznaczone YK032, miało tytuł: „Zaproponuj wykorzystanie dawnych wyświetlaczy: NIXI, VFD, numitronów, żarówek, lamp oscyloskopowych i podobnych”. Poniżej przedstawione jest interesujące rozwiązanie z dużymi lampami, nadesłane przez Konrada Klekota.

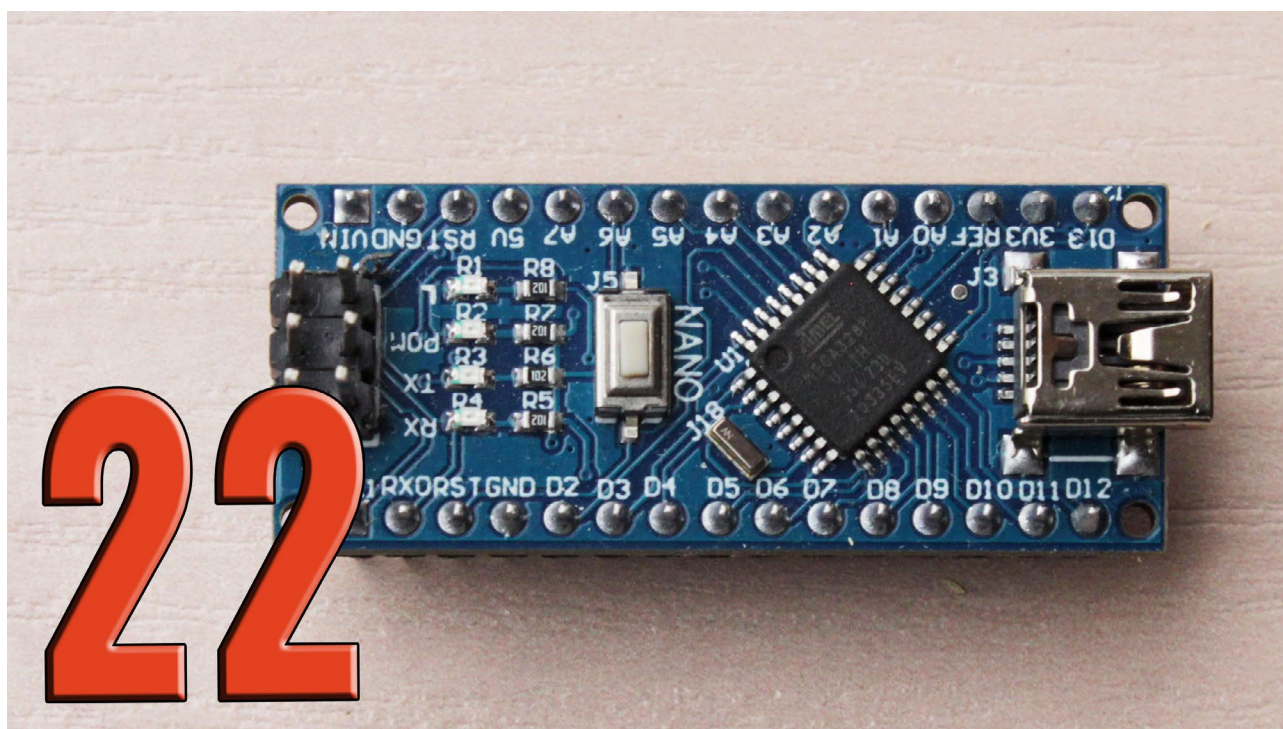
Piszę w związku z artykułem dotyczącym wykorzystania starych wyświetlaczy/lamp. Mam obecnie w trakcie budowy własny zegar na lampach nixie.

Temat, można powiedzieć, wyświechtany do maksimum, chociaż jak przyjdzie co do czego, to niekoniecznie. W założeniach miał to być projekt jak najprostszy, ale bez mikroprocesorów. Wykorzystując to, co mam w szufladzie, dla minimalizacji kosztów. Ponieważ miałem układy 74141, odrzuciłem wszelkie projekty na licznikach 4017. Przeszukując internet, stwierdziłem, że najlepszym wyjściem będzie zastosowanie układów serii HCMOS. Liczniki na

o połowę. Generator kwarcowy 32768 Hz klasycznie na 4060, ale w wersji CMOS, tak wyjątkowo, bo wersja HC siła zakłóceniami aż do fal UKF...



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



Mikroprocesorowa ośła łączka, część 22

Tworzenie oprogramowania dla mikrokontrolerów, jako warunek konieczny wymaga znajomości samego języka jak i posługiwania się oprogramowaniem narzędziowym. W kilku poprzednich częściach było podsumowanie samego języka, ale czy ono wystarczy?

Proste błędy

Zachęcam Czytelników do własnych prób tworzenia oprogramowania. Zdobyta w ten sposób wiedza jest bezcenna. Tworząc program można napotkać wiele problemów, których pokonanie znacząco wpływa na samodzielność. Ogólnie programowanie to sztuka łączenia kilku elementów w jedną całość. Jednym z oczywistych jest sama znajomość języka, ale istnieje coś ważniejszego – umiejętność tworzenia algorytmów. Pod tym pojęciem kryje się przepis na niezbędne działania: co i w jakiej kolejności wykonać, by program realizował założone cele. Skonstruowany algorytm należy zapisać jako program, posługując się możliwościami danego języka (w naszym przypadku języka C).

W dotychczasowych rozważaniach wystarczająco dokładnie zostały opisane dwa główne elementy dotyczące interakcji użytkownika z mikrokontrolerem:

Trudniejsze błędy

możliwość „zadania” mikrokontrolerowi informacji oraz możliwość poinformowania użytkownika o efektach działania programu. Do pierwszej grupy można zaliczyć różnego rodzaju klawiatury. Stosowane rozwiązania dotyczyły wariantów prostych (gdzie piny portów były łączone przez przyciski do masy) oraz klawiatur matrycowych pozwalających na obsługę dużo większej liczby przycisków za pomocą niewielkiej liczby pinów portów (jak przykładowo klawiatury telefoniczne). Wraz z opisem różnych wariantów klawiatur udostępnione były dedykowane im pliki zawierające ich programową obsługę. Po przeciwległej stronie interakcji użytkownika oraz mikrokontrolera występują wyświetlacze. Były używane wyświetlacze 7-segmentowe LED, pozwalające na prezentację danych liczbowych (plus trochę innych znaków dających się utworzyć przy użyciu siedmiu segmentów).

Bardziej uniwersalne są moduły LCD, które pozwalają wyświetlić każdy znak ze zbioru ASCII.

Istniejąca programowa obsługa obu „styków” mikrokontrolera z otoczeniem stanowi bazę do innych ćwiczeń. W tym miejscu mogę zaproponować Czytelnikom przykładowo stoper, który mierzy czas od jednego zdarzenia (naciśnięcia przycisku typu start) do drugiego (naciśnięcia przycisku stop).

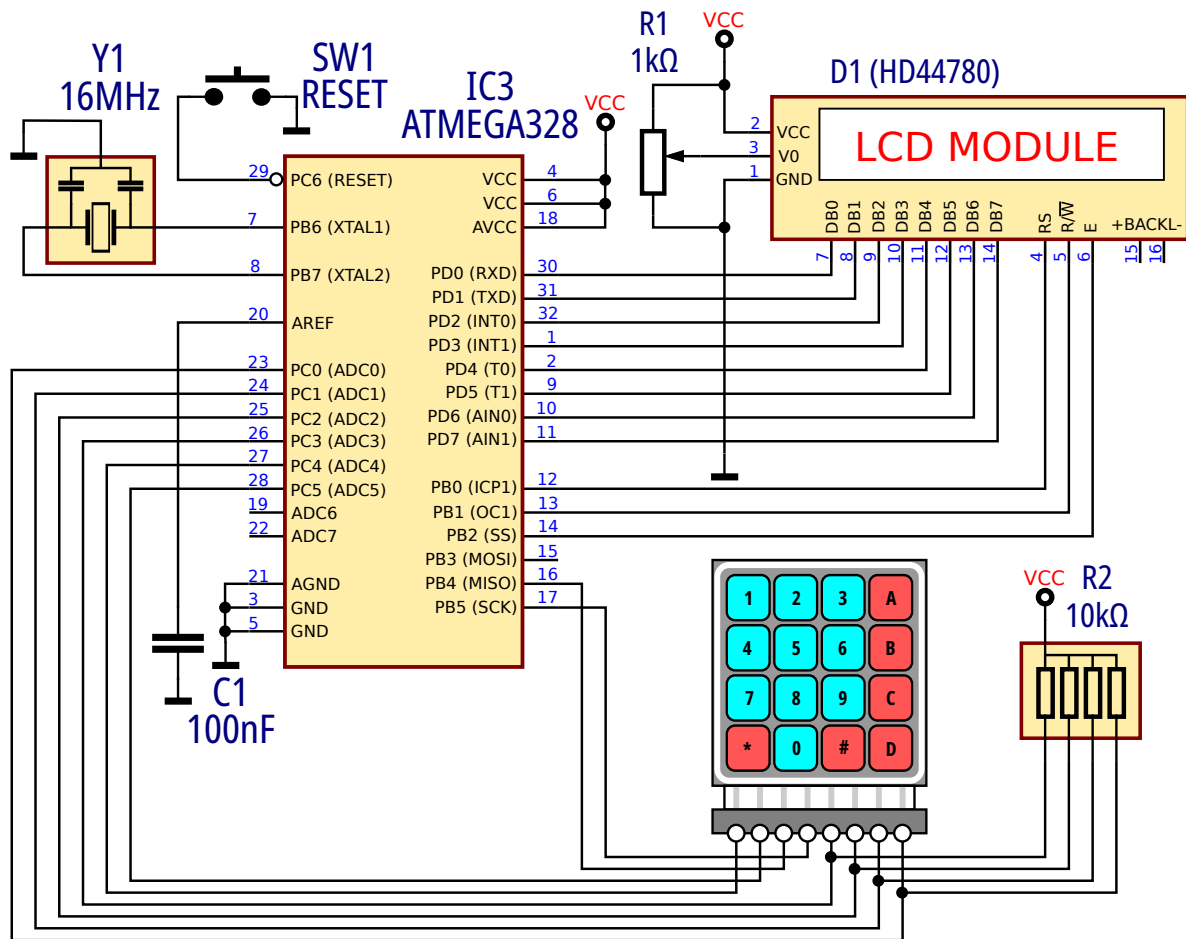
W ten sposób można zbudować „zabawkę” do pomiaru refleksu.

Tworząc jakkolwiek program nie sposób ustrzec się od błędów i pomyłek dotyczących opisu algorytmu w języku C (pomijam błędy samego algorytmu). Wiele z tych błędów daje się łatwo ustalić i można wnieść korektę do programu. Jednak niektóre z nich nie są aż tak oczywiste i trzeba się trochę natrudzić aby zlokalizować „wadliwe” miejsce w programie.

Bazą do tematu będzie obsługa klawiatury matrycowej, opisana w artykule *Mikroprocesorowa ośła łączka*, część 15 (ZE 8/2025), dotycząca układu pokazanego na **rysunku 1**. Materiały dodatkowe związane z tamtym artykułem nie zawierają niżej opisanych błędów. Materiały dodatkowe do aktualnej części są odpowiednio spreparowane i dołączone do bieżącego artykułu.

Proste błędy

Rzeczywista technika pisania tekstu programu nie jest prostym wpisaniem odpowiedniej treści jako tekstu programu. Przy tworzeniu funkcji progra-

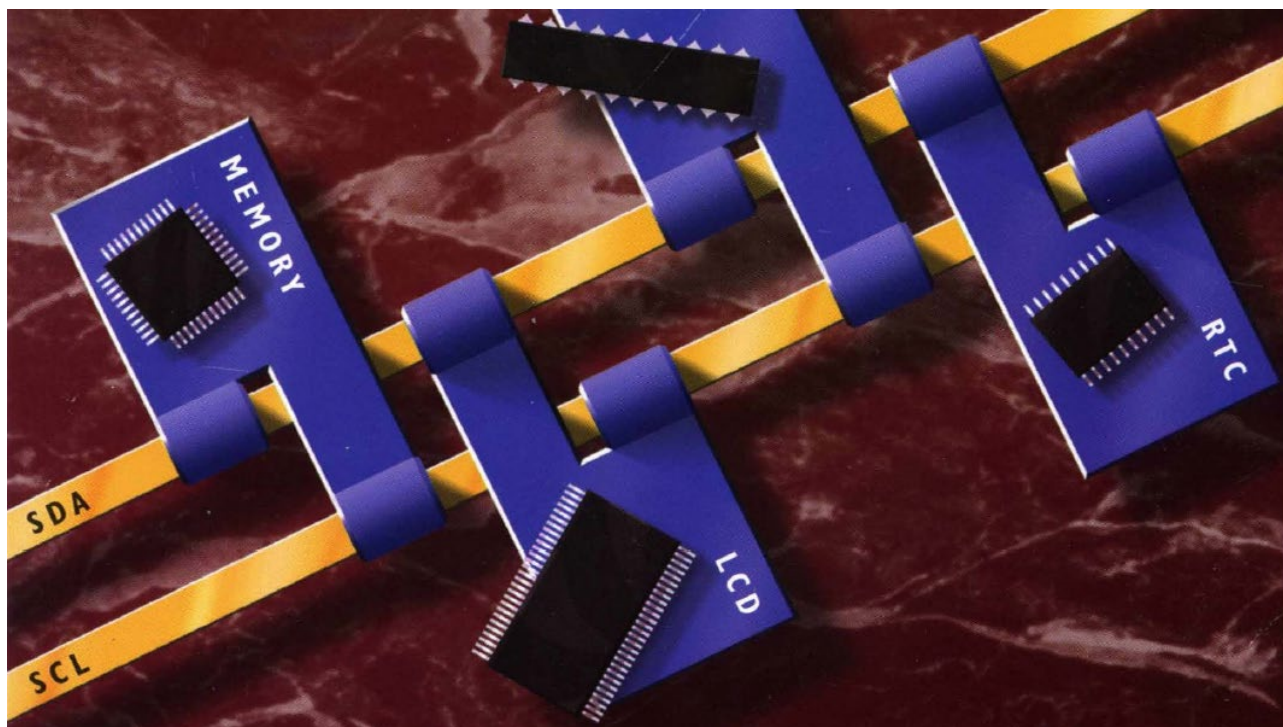


Rysunek 1

LCD jakiegoś znaku. Pamiętamy, że mamy do dyspozycji w module (jako plik nagłówkowy *.h oraz implementacji *.c) odpowiednie już gotowe funkcje. Naturalnym jest, że zostają one użyte. Tak powstał program U107 A (dostępny w materiałach dodatkowych, w rzeczywistości odpowiednio spreparowany program U100 A), którego część pokazuje **rysunek 2**. Mamy tu wybrane słowa podkreślone czerwoną linią falistą. W ten sposób kompilator wskazuje nam potencjalne „problemy gramatyczne” języka.

```
U107A - AtmelStudio
File Edit View VAssistX ASF Project Build Debug Tools Window Help
main.c
main.c
D:\AVR_project\U107A\U107A\main.c
static const KeybEncodeElementT Row0EncodeArray [ 5 ] PROGMEM =
{ /* 00 */ { 0xFE, '1' },
/* 01 */ { 0xFD, '2' },
/* 02 */ { 0xFB, '3' },
/* 03 */ { 0xF7, 'A' },
/* 04 */ { 0x00, 0x00 } };
static const KeybEncodeElementT Row1EncodeArray [ 5 ] PROGMEM =
{ /* 00 */ { 0xFE, '4' },
/* 01 */ { 0xFD, '5' },
/* 02 */ { 0xFB, '6' },
/* 03 */ { 0xF7, 'B' },
```

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



Mikroprocesory i mikrokontrolery – interfejs I²C

Obecnie mikrokontrolery stały się praktycznie niezbędnym elementem każdego urządzenia, gdyż realizują przeróżne operacje. Nowoczesne układy integrują w swojej strukturze wiele przydatnych „zabawek”, jednak niektóre muszą być zewnętrznymi elementami do nich przyłączanymi. Jak to zrobić?

