

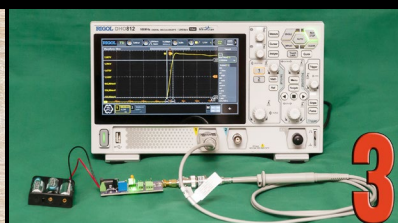
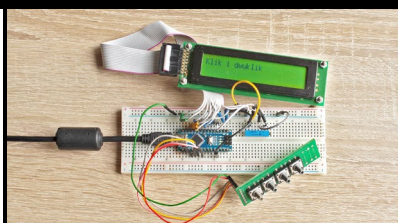
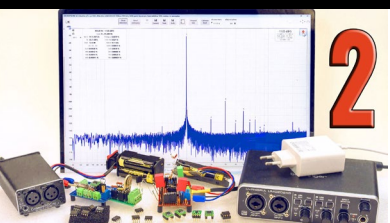
10/2026 Październik (46)

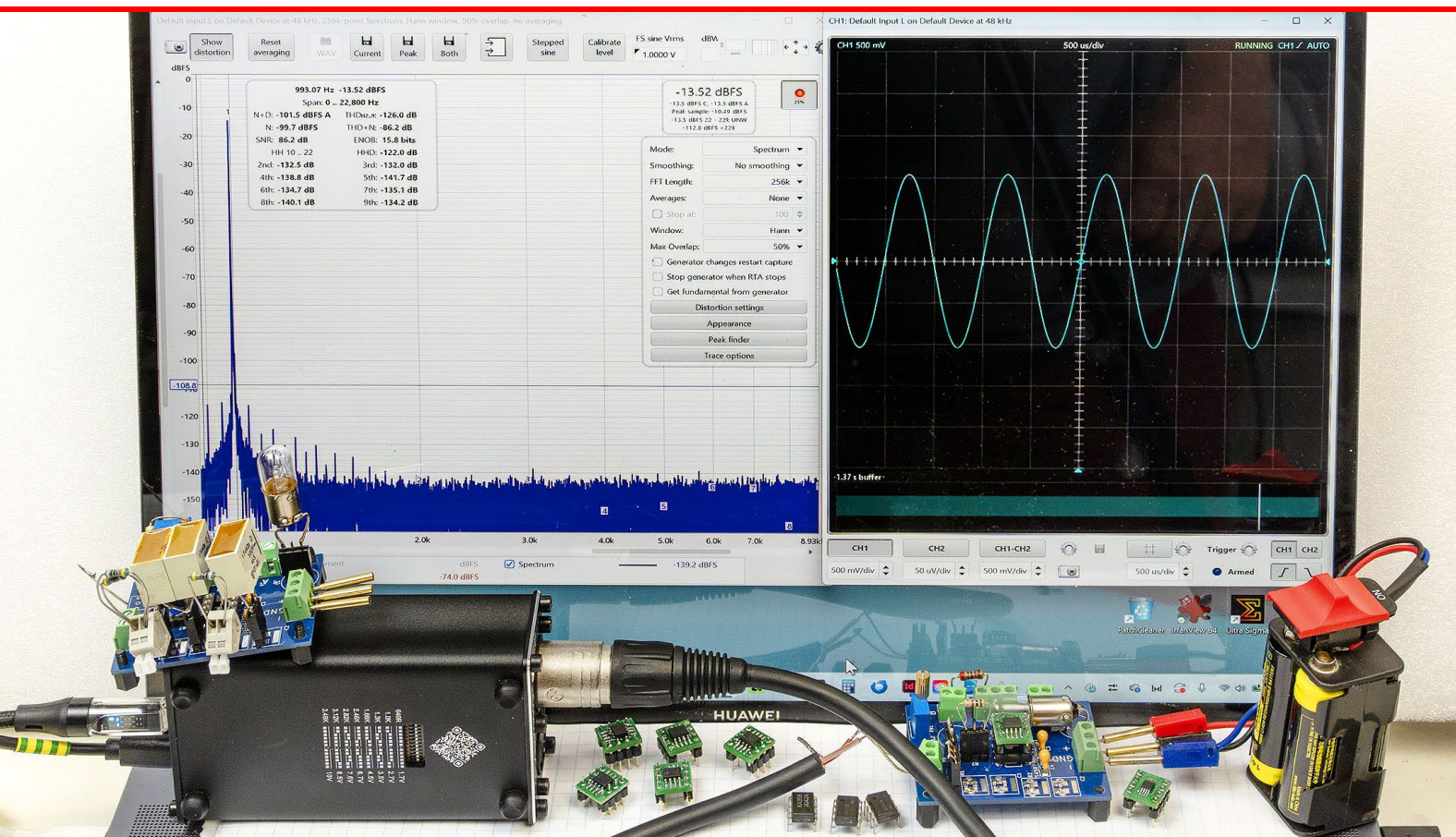
piotr-gorecki.pl



Eksperymentalny generator przebiegu sinusoidalnego

- „Malinka” dla początkujących • Moduł Raspberry Pi Pico 2 W • Klik oraz dwuklik
- Akumulatory litowe – balansery pasywne i aktywne • Szumy, zniekształcenia zakłócenia i decybele
- Obserwacja naturalnego tła promieniowania • Klasyczne konfiguracje zasilaczy stabilizowanych
- Ojciec elektromagnetyzmu: Romagnosi czy Ørsted? • Płyty czołowe do obudów Kradex i Maszcyk





Eksperymentalny generator przebiegu sinusoidalnego

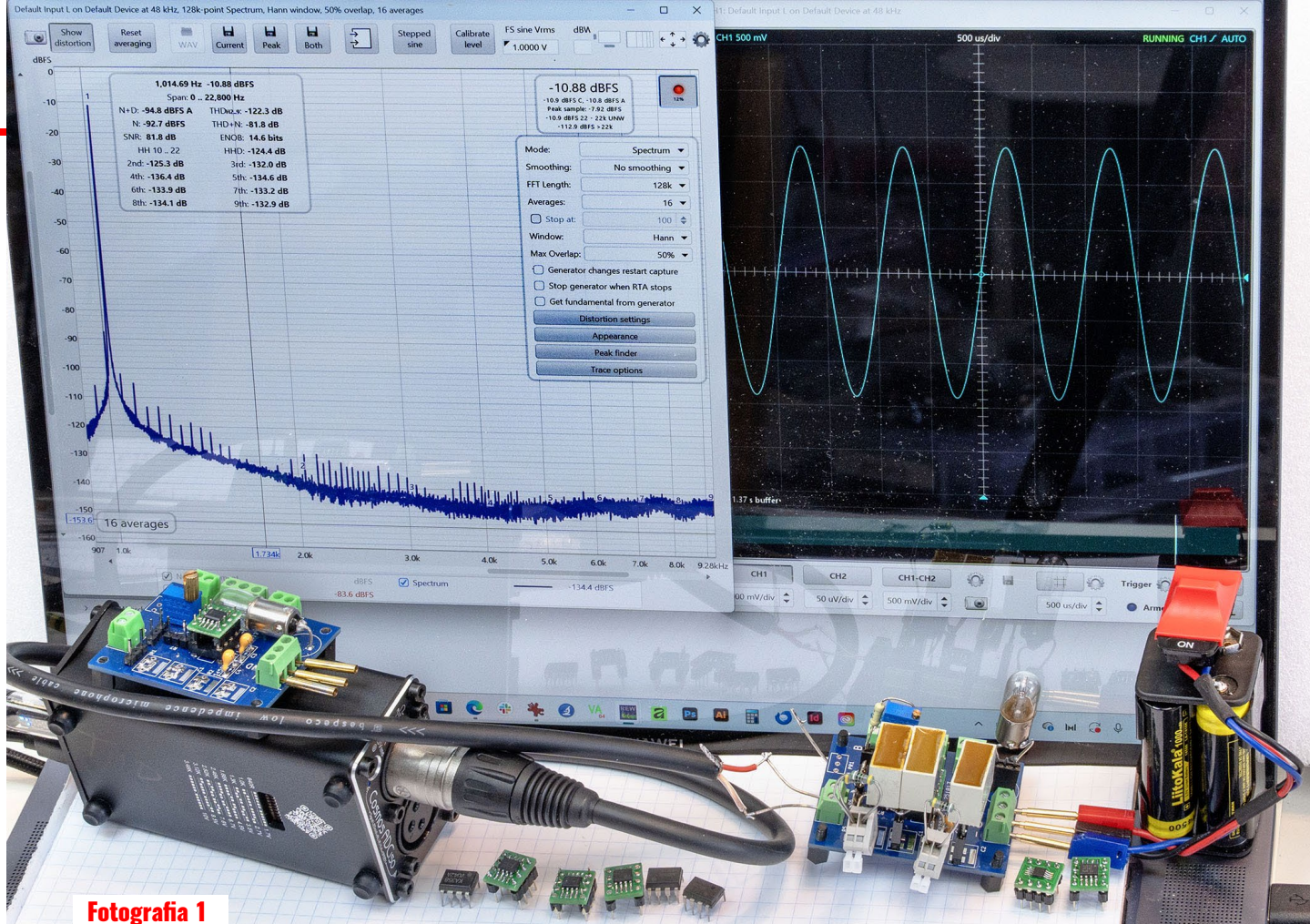
W tym artykule opisuję zaskakująco prosty i zaskakująco dobry eksperymentalny generator sinusoidalny o znikomo małych zniekształceniach. Artykuł przedstawia układ oraz historię projektu. Wstępne wyniki są absolutnie rewelacyjne, ale układ należy dokładniej zbadać i dopracować szczegóły.

Podstawowe rozważania projektowe
Przygotowania – zakupy elementów

Testy modeli

Fotografia tytułowa pokazuje wyniki pomiaru zniekształceń harmoniczných prostego generatora przebiegu sinusoidalnego (993 Hz) z mostkiem Wien'a i żarówką w wersji z taną i nadal ogromnie popularną kostką NE5532. Układ jest śmiesznie prosty i bardzo tani, a parametry okazują się zadziwiająco dobre. Zawartość drugiej harmonicznej to $-132,5$ dB czyli $0,000024\%$, a trzeciej harmonicznej -132 dB, czyli $0,000025\%$. Całkowita zawartość harmoniczných to -126 dB, czyli $0,00005\%$. Wartość THD+N jest dużo gorsza ($-86,2$ dB) z uwagi na obecność szumów, zakłóceń o wyższych częstotliwościach oraz przydźwięku sieci, w tym nieekranowanym modelu.

Widoczne na ekranie wyniki pomiaru harmoniczných wyglądają niewiarygodnie dobrze. Nie jest to wcale jakaś sztuczka, tylko przykład wykorzystania rewelacyjnego, 32-bitowego przetwornika ADC E1DA Cosmos ADCiso 768 kHz, który opisywałem w artykule [Dobre oraz „oszukane” karty audio USB](#) w ZE 5/2025. Zmierzone zniekształcenia harmoniczne są niewiarygodnie małe. Na fotografii widać, że są to zniekształcenia bardzo prostego generatora, zawierającego tylko kilka elementów, w tym małą żarówkę. Po lewej stronie widać też drugi egzemplarz podobnego generatora, gdzie kluczowe elementy dołączane są za pomocą złączy goldpinowych.



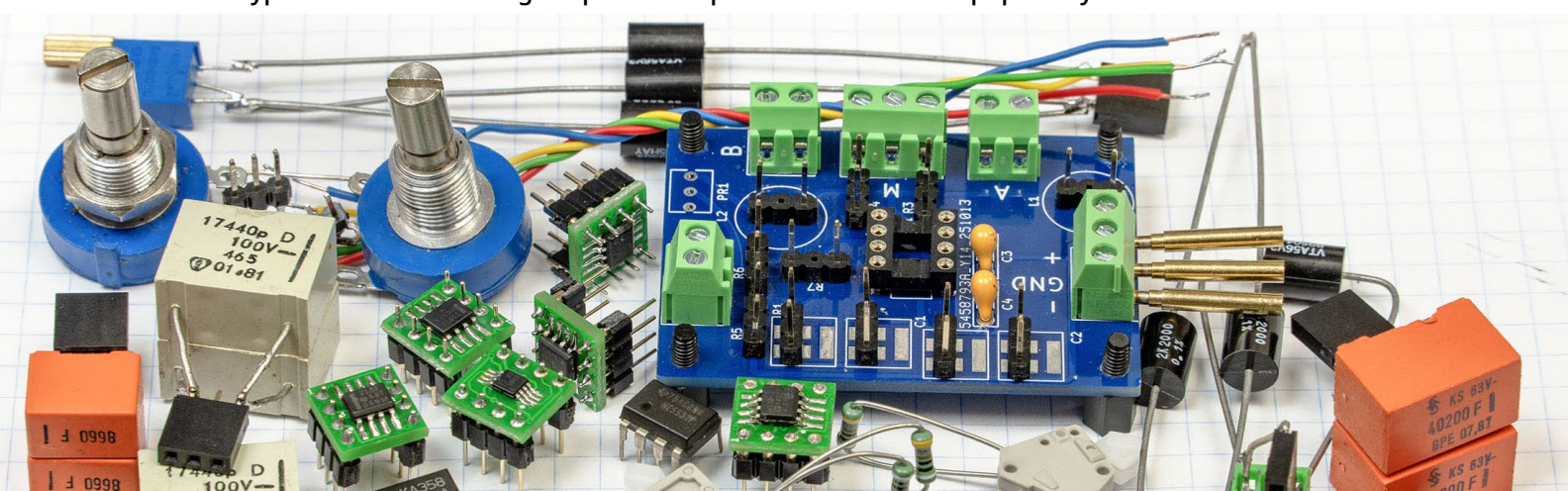
Fotografia 1

Fotografia 1 pokazuje pomiar tego drugiego egzemplarza, przewidzianego do najróżniejszych eksperymentów z różnymi elementami składowymi. Wynik też jest znakomity: druga harmoniczna jest o 125,3 dB mniejsza od podstawowej (1015 Hz), a wartość trzeciej harmonicznej to -132 dB.

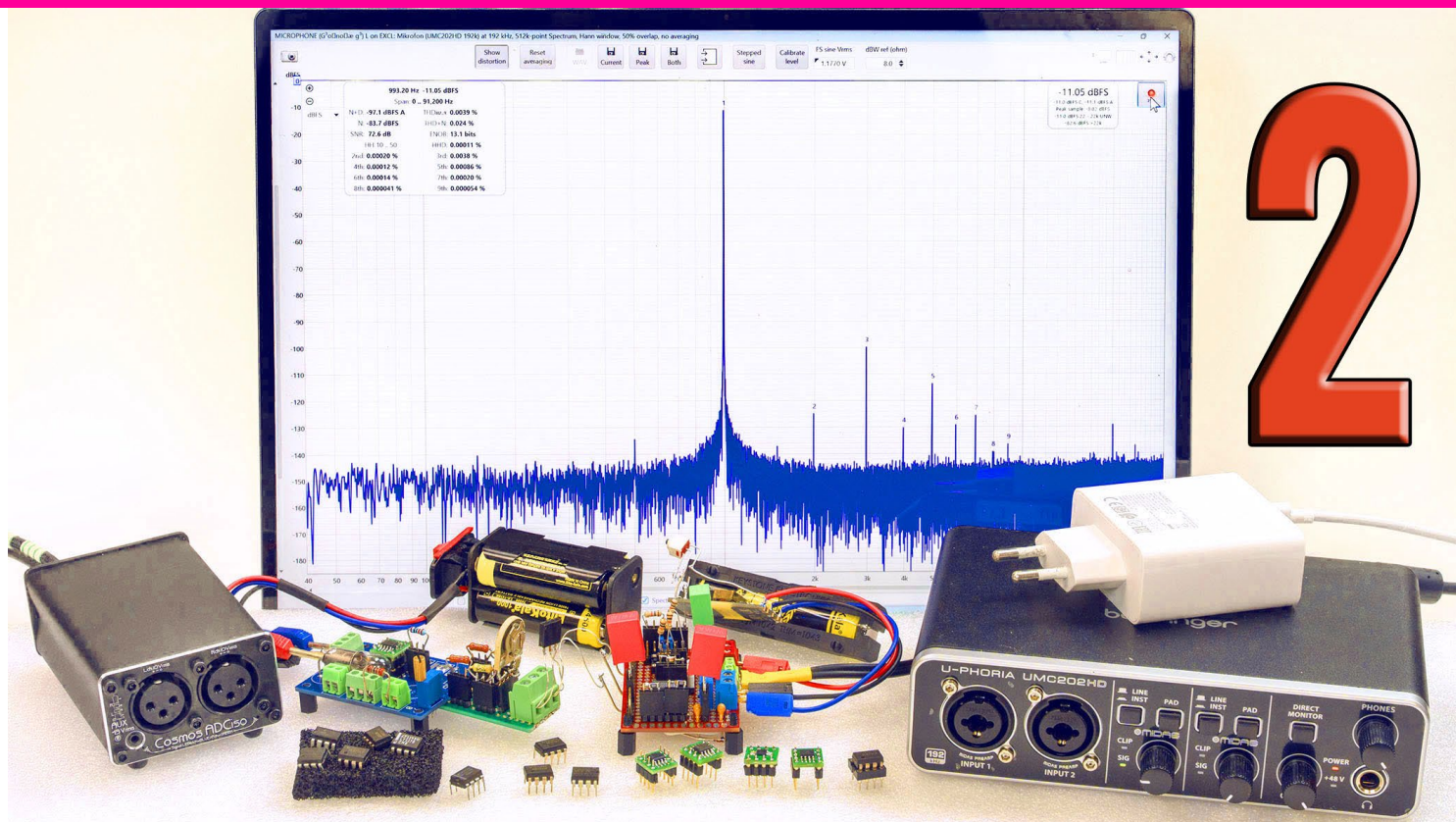
Fotografia 2 przedstawia budowę tego egzemplarza: w płytkę drukowaną wlutowanych jest mnóstwo goldpinów i innych złączy stykowych, a jedynymi wlutowanymi na stałe elementami elektronicznymi są dwa kondensatory ceramiczne filtrujące zasilanie. na goldpiny zakładane są kondensatory i rezystory – jak widać wypróbowałem różne egzemplarze. W pod-

stawkę wkładany jest podwójny wzmacniacz operacyjny. Widać, że miałem do dyspozycji liczne takie wzmacniacze, w tym najlepsze oryginalne wzmacniacze audio, takie jak LM4562, OPA1612, OPA1642, OPA2134, ADA4522-2, a także MC33272, OPA2188, AD823, OPA2134 i LME49860 – te ostatnie kupione w Chinach, duuużo taniej niż oryginały...

Najlepsze wyniki osiągnąłem z kostką LM4562, co widać na fotografiach. Testowałem najlepsze wzmacniacze operacyjne, kosztujące po kilkadziesiąt złotych za sztukę, ale punktem wyjścia był najpopularniejszy klasyk, czyli dostępny dla każdego, nadal bardzo popularny NE5532.



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



Generatory z mostkiem Wiena (2)

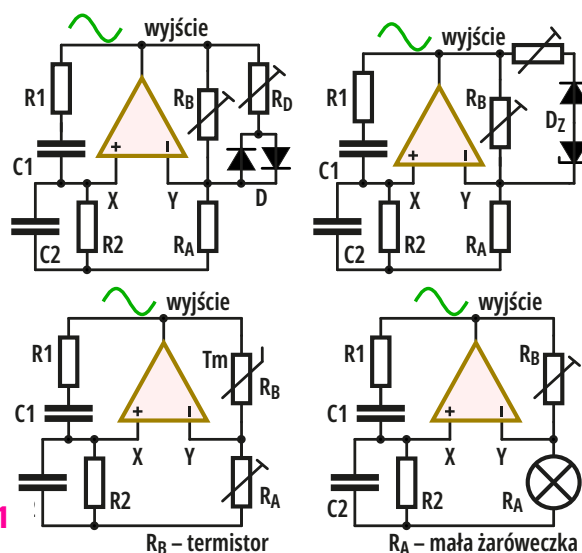
W tym trzyczęściowym artykule opisane są, nadal atrakcyjne i bardzo pożyteczne, generatory analogowe z mostkiem Wiena. Część pierwsza przedstawia podstawowe informacje i proste rozwiązania, druga pokazuje historię i rozwiązania trudniejsze, a trzecia przeznaczona jest głównie dla eksperymentatorów.

