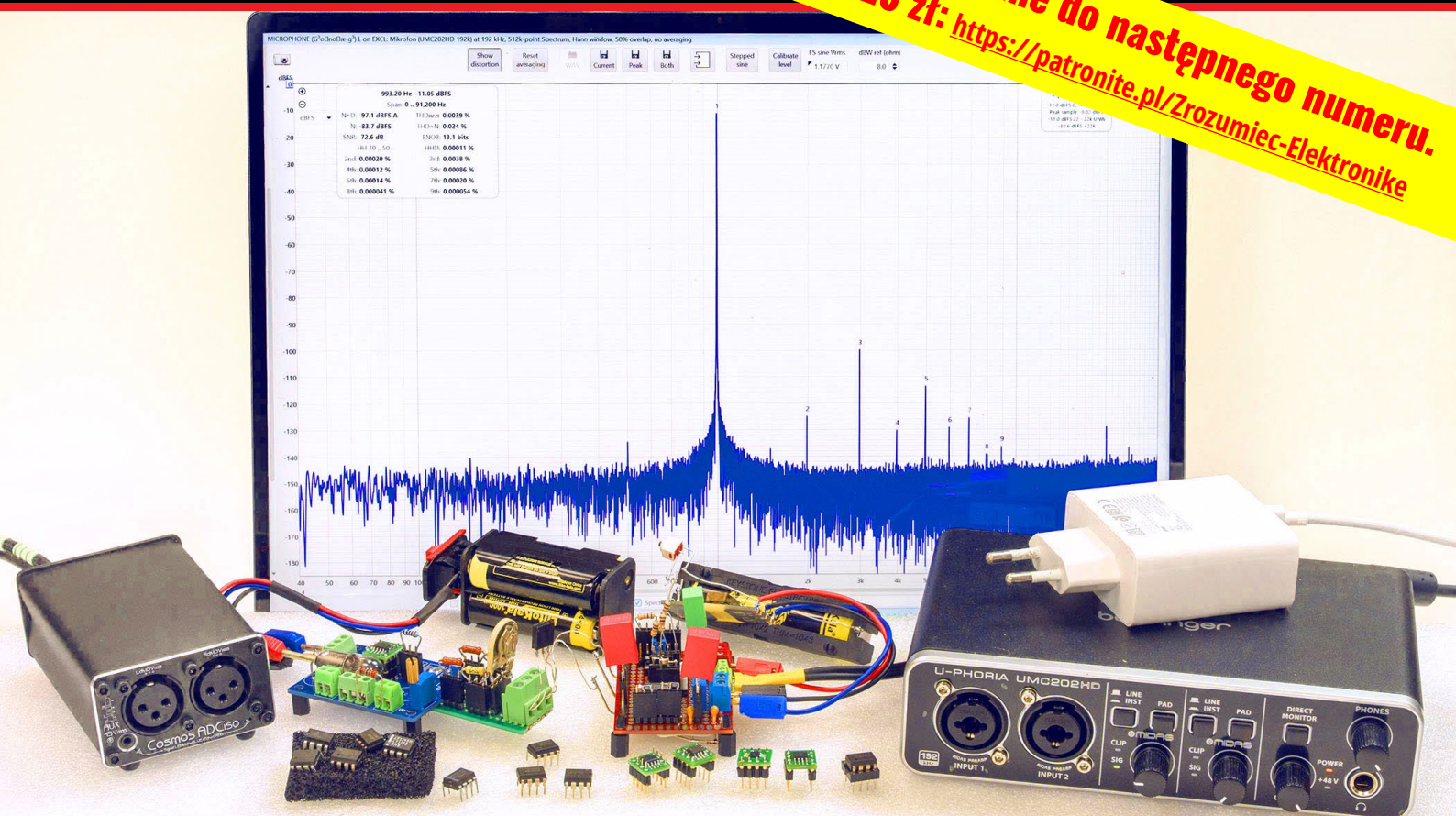


9/2026 Wrzesień (45)