[Geneza interfejsu I²C](#)

[Adresacja w I²C](#)

[Rozpoczęcie transmisji w I²C](#)

[Transmisja w I²C](#)

[Ilustracja działania I²C](#)

Technice mikroprocesorowej zawsze towarzyszyły specjalizowane układy do realizacji określonych funkcji. Tu można wskazać przykładowo przetworniki analogowo-cyfrowe lub cyfrowo-analogowe. Wiele z nich obecnie znajduje się w strukturze mikrokontrolera, jednak początki techniki mikroprocesorowej były trochę trudniejsze. Wykształciło się pojęcie portu mikroprocesora, który umożliwiał mikroprocesorowi interakcję z otoczeniem. Dzisiaj odpowiednik takich portów jest wbudowany w sam mikrokontroler. Przykładem może być choćby przetwornik ADC występujący w mikrokontrolerach AVR. W programie wystarczy wczytać dane z odpowiedniego rejestru mikrokontrolera, by uzyskać cy-

frowy wynik odzwierciedlający wielkość analogową. Na początku burzliwego rozwoju mikroprocesorów również istniała taka potrzeba, tylko wtedy odbywało się to trochę inaczej. Sam mikroprocesor „posiadał” pamięć na program czy pamięć operacyjną, również jako elementy zewnętrzne. Nie ma powodu by inne układy (porty) działały na odmiennych zasadach. Miały (najczęściej) 8-bitową szynę danych, za pomocą której następowała wymiana danych między mikroprocesorem a dołączonym podzespołem. To oczywiście wymaga 8-bitowego interfejsu pozwalającego na te operacje (tego typu połączenie można określić jako 8-bitowe równoległe: następuje jednoczesne przesłanie wszystkich danych).

Te zintegrowane w jedną strukturę mikrokontrolera nadal działają w ten sposób. Różnorodność zastosowań jednak nie pozwala na integrację w jednym układzie wszystkiego, co kiedykolwiek może być potrzebne projektantowi.

Weźmy przykładowo pod uwagę mikrokontroler ATMEGA32 – ma on do dyspozycji cztery porty po osiem bitów (więcej nie ma). Chcąc cokolwiek przyłączyć można wykorzystać do przesyłania danych cały port, ale to utylizuje jego wręcz bezcenne zasoby, piny do interakcji ze środowiskiem. Powstała więc koncepcja zamiany równoległego przesyłania danych na rozwiązanie szeregowe. To pozwala jednocześnie minimalizować liczbę niezbędnych pinów przeznaczonych do komunikacji. Tu można dostrzec jeszcze jedną istotną cechę, jaką jest niezależność od długości przesyłanego słowa przy tej samej liczbie wykorzystanych pinów. Niestety zawsze jest „coś za coś”. Zamiana interfejsu równoległego na szeregowy znacząco komplikuje i zwiększa czas operacji.

Geneza interfejsu I²C

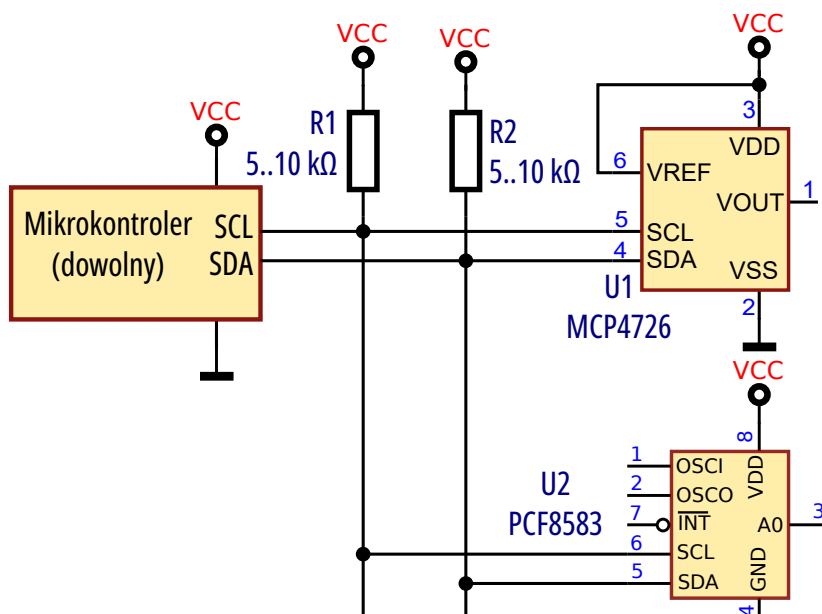
Koncepcja rozwiązania tego interfejsu powstała w firmie Philips Semiconductor i w materiałach tej firmy (a właściwie następników Philips, jak NXP Semiconductors) najlepiej szukać materiałów opisujących ten interfejs (fotografia tytułowa pokazuje okładkę dosyć obszernego katalogu poświęconego układom z interfejsem I²C). Akronym I²C (właściwie IIC pochodzi od ang. Inter-Integrated Circuit) jako dwuprzewodowa, szeregową synchroniczną magistrala komunikacyjna stworzona do przesyłania danych między układami scalonymi na krótkich dystansach używająca tylko dwóch linii: **SDA** (dane) i **SCL** (sygnał taktujący). Krótki dystans należy rozumieć jako taki, który nie wychodzi poza urządzenie, gdyż jego zadaniem jest powiązanie układów w obrębie na przykład telewizora. Dysponując dwoma liniami (SDA, SCL) telewizyjny procesor może przeprogramować głowicę (zmienić program), wnieść korektę do procesora dźwięku (zmienić głośność) i wykonać podobnej klasy działania, wykorzystując do tego celu jedynie dwie linie swoich portów. W każdym przypadku zachodzi potrzeba przesłania kilku/kilku-

Oczywiście różnych układów przyłączonych do mikrokontrolera może być więcej. Powstaje podział na dwie klasy elementów pracujących z interfejsem I²C: strona nadrzędna (master, w tej roli występuje mikrokontroler) oraz podrzędna (slave, jako układ MCP4726 – 12-bitowy przetwornik cyfrowo-analogowy oraz PCF8583 jako RTC – zegar czasu rzeczywistego). Występujące dwa rezystory są niezbędne, gdyż linia danych (SDA) raz jest sterowana przez stronę master (mikrokontroler), a w kilku innych sytuacjach przez stronę slave (przetwornika ADC). To oznacza, przynajmniej po stronie slave, pracę w trybie otwartego drenu (kolektora). Aby dane na linii były stabilne, niezbędny jest rezystor wymuszający w stanach pasywnych wartość logicznej jedynki. W moim przekonaniu rezystor na linii SCL nie jest niezbędny, ale firmowe katalogi dotyczące I²C zalecają taki.

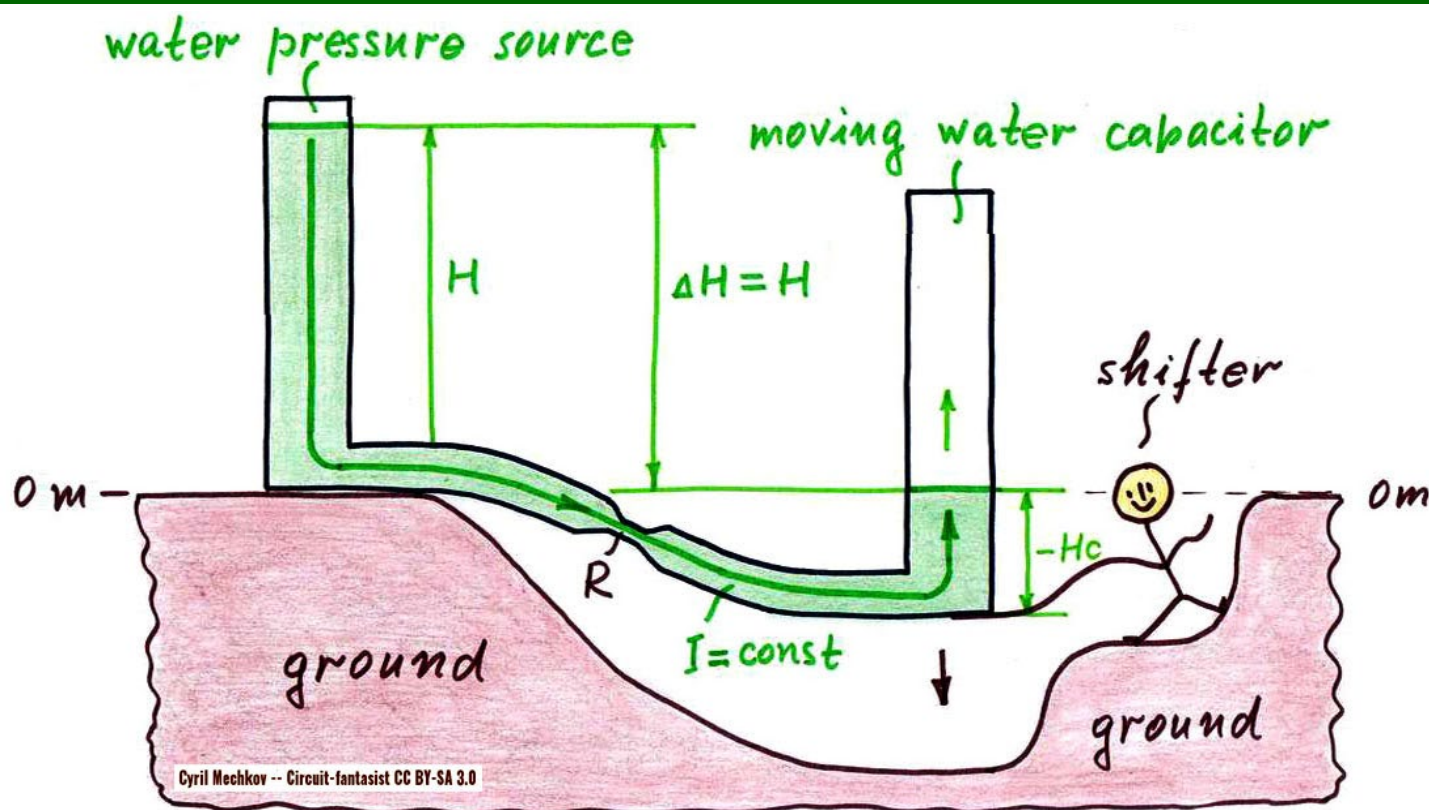
Aby ominąć kwestie patentów (I²C jest opatentowany), niektórzy producenci układów scalonych używają określenia TWI (ang. Two-Wire Interface) dla interfejsu dwuprzewodowego, który jest kompatybilny z I²C.

Adresacja w I²C

Wyobraźmy sobie, że do pary linii I²C jest przyłączonych kilka układów typu slave. Te linie jednocześnie dochodzą do każdego układu. Tu powinno zrodzić się słuszne pytanie: skąd te układy będą „wiedziały”, że dane są przesyłane do nich, skoro wszystkie jednocześnie dostają ten sam sygnał. Specyfiką protokołu I²C jest to, że w pierwszym wysłanym szeregowo bajcie znajduje się adres układu slave. To oczywiście prowadzi do wniosku,



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



Fundamenty elektroniki – analogie hydrauliczne

We wcześniejszych artykułach tego cyklu omówiłem dwa podstawowe wzory, dwa prawa Kirchhoffa oraz bardzo ważną kwestię źródeł napięciowych i źródeł prądowych. W tym artykule zajmujemy się brzemieniami w skutki kwestiami dotyczącymi nauczania elektroniki za pomocą analogii hydraulicznych.

Analogia hydrauliczna – trzy różne wersje
Rezystor w analogii hydraulicznej

Problem ze źródłem prądowym
„Thévenin i Norton” w największym skrócie

Podstawy elektroniki przewodowej, zwłaszcza te dotyczące prądu stałego, są dość proste. Ale elektronika bezprzewodowa, a raczej elektromagnetyzm to bardzo szeroka dziedzina, a szczegóły są skomplikowane i trudne. Tych szczegółów nie można poznać i zrozumieć bez znajomości matematyki i to wyższej matematyki. Naprawdę bez wyższej matematyki się nie obejdzie, jeśli ktoś faktycznie chce zrozumieć elektronikę, zwłaszcza elektronikę bezprzewodową.

Jednak w wielu prostszych „zastosowaniach przewodowych” można te trudniejsze szczegóły pominąć, zaniedbać i przedstawić zagadnienie w sposób uproszczony. I tak robiono to od dawna.

Już w połowie XIX wieku wykorzystywano tak zwaną **analogię hydrauliczną**, która wtedy ułatwiała zrozumienie nowej i ogromnie tajemniczej wówczas elektryczności. Ale wtedy zupełnie nie rozumiano prawdziwej natury elektryczności. W połowie XIX wieku analogia hydrauliczna wydawała się znakomita, pełna i niesprzeczna z prawdą. Ale pod koniec XIX wieku, za sprawą takich tęgich umysłów jak Faraday, Maxwell i Heaviside okazało się, że prawda jest inna, dziwna, nieporównanie trudniejsza. Analogia hydrauliczna na pewno nie jest „całą prawdą i fundamentem elektryczności”, niemniej warto ją wykorzystywać, ale ze świadomością jej wad, ułomności i pułapek.

Analogia hydrauliczna – trzy różne wersje

Analogia hydrauliczna jest bardzo prosta, ale niestety jest tylko częściowa – ułatwia wprowadzenie uchwycenie aspektów najprostszych, ale niestety mocno utrudnia zrozumienie kwestii trudniejszych. W analogii hydraulicznej odpowiednikiem prądu elektrycznego I jest po prostu przepływ wody. Czym więcej (wody) przepływa tym większe natężenie I . **UWAGA!** Wielkość, czyli **natężenie prądu nie jest prędkością ruchu, tylko ilością wody przepływającą w jednostce czasu**. Jeśli rura jest gruba, czyli ma duży przekrój, to dana ilość wody w jednostce czasu będzie przepływać z małą prędkością. A jeżeli rura jest cienka, to prędkość musi być większa, żeby uzyskać przepływ danej ilości wody w jednostce czasu. To prawda, ale nie mieszaj prędkości z natężeniem prądu. W elektronice interesuje nas prąd (natężenie prądu), a nie prędkość ruchu nośników (elektronów), która jest śmiesznie mała.