Ogólne reguły wytwarzania sinusoidy
Generatory z tranzystorem JFET

Generatory z fotorezystorem
Historia generatorów z mostkiem Wiena

W pierwszej części artykułu przedstawiłem elementarne informacje o generatorach przebiegu sinusoidalnego z mostkiem Wiena. Omówiłem trzy sposoby stabilizacji amplitudy wytwarzanego sygnału: z diodami zwykłymi i z diodami Zenera, z termistorem oraz z żarówką – **rysunek 1**. Generatory z małą żaróweczką to znane od dawna, proste, łatwe do realizacji i nadal popularne rozwiązanie. Jednak słusznie czy nie, w literaturze za lepsze uchodzą układy zdecydowanie bardziej skomplikowane, gdzie wartość jednej z rezystancji R_A , R_B zmieniana jest „elektronicznie” przy wykorzystaniu tranzystora JFET lub fotorezystora, zawartego w specyficznym transoptorze. I właśnie takim, bardziej złożonym rozwiązaniom, poświęcony jest ten artykuł.

Rysunek 1



Ogólne reguły wytwarzania sinusoidy

Przypominam, że w generatorach z mostkiem Wiena i ze wzmacniaczem operacyjnym mamy dwie pętle sprzężenia zwrotnego, realizowane przez dwa dzielniki napięcia wyjściowego – **rysunek 2**. Jeden dzielnik z elementami R_1 , C_1 , R_2 , C_2 , daje jakiś współczynnik podziału i drugi dzielnik z rezystorami R_B , R_A musi mieć idealnie taki sam współczynnik podziału, aby uzyskać na wyjściu przebieg sinusoidalny o znikomym zniekształceniu.

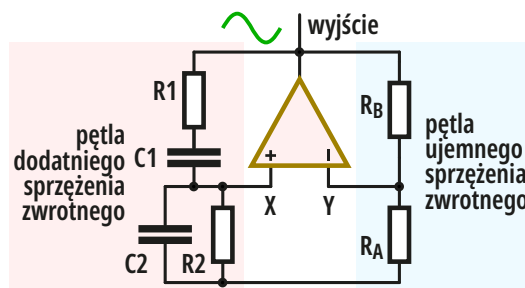
Nie przesadzam ani trochę: podczas pracy generatora sinusoidalnego współczynniki podziału obu dzielników muszą być dobrane wręcz idealnie, żeby wypadkowe wzmocnienie było dokładnie równe jeden, bo to jest konieczne do uzyskania czystutkiej, niezniekształconej sinusoidy. Można to osiągnąć w prosty sposób z wykorzystaniem diod, termistora lub żaróweczki, według rysunku 1. Ale tego wzmocnienia i potrzebnej amplitudy przebiegu wyjściowego może też pilnować układ elektroniczny – kontroler, który w jakiś sposób będzie korygował współczynnik podziału dzielnika, najczęściej wartość rezystancji R_A , według ogólnej idei z **rysunku 3**.

W praktyce funkcję zmiennej rezystancji pełni albo tranzystor polowy złączowy (JFET), albo fotorezystor. Zwykle nie jest to cała rezystancja R_A , a jedynie jej część, według nieco uproszczonej idei z **rysunku 4**. Aby nie zwiększyć nieliniowości rezystancji R_A i poziomu zniekształceń, zmienna rezystancja stanowi tylko małą, a nawet bardzo małą część potrzebnej rezystancji R_A . W przypadku fotorezystora problem nieliniowości jest zdecydowanie mniejszy. Jest ogromnie ważny w przypadku wykorzystania tranzystora JFET, który może pracować w roli regulowanej rezystancji dla przebiegów przemiennych.

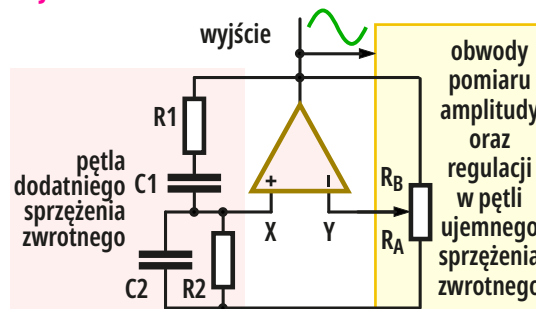
Generatory z tranzystorem JFET

Podkreślam, że w roli potrzebnej tu zmiennej, regulowanej rezystancji nie można wykorzystać obwodu kolektor – emiter tranzystora bipolarnego, bowiem na tej rezystancji występuje napięcie przemiennie i prąd musi płynąć w obu kierunkach. W takiej dziwnej, nietypowej roli może natomiast pracować właśnie tranzystor polowy złączowy (JFET), bo napięcie jego bramki zmienia rezystancję kanału.

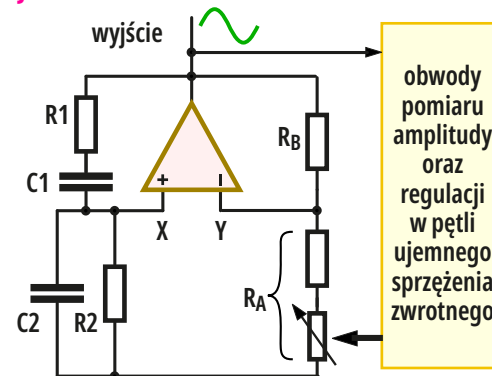
Tak, ale tylko w przypadku przebiegów o małych amplitudach, rzędu miliwoltów. Gdy amplituda rośnie, rezystancja kanału „jotfeta” staje się



Rysunek 2

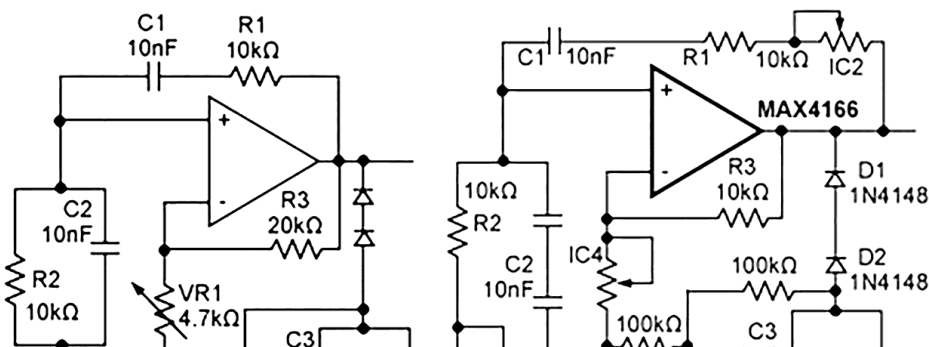


Rysunek 3

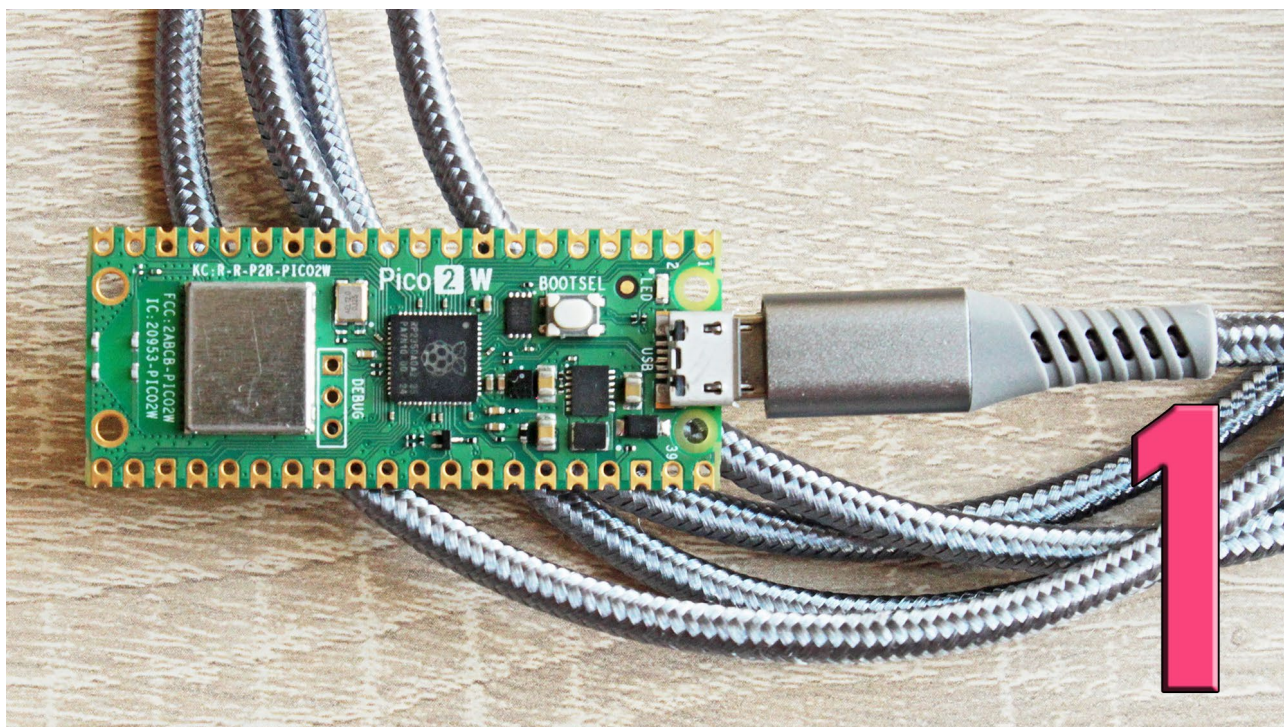


Rysunek 4

udane próby wykorzystania elektronicznych potencjometrów. Tranzystor JFET może pełnić funkcję rezystora o regulowanej wartości, ale tylko wtedy, gdy występują na nim bardzo małe przebiegi zmienne o amplitudach rzędu miliwoltów. Dlatego po pierwsze trzeba pracować z małym napięciem na tranzystorze, a po drugie, w niektórych rozwiązaniach wykorzystuje się dodatkowy obwód linearyzacji, zwierający dwa rezystory 100 k Ω (i ewentualnie kondensator), co wiadać na schemacie z prawej strony rysunku 5.



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



„Malinka” dla początkujących, część 1

Moduł określany jako Raspberry Pi Pico 2W z mikrokontrolerem RP2350 to gotowe środowisko o ogromnych możliwościach. Aby móc go wykorzystać, konieczne jest opanowanie sztuki programowania tego modułu. Okazuje się, że środowisko Arduino daje taką szansę.

Instalacja Arduino

Zainstalowanie kompilatora dla architektury ARM

Pierwszy program

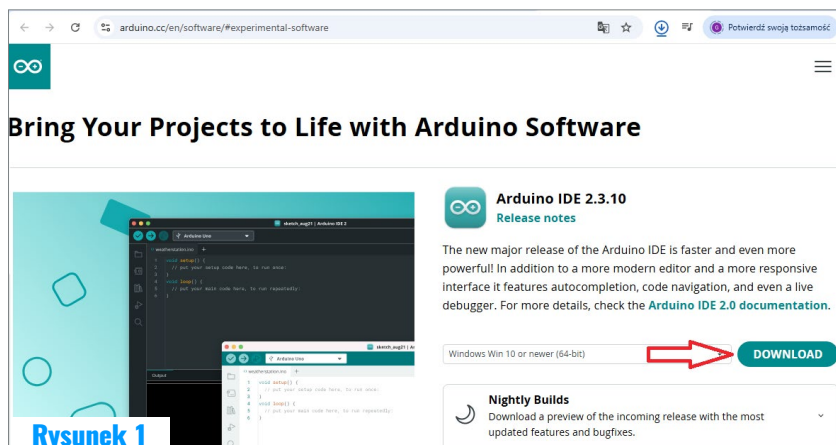
Drugi program

Do tworzenia programów dla tego modułu można wykorzystać kilka narzędzi, od zaawansowanych do zupełnie prostych, takich jak środowisko Arduino. Chociaż pakiet Arduino zwyczajowo jest kojarzony z mikrokontrolerami AVR, obecnie rozszerza swoje możliwości o nowe platformy, do których można zaliczyć tytułowe *Raspberry Pi*. Zanim ogarnie nas radość z uruchomienia swojego pierwszego programu w tym środowisku musimy zacząć od zainstalowania oprogramowania narzędziowego.

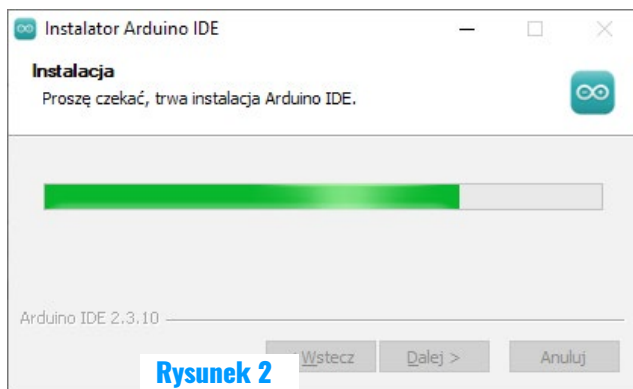
Instalacja Arduino

Do zabawy z Raspberry Pi najlepiej nadaje się oprogramowanie Arduino IDE w wersji min 2,0,

który można pobrać ze strony <https://www.arduino.cc/en/software/#experimental-software> (rysunek 1).



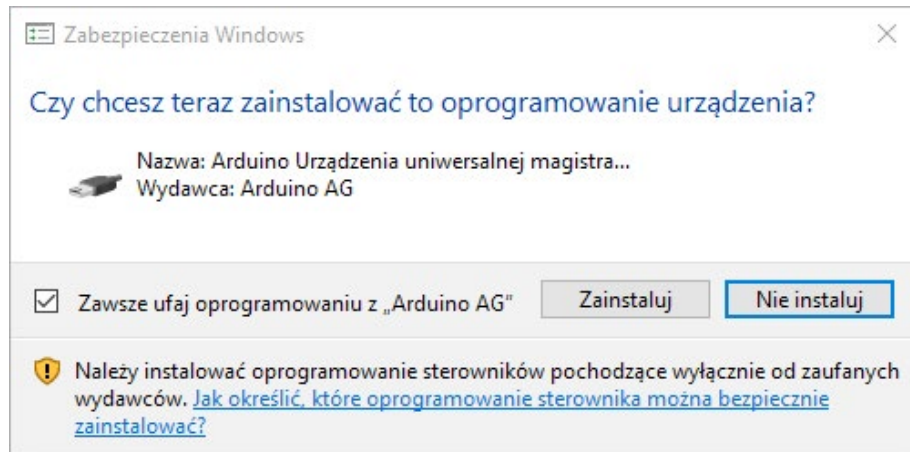
Rysunek 1



Rysunek 2

W przypadku, gdy takie oprogramowanie jest zainstalowane, krok ten można pominąć. Ściągnięty pakiet należy uruchomić (rysunek 2). Sama instalacja nie jest skom-

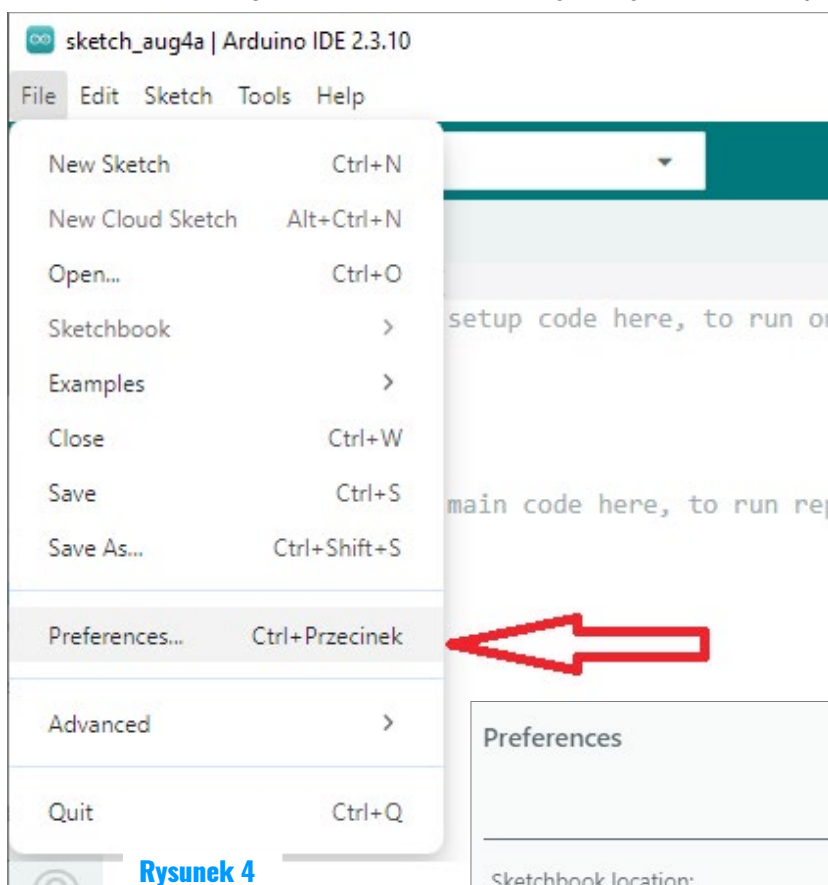
Rysunek 3



la edytować tekst źródłowy programu oraz go kompilować uruchamiając w tle standardowy kompilator GCC języka C/C++, przeznaczony dla mikrokontrolerów AVR. Uruchamiany w tle kompilator jest elementem niezależnym od IDE, więc może być związany z całkowicie inną architekturą.

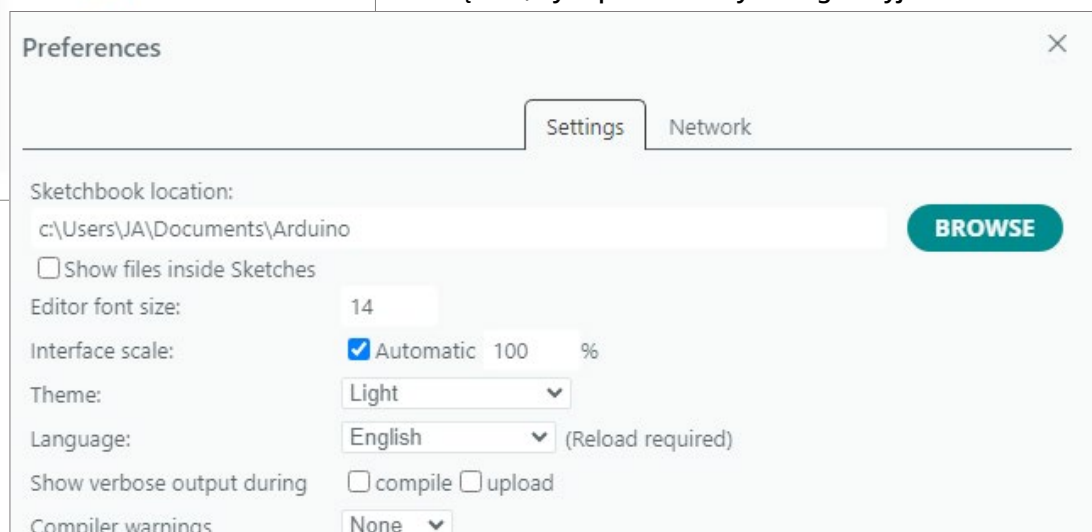
Zainstalowanie kompilatora dla architektury ARM

Po uruchomieniu Arduino należy doinstalować kompilator dla innych architektur mikrokontrolerów wraz ze specyfikacją płytek rozwojowych zawierających mikrokontrolery spod znaku Raspberry Pi. W tym celu należy wejść w opcję *File* → *Preferences* (rysunek 4). W okienku *Preferences* odszukać pole *Additional boards manager* (rysunek 5) i wpisać tam https://github.com/earlephilhower/arduino-pico/releases/download/global/package_rp2040_index.json (jeżeli pole to zawiera już jakieś informacje, dotyczące przykładowo ESP32, oddzielić je przecinkiem). Kliknąć OK, by zapisać zmiany konfiguracyjne.