Odpowiednikiem napięcia elektrycznego U w analogii hydraulicznej jest ciśnienie wody, a także... wysokość. **Napięcie słusznie możemy traktować jako „ciśnienie elektryczne”**. Można je też traktować jako „wysokość elektryczną”, a raczej jako „różnicę wysokości elektrycznych”, co ma związek z *potencjałem elektrycznym*, ale *potencjał* to dość trudne pojęcie, do którego mogę wrócić na życzenie Czytelników.

W analogii hydraulicznej odpowiednikiem przewodów elektrycznych (drutów) są rury, w których płynie woda. Podobnie jak przewody mają różną grubość (średnicę, przekrój poprzeczny), tak i rury w analogii hydraulicznej mają różną średnicę (przekrój), co dość dobrze obrazuje zależność rezystancji przewodów od przekroju i długości.

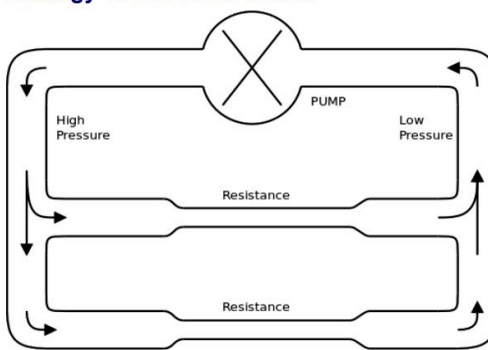
Analogia hydrauliczna ma przynajmniej dwie wersje. Jedna to wersja „ciśnieniowa” „pozioma”, druga – wersja „wysokościowa”, „pionowa”.

Analogia ciśnieniowa – pozioma. Mamy tu zamkniętą instalację, w której woda krąży dzięki obecności pompy – prościutki przykład na **rysunku 1**. Lepszy obrazek pokazujący taką analogię to **rysunek 2**.

Analogia wysokościowa – pionowa. Tutaj mamy układ hydrauliczny z rurami, ale po

<https://pfnicholls.com/Electronics/currentandvoltage.html>

Analogy 1: The water circuit



This analogy considers water flowing through a water circuit in a central heating system.

Analogy: The water being pumped round a circuit. A pump increases the water pressure. Radiators and valves prevent the water from flowing from the pump to the other is analogous to the EMF in the circuit. As water flows through the radiators, the pressure drops, analogous to the P_d across a battery. It divides a circuit into two parts, illustrating how electrical current behaves in a circuit just as the electrical current is measured.

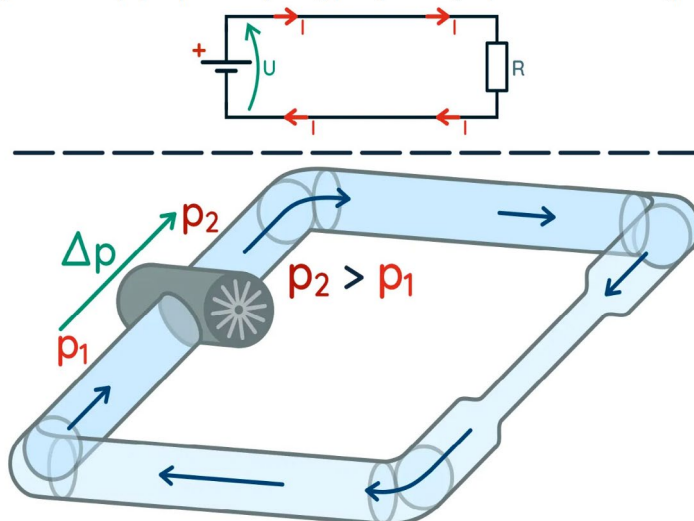
Good For: Understanding current flow in a circuit.

Rysunek 1

<https://forbot.pl/blog/czym-jest-prad-elektryczny-oto-czeste-bledne-wyobrazenia-id51841>

Analogia hydrauliczna (płaska)

Niestety, analogia „wysokościowa” nie pasuje do tego, że prąd płynie zawsze w zamkniętych obwodach – pętach. Odrobinę lepsza pod niektórymi względami jest analogia hydrauliczna ciśnieniowa „płaska”.

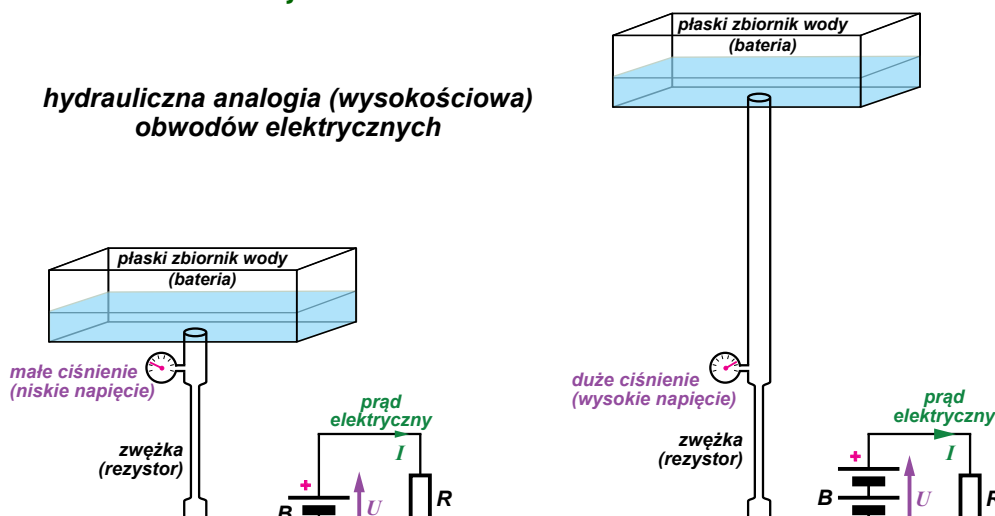


W „płaskiej” analogii hydraulicznej nie mówimy o wysokości, tylko o ciśnieniu i różnicy ciśnień

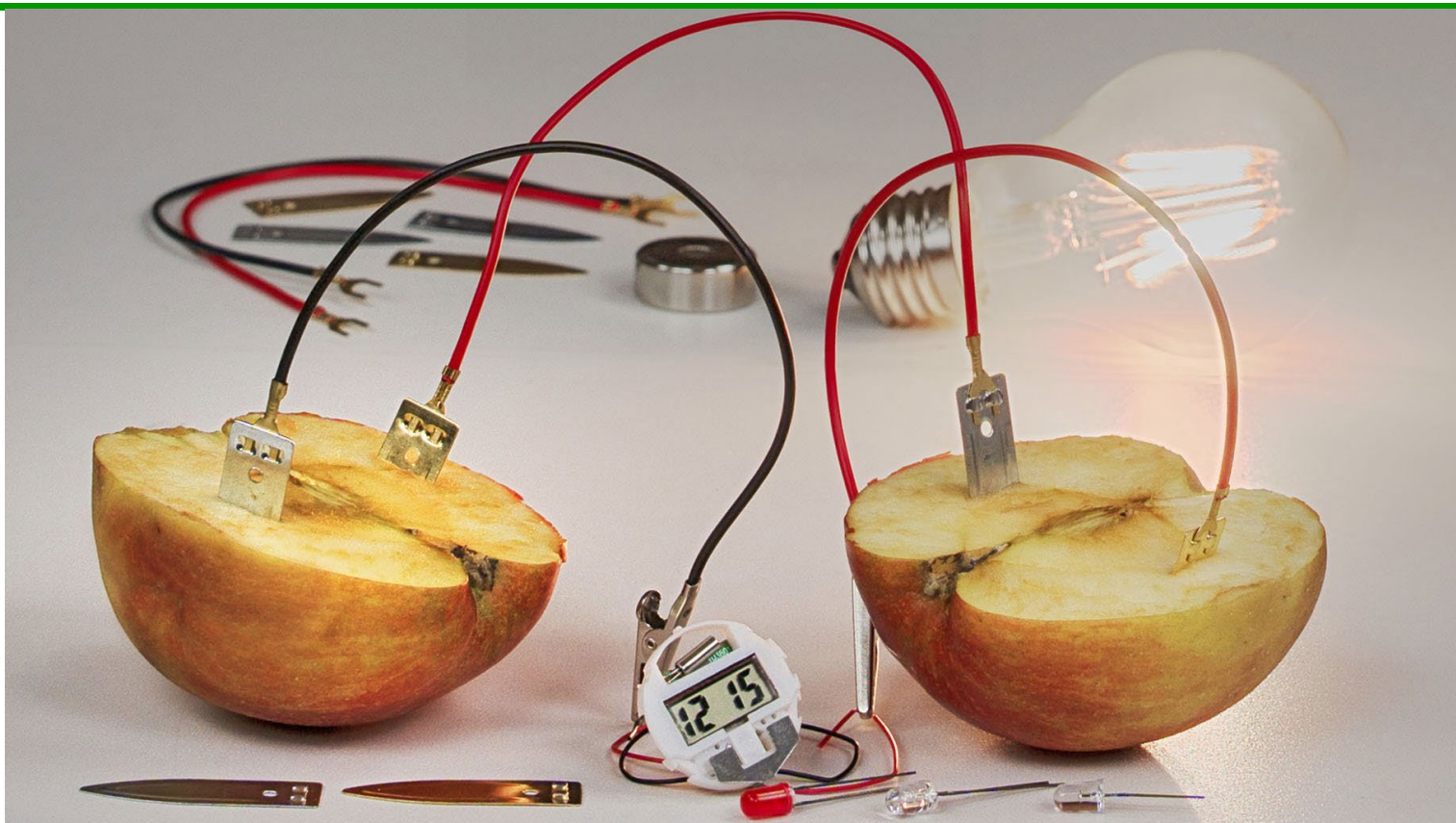
Przewody w takiej analogii to rury, które ułożone są płasko, na jednym poziomie, a odpowiednikiem baterii nie jest zbiornik, lecz pompa o specyficznych właściwościach. Tego typu model oprócz pewnych istotnych zalet ma też poważne ograniczenia.

Rysunek 2

hydrauliczna analogia (wysokościowa) obwodów elektrycznych



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



Fascynujące Przemiany Energii

Proste ogniwa galwaniczne (1)

Własnoręcznie wykonaj rozmaite odmiany ogniw i baterii galwanicznych! Oprócz fascynującej zabawy możesz się wiele nauczyć, i to niezależnie od stopnia technicznego zaawansowania. Artykuł jest przeznaczony także dla początkujących i zawiera wszystkie informacje i wskazówki, które nie zmieściły się w filmach.

A jeśli coś się nie uda? Uniknij błędów!
Moc czyli tempo przemian energii

Proponowane eksperymenty
Pomiary podczas moich eksperymentów

Poniższy artykuł jest rozszerzeniem i uzupełnieniem trzech filmów na moim kanale YT: **B101a**, **B101b**, **B101c**. Jest to też część instrukcji do zestawu **Fascynujące przemiany energii**, zawartego w „pękатыm kuferku” przygotowanym przez sklep Bolland, który genialnie ułatwi przeprowadzenie mnóstwa fascynujących eksperymentów. (**fotografia 1**). Opis „kuferka” jest tutaj.

Nie poprzestań na obejrzeniu i przeczytaniu artykułu! Wykonanie różnych zaskakujących odmian ogniw i baterii daje bardzo dużo satysfakcji! Także zaawansowanym elektronikom, a nie tylko początkującym! Wiem po sobie!



Fotografia 1

Zawartość „**pekatego kuferka**” genialnie ułatwi przeprowadzenie eksperymentów. Do takich eksperymentów potrzebne będą też: kawałek papierowego ręcznika kuchennego (ale nie papieru toaletowego, który się w wodzie rozpada) oraz jakieś owoce lub warzywa. Ja wykorzystałem jabłko i sok z kiszonych ogórków, ale można wypróbować inne, na przykład cytryny, pomarańcze, ziemniaki, itp. Naprawdę warto wykorzystać, pokazane w filmie, zakupione w aptece, tabletki węgla medycznego – pozwolą zbudować stos Volty o zaskakująco dobrych parametrach. Bardzo interesujące są próby budowania ogniw galwanicznych z najróżniejszych metalowych przedmiotów.

Gorąco zachęcam: jeśli tylko masz multimetr, czyli miernik uniwersalny, zbadaj też kwestie napięcia i prądu. A jeszcze lepiej wykorzystać dwa multimetry jako woltomierz i amperomierz, co pozwoli zmierzyć rezystancję, moc, a nawet energię.

A jeśli coś się nie uda? Uniknij błędów!

Ja z powodzeniem przeprowadziłem mnóstwo eksperymentów, co widać w filmie, ale czasem się one nie udają. Prawie zawsze jest to efekt błędu człowieka. Najczęściej jest problem z biegunowością. Akumulatory i baterie mają określoną biegunowość, najprościej „plus” i „minus”. Układy elektroniczne, na przykład zegarek czy brzęczyk oraz liczne elementy, jak dioda LED, nie będą działać przy błędnej biegunowości, czyli gdy zostaną podłączone odwrotnie. Co prawda w prezentowanych w filmie eksperymentach błędna biegunowość nie grozi uszkodzeniem, ale gdybyśmy wykorzystali prawdziwe baterie i akumulatory, to błędna biegunowość może spowodować przepływ dużego prądu i nieodwracalne uszkodzenie.

Najczęstsze błędy to właśnie nieprawidłowa, odwrotna biegunowość. Innym popularnym błędem jest brak styku, brak połączenia elektrycznego. Prąd płynie przez metalowe druty i inne elementy w zamkniętych obwodach – w pętach, więc jeśli pętla będzie przzerwana, to eksperyment się nie uda. Nie zawsze te przerwy (brak dobrego styku) są widoczne.

Jednak czasem zdarzają się najrozmaitsze inne usterki, których nie sposób przewidzieć. Bardzo rzadko, ale

zdarzają się niesprawne „od urodzenia” fabryczne elementy. Znalazienie i wyjaśnienie niektórych usterek sprawia kłopot nawet doświadczonym elektronikom.

Ja przygotowując tę serię ćwiczeń, filmów i artykułów przeprowadziłem mnóstwo testów. Napotkałem niespodzianki, które i mnie zaskoczyły. Przykład na **fotografii 2**. Otóż ogniwa zawierające jako elektrody miedź i cynk bez obciążenia dają napięcie ponad 700 mV (0,7 V). Na fotografii podkładki oraz „ogniwo śrubowe” z dłuższym wkrętem dają prawidłowe napięcie nawet ponad 800 mV. Jedno z krótszych „ogniw śrubowych” też daje prawidłowe napięcie 850,1 mV, natomiast na ostatnim ujęciu widać, że bardzo podobne ogniwo daje tylko 433,6 mV! Długość wkrętów nie ma tu żadnego znaczenia. Najprawdopodobniej w tym ostatnim wkręcie jest jakiś kłopot z niejednorodną lub zanieczyszczoną warstwą cynku, pokrywającego stalowy wkręt.

Nie zdziw się więc, jeśli i Ty napotkasz niespodzianki. Życie jest bogatsze niż podręcznikowa teoria, więc należy nastawić się na niespodzianki i kłopoty. Jeśli wystąpi problem, przede wszystkim trzeba sprawdzić biegunowość oraz upewnić się, że obwód jest zamknięty (nie ma przerw). To najczęściej pozwala znaleźć przyczynę problemu. A jeśli przyczyny nie znajdziesz, radzę nie zniechęcać się, przejść do innych eksperymentów, a do problemu ewentualnie wrócić za jakiś czas.

UWAGA! Absolutnie nie podejmuj żadnych prób z gniazdkiem sieciowym – tam napięcie jest wielokrotnie większe (230 V) i jest śmiertelnie groźne!

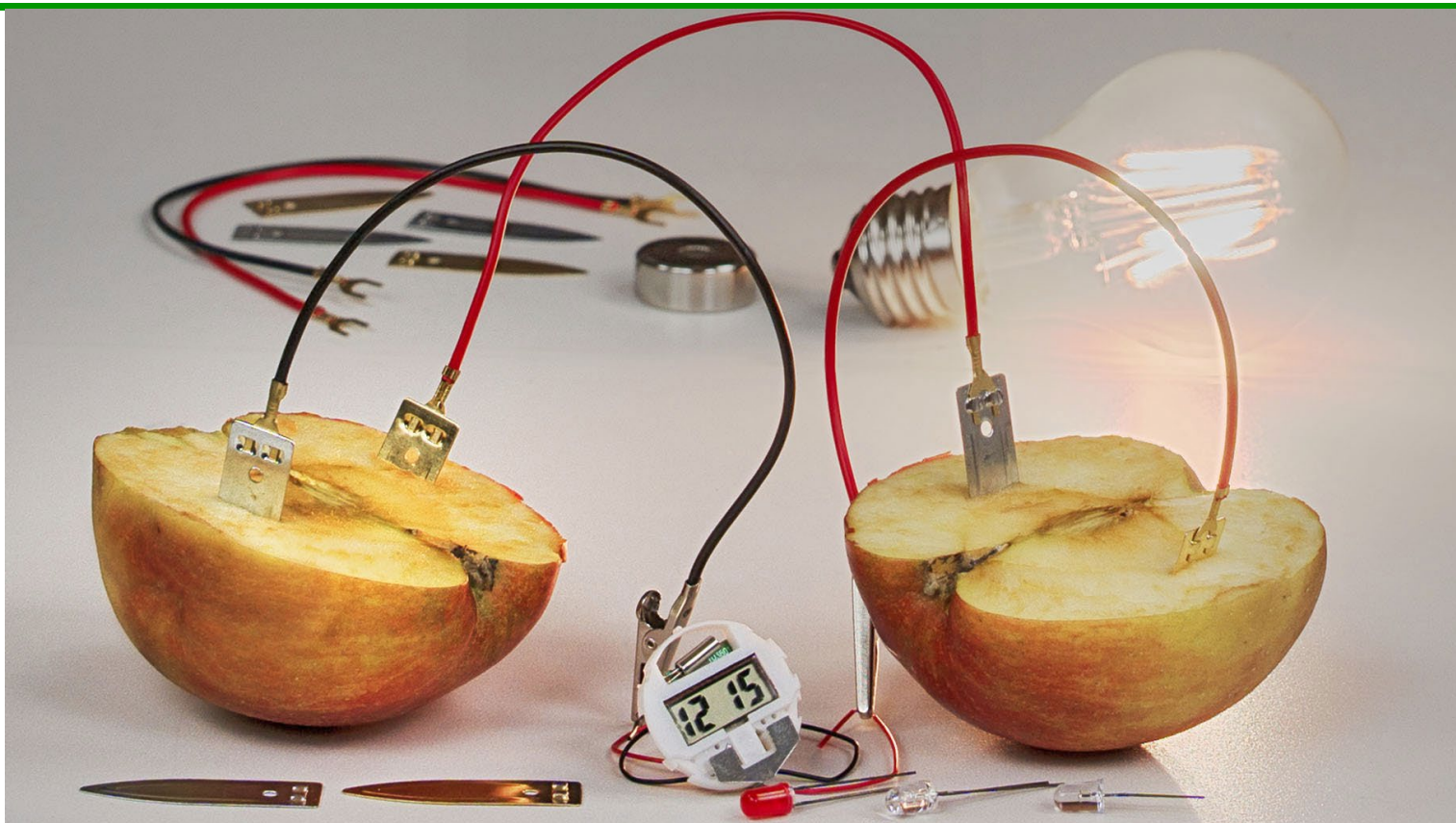
Nie wykorzystuj też popularnych baterii ani zasilacza, bo możesz zepsuć niektóre elementy, np. delikatny zegarek czy diody LED.

Moc czyli tempo przemian energii

To jest cykl o fascynujących przemianach energii, więc bardzo nas interesuje tempo przemian i przekazywania energii, czyli moc. W tym przypadku energię przekazujemy z wykorzystaniem przewodów, więc moc możemy obliczyć bardzo łatwo, mierząc, a potem mnożąc wartości napięcia i prądu. Energia może też być przekazywana w sposób bezprzewodowy, ale wtedy jest zdecydowanie trudniej. A w „przypadkach przewodowych”



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



Fascynujące Przemiany Energii

Proste ogniwa galwaniczne (2)

Własnoręcznie wykonaj rozmaite odmiany ogniw i baterii galwanicznych! W poprzednim artykule tej serii podałem wstępne informacje i wskazówki dotyczące realizacji takich ogniw i baterii, głównie z wykorzystaniem zawartości „pękatego kuferka”. Poniższy artykuł przeznaczony jest dla dociekliwych.

Ogniwa galwaniczne – chemiczne źródła energii
Dla dociekliwych – szereg napięciowy metali

Podsumowanie: Baterie i akumulatory

Poniższy artykuł też jest rozszerzeniem i uzupełnieniem trzech filmów na moim kanale YT: B101a, B101b, B101c. Jest to część instrukcji do zestawu **Fascynujące przemiany energii**, zawartego w „pękatym kuferku” przygotowanym przez sklep Botland, który genialnie ułatwi przeprowadzenie mnóstwa fascynujących eksperymentów. (**fotografia 1**). Opis „kuferka” jest tutaj.

Nie poprzestań na obejrzeniu i przeczytaniu artykułu! Wykonanie różnych zaskakujących odmian ogniw i baterii daje ogromnie dużo satysfakcji! Także zaawansowanym elektronikom, a nie tylko początkującym! Wiem po sobie!



Fotografia 1

Ogniwa galwaniczne – chemiczne źródła energii

Ten cykl dotyczy różnorodnych przemian energii i właśnie zaczęliśmy od **chemicznych źródeł energii, w których energia magazynowana jest we wiąźniach chemicznych**. Budujemy proste (pojedyncze) ogniwa chemiczne, zwane galwanicznymi od nazwiska L. Galvaniego oraz budujemy zestawy ogniw, czyli baterie. Ich działanie można opisać tak: podstawą ogniwa galwanicznego są dwie elektrody wykonane z dwóch różnych metali. Metale te nie stykają się bezpośrednio, gdyż pomiędzy nimi umieszczony jest przewodzący prąd elektrolit – wodny roztwór soli, kwasu lub zasady.

I metal, i elektrolit przewodzą prąd elektryczny, ale na granicy między przewodzącym elektrolitem i metalem tworzy się coś dziwnego. Z uwagi na pewne właściwości chemiczne nie tworzy się tam „przejsię”, tylko coś w rodzaju dziwnej bariery.

Tak, bariery! Mówimy, że na granicy między elektrolitem i jednym metalem tworzy się tak zwane półogniwo, a pomiędzy elektrolitem i drugim metalem – drugie półogniwo. Dwa odpowiednio dobrane półogniwa tworzą (całe) ogniwo i między metalowymi elektrodami występuje niewielkie napięcie elektryczne, które możemy zmierzyć woltomierzem. Napięcie ogniwa zależy głównie od tego, z jakich metali zbudowane są elektrody (jedną z elektrod może być węgiel, który metalem nie jest).

Nasze najprostsze ogniwa mają małe napięcie poniżej 1 wolta. Ale powszechnie wykorzystujemy lepsze ogniwa, które działają na tej samej zasadzie, tylko zawierają inne materiały i dlatego wytwarzają napięcie mniej więcej 1,5 wolta.

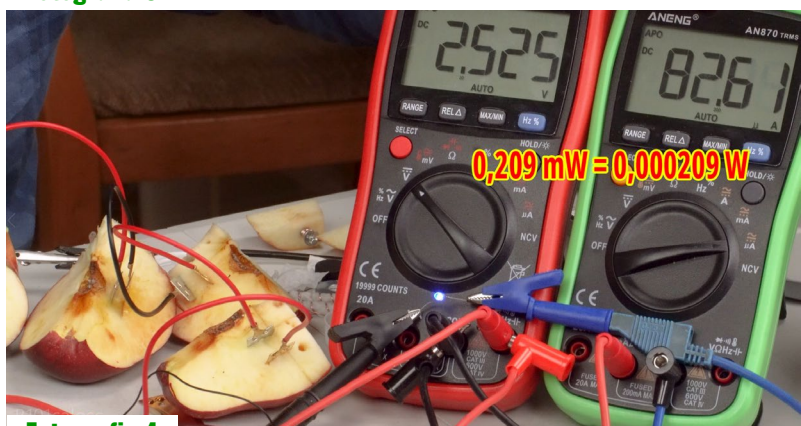
Napięcie (które możemy nazwać „ciśnieniem elektrycznym”) to jeden ważny parametr, a drugi główny parametr to zawartość energii, co nazywamy pojemnością. Ogólnie biorąc, gdy ogniwo pracuje, gdy płynie prąd, wtedy na obu granicach między elektrolitem i metalami zachodzą reakcje chemiczne, zwane **redox** (redukcja i utlenianie). Ścisłej biorąc, energia wiązań chemicznych wykorzystuje oddziaływanie elektryczne, więc energia chemiczna w istocie też jest energią elektryczną, tylko „zaklętą w powłokach atomowych”. Najprościej biorąc, w czasie pracy ogniwa następują takie przemiany chemiczne polegające na pobieraniu oraz oddawaniu elektronów, w których powstają związki chemiczne „o mniejszej energii”, a „różnicę energii chemicznej” ogniwo przekazuje na zewnątrz w postaci energii elektrycznej. W wyniku tych reakcji z obu metali powstają jakieś sole tych metali, a my mówimy że



Fotografia 2



Fotografia 3

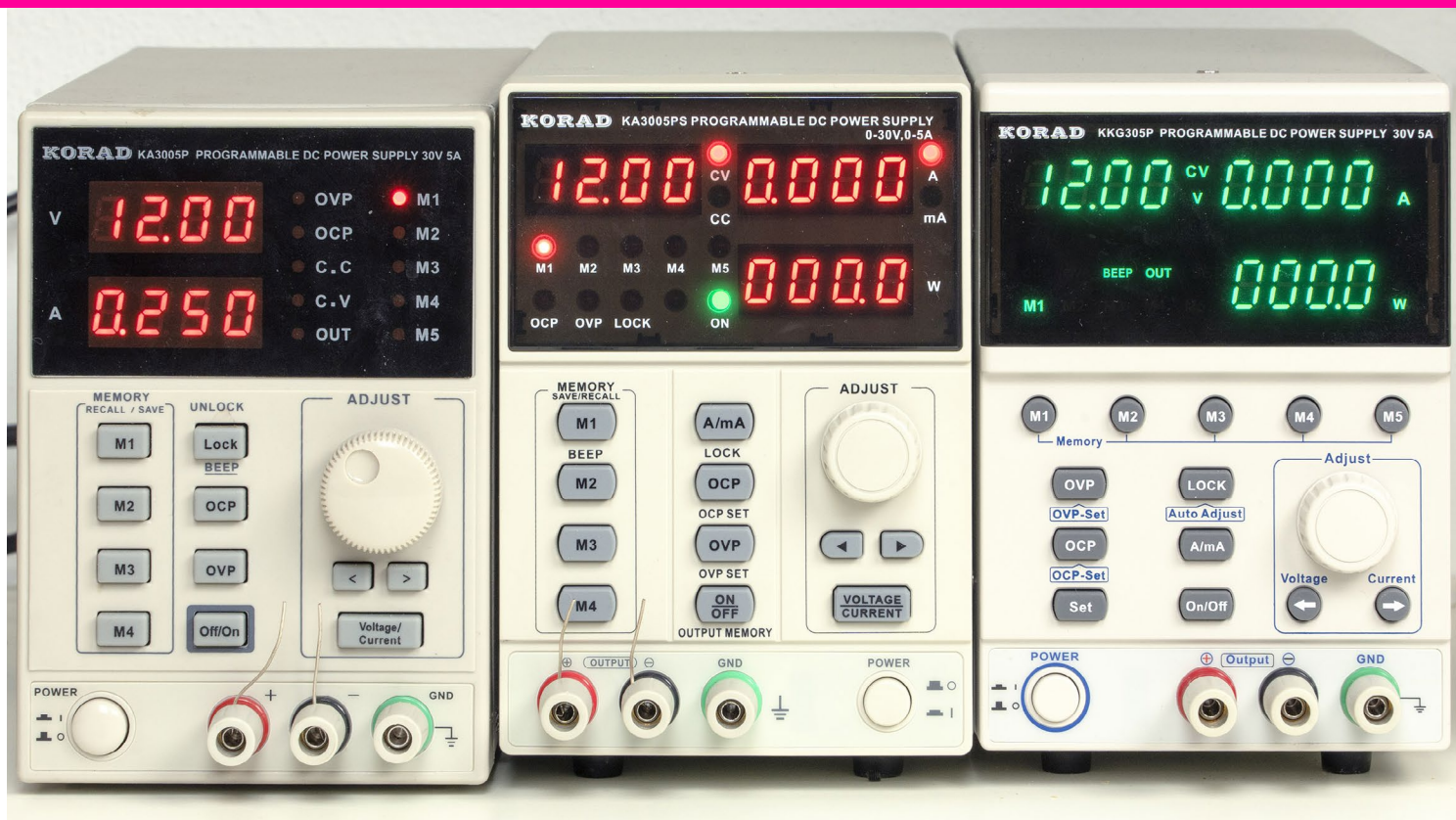


Fotografia 4

(masy) reagujących substancji chemicznych, w tym masy elektrod. Po drugie podczas pracy naszych prymitywnych ogniw przemiany chemiczne pogarszają kontakt między elektrolitem i elektrodami, co skutkuje zmniejszaniem napięcia (i prądu).