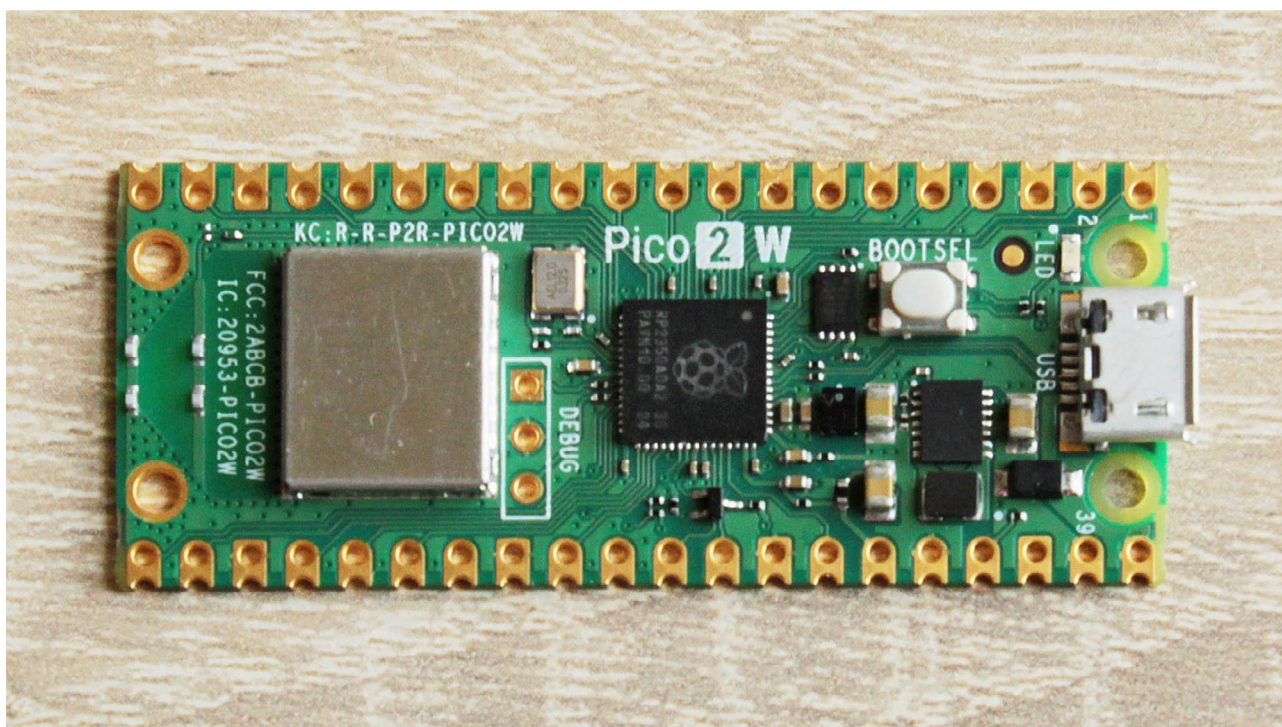
Rysunek 4

plikowana, czasami będzie wymagać zgody użytkownika (rysunek 3). Po zakończeniu tej operacji należy uruchomić samo Arduino. Domyślnie ten pakiet jest przygotowany do „obróbki” mikrokontrolerów z ro-



Rysunek 5

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



Moduł Raspberry Pi Pico 2 W

Nowoczesne mikrokontrolery z punktu widzenia hobbysty miały poważną wadę – trudno było je aplikować we własnych konstrukcjach. Ta bariera została przełamana, przemysł elektroniczny produkuje gotowe do użycia moduły z nowoczesnymi mikrokontrolerami.

Mikrokontroler RP2350

Moduł z mikrokontrolerem RP2350

Ładowanie kodu programu do Flash

Narzędzia do programowania Raspberry Pi Pico 2 W

Obecnie przemysł elektroniczny daje bardzo szeroką ofertę mikrokontrolerów, od najprostszych (jak choćby popularne AVR-y) do bardzo zaawansowanych, do których bez wątpienia należą ARM-y. Mikrokontrolery ARM (ang. Advanced RISC Machine) to rodzina architektur (modeli programowych) procesorów 32-bitowych oraz 64-bitowych, typu RISC (ang. reduced instruction set computing) – procesorów o zredukowanej liście instrukcji a jednocześnie o dużej wydajności przetwarzania. Procesory z architekturą ARM są jednymi z najczęściej stosowanych na świecie. Używa się ich między innymi w sprzęcie komputerowym, telefonach komórkowych, routerach, kalkulatorach. Obecnie wiele firm półprzewodnikowych produkuje te układy, należące do rodziny ARM, jako produkty licencyjne. Brytyjskie przedsiębiorstwo z siedzibą w Cambridge, ARM Holdings, zajmuje się projektowaniem mikropro-

cesorów. W przeciwieństwie do innych potentatów mikroprocesorowych (np. Intel), ARM Holdings sprzedaje jedynie licencje na zaprojektowane architektury, nie produkując swoich własnych. Warunki licencyjne oznaczają, że mają ujednoliconą architekturę bez względu na końcowego producenta.

Mikrokontroler RP2350

Jednym z oferentów procesorów ARM jest również brytyjska firma Raspberry Pi Ltd z siedzibą w Cambridge (ciekawa zbieżność), która zaprojektowała mikrokontroler RP2350. Z kilku względów jest to dosyć niezwykły układ. Mikrokontroler ten, zaprojektowany przez Raspberry Pi Ltd, zawiera w sobie dwie różne architektury sprzętowe jednocześnie: ma dwurdzeniowy procesor ARM Cortex-M33 (licencyjny z ARM Holdings) oraz, również dwurdzeniowy, procesor Hazard3 RISC-V (własne

opracowanie jako rozwiązanie otwarte). Podczas uruchamiania urządzenia można wybrać, na których rdzeniach ma pracować mikrokontroler. Oznacza to, że układ może działać jako procesor ARM Cortex-M33, oferując pełne wsparcie standardowych narzędzi programistycznych ARM, a w razie potrzeby pozwala na przełączenie się na otwartą architekturę RISC-V.

Procesor dwurdzeniowy (ang. dual-core) to jeden układ scalony, który integruje w sobie dwa niezależne, fizyczne rdzenie przetwarzające instrukcje. Można to sobie wyobrazić jako jeden komputer, który posiada dwa osobne „mózgi”, pracujące jednocześnie nad powierzonymi zadaniami. System operacyjny kontrolujący pracę RP2350 może zlecić jednemu rdzeniowi obsługę programu (np. gry), a drugiemu procesy w tle (np. odtwarzanie muzyki). Innym wariantem jest praca zespołowa, gdzie oba rdzenie mogą współpracować nad jednym ciężkim zadaniem, dzieląc je na mniejsze części (tzw. wielowątkowość).

Wspomniany wyżej procesor ARM Cortex-M33 jest rozwinięciem wariantu Cortex-M4 (to 32-bitowy rdzeń procesora stworzony z myślą o wydajnym przetwarzaniu sygnałów w systemach wbudowanych o niskim poborze prądu, mający dużą wydajność w zakresie przetwarzania sygnałów DSP oraz obliczeń zmiennoprzecinkowych). Seria Cortex to grupa nowoczesnych rdzeni procesorów firmy ARM Holdings dzieląca się na trzy główne linie: A – aplikacyjne, R – czasu rzeczywistego oraz M – mikrokontrolery. M33 to oznaczenie konkretnego modelu rdzenia, bazującego na nowoczesnej architekturze ARMv8-M.

Skoro dwa rdzenie w mikrokontrolerze pracują jednocześnie (sięgają do pamięci z programem, pamięci operacyjnej RAM), musi istnieć jakiś arbiter synchronizujący ich pracę. Tę funkcję pełnią magistrala systemowa (Bus Matrix) i kontroler dostępu do pamięci (Memory Controller). W RP2350 pamięć RAM jest podzielona na banki, możliwy jest jednocześnie dostęp do różnych banków. Dodatkowo występuje tutaj pamięć typu Cache (szybka pamięć podręczna). Jest to niewielka pamięć wbudowana bezpośrednio w procesor. Jej główną rolą jest pośredniczenie między szybkim rdzeniem procesora a znacznie wolniejszą pamięcią główną (np. Flash) w celu likwidacji tzw. wąskiego gardła. Procesor wykonuje operacje w ułamkach nanosekund, podczas gdy pobranie danych z zewnętrznej pamięci Flash trwa wielokrotnie dłużej, więc rdzeń musiałby

Przyglądając się wewnętrznej strukturze mikrokontrolera (**rysunek 1**, na następnej stronie) można dostrzec, że nie zawiera on wewnętrznej pamięci Flash przeznaczonej generalnie do przechowywania kodu programu. To kolejne cecha, która odróżnia „malinkę” od wielu innych mikrokontrolerów (przykładowo AVR mają wewnętrzną pamięć Flash na kod programu). Analizując różne rozwiązania przyłączenia pamięci Flash do procesorów, zwyczajowo stosowane są rozwiązania z interfejsem równoległym (przesyłane są jednocześnie wszystkie bity, przykładowo 8). W mikrokontrolerze RP2350 stosowana pamięć Flash jest rozwiązaniem szeregowym. Mikrokontroler komunikuje się z pamięcią za pomocą interfejsu QSPI (Quad SPI – czterobitowej szeregowej magistrali SPI). Pozwala to znacząco zredukować liczbę pinów do obsługi pamięci (magistrala QSPI potrzebuje 6 linii). Dzięki temu większość fizycznych pinów mikrokontrolera RP2350 pozostaje wolna do dyspozycji użytkownika. Choć wykorzystana jest pamięć szeregową, tryb „Quad” pozwala na przesyłanie 4 bitów danych jednocześnie w każdym cyklu zegara, dodatkowo w połączeniu z technologią DDR (Double Data Rate – odczyt następuje na narastającym i opadającym zboczach sygnału zegarowego) zapewnia błyskawiczny dostęp do danych.

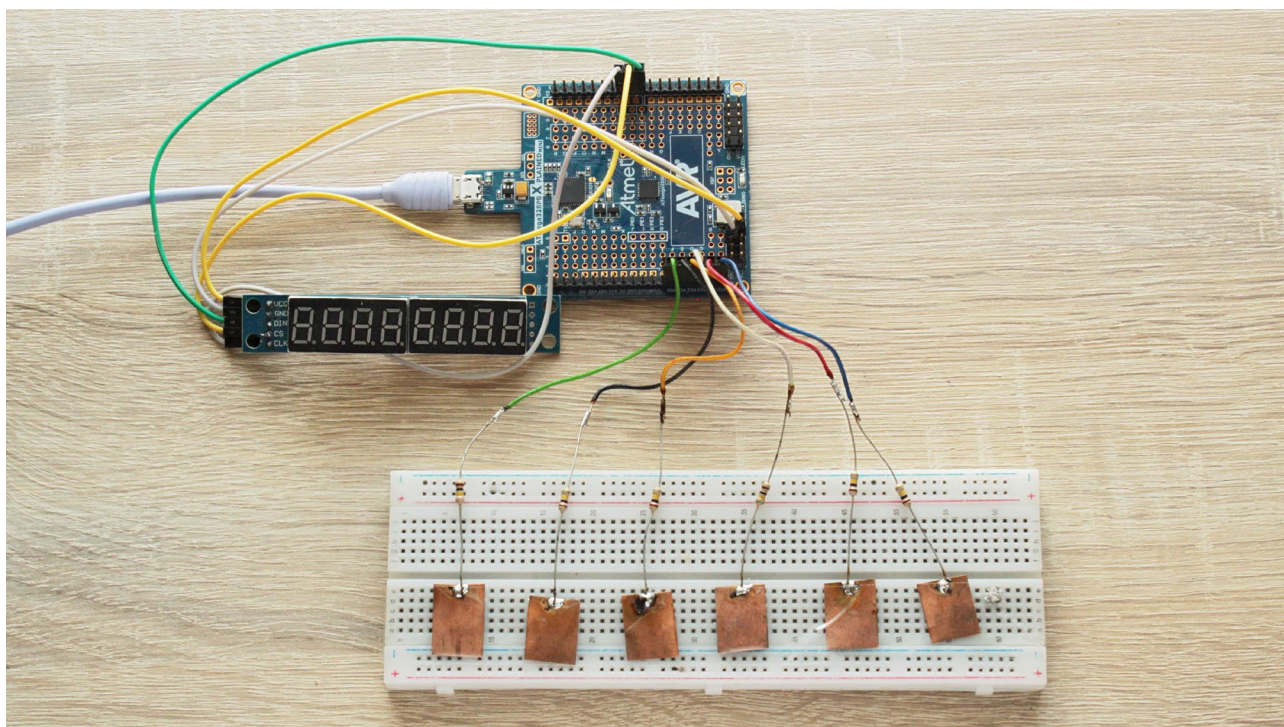
Wewnątrz mikrokontrolera RP2350 znajduje się dedykowany, sprzętowy kontroler XIP (Execute-in-Place) z dedykowaną pamięcią podręczną (Cache). Sprawia on, że z punktu widzenia rdzeni ARM czy RISC-V ta szeregową pamięć Flash jest „widziana” tak, jakby była zwykłą, szybką, równoległą pamięcią wewnętrzną przypisaną do konkretnego adresu w przestrzeni adresowej procesora.

Moduł z mikrokontrolerem RP2350

Firma Raspberry Pi Ltd oferuje również moduł z wlutowanym mikrokontrolerem (właściwie serię modułów o różnym wyposażeniu) o ogólnej nazwie *Raspberry Pi Pico 2*. Seria *Pico 2* nie jest jedyną, Raspberry Pi Ltd oferuje również inne, o bardziej lub mniej rozbudowanym wyposażeniu. Nawet krótkie przedstawienie wszystkich wykracza poza jeden artykuł.

Ograniczając się do serii *Pico 2* – występuje tu kilka wariantów sprzętowych modułu. Popularnym wariantem jest *Pico 2* (**fotografia 2**, którą udostępnił Thomas Amberg – na kolejnej stronie). Artykuł poświęcony jest wersji *Pico 2 W*, którą pokazuje fotografia tytułowa. Główna różnica między płytkami *Raspberry Pi Pico 2* a *Raspberry Pi Pico 2 W* sprowadza się do obecności modułu łączności bezprze-

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

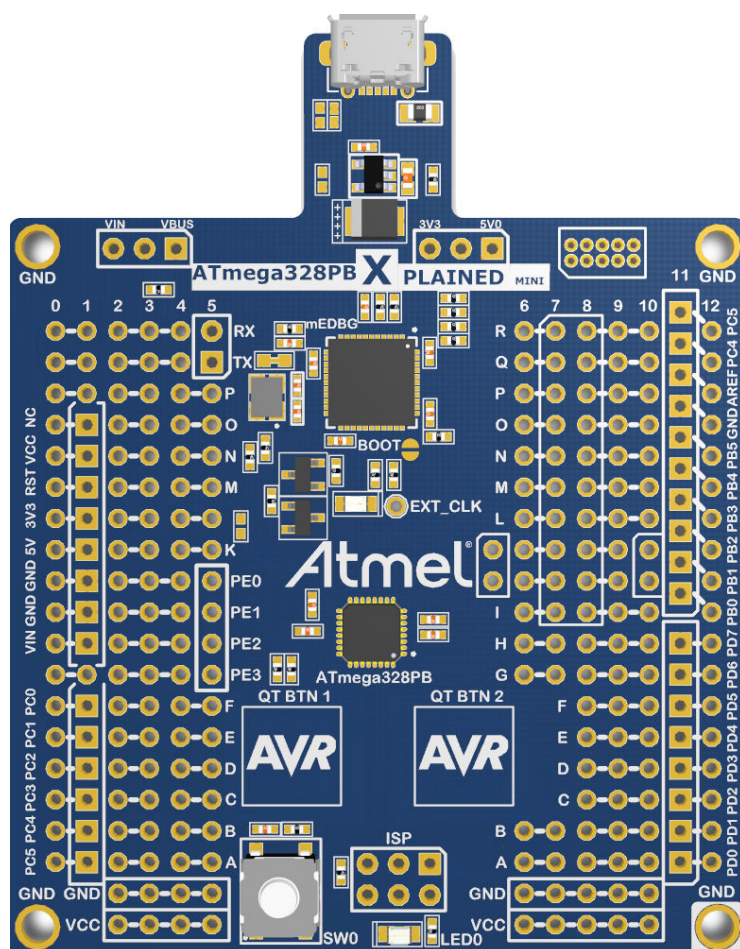


Przyciski dotykowe w Xplained

Firma Microchip (obecny producent mikrokontrolerów AVR) udostępniła użytkownikom moduł Xplained Mini zawierający na pokładzie mikrokontroler ATMEGA328PB. Jednocześnie dla tego modułu udostępnia kilka gotowych programów. Jeden z nich prezentuje obsługę pól dotykowych.

Idea wykrywania dotknięć
Oryginalny programu obsługi dotyku
Zmiany w konfiguracji projektu
Praktyczne rozwiązanie
Inne parametry konfigurujące PTC
Ograniczenia mikrokontrolera

Według oryginalnej dokumentacji (Microchip Technology) sam moduł *Xplained Mini* zawierający mikrokontroler ATMEGA328PB, jest dostępny w dwóch wariantach, które różnią się sposobem rozwiązania (oraz liczbą) pól dotykowych stanowiących bazę do programowego rozpoznawania ich dotknięcia. Jedna wersja (**rysunek 1**), zawiera dwa pola dotykowe (na PCB opisane jako QT BTN 1 oraz QT NTN 2).



Rysunek 1

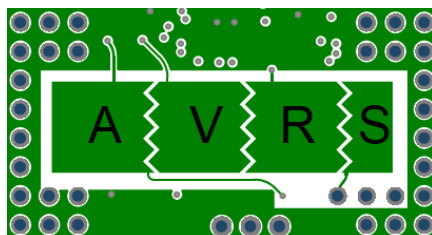
Inna wersja (**rysunek 2**) w tym miejscu oferuje cztery pola dotykowe (gdy popatrzyć na płytkę drukowaną „pod światło”). W rzeczywistości jest to bardziej złożone, dwa moduły stanowią bazę dla różnych wariantów obsługi pól dotykowych. W tej samej dokumentacji jest informacja, że pole S (**rysunek 3**) jest połączone z polem A przez zwórkę $0\ \Omega$. Takie rozwiązanie umożliwia zbudowanie rozwiązań typu suwak (ang. slider).

Idea wykrywania dotknięć

Współczesne interfejsy użytkownika coraz częściej rezygnują z mechanicznych przełączników na rzecz paneli dotykowych. W mikrokontrolerze ATmega-328PB zaimplementowano dedykowany, sprzętowy podzespół – *Peripheral Touch Controller* (PTC). Jest to wyspecjalizowany podsystem, który całkowicie odciąża jednostkę centralną od procesów generowania sygnałów, próbkowania i wstępnej filtracji danych. Idea działania kontrolera PTC jest technologią transferu ładunku, realizowana w dwóch zasadniczo odmiennych wariantach architektonicznych: *Self-Capacitance* (pomiar pojemności własnej) oraz *Mutual-Capacitance* (pomiar pojemności wzajemnej).