Trzecia ważna sprawa to właśnie kwestia prądu. Fabryczne ogniwa i baterie mają zoptymalizowane wszystkie parametry i można z nich pobrać prąd elektryczny o dużej wartości, nawet kilku amperów. A w naszych eksperymentach „wydajność prądowa” jest generalnie bardzo mała, rzędu mikroamperów i pojedynczych miliamperów (**fotografie 2...4**). Głównie z powodu małej powierzchni elektrod. Jeże-

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



Zasilacze KORAD, czyli nie tylko KA3005D

Artykuł powstał przy współpracy firmy i sklepu internetowego **botland.pl**, autoryzowanego dystrybutora wyrobów chińskiej firmy KORAD. Zasilacze KORAD oznaczone KA3005, w ostatnich latach nieprzypadkowo zdobyły ogromną popularność na całym świecie, a dziś dostępne są ich różne wersje, które warto poznać.

Dlaczego akurat KORAD?

Klasyczne liniowe czy impulsowe?

Podstawowe właściwości zasilaczy KORAD

W pracowniach zaawansowanych hobbystów a także wielu profesjonalistów, najpopularniejsze są dziś zasilacze KORAD serii KA3005. Nieprzypadkowo! Naprawdę warto bliżej poznać tę rodzinę zasilaczy, ale nie tylko tę. Otóż założona w roku 2011 chińska firma Dongguan KORAD Technology Co., Ltd. produkuje sporo odmian zasilaczy regulowanych, w tym rodziny U200, KWR100, KKG, KD3000, KD3300, KA3300, czterokanałowe KC i potężne KS. Trzeba również wspomnieć, że KORAD produkuje też „odwrotność zasilaczy”, a mianowicie aktywne elektroniczne obciążenia o regulowanych parametrach – KEL100 oraz KEL2000.

Którą wersję zasilacza KA3005 warto kupić?

Liczne wersje zasilaczy KA3005

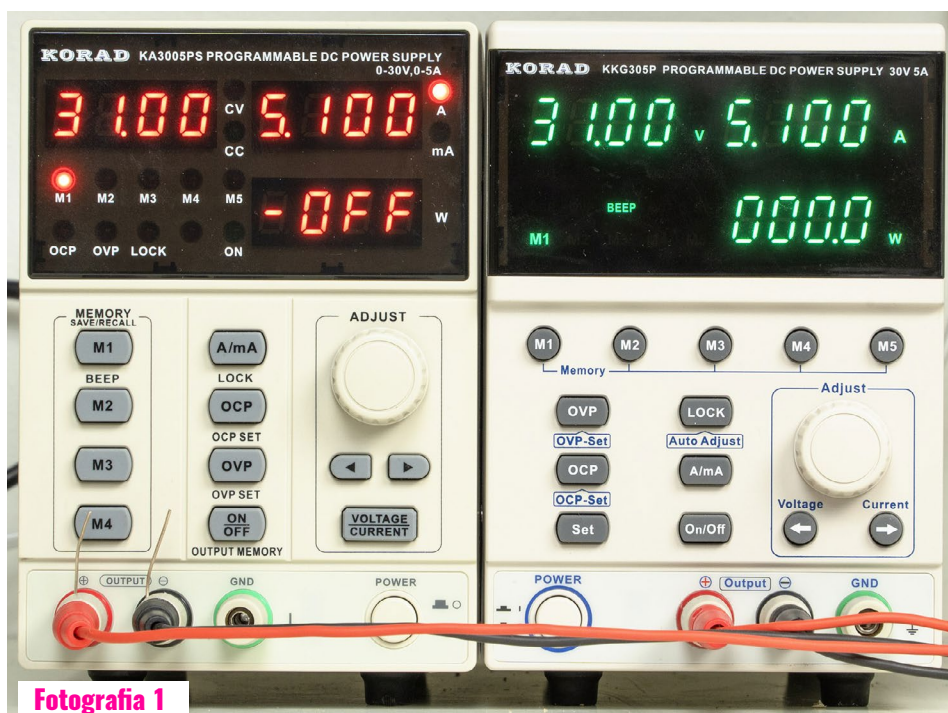
Rodziny i odmiany zasilaczy KORAD

Dlaczego akurat KORAD?

W Internecie można znaleźć nieprzeliczone mnóstwo ofert najróżniejszych zasilaczy, w większości oczywiście chińskich producentów. Mają one różną jakość. Zgodnie z ogólną zasadą, że „nie ma darmowych obiadów”, najczęściej niska cena oznacza też niską jakość, słabe, mocno naciągane parametry oraz kiepską trwałość. KORAD kilkanaście lat temu wszedł na rynek i oferuje przede wszystkim zasilacze KA3005. Wprawdzie nadal mają one pewne wady, ale usunięto „wady wieku dziecięcego”. Stosunek możliwości i właściwości do ceny współczesnych ulepszonych modeli jest naprawdę bardzo dobry.

Dlaczego jednak akurat KORAD KA3005, jeżeli podobne zasilacze można kupić o połowę taniej?

Dlatego, że to naprawdę przyzwolite **zasilacze liniowe**. Są już kilkanaście lat na rynku, zostały gruntownie przetestowane przez użytkowników, którzy wykryli rozmaite błędy i słabości. W kolejnych wersjach usuwano te wady i poprawiano właściwości. Właśnie te zasilacze są powszechnie uznawane za wartę swojej ceny. Co ważne, są dużo tańsze od profesjonalnych zasilaczy renomowanych wytwórców. A **nowe wersje (fotografia 1) mają właściwości i parametry, pozwalających na zaliczenie ich do zasilaczy laboratoryjnych**.



Fotografia 1

Klasyczne liniowe czy impulsowe?

Ale najpierw ważne przypomnienie: dziś dominują zasilacze i ładowarki **impulsowe**. Dostępne są też impulsowe zasilacze, zupełnie niesłusznie nazywane laboratoryjnymi. Przykład (z oferty KORAD) na **fotografii 2**. Są atrakcyjne z wyglądu, lekkie i tanie, natomiast mają istotne wady: jako zasilacze impulsowe wytwarzają zakłócenia radiowe i „przewodowe”, a napięcie wyjściowe nie jest „ładne”, tylko „zaśmieczone” zakłóceniami i szumami.

Owszem, w niektórych zastosowaniach nie ma to znaczenia, ale w pracowni elektronika często potrzebne jest „czyste” źródło zasilania. „Najczystsze” są akumulatory i baterie, ale „czyste” napięcie dają też zasilacze liniowe, które nie zawierają przetwornicy impulsowej, tylko klasyczny, ciężki transformator sieciowy.

Właśnie dlatego **w laboratoriach oraz w pracowniach zaawansowanych i świadomych elektroników nadal dominują klasyczne, ciężkie zasilacze z klasycznym transformatorem sieciowym**. Wśród nich przede wszystkim chińskie zasilacze KORAD rodziny **KA3005**, czyli o napięciu regulowanym w zakresie 0...**30 V** i maksymalnym prądzie nastawianym w zakresie 0...**5 A**. Zdecydowanie mniej popularne są wersje mniejsza: **KA3003** (30 V, 3 A) oraz większe **KA3010** czy **KA6005**. KORAD produkuje też podobne, nieco lepsze zasilacze KKG, w szczególności **KKG305D** oraz **KKG6005P**.

Nie wchodząc w szczegóły: zasilacze impulsowe są dużo gorsze od liniowych. Zgrubna reguła jest taka: zasilacz liniowy **KA3005** (30 V, 5 A) z klasycznym trans-

Podstawowe właściwości zasilaczy KORAD

Dla tych, którzy jeszcze nie znają nowych zasilaczy KORAD **KA3005**, oraz **KKG305** (fotografia 1) przedstawię ich ogólne cechy.

Otóż mają one jedno pokrętło – nie jest to potencjometr, tylko enkoder cyfrowy. Pozwala on na ustawić zarówno napięcie, jak i prąd maksymalny, a potem odczytywać ich wartości podczas pracy. Trzy wyświetlacze pokazują napięcie, prąd oraz moc pobieraną z zasilacza.

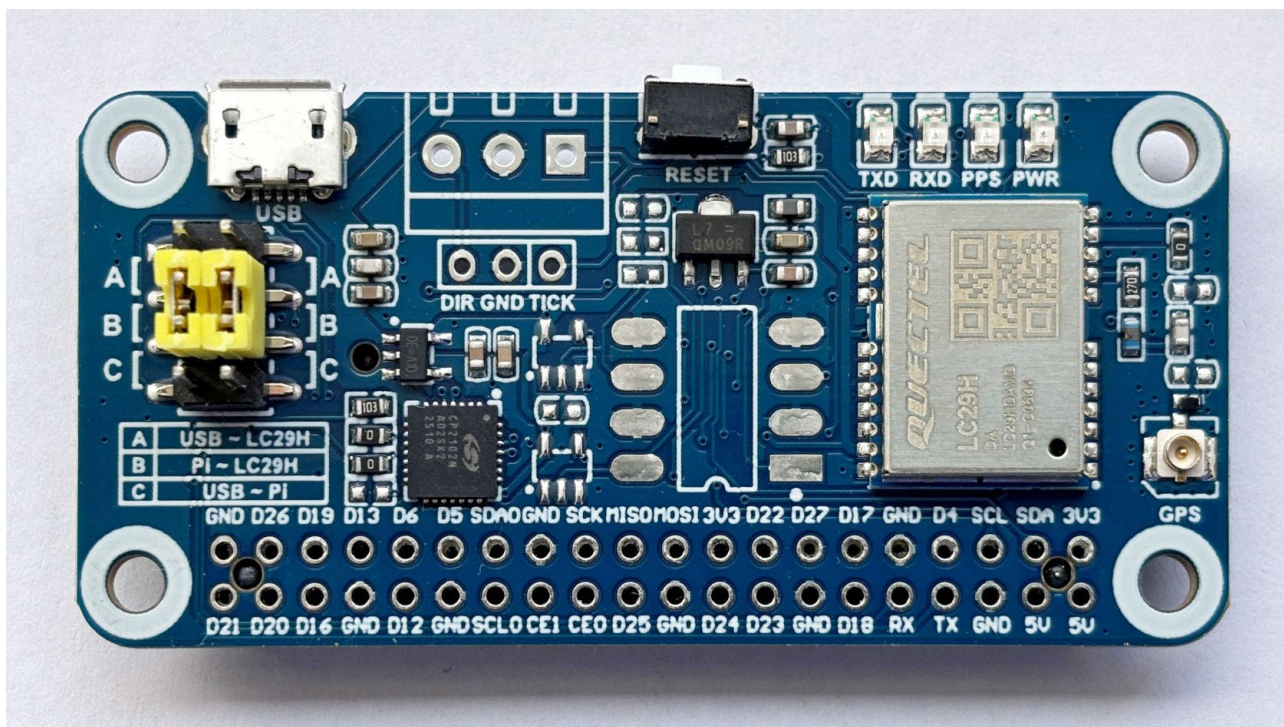
Gdy włączymy przyrząd za pomocą przycisku POWER, wskaźniki pokazują nastawione wartości napięcia i prądu maksymalnego (zapamiętane w chwili ostatniego wyłączenia zasilania), ale **na wyjściu nie ma napięcia i wyświetlony jest napis OFF** – fotografia 1.

Później napięcie wyjściowe włącza i wyłącza w sposób elektroniczny klawisz ON/OFF.

I to jest typowe dla wszystkich dobrych zasilaczy tego typu! Otóż po włączeniu na zaciskach nie ma jeszcze napięcia, bo najpierw sprawdzamy lub ustawiamy potrzebne wartości napięcia i prądu maksymalnego, a potem naciskamy przycisk



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



GNSS RTK, czyli nawigacja o dokładności centymetrowej

Współczesna nawigacja to nie tylko GPS, ale cała rodzina systemów GNSS. Jeszcze niedawno centymetrowa dokładność wiązała się z ogromnymi kosztami i dostępna była tylko dla profesjonalistów. Dziś, dzięki modułowi RTK dla Raspberry Pi, niemal każdy może sam spróbować osiągnąć tę niesamowitą precyzję lokalizacji.

Kilka słów o nawigacji satelitarnej

W poszukiwaniu dokładności

RTK w życiu codziennym

Czy można samemu zrealizować nawigację RTK?

Pierwsze eksperymenty pod Windows 10

W kierunku Raspberry Pi

Uruchomienie trybu RTK

Weryfikacja dokładności lokalizacji

Łyżka dziegciu w beczce miodu

Podsumowanie

Kilka słów o nawigacji satelitarnej

Nawigacja satelitarna stała się dziś na tyle powszechna, że niemal każdy z nas ma codziennie do czynienia z urządzeniami umożliwiającymi określanie położenia na podstawie precyzyjnych sygnałów wysyłanych przez satelity. Rzadko jednak zastanawiamy się nad zasadami ich działania.

Choć celem niniejszego tekstu nie jest szczegółowe omawianie tego systemu, warto w kilku słowach zwrócić uwagę na kwestię jego dokładności. Pierwszym systemem nawigacji satelitarnej był amerykański NAVSTAR GPS (*Navigation Signal Timing and*

Ranging – Global Positioning System), który powstał jako projekt militarny i początkowo był dostępny wyłącznie dla wojska. Dopiero w latach 80. ubiegłego wieku prezydent Stanów Zjednoczonych zezwolił na jego cywilne wykorzystanie. Przez wiele kolejnych lat w systemie tym stosowano jednak tzw. selektywną dostępność (*Selective Availability*), która ograniczała dokładność pozycjonowania dla użytkowników cywilnych do około 100 metrów.

Na początku 2000 roku wyłączono celowe zakłócanie sygnału, dzięki czemu każdy odbiornik GPS mógł określać swoją pozycję z dokładnością do kilku metrów –

był to prawdziwy przełom. To właśnie, w połączeniu ze znacznym spadkiem cen urządzeń, przyczyniło się do masowego upowszechnienia nawigacji satelitarnej.

Choć GPS jest powszechnie wykorzystywany cywilnie, pozostaje wciąż systemem wojskowym. Departament Obrony USA w każdej chwili może wyłączyć jego działanie lub ponownie włączyć celowe zakłócanie sygnału w wybranych rejonach świata. Dlatego inne państwa postanowiły stworzyć własne systemy nawigacji satelitarnej. Obecnie na orbicie Ziemi krążą satelity przesyłające sygnały dla rosyjskiego GLONASS, europejskiego Galileo, chińskiego BeiDou oraz japońskiego QZSS. Wszystkie te systemy razem tworzą Globalny System Nawigacji Satelitarnej, znany pod angielskim skrótem GNSS (*Global Navigation Satellite System*).

W poszukiwaniu dokładności

Nowsze odbiorniki mogą jednocześnie korzystać ze wszystkich systemów, co pozwala zwiększyć dokładność pomiarów. Jednakże prawa fizyki, ograniczenia sprzętowe oraz wszechobecne zakłócenia sprawiają, że typowe urządzenie nie jest w stanie osiągnąć precyzji lepszej niż 1 metr.

Od wielu lat znane jest jednak rozwiązanie pozwalające na wielokrotne zwiększenie dokładności pomiarów: RTK (*Real Time Kinematic*). W dużym skrócie polega ono na użyciu dwóch odbiorników: jeden znajduje się w dokładnie znanej pozycji i przesyła dane korekcyjne do drugiego, którego lokalizację chcemy określić. Mogłoby się wydawać, że skoro znamy pozycję pierwszego odbiornika, możemy w każdej chwili określić błąd systemu i po prostu odjąć go od wskazań drugiego odbiornika, aby uzyskać właściwą lokalizację. W praktyce jednak takie rozwiązanie jest skuteczne głównie w przypadku usuwania wspomnianych wcześniej celowych zakłóceń; dla niezakłóconych sygnałów poprawa dokładności nie byłaby znacząca.

W systemie RTK, jak podaje Wikipedia, stosuje się metody polegające na „szybkim rozwiązaniu problemu nieoznaczoności pomiarów fazowych przez odbiornik ruchomy na podstawie przesłanych ze stacji referencyjnej poprawek do pseudoodległości oraz surowych danych pomiarów fazy sygnałów”. Choć zagadnienie jest skomplikowane, najważniejszą praktyczną informacją jest to, że pomiary RTK wymagają użycia nie jakichkolwiek dwóch odbiorników, lecz dwóch specjalnych urządzeń, które dodatkowo muszą się ze sobą komunikować.

RTK w życiu codziennym

oraz koszt niezbędnych urządzeń – z tego powodu początkowo używana była jedynie w geodezji. Jednak zalety tej metody sprawiły, że z czasem zyskiwała na popularności, co przełożyło się również na stopniowy spadek cen odbiorników.

Dodatkowo okazało się, że choć najkorzystniejsze jest, aby stacja referencyjna znajdowała się jak najbliżej miejsca pomiarów, to akceptowalne wyniki można uzyskać także przy odległościach sięgających wielu kilometrów. Wraz z powstaniem zestandaryzowanych formatów wymiany danych korekcyjnych utworzono całe sieci stacji referencyjnych prowadzących ciągłe obserwacje, a rozwój Internetu sprawił, że dostęp do nich stał się możliwy praktycznie w każdym miejscu.

Dziś odbiorniki GNSS RTK wykorzystywane są już nie tylko w geodezji. Znajdziemy je np. w nowoczesnych ciągnikach rolniczych, które są w stanie precyzyjnie realizować takie prace polowe jak siew czy oprysk, albo też w robotach koszących przydomowe trawniki.

Czy można samemu zrealizować nawigację RTK?

Ceny odbiorników GNSS RTK wykorzystywanych w geodezji to poziom przynajmniej kilku lub kilkunastu tysięcy złotych. Jednak ostatnio na rynku pojawiają się też tańsze konstrukcje, stosowane choćby w autonomicznych kosiarkach. Z kolei w kierunku amatorów chcących spróbować swoich sił pomocną dłoń wyciągają Chińczycy, którzy udostępniają kompletne płytki z już wlutowanymi podzespołami. Eliminuje to konieczność samodzielnego montażu, co w przypadku wielu nowoczesnych układów scalonych jest, bez doświadczenia i bez specjalistycznego wyposażenia, niemalże niemożliwe.

Przykładem takiego produktu jest oferowany przez firmę Waveshare moduł „LC29H(XX) GPS/RTK HAT”. Ten odbiornik GPS/RTK bazuje na układzie Quectel LC29H, a widoczne w nazwie litery HAT (*Hardware Attached on Top*) informują, że jest to przeznaczona dla Raspberry Pi płytka, którą montuje się bezpośrednio na złączu rozszerzeń tego urządzenia.

Moduł dostępny jest w trzech odmianach. Najtańsza z nich, oznaczona literami AA, to jedynie odbiornik GNSS, który nie ma możliwości pracy w trybie RTK. W trybie RTK pracują moduły oznaczone literami DA i BS, przy czym pierwszy z nich umożliwia lokalizację urządzenia, drugi natomiast pozwala stworzyć własną stację referencyjną.

Ja zdecydowałem się na moduł DA, widoczny na

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



Jak działa smartfon? 3G – wspaniała pomyłka!

To jest piąty 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 0G do 2G, a w tym artykule omawiam zadziwiającą trzecią generację telefonii komórkowej, oznaczaną skrótem 3G.

Dziwna, „wojskowa” droga do 3G
Na czym polega specyfika telefonii 3G?

Oczekiwania dotyczące technologii 3G

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 obok**). 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.

Nieprzypadkowo zacząłem od telegrafu, telegrafu bezprzewodowego i omówiłem klasyczne stare telefony oraz stare radiotelefony. W drugim artykule doszliśmy do telefonii mobilnej zerowej generacji (0G), która była bardzo kosztowna, mało popularna. Potem przedstawiłem pierwszą, prawdziwą, analogową generację telefonii komórkowej (1G). Ostatnio omówiłem generację 2G, czyli pierwszą generację cyfrową, utożsamianą u nas z GSM.



Pora na omówienie następnej generacji. Dziś sporo mówi się o telefonii 3G, a konkretnie o jej ostatecznym zamknięciu, wyłączeniu! Paradoksalne jest to, że gorsza, słabsza, wcześniejsza telefon 2G nadal jest w powszechnym użyciu. Dlaczego więc lepsza telefon 3G, oparta na znakomitych, sprawdzonych rozwiązaniach wojskowych, przegrała z gorszą 2G? I jak to się ma do telefonii 4G, której tak naprawdę nie było i nie będzie? I jaki to wszystko ma związek z telefonią 5G, o której mówi się dziś bardzo dużo? Stopniowo to wyjaśnię.

Wcześniej wspominałem, że lubimy „okrągłe” daty i że pierwsza „prawdziwa” telefon 2G pojawiła się na początku lat 80. Telefon 2G zaczęła się na świecie upowszechniać od wczesnych lat 90. Natomiast bohaterka tego artykułu – cyfrowa telefon 3G to (pierwsze) lata dwutysięczne.

Dziwna, „wojskowa” droga do 3G

Telefon GSM (2G) całkowicie wystarczała do rozmów telefonicznych, ale szybko okazało się, że jest fatalnie słaba jeśli chodzi o dostęp do Internetu, czyli o przesyłanie dowolnych danych cyfrowych. Dlatego już dawno, w ramach generacji 2G, dodano do GSM swego rodzaju protezy, pozwalające zwiększyć szybkość przesyłania cyfrowych informacji „internetowych”. Były to GPRS i EDGE, czasem nazywane telefonami 2.5G i 2.75G.

Dość szybko, bo na początku lat dwutysięcznych, w różnych częściach świata zaczęto też wprowadzać telefon trzeciej generacji (3G). **Fotografia tytułowa** pokazuje pierwszy telefon 3G produkcji Nokii – model 6650 z roku 2002.

Podkreślam: zaczęto ją wtedy wprowadzać do powszechnego użytku, natomiast koniecznie trzeba pamiętać, że badania i ustalanie standardów 3G zaczęły się dużo wcześniej, jeszcze w latach 80., a nawet wcześniej. Otóż wkład w technikę wykorzystaną w telefonach 3. generacji ma między innymi pokazana na **fotografii 1**, urodzona w roku 1914, Hedwig Eva Maria Kiesler, córka austriacko-żydowskiego bankiera, która stała się znana jako amerykańska aktorka Hedy Lamarr. Fotografia pochodzi z roku 1941, w którym wraz z innym artystą złożyła zgłoszenie patentowe). Fragment patentu (2292387) udzielonego w roku 1942 na system tajnej komunikacji pokazany jest na **rysunku 2**.

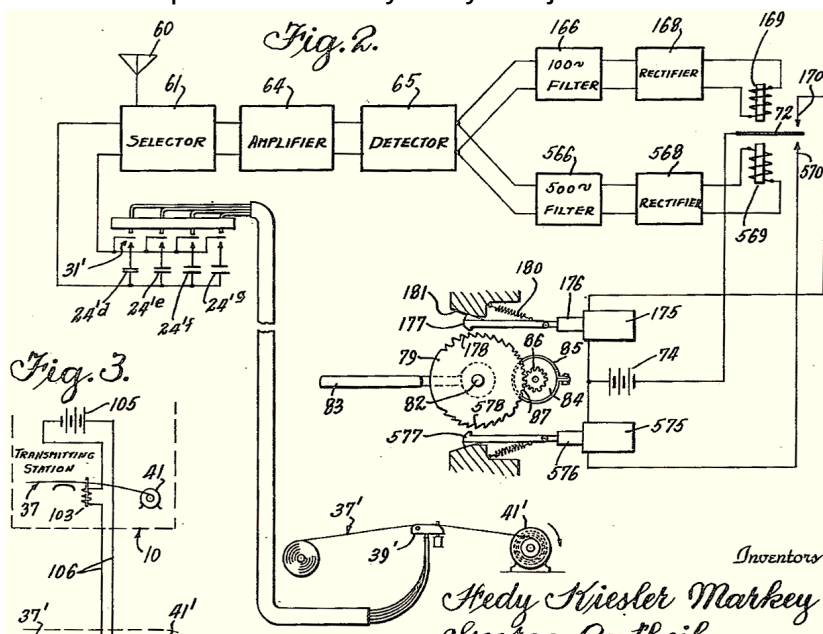
To był drobny wkład prowadzący do telefonii 3G. Na jej kolejne generacje nie powinniśmy patrzeć z dzisiejszej perspektywy. A spojrzenie



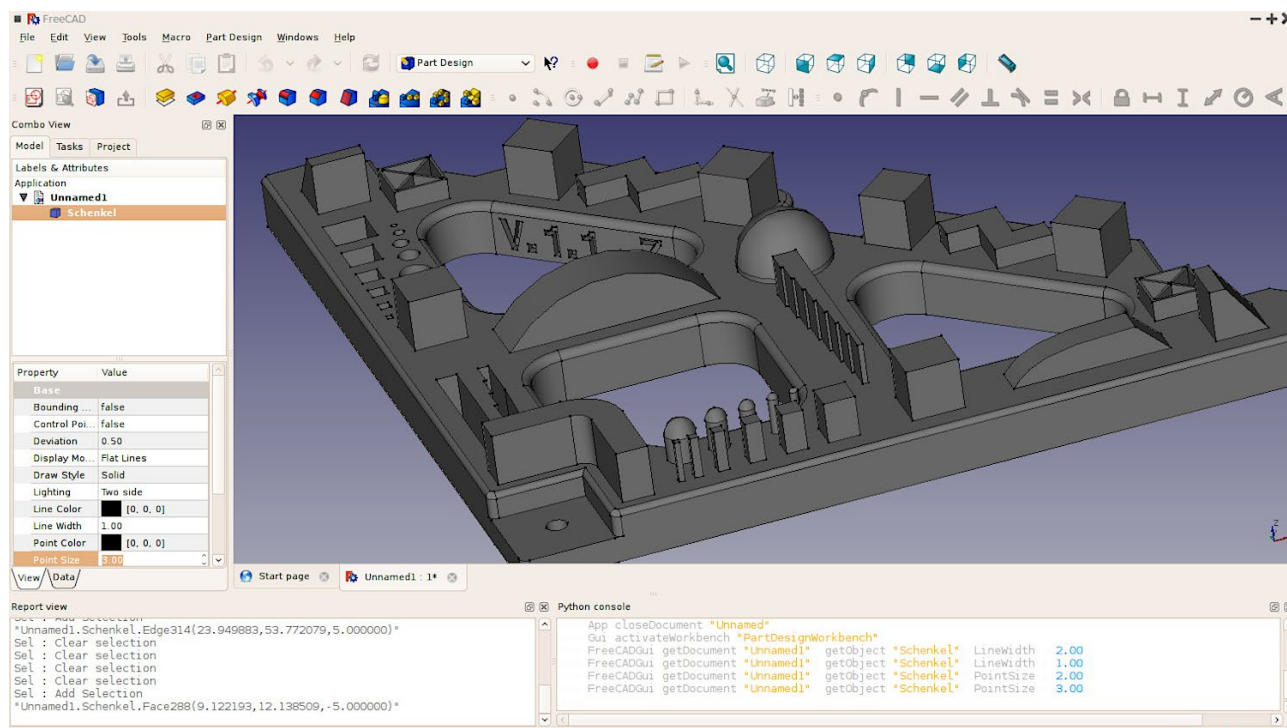
Fotografia 1

rzystania telefonii komórkowej jako drogi dostępu do Internetu, czyli globalnej sieci komputerowej – do tego wątku wrócę w następnym artykule serii. Druga kwestia to ówczesne możliwości i dostępne rozwiązania techniczne oraz najbardziej obiecujące wtedy kierunki rozwoju.

Otóż w latach 80., a nawet w latach 90. elektronika w dużym stopniu była analogowa. Opracowanie i wprowadzenie na rynek cyfrowej telefonii 2G



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



Ten cały druk 3D...

...to nadaje się tylko do druku zabawek. Bardzo często słyszę takie stwierdzenie. I ponieważ jest w nim trochę racji. W końcu najwięcej gotowych modeli, które można znaleźć w Internecie, to właśnie wszelkiego rodzaju zabawki, figurki i gadżety. Ale jak zwykle jest to tylko część prawdy.

Repozytoria modeli 3D

Do drukowania 3D potrzebne są nam trzy elementy – drukarka 3D, materiał i model. Pierwsze dwa już wstępnie omówiliśmy, pora więc zabrać się za trzeci czyli model. Model, lub inaczej mówiąc, projekt, rysunek, przedstawia to, co chcemy wydrukować, czyli uzyskać na końcu całego procesu. W zasadzie cała praca z drukiem 3D sprowadza się właśnie do efektu końcowego – fizycznego elementu o określonym kształcie i właściwościach. Do jego wykonania drukarka potrzebuje przepisu – modelu, rysunku, projektu, nazwa jest nieistotna – który przedstawia to, co chcemy wydrukować. Skąd więc wziąć takie modele? Pierwszym i najprostszym sposobem jest pobranie ich z Internetu.