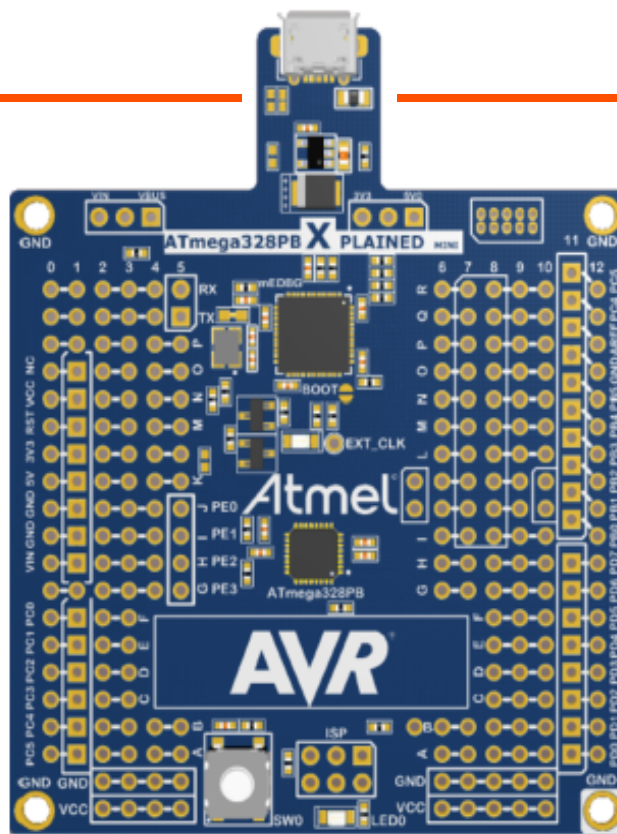
W wariantcie *Self-Capacitance* sensor stanowi pojedynczą, odizolowaną elektrodę przewodzącą, połączoną z jednym pinem wejściowym mikrokontrolera (linią Y). W ujęciu elektrostatycznym, elektroda ta tworzy okładzinę kondensatora, którego drugą, wirtualną okładziną jest masa układu (GND). Przed interakcją z użytkownikiem, układ ma stałą, bazową pojemność pasożytniczą. Zależy ona bezpośrednio od geometrii ścieżek, grubości laminatu oraz odległości od wewnętrznych płaszczyzn masy. Wprowadzenie ludzkiego palca w bliskie sąsiedztwo elektrody tworzy nowy obwód elektryczny. Ciało człowieka, reprezentujące relatywnie dużą pojemność, łączy pole dotykowe z potencjałem ziemi, wnosi do układu dodatkową pojemność. Można to traktować jako dołączenie do istniejącego obwodu kolejnego kondensatora w sposób równoległy. Z punktu widzenia mikrokontrolera całkowita pojemność widziana przez pin mikrokontrolera wzrasta. Sprzętowy kontroler PTC realizuje pomiar poprzez cykliczne i niezwykle precyzyjne przesyłanie

Pola dotykowe są połączone z mikrokontrolerem przez szeregowo włączony rezystor o wartości $100\text{k}\Omega$



Pole A jest połączone z PE2
Pole V jest połączone z PE3
Pole R jest połączone z PC3
Pole S jest połączone zworką $0\ \Omega$ z polem A

Rysunek 3

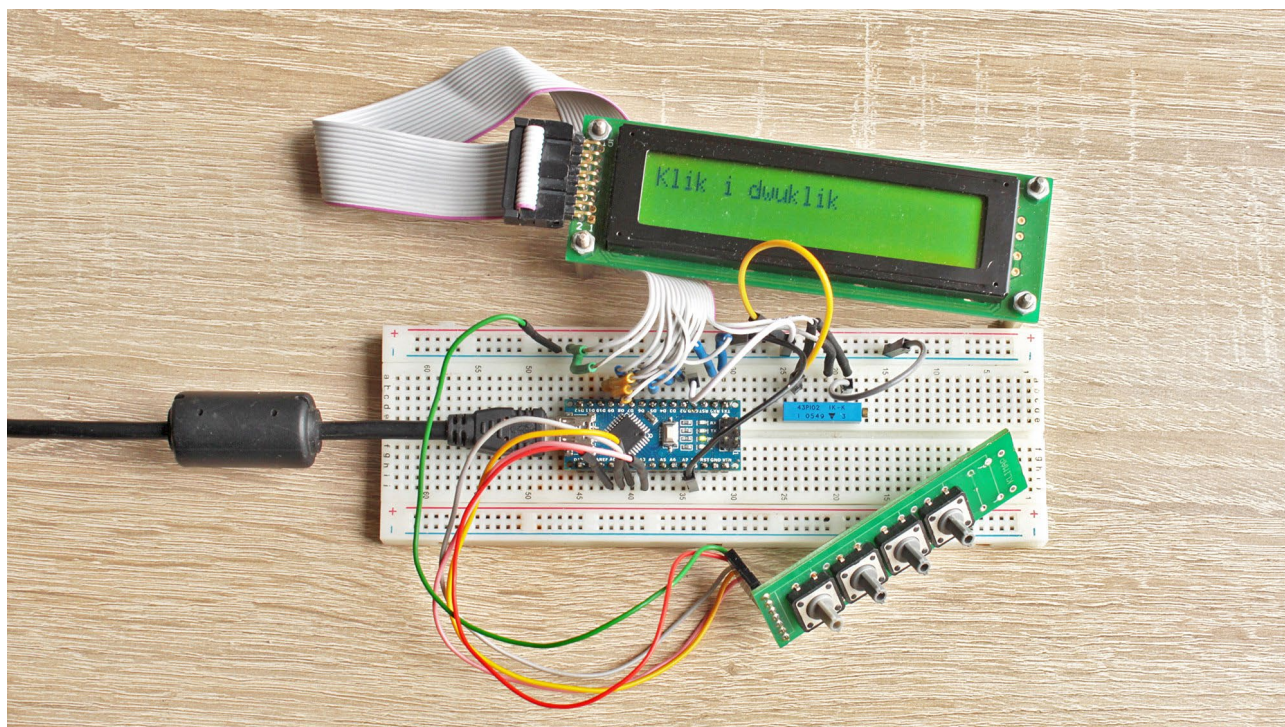


Rysunek 2

PTC. Proces ten przebiega w ściśle zsynchronizowanych fazach. Pierwszą jest zerowanie, gdzie wewnętrzny kondensator oraz zewnętrzna elektroda sensora są całkowicie rozładowywane do potencjału masy (GND), aby usunąć wszelkie ładunki szcztkowe. W drugim kroku następuje ładowanie kondensatora. Zewnętrzna elektroda (pole dotykowe) zostaje podłączona do wewnętrznej szyny zasilania na ściśle określony czas, przez co gromadzi ładunek proporcjonalny do swojej aktualnej pojemności.

W kolejnym cyklu elektroda zostaje odłączona od zasilania i połączona bezpośrednio z wewnętrznym kondensatorem. Zgromadzony na elektrodzie ładunek przepływa i rozdziela się pomiędzy obie pojemności: pola dotykowego oraz kondensatora wzorcowego. Cykl ładowania i przesyłania ładunku jest powtarzany wielokrotnie w ramach jednego procesu pomiarowego. Każde powtórzenie powoduje stopniowe „pompowanie” napięcia na kondensatorze. Po zakończeniu serii transferów, napięcie na kondensatorze próbkowania jest poddawane konwersji przez dedykowany komparator analogowo-cyfrowy lub przetwornik ADC powiązany z PTC. Im większa pojemność sensora, która wzrosła przez obecność

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



Klik oraz dwuklik

Wszyscy znamy pojęcie kliku oraz dwukliku, choćby dlatego, że tego wynalazku używa każdy system operacyjny z graficznym interfejsem. Czy coś takiego da się zrealizować w środowisku prostych mikrokontrolerów? Da się i nie jest to skomplikowane.

[Standardowa obsługa klawiatury o 4 przyciskach](#)
[Nowa funkcja dekodowania klawiatury](#)

[Przykład zastosowania](#)

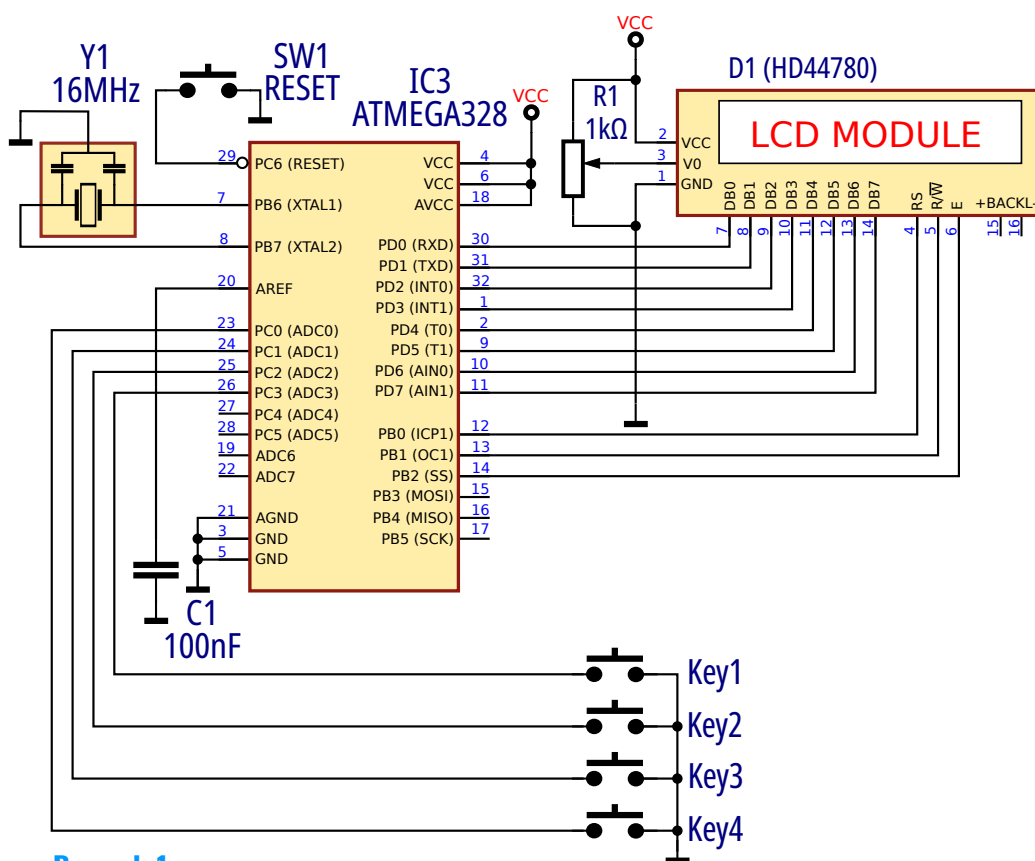
Standardowa myszka komputerowa ma dwa przyciski określone jako przycisk lewy i prawy. Pojedyncze kliknięcie generuje określone polecenie dla systemu komputerowego. Szybkie dwa kliknięcia, określone jako dwuklik, komputer interpretuje jako odmienne polecenie. A gdyby tak zrobić „dziwną” myszkę ograniczoną jedynie do przycisków, o większej liczbie przycisków (przykładowo 4 przyciski), gdzie można by każdym kliknąć, czyli wykonać pojedyncze naciśnięcie oraz wykonać dwuklik, ale w dowolnej kombinacji (klasyczny dwuklik dotyczy tego samego przycisku). Do realizacji tego projektu jest użyty moduł obsługi klawiatury (opisany w ramach *Mikroprocesorowej oślej łączki* w ZE 9/2025). To rozwiązanie zostało celowo zaprojektowane w taki sposób, by było elastyczne i żeby łatwo można je

było adaptować do różnorodnych zastosowań, bez jakichkolwiek modyfikacji samego modułu obsługi klawiatury. W projekcie został również wykorzystany moduł po raz pierwszy opisany w artykule *Generowanie dźwięków* w ZE 8/2025 przeznaczony do obsługi zdarzeń zakolejkowanych w czasie.

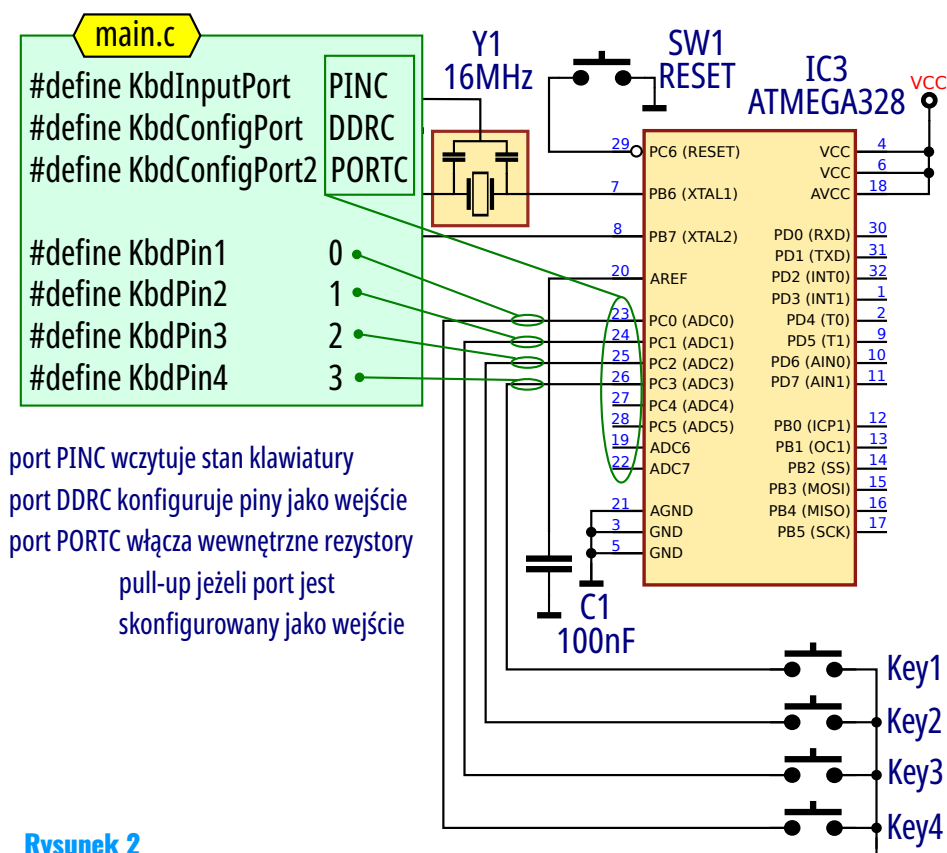
Standardowa obsługa klawiatury o 4 przyciskach

Do tego projektu oprócz samej klawiatury niezbędny jest moduł alfanumerycznego wyświetlacza LCD. Cały schemat pokazuje **rysunek 1** (na następnej stronie). W programie przyłączenie klawiatury jest określone w elastyczny sposób, z wykorzystaniem mechanizmu `#define` – szczegóły przedstawia **rysunek 2** (na następnej stronie).

Istotne elementy związane z obsługą klawiatury prezentuje **listing 1** (na następnej stronie). Znaczenie poszczególnych instrukcji wyjaśnia **rysunek 3** (na następnej stronie). Konfiguracja środowiska mikrokontrolera, gdzie jest przyłączona klawiatura (*EnvirInit*) określa odpowiedni port jako wejście bez modyfikacji kierunku pracy portu dla pozostałych jego wyprowadzeń z jednoczesnym włączeniem wbudowanych w mikrokontroler rezystorów pull-up. W sytuacji, gdy przycisk nie jest naciśnięty, wymuszany jest stan logicznej jedynki (zwalnia to z użycia zewnętrznych rezystorów podciągających). W funkcji *main* występują wywołania funkcji określających szczególnie związane z obsługą klawiatury. Przede wszystkim do modułu jej obsługi dołączona jest tabela przekodowań (zamieniająca bitowe kombinacje generowane przez naciśnięte przyciski na użyteczne w programie kody znaków ASCII, choć w tym przypadku mogą to być dowolne liczby 8-bitowe).



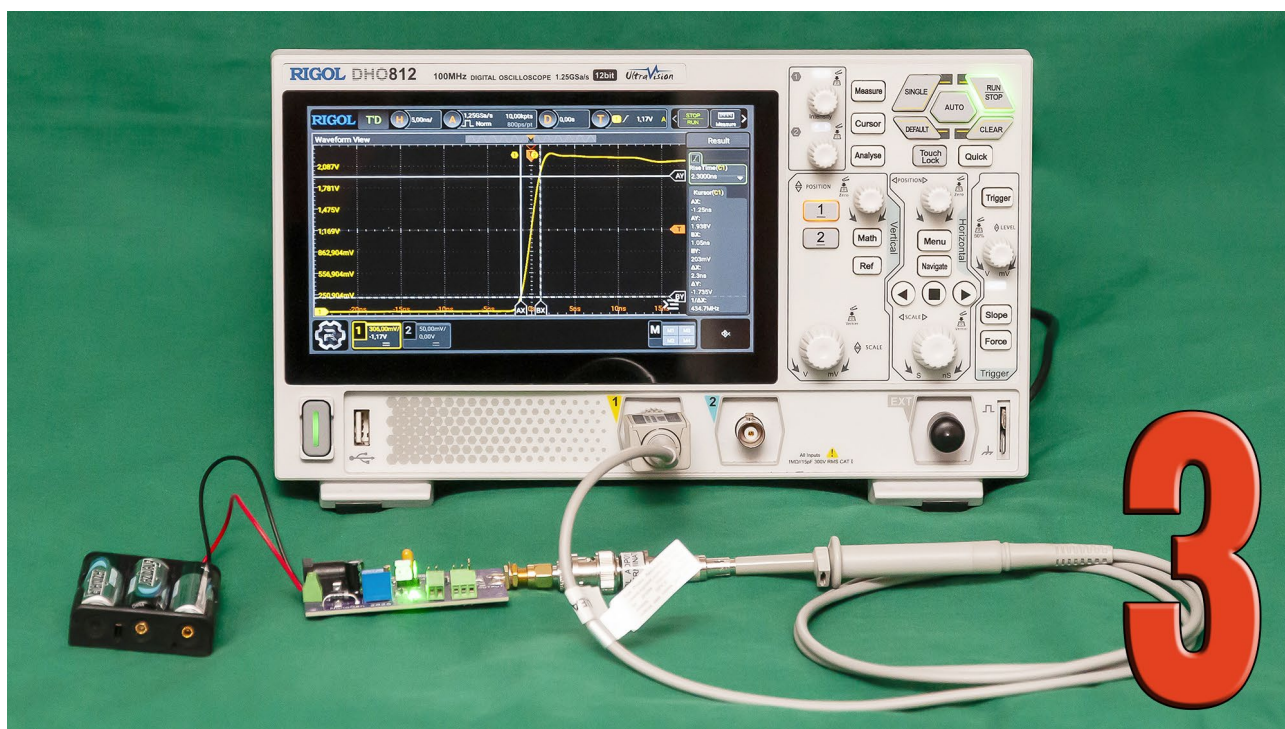
Rysunek 1



port PINC wczytuje stan klawiatury
port DDRC konfiguruje piny jako wejście
port PORTC włącza wewnętrzne rezystory pull-up jeżeli port jest skonfigurowany jako wejście

Rysunek 2

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



Impulsowy test oscyloskopów i sond pomiarowych, część 3

Zmontowałem i przetestowałem „NanoGen – generator nanosekundowy” opisany w ZE 3/2026. Było to bardzo interesujące, wzbogacające wiedzę, ale też zaskakujące doświadczenie. W czterech artykułach chciałbym podzielić się swoimi spostrzeżeniami. Mam nadzieję, że moje uwagi przydadzą się innym Czytelnikom.

Charakterystyka impulsowa badanych sond oscyloskopowych
Sondy pasywne
Sondy budżetowe

Aktywna sonda różnicowa Hantek HT8100
Porównanie wyników badań
Wnioski z porównania sond

Charakterystyka impulsowa badanych sond oscyloskopowych

Po wyznaczeniu referencyjnej odpowiedzi impulsowej oscyloskopu Rigol DHO812, przystąpiłem do oceny wpływu różnych sond pomiarowych na odpowiedź przejściową całego toru pomiarowego. Celem eksperymentu nie było bezpośrednio wyznaczenie rzeczywistego pasma poszczególnych sond, lecz porównanie ich zachowania w identycznych warunkach pomiarowych przy pobudzeniu bardzo szybkim impulsem nanosekundowym, generowanym przez opisany wcześniej nanogenerator.

Wszystkie pomiary wykonałem przy użyciu tego

samego stanowiska oraz tej samej konfiguracji wejściowej oscyloskopu. Sondy pracowały w trybie tłumienia, zgodnym z ich przeznaczeniem, a połączenie realizowałem za pomocą adaptera BNC umożliwiającego bezpośrednie przyłączenie końcówki pomiarowej do wyjścia generatora. Dzięki temu ograniczyłem wpływ dodatkowych przewodów, nieciągłości impedancyjnych oraz pasożytniczych elementów RLC.

Ponieważ referencyjny czas narastania oscyloskopu Rigol DHO812 wynosi około 2,65 ns, uzyskane wyniki należy interpretować przede wszystkim jako ocenę wpływu sond na odpowiedź impulsową całego systemu pomiarowego. W badanym układzie

różnice pomiędzy sondami ujawniają się głównie poprzez poziom przeregulowania, charakter oscylacji przejściowych, czas stabilizacji poziomu ustalonego oraz wpływ na automatycznie wyznaczany czas narastania.

W praktyce parametry te dostarczają cennych informacji o jakości kompensacji częstotliwościowej sondy, jej efektywnej pojemności wejściowej oraz podatności na rezonanse pasożytnicze występujące w rzeczywistym torze pomiarowym.

Sondy pasywne

Rigol PVP3150 (150 MHz)

Specyfikacja fabryczna:

- pasmo przenoszenia: 150 MHz,
- deklarowany czas narastania: 2,3 ns,
- pojemność wejściowa: 10 pF ± 5 pF.

Badaniom poddałem dwa egzemplarze sondy Rigol PVP3150. Pomiar wykonałem w skróconym torze pomiarowym, po wcześniejszej eliminacji dodatkowych adapterów i złączy mogących wpływać na odpowiedź wysokoczęstotliwościową układu. Oba egzemplarze wykazały bardzo dobrą powtarzalność wyników (**rysunek 11**).

Automatyczny pomiar czasu narastania wskazał, $\tau_{\text{trzm}} = 2,60$ ns, natomiast pomiar kursorowy dał wartość $\Delta X \approx 2,45$ ns.

Uzyskany wynik jest nieznacznie krótszy od referencyjnego czasu narastania oscyloskopu, wynoszącego około 2,65 ns. Formalnie prowadzi to do sprzeczności przy próbie zastosowania klasycznego modelu RSS, ponieważ po podstawieniu wartości do równania, pod pierwiastkiem pojawia się wartość ujemna. Zjawisko to nie oznacza jednak rzeczywistego przekroczenia możliwości oscyloskopu, lecz jest skutkiem sposobu kształtowania odpowiedzi impulsowej przez sondę i tor pomiarowy.

Analiza przebiegu wskazuje na występowanie umiarkowanego przeregulowania. Dla poziomu ustalonego, wynoszącego około 2,18 V obserwowałem maksymalną wartość około 2,37 V, co odpowiada overshootowi rzędu 8,7%.

Po wystąpieniu przeregulowania oscylacje przejściowe szybko wygasają, a przebieg osiąga stabilny poziom po około 16–18 ns. Odpowiedź impulsowa



Rysunek 11

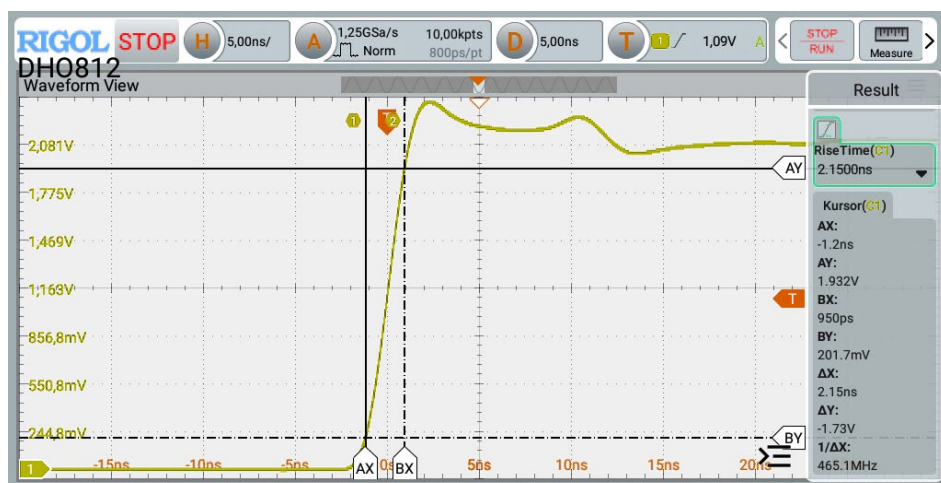
cechą sondy jest niewielkie wyostrzenie odpowiedzi impulsowej, które może wpływać na sposób wyznaczania poziomów 10–90% przez automatykę oscyloskopu. Pomimo umiarkowanego overshootu oba egzemplarze wykazały bardzo dobrą powtarzalność parametrów, potwierdzając zgodność z klasą sond 150 MHz i dobrą jakość wykonania.

Rigol PVP2350 (350 MHz)

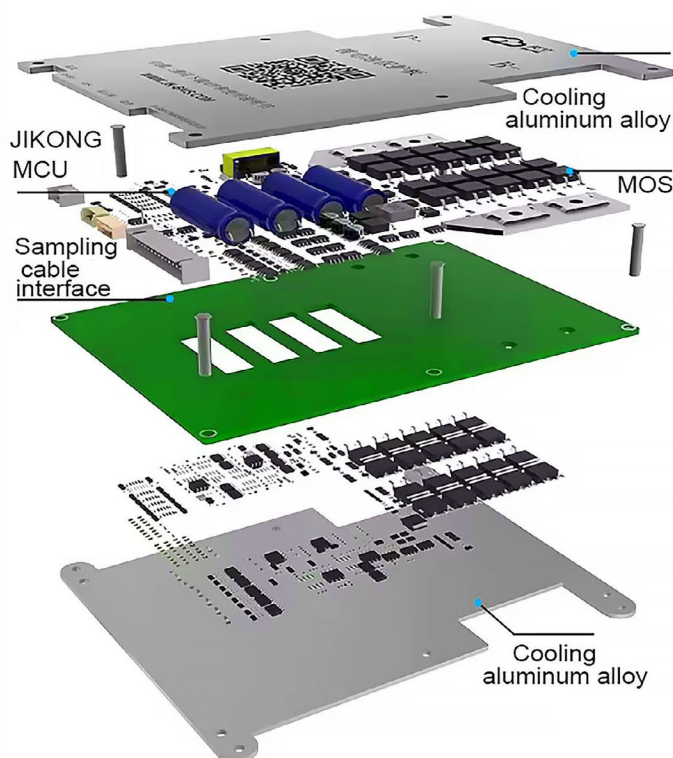
Specyfikacja fabryczna:

- pasmo przenoszenia: 350 MHz,
- deklarowany czas narastania: 1,0 ns,
- pojemność wejściowa: 10 pF ± 5 pF.

Rigol PVP2350 jest najszybszą z badanych sond pasywnych producenta. W badanym układzie uzyskałem najkrótszy czas narastania spośród wszystkich sond pasywnych, co wskazuje na bardzo dobrą odpowiedź wysokoczęstotliwościową sondy. Pomiar wykonałem już w zoptymalizowanej konfiguracji stanowiska pomiarowego, po wcześniejszym wyeliminowaniu dodatkowego łącznika BNC-BNC opisanego w części 2 (**rysunek 12**).



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



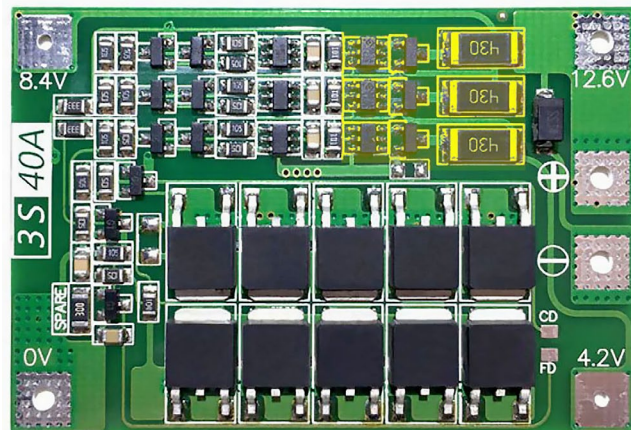
Akumulatory litowe – balansery pasywne i aktywne

We wcześniejszych artykułach cyklu omówione były zasady ładowania i rozładowania. W ostatnich opisane były zasady i sposoby ochrony pojedynczych ogniw i baterii ogniw – akumulatorów litowych. W poniższym artykule, czwartym, w przystępny sposób omawiam niełatwy, złożony temat balansowania.

Balansowanie podczas ładowania Balansery „skokowe”

Typowe balansery pasywne Balansery aktywne

W poprzednich artykułach tej serii omówiłem specyfikę akumulatorów litowych, podkreślałem ich delikatność, konieczność stosowania różnych zabezpieczeń. Omówiłem kostki DW01 i pokrewne, służące do ochrony pojedynczych ogniw. W poprzednim artykule doszliśmy do BMS-ów, czyli systemów ochrony baterii ogniw. Doszliśmy do popularnych rozwiązań i ich realizacji w postaci charakterystycznych modułów. Jeden z takich modułów, czterokanałowy BMS pokazany jest na **fotografii 1**. Wiemy, że moduł zawiera cztery sześcionóżkowe układy DW01 lub odpowiedniki, rodzaj tranzystorowych bramek AND oraz rezystory do ochrony nadprądowej. Intryguje obecność wyróżnionych żółtym kolorem, trzech dodatkowych układów scalonych i czterech rezystorów oznaczonych 430.



100mA Balanced current
40A continuous current

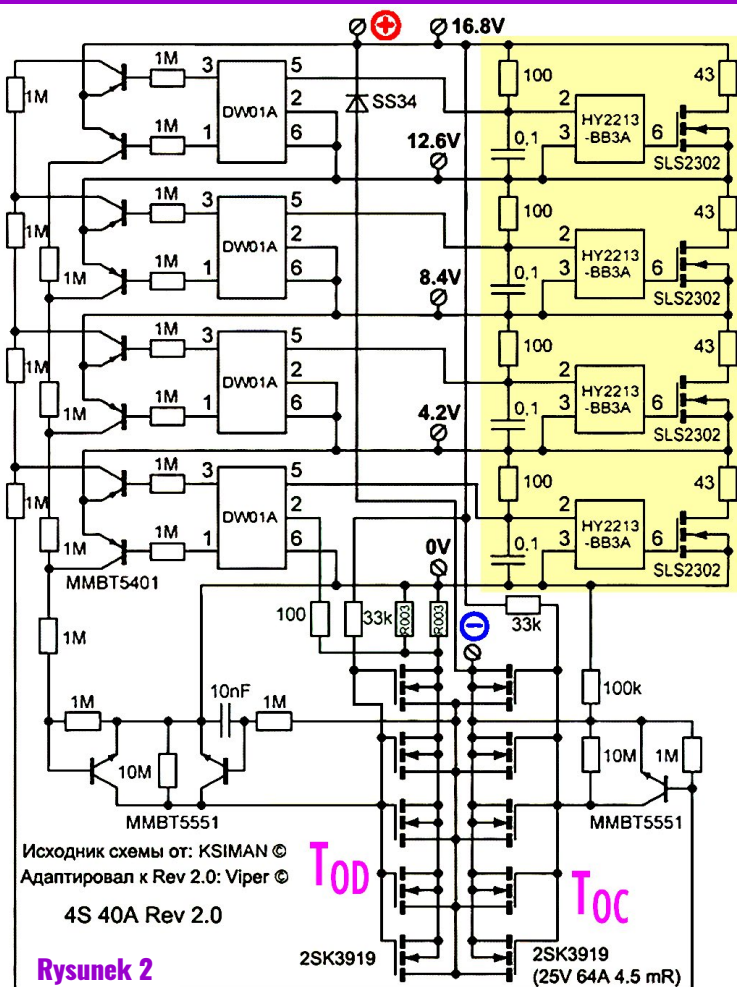
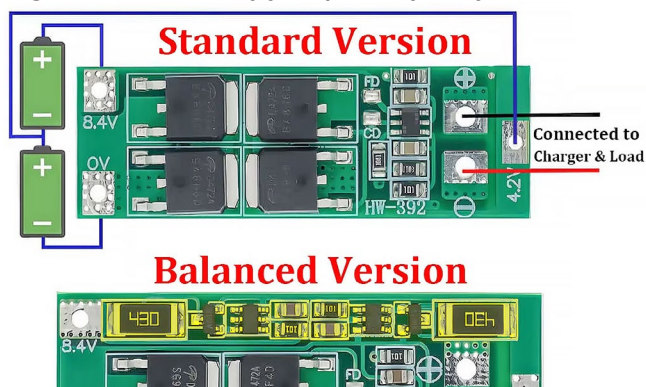
Fotografia 1 Can be enabled drill motor

Tego rodzaju moduły realizowane są według schematów takich lub podobnych, jak pokazuje **rysunek 2**. Schemat tych intrygujących obwodów widoczny jest w prawej górnej części rysunku, wyróżnionej kolorem żółtym. To są obwody balansera. Dostępne są BMS-y bez balansera oraz z balanserem – kolejny przykład na **fotografii 3**. Słowo **balanser** (ang. *balancer*) pochodzi od słowa **balans**, a z nim nieodłącznie kojarzy się pojęcie równowagi i szukania równowagi. Czyli nazwa balanser sugeruje zrównoważenie, wyrównanie – zapewne wyrównanie napięcia, a jeszcze lepiej ładunku czy energii poszczególnych ogniw.

Nazwa coś obiecuje, a raczej dużo obiecuje: że BMS z balanserem spowoduje jakiegoś rodzaju wyrównanie, a to nasuwa wniosek, że pozwoli uzyskać z akumulatora więcej energii, a raczej całą energię ze wszystkich ogniw baterii.

Prawda jest dużo bardziej skomplikowana, a zdecydowana większość użytkowników ma nadmierne oczekiwania i całkowicie błędne wyobrażenia na temat roli i realnych korzyści z obecności balanserów. Te nadmierne oczekiwania powodują i utwierdzają nieprecyzyjne i błędne wpisy na stronach internetowych. Przykładem może być fragment artykułu z pewnej, wartościowej skądinąd, strony internetowej: *BMS monitoruje napięcia ogniw, czy to podczas ładowania, czy rozładowywania po wykryciu np. spadku napięcia, na którymś ogniwie odłącza układ i analogicznie podczas ładowania balansuje cele i jeżeli ktoś przekroczy górny próg napięcia (spadek natężenia prądu poniżej 0,05C ok 130 mA), to zostaje odłączona, a reszta ładuje się do tego samego poziomu.*

To prawda, że każdy BMS nieustannie monitoruje napięcia ogniw. Tylko jest poważny problem z tym „odłączaniem po przekroczeniu górnego progu napięcia”, a zupełnie fałszywe nadzieje budzi stwierdzenie, że „reszta ładuje się do tego samego poziomu”. Jeżeli ktoś nie rozumie szczegółów, to może odnieść wrażenie, że każdy balanser skutecznie wspomaga ogniwa słabsze, żeby je w pełni wykorzystać.



Rysunek 2

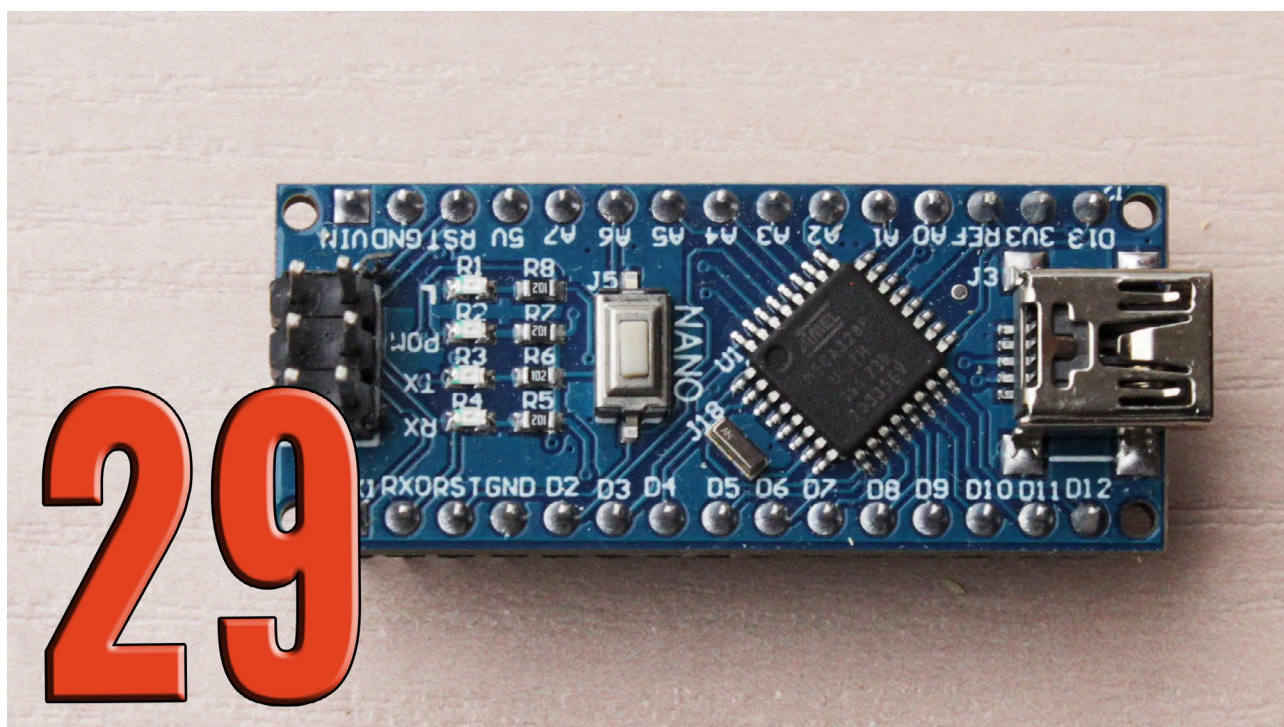
Tymczasem rzeczywistość jest inna, bardzo skomplikowana, a najprostsze balansery pasywne w wielu przypadkach tak naprawdę nie spełniają swojej roli i nic nie pomagają – to zależy też od właściwości użytej ładowarki. Tylko bardziej skomplikowane i droższe balansery aktywne pozwalają w pełni wykorzystać możliwości wszystkich ogniw baterii. Najprostsze i najpopularniejsze balansery (tzw. pasywne) mogą trochę pomóc, ale tylko w niektórych przypadkach. Zagadnienie jest dość skomplikowane i obszerne, dlatego trzeba zacząć od przypomnienia podstaw.

Balansowanie podczas ładowania

Ogólnie biorąc, nowoczesne akumulatory litowe są prawie pod wszystkimi względami dużo lepsze od kwasowych i zasadowych, ale są też dużo delikatniejsze. Bardzo łatwo je zepsuć. Już jednokrotne silne przeładowanie lub jednorazowe rozładowanie „do zera” może spowodować dużą, nieodwracalną utratę pojemności, a nawet je całkowicie uszkodzić.

Żeby nie dopuścić do przetładowania, **należy kontrolować napięcie wszystkich ładowanych ogniw i nie dopuścić, żeby napięcie któregośkol-**

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



Mikroprocesorowa ośła łączka, część 29

Istotnym elementem w oprogramowaniu jest obsługa przerwań. Ta operacja w języku programowania C jest prosta i realizuje się wręcz naturalnie. Trochę gorzej to wygląda przy tworzeniu programu w języku C++ (pomijając oczywiście rozwiązanie zaczerpnięte z języka C, gdyż język C jest podzbiorem C++).