Repozytoria modeli 3D

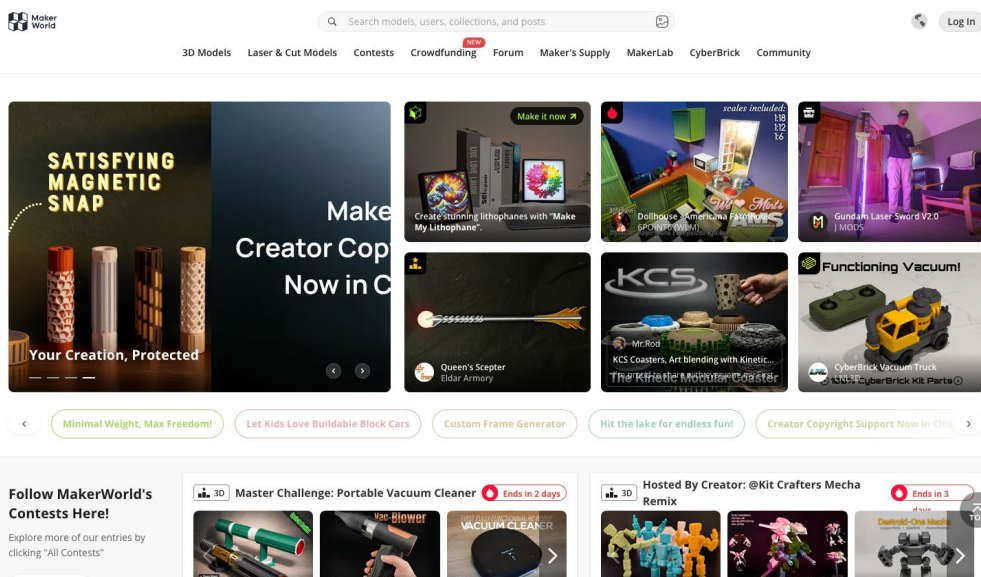
W Internecie dostępnych jest co najmniej kilkanaście tzw. **repozytoriów** modeli 3D. **Repozytorium**

Programy do projektowania 3D

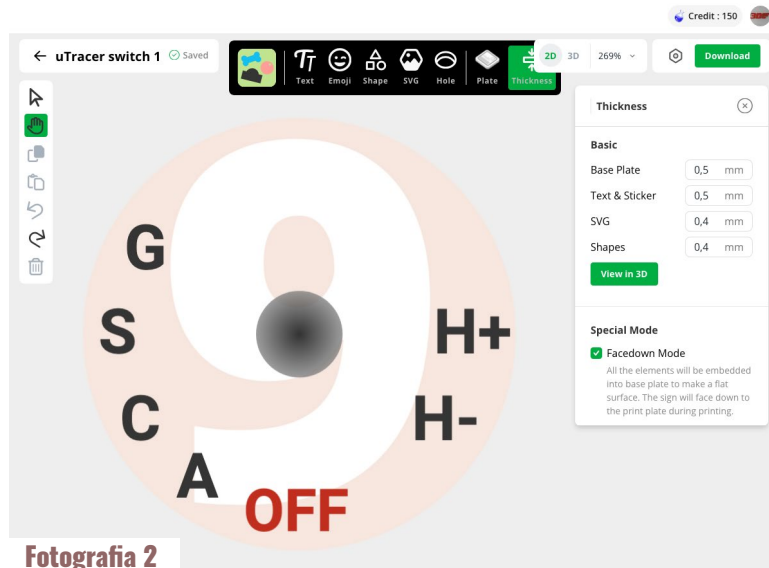
to takie mądre słowo które w zasadzie oznacza dedykowany serwis z gotowymi modelami 3D przeróżnych elementów (m.in. zabawek i figurek :) w postaci plików w formacie rozumianym przez każdy **Slicer** (przypominam – Slicer to program w którym przygotowujemy model do druku na konkretnej drukarce). Modele są dostępne na różnych zasadach i na różnej licencji, o czym należy pamiętać. Czasem są po prostu darmowe, do dowolnego wykorzystania, czasem darmowe, o ile nie chcemy ich drukować i sprzedawać, a czasem są płatne. Dlatego zawsze, a w szczególności dotyczy to plików dostępnych za darmo, trzeba sprawdzić licencję na jakiej autor udostępnia dany plik. Niestety wielu nieuczciwych „przedsiębiorców” pobiera pliki z licencją do wykorzystania domowego, indywidualnego, po czym masowo drukuje modele i próbuje je sprzedawać. No, ale to zupełnie inna sprawa.

Trzy największe i najpopularniejsze serwisy tego typu prowadzone są przez producentów drukarek 3D, którzy niejako w ten sposób kreują sobie popyt na urządzenia, bo najczęściej serwis taki oferuje dodatkowe ułatwienia dla posiadaczy drukarki danej firmy. Ale nie oznacza to, że modele są „zamknięte” – jak najbardziej można je drukować na dowolnej innej drukarce 3D, po prostu będzie to wymagało minimalnie więcej pracy – modele najczęściej dostępne są w uniwersalnym formacie STL rozumianym przez każdą drukarkę. Np. serwis **MakerWorld** (<https://makerworld.com/>) pokazany na **fotografii 1**, prowadzony przez producenta drukarek Bambu Lab, oferuje maksymalne ułatwienia dla swoich klientów – jeśli masz odpowiednio skonfigurowaną drukarkę Bambu Lab i konto w serwisie, to z poziomu strony internetowej serwisu wystarczy często kliknąć w klawisz „Drukuj” i gotowy, już pocięty plik z wybranym modelem automatycznie zostanie wysłany do Twojego urządzenia. Potem wystarczy tylko potwierdzić chęć drukowania i za jakiś czas mamy gotowy, wydrukowany element. Dla innych drukarek musisz samodzielnie pobrać plik STL do komputera, otworzyć go w slicerze, wybrać parametry drukowania i uruchomić druk. Z racji możliwości łatwego drukowania wielokolorowego na drukarkach Bambu Lab oraz niskiej ceny urządzeń (za czym idzie duża popularność), w tym serwisie znajdziesz najwięcej modeli figurkowo – zabawkowych, kolorowych i hobbystycznych. Dostępne tam są także działy z modelami do wycinania laserem, nożem, rysowania i grawerowania – wykorzystują one możliwości najbardziej zaawansowanych drukarek, a właściwie już „kombajnów” druku 3D, wyposażonych w głowice laserowe. Nie oznacza to jednak, że nie ma tam modeli technicznych – jak najbardziej jest ich sporo i to naprawdę całkiem sensownych. Tym bardziej, że w portfolio materiałów Bambu Lab znajduje się całkiem sporo technicznych, wytrzymałych filamentów.

Ciekawostką tego serwisu jest to, że w dziale MakerLab oferuje on kilkanaście narzędzi przeglądarkowych do projektowania 3D. Np. Make My Sign to prościutki, ale bardzo użyteczny program, w którym

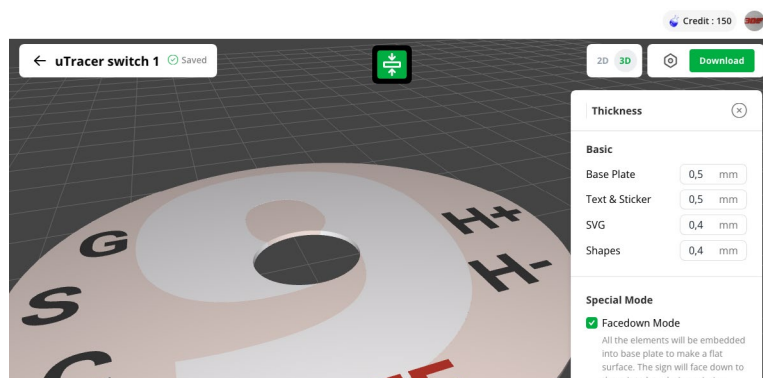


Fotografia 1

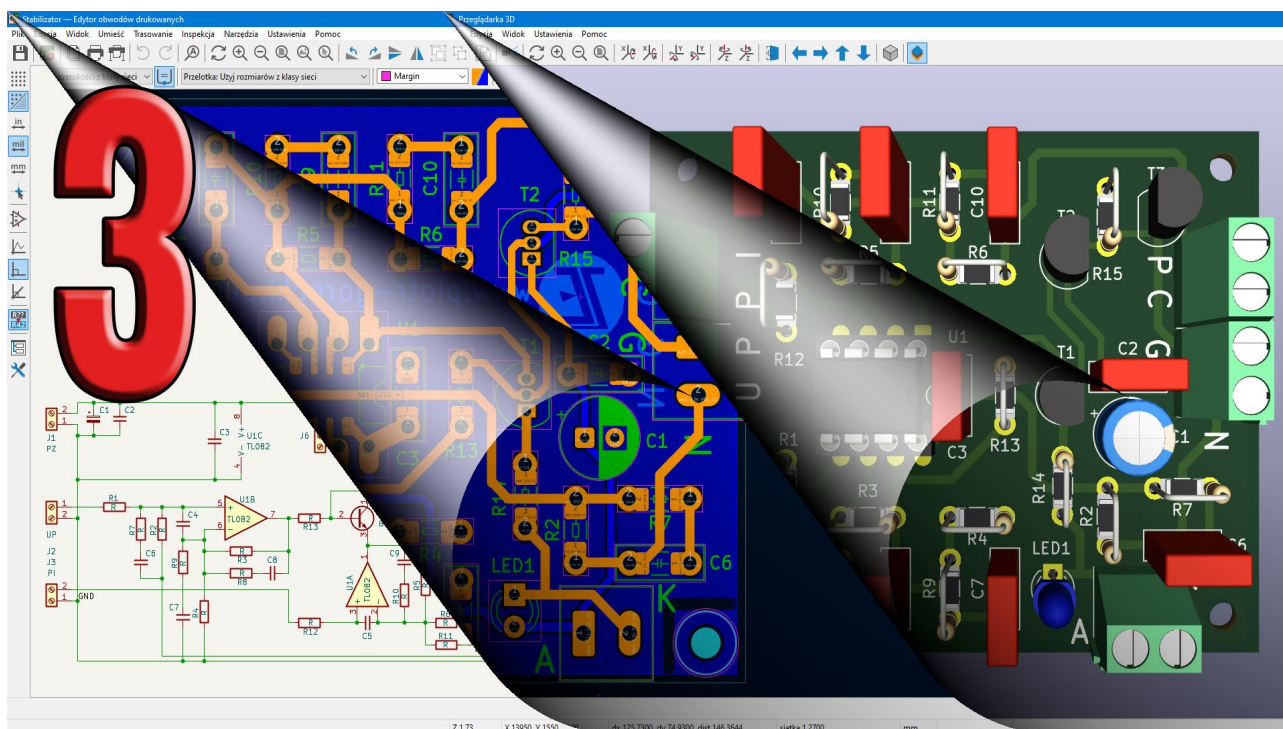


Fotografia 2

rzysłałem go do stworzenia opisu pozycji przełącznika obrotowego w postaci cienkiej, okrągłej płytki z napisami i środkowym otworem – podłożona pod przełącznik i przymocowana jego nakrętką idealnie spełnia swoją funkcję. Na **fotografiach 2 i 3** pokazany jest projekt tej płytki w 2D oraz jej widok w 3D.



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



Projektowanie płytek drukowanych za pomocą KiCad (3)

W dwóch pierwszych artykułach serii zapoznaliśmy się z pakietem KiCad oraz zrealizowaliśmy schemat. W artykule trzecim zrobimy najważniejszy krok: zaprojektujemy płytkę drukowaną na podstawie wcześniej utworzonego schematu. To krok najtrudniejszy, ale artykuł szczegółowo pokazuje, jak to zrobić.

Ustawienia edytora PCB

Ustawienia projektowe płytki

Paski narzędzi

Krawędzie płytki

Umieszczenie elementów na płytce

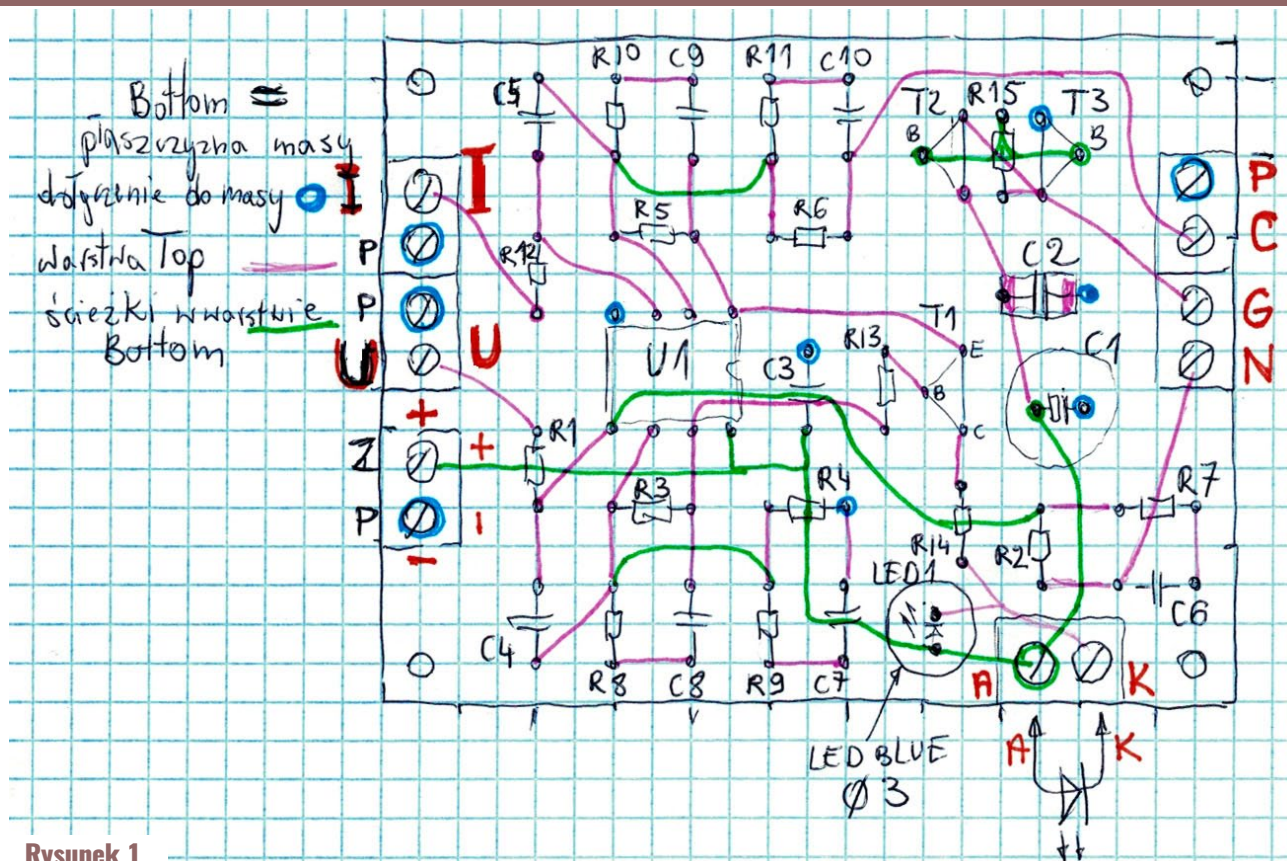
Pakiet KiCad znakomicie ułatwia projektowanie płytki. Po pierwsze dlatego, że już na etapie rysowania schematu możemy określić, jakie elementy płytkowe (footprinty) będą wykorzystane na płytce. Po drugie, program ułatwia zarówno rozmieszczanie elementów na płytce, jak też prowadzenie ścieżek. Nie ma wprawdzie dobrych narzędzi do automatycznego rozmieszczania elementów i automatycznego projektowania przebiegu ścieżek, ale to nie szkodzi, bo ich praktyczna przydatność jest bardzo ograniczona. Lepiej zaczynać od projektowania ręcznego.