[Obsługa przerwania w języku C++](#)

[Środowisko do eksperymentów](#)

Obsługa przerwania w jakichkolwiek mikroprocesorach ma pewną specyfikę: zawsze jest uruchamiana w sposób sprzętowy (zaistnienie potrzeby obsługi przerwania prowadzi do wywołania odpowiedniej funkcji, które nie pochodzi z kodu programu). Występuje tu kilka możliwych wariantów uzależnionych od wyposażenia sprzętowego. Niektóre mikroprocesory są wyposażone w tryb wektorowy, odpowiedni kontroler obsługi przerwania skieruje wykonanie programu w docelowe miejsce w przestrzeni pamięci zawierającej kod programu (mówiąc w uproszczeniu, wykona wywołanie dedykowanej funkcji do obsługi). W takim przypadku na początku programu koniecznością jest odpowiednie zaprogramowanie kontrolera przerwania przede wszystkim w zakresie podania adresów funkcji do

obsługi przerwania. Taka funkcja może znajdować się w dowolnym miejscu przestrzeni programu, więc kontroler przerwania musi wcześniej „wiedzieć” gdzie one się znajdują. Innym wariantem jest przydzielenie każdej funkcji obsługi przerwania stałego miejsca w przestrzeni adresowej i taki wariant występuje w mikrokontrolerach AVR. W odpowiedniej dokumentacji można znaleźć szczegółowe informacje dotyczące lokacji funkcji obsługi przerwania w przestrzeni adresowej (**rysunek 1** na następnej stronie pokazuje taką tabelę dla mikrokontrolera ATMEGA328P, gdzie jest zaznaczone miejsce obsługi przerwania przetwornika analogowo-cyfrowego). Uruchomienie pomiaru wartości analogowej trwa ileś czasu, więc program musi poczekać aż ten pomiar zostanie zakończony.

Możliwym rozwiązaniem jest odczyt rejestru statusowego związanego z przetwornikiem ADC i sprawdzenie odpowiedniego bitu, który zostanie ustawiony przez sam przetwornik w chwili zakończenia konwersji napięcia analogowego na wartość liczbowa. Innym wariantem jest takie skonfigurowanie przetwornika, by po zakończeniu działania zgłosił przerwanie, w obsłudze którego można odczytać liczbę wynikającą z konwersji.

Obsługa każdego przerwania finalnie trafia do odpowiedniej funkcji w programie (czy to w wyniku wektoryzowanego adresu, czy stałej lokacji w przestrzeni adresowej) jako jej wywołanie bez podania parametrów (każda obsługa przerwania jest funkcją bezparametrową). W przypadku klasycznego języka C nie stanowi to problemu, wystarczy zgłosić odpowiednią funkcję do obsługi przerwania. Taka funkcja ma specyficzny nagłówek. W rzeczywistości używa się odpowiedniej konstrukcji utworzonej przez `#define`. W dotychczasowych przykładach było używane rozwiązanie z użyciem `SIGNAL`, jak pokazuje **listing 1**:

```
SIGNAL ( TIMER0_OVF_vect )
{
    ( . . . )
} /* TIMER0_OVF_vect */
```

W nawiasach umieszczana jest informacja identyfikująca przerwanie, również utworzona przez `#define` (aby kompilator wraz z linkerem poprawnie skojarzył funkcję obsługi z odpowiednim przerwaniem). Podobnie w przypadku programów w języku C++ tworzone są funkcje obsługi przerwania, z tym, że do jej konstrukcji zamiast wyrażenia `SIGNAL` używa się `ISR`

12.4 Interrupt Vectors in ATmega328 and ATmega328P

Table 12-6. Reset and Interrupt Vectors in ATmega328 and ATmega328P

VectorNo.	Program Address ⁽²⁾	Source	Interrupt Definition
1	0x0000 ⁽¹⁾	RESET	External Pin, Power-on Reset, Brown-out
2	0x0002	INT0	External Interrupt Request 0
3	0x0004	INT1	External Interrupt Request 1
4	0x0006	PCINT0	Pin Change Interrupt Request 0
5	0x0008	PCINT1	Pin Change Interrupt Request 1
6	0x000A	PCINT2	Pin Change Interrupt Request 2
7	0x000C	WDT	Watchdog Time-out Interrupt
8	0x000E	TIMER2 COMPA	Timer/Counter2 Compare Match A
9	0x0010	TIMER2 COMPB	Timer/Counter2 Compare Match B
10	0x0012	TIMER2 OVF	Timer/Counter2 Overflow
11	0x0014	TIMER1 CAPT	Timer/Counter1 Capture Event
12	0x0016	TIMER1 COMPA	Timer/Counter1 Compare Match A
13	0x0018	TIMER1 COMPB	Timer/Counter1 Compare Match B
14	0x001A	TIMER1 OVF	Timer/Counter1 Overflow
15	0x001C	TIMER0 COMPA	Timer/Counter0 Compare Match A
16	0x001E	TIMER0 COMPB	Timer/Counter0 Compare Match B
17	0x0020	TIMER0 OVF	Timer/Counter0 Overflow
18	0x0022	SPI, STC	SPI Serial Transfer Complete
19	0x0024	USART, RX	USART Rx Complete
20	0x0026	USART, UDRE	USART, Data Register Empty
21	0x0028	USART, TX	USART, Tx Complete
22	0x002A	ADC	ADC Conversion Complete

Rysunek 1

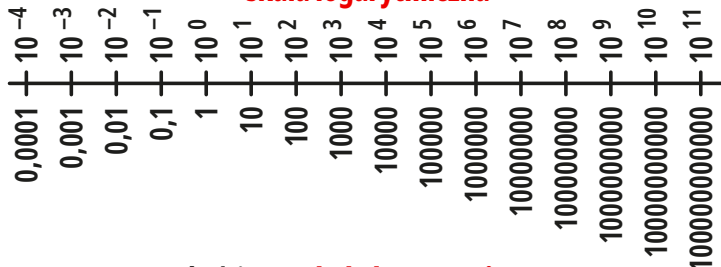
za przestarzałą (ale poprawną). Pod kątem technicznym oba makra generują funkcję obsługi przerwania, która domyślnie blokuje wykonywanie innych przerwania podczas swojego działania (nie ma obsługi przerwania zagnieżdżonych). Nowoczesność `ISR` w stosunku do `SIGNAL` polega na tym, że istnieje możliwość specyfikowania dodatkowych informacji pozwalających przykładowo na obsługę przerwania zagnieżdżonych jak również powiązanie dwóch różnych przerwania z tą samą funkcją obsługi (to są już bardziej zaawansowane rozwiązania, wykraczające poza zakres tego cyklu artykułów). Typowym elementem w programach w języku C++ jest korzystanie z obiektów i klas. Ich składnikami są między innymi metody (jako funkcje działające na danych zawartych w obiektach). Specyfiką metod jest istnienie jednego ukrytego parametru wywołania (identyfikowanego słowem kluczowym `this`). Nawet bezparametrowa funkcja będąca

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

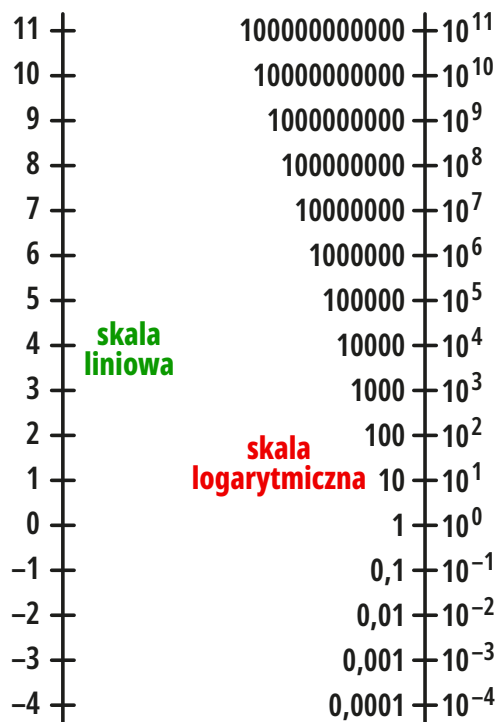
„zwyczajna” **skala liniowa** jest dobra do wielu zastosowań, ale jest kłopot z jednoczesnym przedstawieniem bardzo dużych i bardzo małych liczb



skala logarytmiczna



na pozór dziwna **skala logarytmiczna** pozwala w prosty sposób przedstawić liczby dodatnie z bardzo szerokiego zakresu



Szumy, zniekształcenia zakłócenia i decybele (1)

To jest artykuł wstępny cyklu, który w przystępny sposób przedstawia tytułowe niełatwe zagadnienia. Dwa pierwsze artykuły są wprowadzeniem i dotyczą głównie wykorzystania decybeli w elektronice. Poniższy materiał nakreśla tło oraz wprowadza w tematykę decybeli, procentów, a także „papiemów”.

Zakłócenia, szumy i zniekształcenia
Dlaczego wykorzystujemy decybele?
Co to jest decybel?

Decybele w elektronice
Procenty, ppm i ppb

Temat szumów, zakłóceń oraz zniekształceń w urządzeniach elektronicznych jest bardzo obszerny, skomplikowany, a szczegóły są trudne. Jednak podstawy tego zagadnienia są łatwe. W urządzeniach elektronicznych **przetwarzamy rozmaite sygnały elektryczne**, na przykład wzmacniamy, i **są to sygnały użyteczne** – jak najbardziej pożądane. Niestety, obok przetwarzanego sygnału użytecznego, w urządzeniach pojawiają się rozmaite sygnały, składniki niepożądane. I właśnie **szumy, zakłócenia i zniekształcenia to wszelkie niepożądane sygnały, przebiegi i zmiany**, to zło konieczne, którego nie można wyeliminować, a z którym trzeba walczyć.

Zakłócenia, szumy i zniekształcenia

Łatwo zrozumieć co to są **zakłócenia** – tak nazywamy **niepożądane sygnały przychodzące do urządzenia z zewnątrz**. Najprostszym przykładem zakłóceń są sygnały elektryczne pojawiające się w urządzeniu, a raczej przenikające do urządzenia elektronicznego jako skutek wyładowania atmosferycznego – pioruna. Inne popularne źródło zakłóceń to iskrzące silniki komutatorowe. Ogólnie biorąc, zakłócenia są bardzo mało przewidywalne, bo ich pojawianie się zależy od różnych czynników losowych. Nawet jeśli tego nie mówimy, przede wszystkim kojarzymy zwrot: „zakłócenia zewnętrzne”. I słusznie.

A teraz szumy. Najprościej, ale nie do końca ściśle biorąc, za **szumy** uważamy te **niepożądane sygnały, których źródła są wewnątrz urządzenia** i które występują zawsze, także wtedy gdy włączone urządzenie nie przetwarza sygnału użytecznego. Sygnału użytecznego nie ma, ale urządzenie elektroniczne, np. wzmacniacz, jest źródłem szumów. Dlatego obok „zakłóceń zewnętrznych” nierozłącznie kojarzymy popularne określenie „szumy własne”.

Tu trzeba podkreślić, że w języku angielskim zarówno zakłócenia jak i szumy są określane słowem *noise*. Nie warto tego przenosić do języka polskiego, co czynią niektórzy autorzy, tylko w miarę możliwości pozostać przy prostym rozróżnieniu, że **zakłócenia** są zewnętrzne, a **szumy** generalnie są wewnętrzne (choć są pewne wyjątki).

Jeszcze inna historia jest ze zniekształceniami – **zniekształceniami przetwarzanego sygnału użytecznego**. Jeżeli nie ma sygnału użytecznego – nie ma też zniekształceń (są tylko szumy i mogą być zakłócenia). Zniekształcenia pojawiają się tylko podczas przetwarzania sygnału i zgodnie z nazwą są to zniekształcenia, czyli deformacje tego sygnału użytecznego. Słowa „deformacje” i „zniekształcenia” w zasadzie nie kojarzą się z niepożądanymi przebiegami, ale w praktyce sygnały użyteczne zwykle są sygnałami, przebiegami zmiennymi, a wtedy w grę wchodzi coś takiego jak transformacja i szeregi Fouriera.

Najprościej biorąc, **fakt wystąpienia zniekształceń, deformacji sygnału użytecznego oznacza też pojawienie się sygnałów, składników, których nie było w wejściowym sygnale użytecznym**. Te dodatkowe składniki pojawiają się podczas przetwarzania sygnału użytecznego jako skutek różnych niedoskonałości. Najbardziej znane są **zniekształcenia nieliniowe**, wynikające z nieliniowości, skutkujące pojawieniem się tzw. harmonicznych (*THD*). Mniej znane są tak zwane **zniekształcenia intermodulacyjne (IMD)**, a jeszcze mniej znane tak zwane **zniekształcenia TIM**, związane z ograniczoną szybkością układu elektronicznego (wzmacniacza).

Ogólnie biorąc mamy więc trzy grupy **niepożądanych sygnałów**, składników: zakłócenia, szumy i zniekształcenia. A jaki mają z tym związek **decybele**, które większość ludzi kojarzy z głośnymi dźwiękami i hałasem? Odpowiedź jest dość prosta.

Dlaczego wykorzystujemy decybele?

Mówimy o przetwarzaniu najrozmaitszych (elektrycznych) sygnałów użytecznych, w tym sygnałów audio, czemu niestety towarzyszą różne „śmieci”: zakłócenia, szumy i zniekształcenia. Bardzo często te „śmieci” są bardzo małe, wielokrotnie mniejsze, niż przetwarzany sygnał użyteczny.

Przykładowo szumy mogą być 10 000 000 razy mniejsze, niż sygnał użyteczny, a zniekształcenia harmoniczne (*THD*) mogą być mniejsze niż 0,0001% przetwarzanego sygnału użytecznego. Przy takim zapisie z wieloma zerami można się łatwo pomylić i nie dostrzec znacznej, dziesięciokrotnej różnicy między na przykład między 42300000 a 423000000. Podobnie między 0,00012% i 0,000012%.

Jeżeli liczba zawiera kilka zer, to mamy kłopot. Niezbyt dobrym rozwiązaniem jest „zapis z zerami”, czyli 42 300 000 i 423 000 000, bo jest niedopuszczalny w przypadku liczb bardzo małych (nie można napisać 0,000 12% i 0,000 012%).

Jednym ze skutecznych rozwiązań „problemu wielu zer” jest sposób zapisu, zwany **notacją naukową**. W notacji naukowej liczby 42300000 oraz 423000000 zapiszemy jako $4,23 \cdot 10^7$ oraz $4,23 \cdot 10^8$, w skrócie 4,23E7, 4,23E8. Z kolei wyrażone w procentach wartości 0,00012% oraz 0,000012%, to ułamki $12/10000000$ oraz $12/100000000$, które w notacji naukowej zapiszemy w postaci $1,2 \cdot 10^{-6}$ (1,2E-6) oraz $1,2 \cdot 10^{-7}$ (1,2E-7).

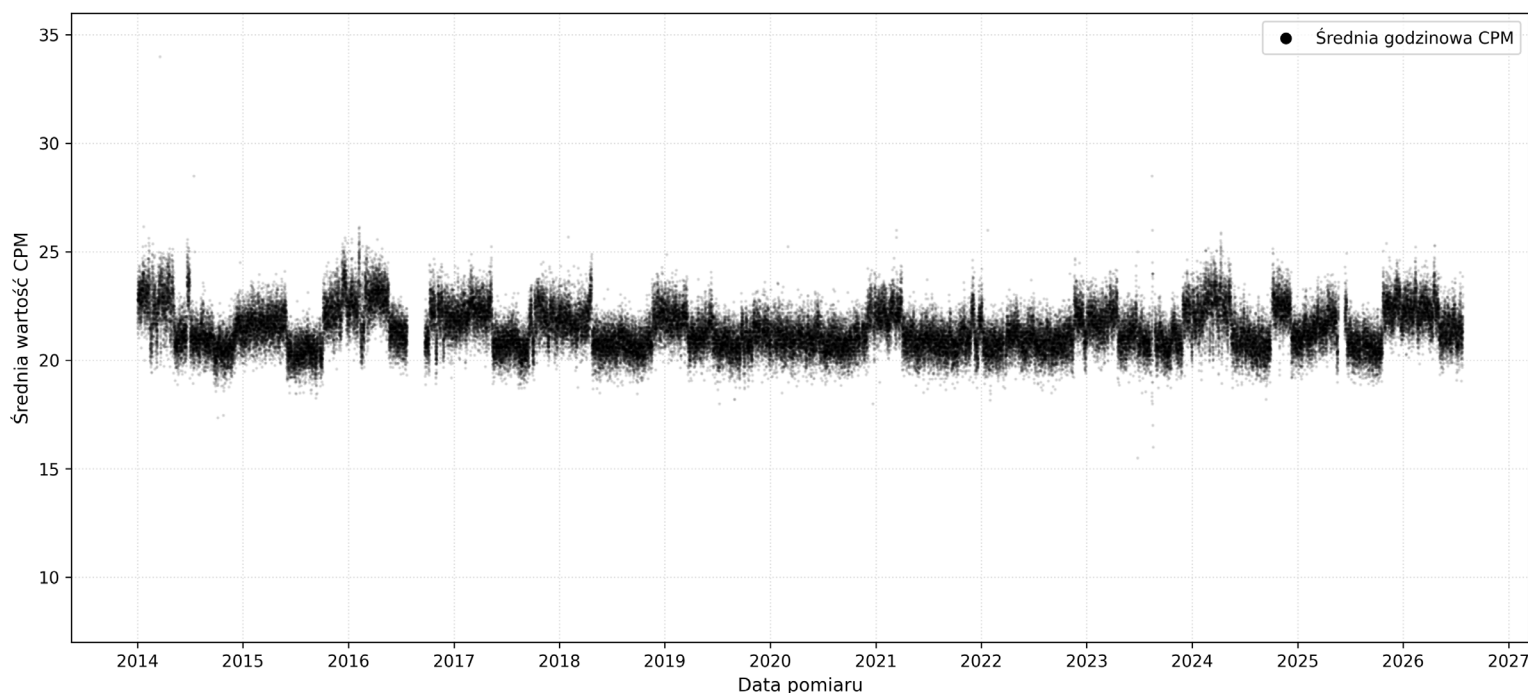
Notacja naukowa skutecznie rozwiązuje „problem wielu zer”, ale wielu praktyków jej nie lubi i nie stosuje. Praktycy zdecydowanie częściej wykorzystują decybele, i to w mnóstwie odmian.

Co to jest decybel?

Decybel to uniwersalna logarytmiczna jednostka miary, która wielu osobom kojarzy się przede wszystkim z dźwiękiem, hałasem i głośnością. Owszem decybele są wykorzystywane między innymi w akustyce, ale też w innych dziedzinach, w szczególności w elektronice, gdzie często nie mają żadnego związku z dźwiękiem. Wprowadzenia w temat nie warto jednak zaczynać od *decy-bela*, czyli od dziesiątej części bela. Warto zacząć od miar.

Na co dzień mamy do czynienia głównie z miarami liniowymi – przykład z mojej pracowni na **fotografii 1**.

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



Obserwacja naturalnego tła promieniowania

Kilkanaście lat ciągłych pomiarów wykonanych prostym miernikiem Geigera-Müllera pozwoliło zaobserwować niewielkie, lecz wyraźne zmiany naturalnego tła promieniowania. Ich najbardziej prawdopodobną przyczyną okazały się zmiany stężenia radonu.

Naturalne tło promieniowania

Mierniki promieniowania

Budowa mojego miernika

Pomiar promieniowania tła

Wieloletnie zbieranie danych

Zmiany poziomu promieniowania tła

Obecność radonu w piwnicy

Jak interpretować uzyskany wynik?