Projekt płytki powstanie na podstawie szkicu z **rysunku 1** (na następnej stronie) z naniesionymi zaleceniami projektowymi dla tej płytki.

Ustawienia edytora PCB

W zasadzie można pozostać przy ustawieniach domyślnych, ale warto troszkę poznać możliwości i kluczowe opcje. Otóż podobnie jak w edytorze schematów, również w edytorze płytek drukowanych mamy podział na ustawienia samego edytora płytek (programu) oraz ustawienia projektowe konkretnej płytki drukowanej. Trzeba wiedzieć o takim podziale ustawień.

Część ustawień edytora płytek, jak właściwości kursora i siatki, jest podobna jak w przypadku edytora schematów. Podobna, ale ustawień rozmiarów siatki jest dużo więcej, niż w przypadku edytora schematów. Spowodowane jest to tym,



Rysunek 1

że niektóre płytki drukowane są projektowane z dużą precyzją. W ustawieniach edytora płytek drukowanych warto zwrócić uwagę na opcje wyświetlania, które widzimy na skadrowanym **rysunku 2**. W czerwonej ramce możemy wybrać sposób wyświetlania nazw sieci połączeń oraz numery pól lutowniczych. Niebieska ramka pozwala z kolei wybrać czy i w jaki sposób ma być wyświetlany obrys prześwi-

tu ścieżek i pól lutowniczych. Natomiast w zielonej ramce możemy wybrać czy zaznaczenie footprintu ma podświetlać powiązane z nim pola tekstowe, takie jak na przykład oznaczenie lub wartość. Ostatnia, brązowa ramka zawiera ustawienia zaznaczania obiektów na schemacie, czemu odpowiada podświetlenie odpowiadającego obiektu na płycie oraz przypisanie temu działania, jak skoncentrowanie widoku i jego powiększenie. Można też włączyć podświetlenie zaznaczonych sieci. Opcje te stanowią ułatwienie przy dużych projektach. Dostępne jest też włączenie automatycznego odświeżania widoku 3D. Na **rysunku 3** mamy możliwość zmiany ustawień punktów bazowych oraz kierunku zwiększania osi X i Y.

Numeracje
Nazwy sieci: Pokaż na polach
☒ Pokaż numery pola

Obrysy prześwitu
Ścieżki: Pokaż podczas trasowania przelotką na końcu
☒ Pokaż prześwit pola lutowniczego

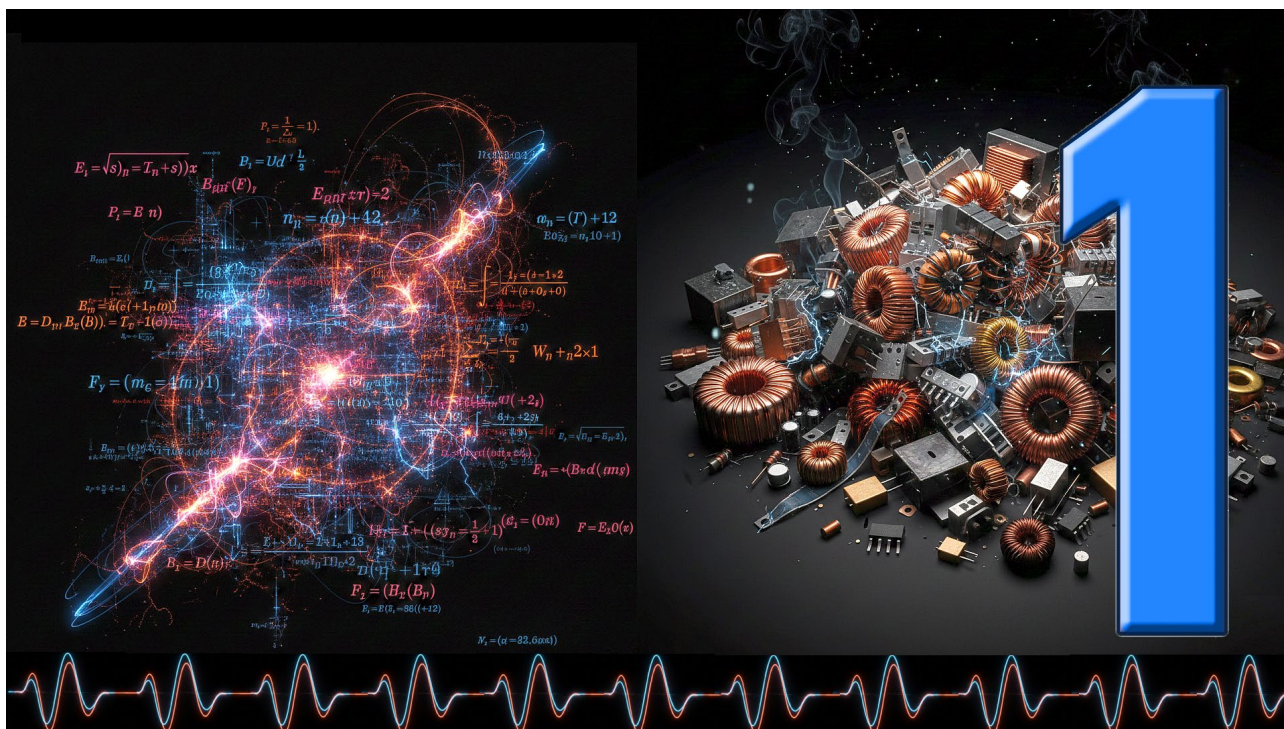
Zaznaczenie oraz podświetlanie
☒ Podświetlaj pola gdy wybrane są ich nadrzędne footprinty

Ustawienia zaznaczania
☒ Wybierz/podświetl obiekty odpowiadające wyborowi na schemacie
☒ Skoncentruj widok na zaznaczonych elementach
☒ Powiększ zaznaczone elementy

Preferencje
Wspólne
Mysz i płytki dotykowa
Skróty klawiszowe
Kontrola wersji
Zbieranie danych
Edytor bibliotek symboli
Edytor Schematów
Edytor Footprintów
Edytor obwodów drukowanych
Opcje wyświetlania
Siatki
Punkty bazowe oraz osie

Wyświetlaj punkt odniesienia
☒ Punkt odniesienia strony
☐ Punkt bazowy pliku wierceń/położeń
☐ Punkt bazowy siatki
Oś X
☒ Zwiększa się w prawo
☐ Zwiększa się w lewo
Oś Y

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



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”.

Główny obiekt naszych rozważań
Podstawowe, jakże oczywiste prawa

Pytania i odpowiedzi dotyczące natury zjawisk

Wydawałoby się, że kawałek drutu, przewodu lub ścieżki na płytce PCB nie ma praktycznie żadnego znaczenia, jeśli chodzi o analizę zjawisk towarzyszących przepływowi prądu przez taki twór. Niestety sprawa wygląda troszkę bardziej zawile. Ogólnie, pojęcie indukcyjności jest słabo rozumiane, co często powoduje błędną jej interpretację oraz błędne zrozumienie zjawisk, których ona dotyczy. Postaram się to przystępnie wyjaśnić, nie wnikając w skomplikowaną analizę tego zjawiska przy pomocy wyższej matematyki. **Lecz docelowo skupię się wyłącznie** na indukcyjności pętli **L_{LOOP}**, gdyż tego typu indukcyjne pętle istnieją w realnych obwodach z szybkimi układami cyfrowymi i niestety mają na nie niekorzystne oddziały-

wanie. Gdy chcemy prawidłowo wykonać urządzenie zawierające szybkie układy cyfrowe – musimy być świadomi tych pętli i umieć je minimalizować jak tylko się da. Chciałbym tu również wspomnieć o cyklu dotyczącym podstaw indukcyjności, którego autorem jest **Piotr Górecki**, a wszystkie te darmowe odcinki można znaleźć w **Kopalni Skarbów** zamieszczonej na stronie ZE. Odcinki te umieszczone są w cyklach: „**Poznajemy elementy indukcyjne**” oraz „**Elektronika (nie tylko) dla informatyków**”. Zachęcam do zapoznania się z nimi i odświeżenia swojej wiedzy przed kontynuacją wszystkich wykładów tego minicyklu. Postanowiłem również skorzystać z uproszczonego, niedoskonałego modelu w postaci okręgów

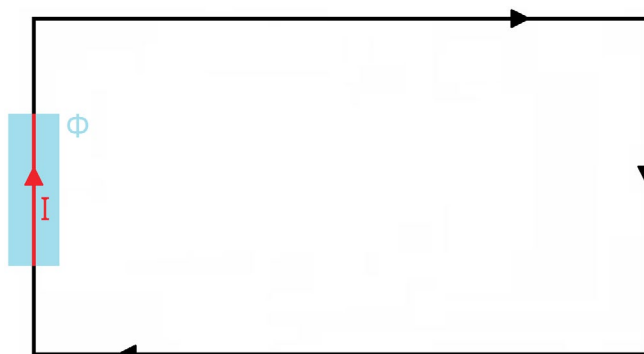
reprezentujących pole magnetyczne wytwarzane wokół przewodnika z prądem. Choć model ten jest niepełny i ułomny, a nawet wprowadzający w błąd w wielu kwestiach, to mimo wszystko wręcz genialnie wizualizuje niektóre cechy magnetyzmu.

Główny obiekt naszych rozważań

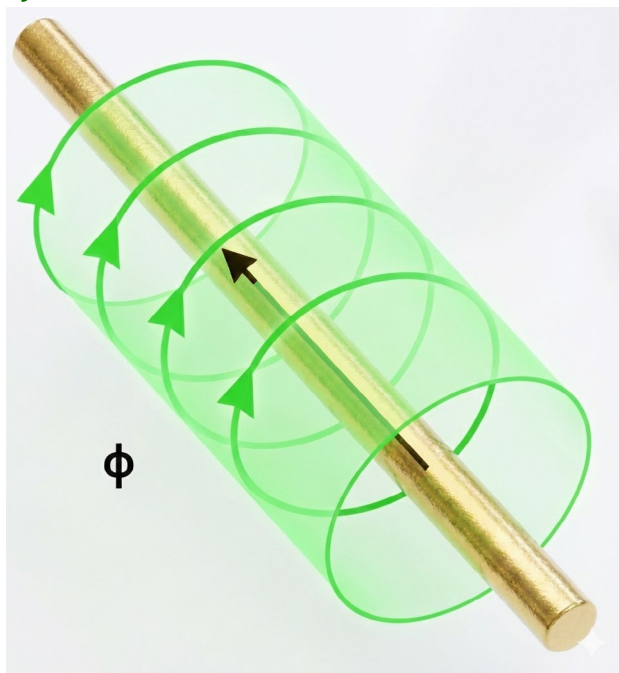
Dla ułatwienia dalszej analizy materiału, szczególnie, jeśli chodzi o wyobrażenie pola magnetycznego, przyjmijmy, że najpierw badamy **wycinek prostoliniowego przewodnika, z którego zbudowana jest pętla, ponadto założmy, że pozostałe odcinki pętli są tak daleko, że nie mają realnego wpływu na wspomniany wycinek. Przez pętlę płynie prąd stały i znajduje się ona w próżni.** Pętla ta wraz z wydzielonym odcinkiem przedstawiona jest na **rysunku 1**.

Podstawowe, jakże oczywiste prawa

Gdy prąd I płynie w przewodniku, wokół tego przewodnika wytwarzane jest pole magnetyczne (**rysunek 2**), a miarą całkowitej ilości tego pola przenikającego przez pętlę jest strumień magnetyczny Φ . Ważną zależnością jest fakt, że wraz ze wzrostem prądu w przewodniku – proporcjonalnie rośnie też wytwarzany przez niego strumień magnetyczny Φ , co w nazbyt uproszczony sposób przedstawia **rysunek 3a** i **3b**. Koniecznie trzeba mieć tu świadomość, że rysunek ten jest bardzo uogólniony, i de facto nie przedstawia gęstości strumienia magnetycznego Φ , czyli indukcji magnetycznej B dla dwóch różnych przypadków związanych z wartością płynącego prądu, a tak naprawdę wprowadza w błąd i buduje błędne wyobrażenie o naturze tego zjawiska. Tę kontrowersyjną kwestię wyjaśnię w dalszej części tego artykułu.

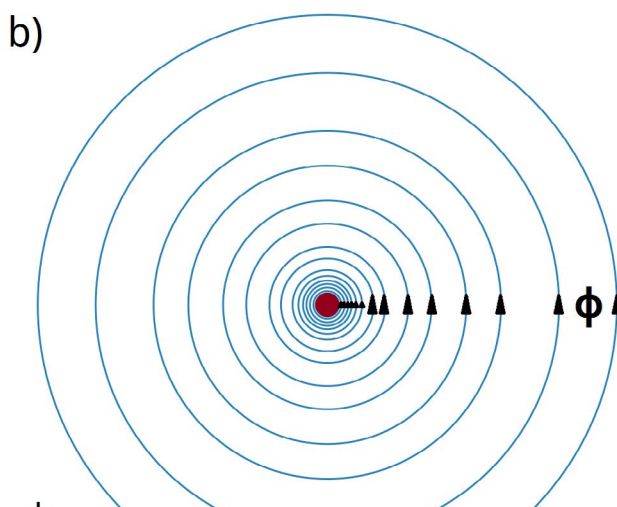
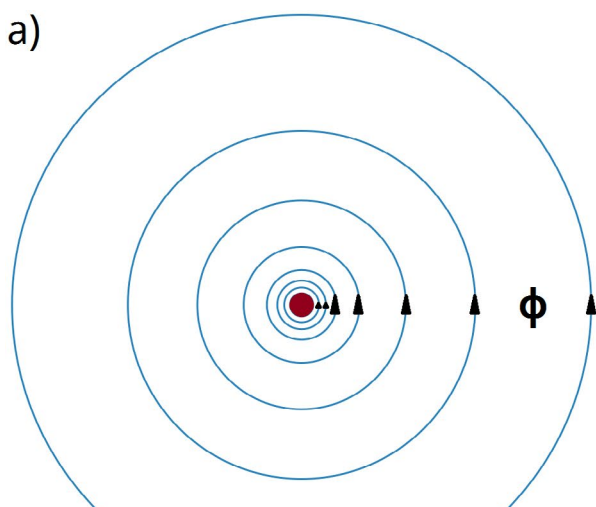


Rysunek 1



Rysunek 2

Wraz ze wzrostem prądu w przewodniku, proporcjonalnie rośnie też wytwarzany przez niego strumień magnetyczny Φ . Należy tu jednak koniecznie zaznaczyć, że wniosek ten jest prawdziwy tylko dla materiałów niemagnetycznych otaczających ten przewodnik



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

ZROZUMIEĆ **ELEKTRONIKĘ** z Piotrem Góreckim

ZE 3/2026

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.