Podsumowanie

Naturalne tło promieniowania

W życiu nieustannie towarzyszą nam różne rodzaje promieniowania. Zdajemy sobie sprawę, że wokół nas występuje wiele jego źródeł, większość z nas słyszała np. o promieniowaniu cieplnym, radiowym, laserowym czy ultrafioletowym. Wiemy również, że nadmierna ekspozycja na każde z nich może być szkodliwa dla organizmu. Mimo to największy niepokój zwykle wzbudza promieniowanie jonizujące.

Oznaczane jest ono greckimi literami α , β i γ ,

a nam najczęściej kojarzy się z bronią jądrową lub awariami elektrowni atomowych. Nie wszyscy jednak zdają sobie sprawę, że promieniowanie jonizujące towarzyszy nam nieustannie. Jest ono naturalnym elementem naszego otoczenia, ponieważ jego źródłem są zarówno występujące w przyrodzie radioizotopy wielu pierwiastków, jak i promieniowanie kosmiczne docierające do Ziemi z przestrzeni kosmicznej. Ten wszechobecny poziom określamy mianem naturalnego tła promieniowania jonizującego.

Mierniki promieniowania

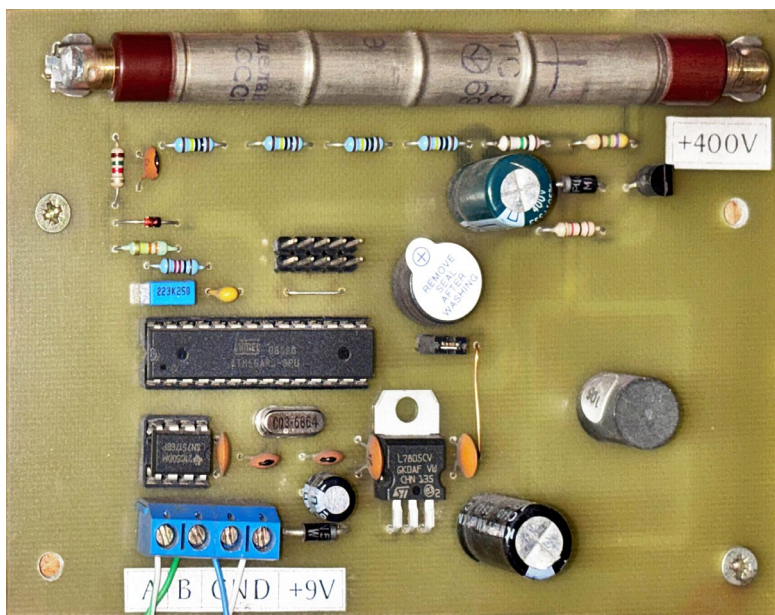
Wielu elektroników konstruuje różnego rodzaju przyrządy pomiarowe służące do obserwacji otaczającego nas środowiska. Termometry, higrometry czy barometry umożliwiają śledzenie zmian zachodzących w przyrodzie. Równie interesującym, a przy tym nieco mniej typowym projektem, jest budowa miernika promieniowania jonizującego.

Wbrew pozorom wykonanie takiego urządzenia nie jest szczególnie trudne, ponieważ bez większego problemu można kupić jego najważniejszy element: licznik Geigera-Müllera. Konstrukcja tego czujnika została opracowana jeszcze w latach międzywojennych. Z elektronicznego punktu widzenia licznik Geigera-Müllera jest lampą wypełnioną gazem, której elektrody spolaryzowane są napięciem rzędu kilkuset woltów. Gdy do jej wnętrza przedostanie się cząstka promieniowania, powoduje jonizację znajdującego się tam gazu. Powstałe elektrony i jony są następnie przyspieszane w polu elektrycznym, wywołując kolejne zderzenia i jonizację następnych atomów. W efekcie dochodzi do lawinowego wyładowania elektrycznego, które na wyprowadzeniach detektora pojawia się jako krótki impuls. Każdy taki impuls odpowiada zarejestrowaniu pojedynczego zdarzenia promieniotwórczego.

Budowa mojego miernika

Amatorzy w swoich konstrukcjach często wykorzystują czujniki wyprodukowane jeszcze w czasach Związku Radzieckiego. Również i ja, kilkanaście lat temu użyłem elementu, na którego obudowie widnieje napis „Сделано в СССР”.

Wykonany przeze mnie miernik bazuje na tubie „CTC-5”, której nazwa w polskich publikacjach, po transkrypcji z cyrylicy, zapisywana jest jako STS-5. Nadzór nad pracą urządzenia sprawuje popularny mikrokontroler ATmega8. Odpowiada on zarówno za sterowanie przetwornicą wytwarzającą napięcie 400 V, jak i za zliczanie impulsów rejestrowanych przez czujnik. Każdy zarejestrowany impuls powoduje krótkotrwałe włączenie głośniczka piezoelektrycznego, emitującego charakterystyczne dla dozymetrów „tyrkotanie”, a jednocześnie zwiększa wartość wewnętrznego licznika impulsów. Co minutę dane są odczytywane za pośrednictwem magistrali RS-485 przez miniaturowy komputer Raspberry Pi, od-



Fotografia 1

na zostać usunięta, jednak tu pozostała ona celowo, w celu zmniejszenia głośności sygnalizacji.

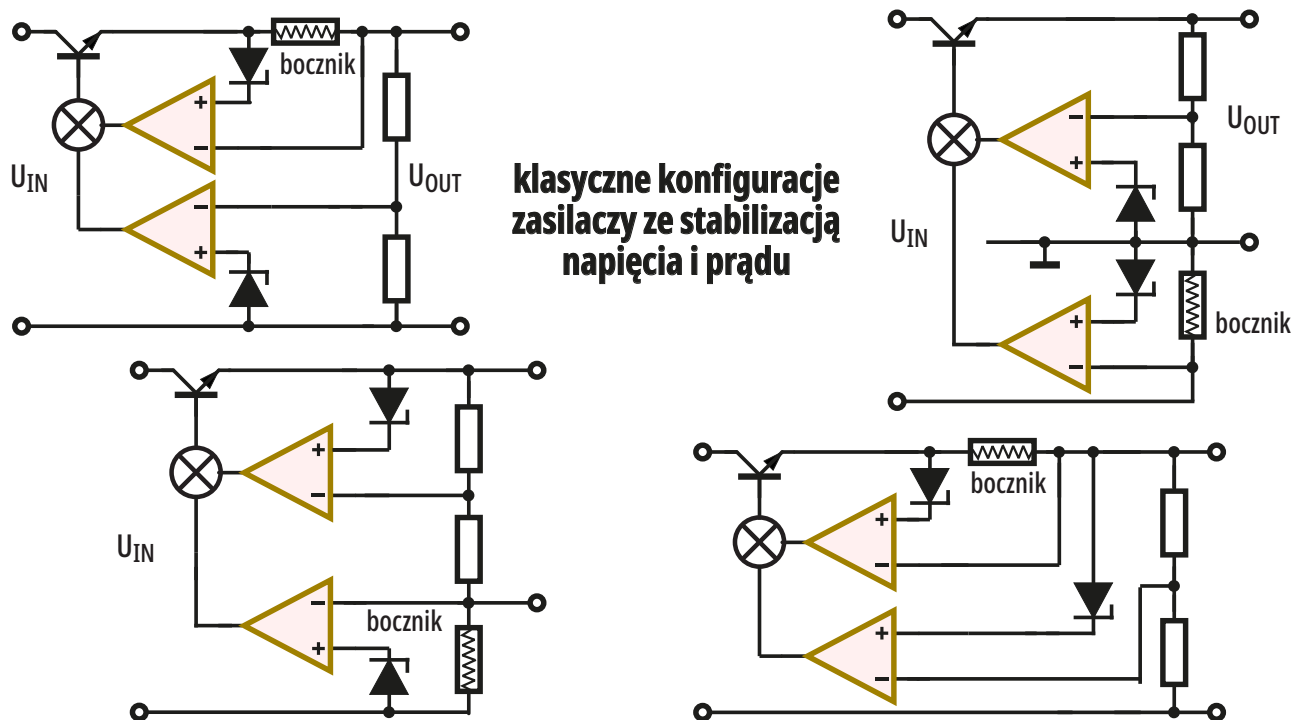
Wyniki pomiarów zapisywane są w jednostce CPM (ang. *counts per minute*), oznaczającej liczbę impulsów zarejestrowanych przez licznik Geigera-Müllera w ciągu jednej minuty. Należy jednak pamiętać, że CPM nie jest jednostką dawki promieniowania, lecz jedynie informacją o liczbie zarejestrowanych cząstek.

Dla zastosowanego w opisywanym mierniku czujnika STS-5 wartość 20 CPM odpowiada w przybliżeniu mocy dawki 0,11 $\mu\text{Sv/h}$, czyli poziomowi typowemu dla naturalnego tła promieniowania w wielu rejonach Polski. Nie jest to jednak przeliczenie dokładne, gdyż zależy ono od rodzaju promieniowania (α , β lub γ), jego energii, napięcia pracy tuby oraz sposobu kalibracji miernika. Dlatego jednostka CPM jest bardziej uniwersalna, niż przeliczanie jej na $\mu\text{Sv/h}$. Warto również pamiętać, że czujnik STS-5 praktycznie nie rejestruje promieniowania α . Ma ono bardzo małą zdolność przenikania i jest zatrzymywane nawet przez zwykłą kartkę papieru, dlatego nie jest w stanie przedostać się przez metalową obudowę czujnika.

Pomiar promieniowania tła

Licznik Geigera-Müllera wykorzystuje się przede wszystkim do monitorowania miejsc oraz obiektów o potencjalnie podwyższonym poziomie promieniowania. Urządzenie zamontowane na stałe w warunkach domowych powinno rejestrować stałe tło. Zmiana odczytów, przy założeniu braku obecności

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



Klasyczne konfiguracje zasilaczy stabilizowanych

W ZE ukazała się seria dziesięciu artykułów o numerach od [A050](#) do [A059](#), która pokazała problemy i możliwości ich pokonania przy projektowaniu i realizacji „małego” zasilacza liniowego ze stabilizatorem LM317. Tym więcej wiedzy wymaga planowana budowa zasilacza liniowego o większej mocy.

Dwie główne odmiany stabilizatorów napięcia
Ogranicznik – stabilizator prądu

Klasyczne konfiguracje stabilizatorów
Zalety i wady klasycznych konfiguracji

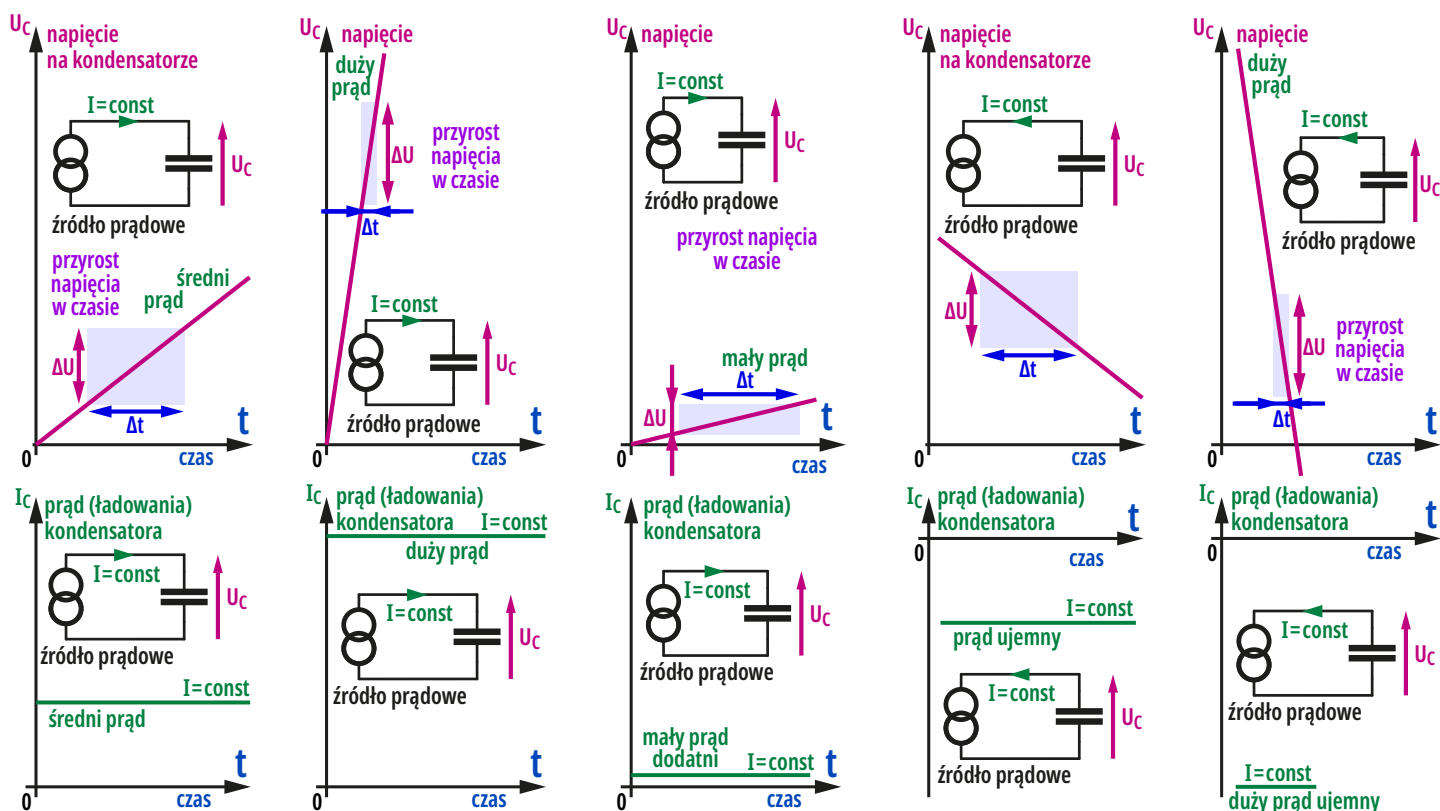
Gdy po serii artykułów dotyczących budowy „małego” zasilacza warsztatowego opartego na układzie LM317 przyszła pora na budowę nieco wydajniejszego urządzenia, postanowiłem i ja podzielić się swoimi przemyśleniami i doświadczeniami w tym zakresie.

Chciałbym zacząć, jak na fizyka teoretyka przystało, od pewnych rozważań teoretycznych. Zaczniemy więc od podstawowej oczekiwanej funkcjonalności budowanego zasilacza, czyli dostarczania zasilania o stałym napięciu wyjściowym. W tym momencie napotykamy pierwsze, podstawowe pytanie: w jaki

Dwie główne odmiany stabilizatorów napięcia

Po pierwsze, możemy wykorzystać jakieś zjawiska (procesy) fizyczne lub chemiczne, które zagwarantują nam stałość napięcia zasilania. Najprostszym przykładem takiego podejścia jest użycie zasilania baterijnego (baterii jednorazowej lub akumulatora) – podejście sensowne i wykorzystywane, ale raczej mało praktyczne. Owszem, czasem można wykorzystać akumulator o jak największej pojemności i znikomej impedancji wewnętrznej i uznać, że przy niewielkich zmianach poboru prądu jego napięcie jest nie-

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



Kondensator – zależność napięcia i prądu

W dwóch poprzednich artykułach serii omawiałem kondensatory. Koncentrowałem się na energii, bo ona jest najważniejsza. Teraz pokażę, że niedoskonała analogia hydrauliczna kondensatora nie tylko dobrze ilustruje jego elementarne właściwości, ale też pomaga zrozumieć bardzo ważne, trudniejsze kwestie.

Zależność napięcia i prądu w kondensatorze
Kondensator – zależność napięcia i prądu
Ładowanie czy rozładowanie kondensatora?

Proste przypadki – stała szybkość zmian
Trudne przypadki – zmienna szybkość zmian...

Poprzednie artykuły tej serii pokazały, że **niepełna i niedoskonała analogia hydrauliczna okazuje się pomocna do zrozumienia podstaw**. I nie tylko podstaw. Zgodnie z obietnicą z poprzedniego artykułu pokażę teraz, jak analogia wysokościowa pomaga też zrozumieć niektóre dużo trudniejsze aspekty elektroniki i matematyki. Jest znakomitą ilustracją i wytłumaczeniem, **czym w praktyce jest całkowanie oraz różniczkowanie**, a te pojęcia dla mnóstwa osób są dosłownie czarną magią. Elektronik nie musi być dobrym matematykiem czy fizykiem, ale bez zrozumienia sensu pojęcia pochodnej oraz różniczkowania i całkowania będzie mieć kłopoty ze zrozumieniem „trudniejszej elektroniki”.

Zacznę o przypomnieniu, że **kondensator jest zbiornikiem na energię elektryczną, który magazynuje energię $E = 0,5 \times C \times U^2$ w polu elektrycznym** (lub jak kto woli w postaci pola elektrycznego). I to jest jasne, jeżeli w stanie ustalonym na kondensatorze o pojemności C występuje jakieś niezmiennie, stałe napięcie U . Trudniej wytłumaczyć i zrozumieć wzajemną „dwukierunkową” zależność napięcia U oraz prądu I w kondensatorze, a to dla praktyka jest ogromnie ważne. W przypadku (idealnego) rezystora zależność między napięciem i prądem jest oczywista (prawo Ohma). A jaka jest **zależność prądu i napięcia w kondensatorze**? Spróbuję to przedstawić w sposób maksymalnie przystępny.

Zależność napięcia i prądu w kondensatorze

Dla (idealnego) **rezystora** zależność napięcia U i prądu I jest trywialna, śmiesznie łatwa – wiąże je podstawowy parametr rezystora – rezystancja R tego rezystora definiowana $R = U / I$. W przypadku kondensatora i cewki aż tak łatwo nie jest, ale podstawowe zasady też są zaskakująco proste.

Powszechnie wiadomo też, że **przez kondensator płynie prąd, gdy zmienia się na nim napięcie**. Dlatego **dla prądu stałego kondensator stanowi przerwę, rozwarcie**. Bo jeżeli napięcie na kondensatorze jest, ale się nie zmienia, to nie płynie prąd, a wtedy niezależnie od napięcia na kondensatorze jego oporność jest nieskończenie wielka, bo przecież oporność to po prostu stosunek napięcia i prądu (U / I). Kwestię oporności kondensatora dla przebiegów zmiennych, w szczególności sinusoidalnych omówimy później. A teraz przedstawię podstawowe zależności, które pokażą też praktyczny sens matematycznego pojęcia funkcji oraz praktyczny sens pochodnej, a potem też całki.

Kondensator – zależność napięcia i prądu

Zaczynamy od oczywistej sytuacji, gdy kondensator jest pusty – napięcie na nim jest równe zeru i nie ma jeszcze w nim energii, co ilustruje **rysunek 1a**. Potem według **rysunku 1b** dołączamy do kondensatora źródło prądowe (tak, prądowe, dające niezmienny prąd I , a nie źródło napięciowe o niezmiennym napięciu) i w ten sposób pomału zaczniemy dostarczać doń energię. Co się dzieje z napięciem?

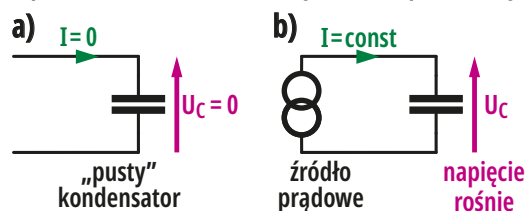
Właśnie tę kwestię znakomicie, genialnie ilustruje, skądinąd niedoskonały, pionowy model hydrauliczny. Jeżeli według **rysunku 2** do kondensatora w postaci rury wpychamy wodę w jakimś określonym, niezmiennym tempie, to...

... oczywiście poziom wody w takim cylindrze rośnie. Rośnie w jakimś określonym, niezmiennym tempie, rośnie liniowo. Możemy to przedstawić jak na **rysunku 3a**, ale od razu zaznaczam, że podobieństwo wykresu do **rysunku 3b**, pokazującego charakterystykę rezystora jest pozorne! Otóż dla rezystora mamy prostą i oczywistą zależność napięcia U i prądu I , niezależną od czasu!

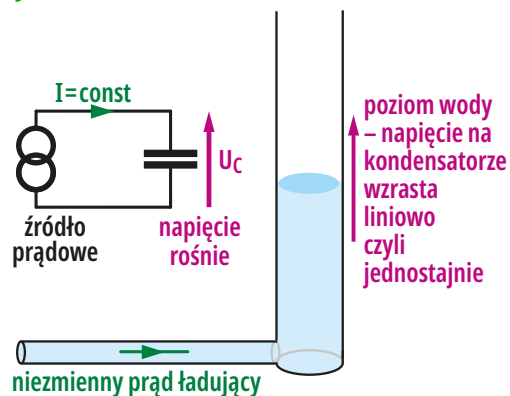
Napięcie na kondensatorze w jakiś sposób zależy od prądu, ale zależy też od czasu!

Jeżeli kondensator ładowany jest prądem I o *niezmiennej, ustalonej* wartości (co zapisujemy jako $I = \text{const}$), to napięcie U na nim rośnie, zmienia się liniowo – w *niezmiennym, ustalonym* tempie. Tempo, szybkość zmian napięcia będzie oczywiście zależeć od wartości prądu: **czym większy prąd, tym szybsze tempo zmian napięcia**, co ilustruje **rysunek 4**.

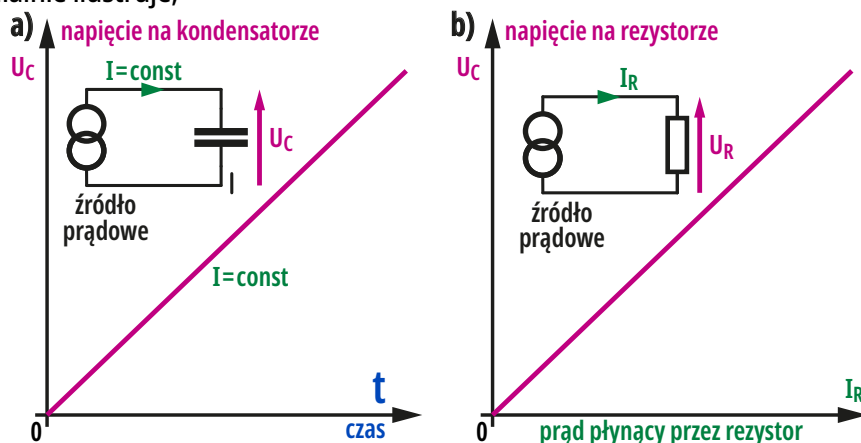
Analogia hydrauliczna znakomicie pokazuje też drugi ważny wniosek: przy niezmiennym, ustalonym tempie wtłaczania wody do pionowej rury,



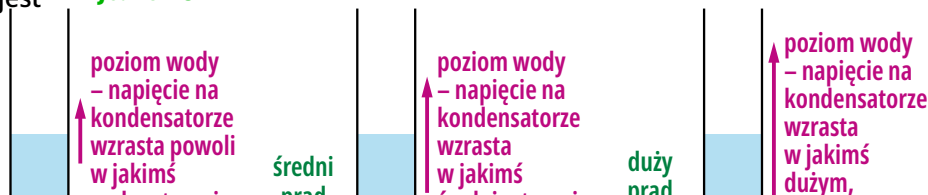
Rysunek 1



Rysunek 2



Rysunek 3



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



Ojciec elektromagnetyzmu: Romagnosi czy Ørsted?

W poprzednim artykule [Czy to Gian Romagnosi jest ojcem elektromagnetyzmu?](#) przedstawiłem informacje o dwóch krótkich artykułach z roku 1802, które zawierały sugestię, że elektryczność jest związana z magnetyzmem. Informacje te nie wzbudziły zainteresowania. Zainteresowanie pojawiło się 18 lat później.

Hans Christian Ørsted

Czy odkrycie zainteresowało świat naukowy?

Przypomnę, że przez stulecia, aż do początków XIX wieku za oczywiste uważano, że elektryczność oraz magnetyzm to dwa zupełnie oddzielne, niepowiązane zjawiska. Pierwsze wzmianki sugerujące związek elektryczności i magnetyzmu pojawiły się w roku 1802. Fizyk amator **Gian Domenico Romagnosi** w dwóch codziennych gazetach opublikował krótkie artykuły sugerujące taki związek.

Choć jego spostrzeżenia dotarły do ówczesnych naukowców zajmujących się tymi zagadnieniami, przez kilkanaście lat były ignorowane. Ignorowane do czasu publikacji zdecydowanie bardziej konkretnych informacji w lipcu roku 1820.

Ataki na Ørsteda

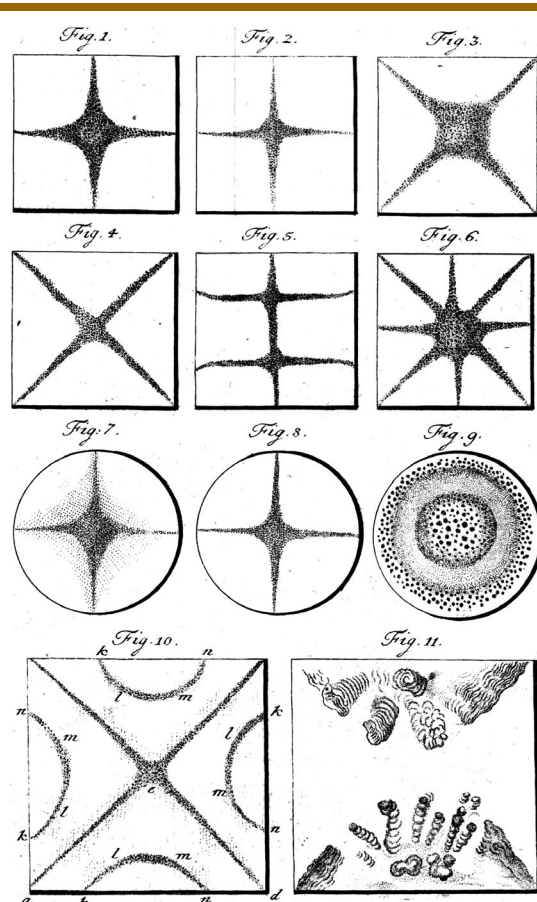
Znaczenie odkryć Ørsteda

Hans Christian Ørsted

Niektóre szkolne podręczniki i inne źródła informują, że duński profesor z Kopenhagi, **Hans Christian Ørsted** (Oersted), zupełnie przypadkowo, przygotowując doświadczenie dla studentów, zauważył reakcję igły magnesu umieszczonego w pobliżu przewodu, przez który popłynął prąd elektryczny i że w roku 1820 opublikował informację o tym zdarzeniu, co udowodniło związek elektryczności i magnetyzmu i stało się początkiem elektromagnetyzmu. To jest drastycznie uproszczony i niepełny obraz. Bardzo interesujące są szczegóły oraz szerszy kontekst tego zdarzenia.



Rysunek 1



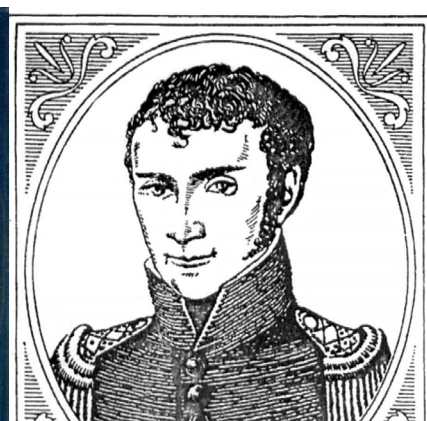
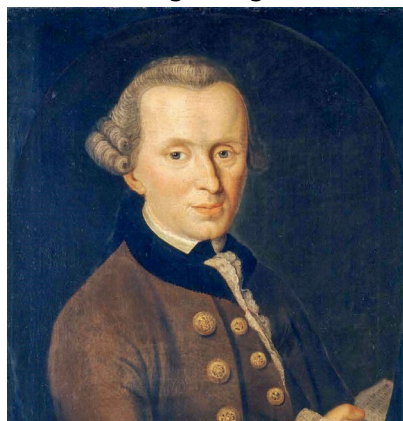
Rysunek 2

Otóż Hans Christian Ørsted (**rysunek 1**) miał wyjątkowo szerokie zainteresowania, między innymi był... poetą, co zresztą wiązało się z otrzymaniem przez niego pierwszego prestiżowego medalu i uzyskaniem stanowiska profesora uniwersyteckiego. Zajmował się nie tylko zagadnieniami fizyki, w tym akustyki i wibracji, ale też chemią, medycyną i innymi dziedzinami. **Rysunek 2** pochodzi z jego pracy z roku 1810 i przedstawia „akustyczne” figury Chladniego (figurę taką można też zobaczyć na płytce wiodocznej na portrecie, na rysunku 1).

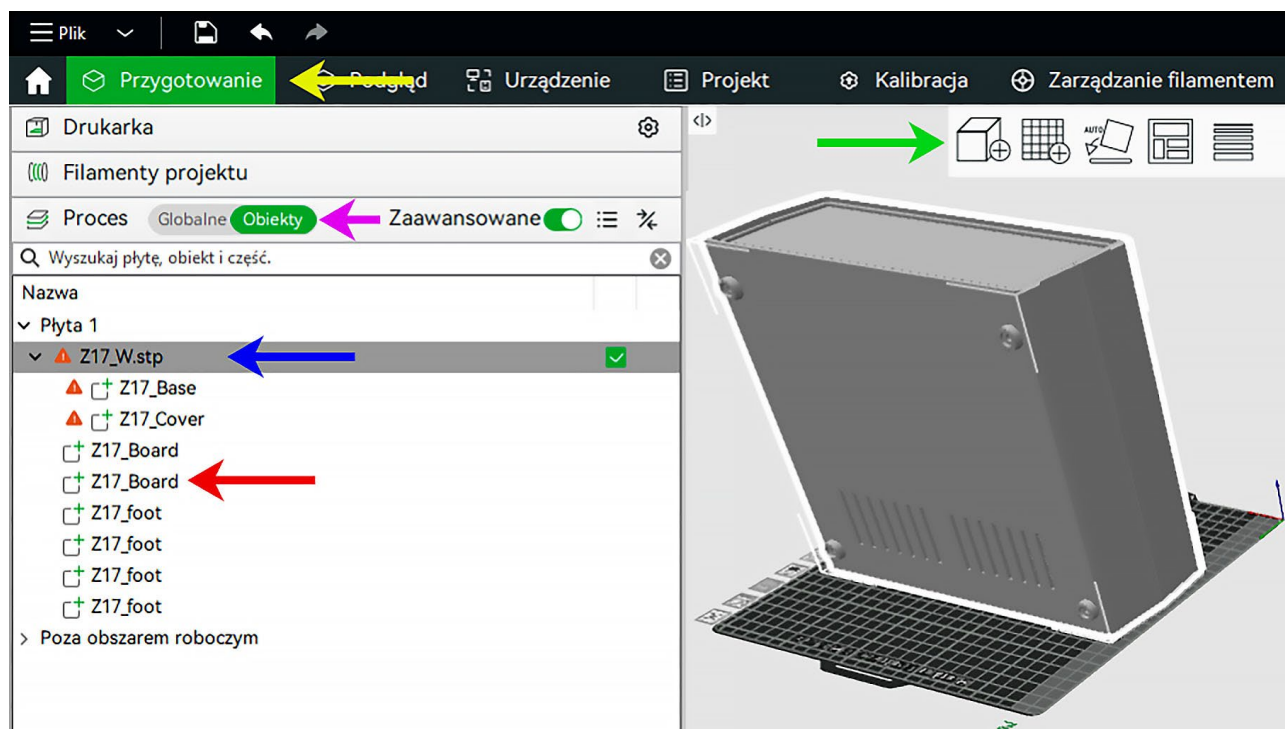
Co ważne, Ørsted (1777 – 1851) był zafascynowany koncepcjami słynnego oświeceniowego filozofa Immanuela Kanta (1724 – 1804) – **rysunek 3** – zajmującego się między innymi filozofią przyrody, w tym fizyki, jednością przyrody (Einheit der Natur), zawartymi w dziele zwanym MAN (Metaphysische Anfangsgründe der Naturwissenschaft).

W szkolnych podręcznikach przewija się informacja, iż Ørsted dokonał swego odkrycia przypadkowo, a to jest co najmniej uproszczeniem, a wręcz wprowadza w błąd. Owszem, w większości odkryć naukowych występuje ja-

Otóż duże znaczenie miał wpływ filozofii Kanta, a także kontakty z licznymi ówczesnymi naukowcami, wśród których byli jego rówieśnik, przyjaciel, wspomniany w moim artykule o akumulatorach, dziś niemal zapomniany „romantyczny naukowiec – samouk” Johann Wilhelm Ritter (1776 – 1810) – **rysunek 4** oraz słynny wówczas Giovanni Aldini. Ożywione kontakty, w tym korespondencja z Aldinim, poznanym w Paryżu podczas studenckiej podróży po Europie (1801 – 1804), sugerują, że Ørsted z niewątpliwie znanych mu prac Aldiniego mógł uzyskać informacje o eksperymentach i wnioskach Romagnosiego.



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



Płyty czołowe do obudów Kradex i Maszczyk, część 1

Wcześniej za pomocą kreatora z platformy <https://makerworld.com> utworzyliśmy prostą obudowę do regulatora PWM. Zaprojektowaliśmy też do niej płytę czołową, używając narzędzi z programu Bambu Studio. W tym artykule zajmiemy się projektowaniem i wydrukiem płyt czołowych do gotowych obudów.

Obudowy

Obudowy firmy Maszczyk

Obudowy firmy Kradex

Import do programów CAD

Import obudowy do Bambu Studio

Szczegóły

Można wydrukować kompletną obudowę, jednak ma to pewne wady. Atrakcyjna jest inna możliwość, którą koniecznie trzeba wziąć pod uwagę. Mianowicie na rynku jest mnóstwo niedrogich obudów, które mają oddzielne płyty czołowe. Nie drukujemy całej obudowy, a jedynie płytę czołową.

Obudowy

Dla elektroników produkowane są gotowe obudowy z polistyrenu, ABS-u oraz metalu. Najbardziej popularne i tanie są obudowy z tworzyw sztucznych, firmy Maszczyk i Kradex. Wiele z nich składa się z pokrywy górnej, dolnej, płyty czołowej przedniej

i tylnej, nóżek i wkrętów mocujących. Stosowane są w nich różne rozwiązania konstrukcyjne i obudowa tego typu może składać się też z dwóch części.

Wśród hobbystów popularne jest wykonywanie dedykowanych płyt czołowych w postaci wydruku na papierze lub folii i naklejenie ich na oryginalną, fabryczną płytę, w której muszą być wykonane potrzebne otwory. Jednak w wielu przypadkach zachodzi konieczność wykonania dużych i nietypowych otworów. W takim przypadku zwykła wiertarka może się nie sprawdzić. Piłowanie pilnikiem takich otworów jest pracochłonne. Lepszym rozwiązaniem będzie użycie małej frezarki CNC.

Jednak nie każdy ma dostęp do takiej frezarki. Również ceny sensownych frezarek CNC „na biurko” są wyższe niż drukarek 3D. Czy zatem możliwe jest wydrukowanie panelu czołowego do popularnych obudów? Tak, choć napotkamy pewne problemy i ograniczenia. Największym problemem jest maksymalny rozmiar wydruku osi X i Y na powierzchni stołu drukarki 3D.

W niektórych przypadkach można próbować ustawić drukowaną płytę czołową po przekątnej stołu drukarki 3D. Dzięki temu można uzyskać nieznacznie większy obszar roboczy w osiach X i Y. Musimy też posiadać model 3D panelu obudowy z jakiego wykonamy płytę czołową. Tutaj, choć nie zawsze, z pomocą przychodzą producenci obudów dla elektroników udostępniając modele 3D produkowanych obudów. Dostępność modelu nasuwa myśl samodzielnego wydrukowania całej obudowy. Jednak jest to nieopłacalne z uwagi na czasochłonność i koszt filamentu potrzebnego do obudowy w porównaniu z ceną gotowej obudowy. Ograniczeniem jest też rozmiar wydruku posiadanej drukarki 3D.

Obudowy firmy Maszczyk

Producenci obudów dla elektroników do swoich obudów udostępniają na stronach internetowych dokumentację, w tym modele 3D. Jednak ta sytuacja nie jest idealna. W przypadku obudów firmy Maszczyk do niektórych obudów nie jest udostępniana żadna dokumentacja. Przykładem jest obudowa o oznaczeniu KM-136. Dla innych obudów udostępniana jest dokumentacja w postaci plików PDF, IGS i DWG. Tutaj przykładem jest obudowa o oznaczeniu KM-85. Z kolei dla obudowy o oznaczeniu KM-5 udostępniona jest dokumentacja w postaci plików PDF, DXF, IGS i STEP. Sekcję z plikami dokumentacji obudowy KM-5 widzimy w czerwonej ramce na **rysunku 1**. Dla naszych potrzeb najbardziej przydatny jest plik STEP (×.step). W przypadku przeglądarki Mozilla Firefox klikamy ikonkę pliku STEP prawym klawiszem myszki i z menu kontekstowego wybieramy **Zapisz element docelowy jako...**



OBUDOWY STANDARDOWE



KM-5

6,56 zł (netto)
8,07 zł (brutto)
Cena detaliczna

Opcja KM-5BK

Ilość sztuk

1

DO KOSZYKA

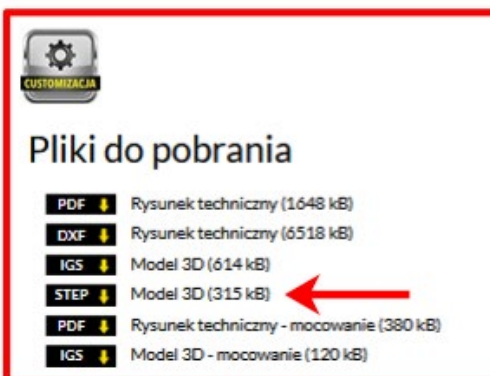
Standardowa obudowa dwuczęściowa o klasycznym wyglądzie. Doskonale nadaje się do zalania żywicą lub silikonem. Brak jakichkolwiek otworów, umożliwia frezowanie w dowolnym miejscu.

Obudowa zamykana za pomocą pokryw, którą trzeba przykleić we własnym zakresie.

Przy dokładnym klejeniu oraz odpowiednim spoiwie można uzyskać stopień szczelności IP67.

Materiał z którego została wyprodukowana obudowa to ABS. Istnieje możliwość wykonania obudowy z innego materiału lub w niestandardowym kolorze.

Długość:	83,6
Szerokość:	70
Wysokość:	23,2
Materiał:	ABS
Dostępne opcje:	KM-5BK KM-5G KM-5P/BK KM-5P/G



Rysunek 1

Obudowy firmy Kradex

W przypadku obudów firmy Kradex jest nieco lepiej niż w przypadku obudów firmy Maszczyk. Udostępniana jest obszerniejsza dokumentacja obudów tej firmy. Jednak nie w każdym przypadku dostępny jest plik STEP (×.stp). Przykładowo dla obudowy Z17W plik STEP jest udostępniony. Natomiast dla obudowy o oznaczeniu Z4A dostępne są pliki dokumentacji, lecz bez pliku STEP. Sprawdziłem wrywkowo kilka modeli obudów firmy Kradex i nie spotkałem się z brakiem dokumentacji dla tych obudów. Jak ma to miejsce w przypadku niektórych obudów

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

ZROZUMIEĆ **E**LEKTRONIKĘ z Piotrem Góreckim

ZE 10/2026

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
będące 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.**

Inicjatywa Zrozumieć Elektronikę realizowana jest dzięki wsparciu Patronów i Mecenasów poprzez [Patronite.pl](https://patronite.pl)