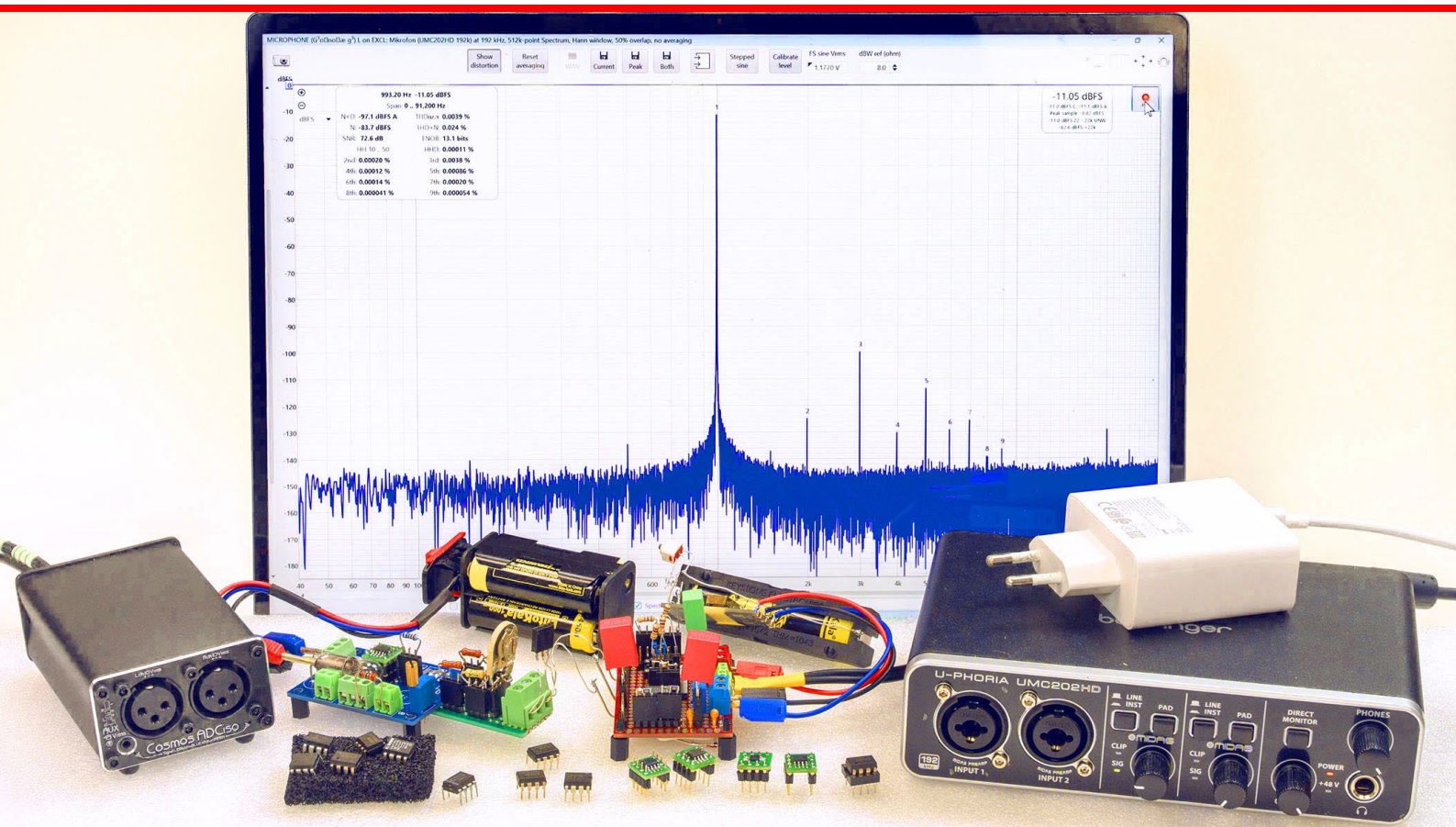
piotr-gorecki.pl



Eksperymentalna przystawka pomiarowa do kart audio

- Programator i debugger w Xplained Mini • Czy można zrobić MOSFET-a w warunkach domowych?
- Moduł wyświetlacza z układem TM1637 • Impulsowy test oscyloskopów i sond pomiarowych
- Generatory z mostkiem Wiena • BMS w akumulatorze – co to i po co? • Zakup przekładników „ręcznych”
- Odbiornik detektorowy i nie tylko detektorowy • Kondensator – zdolność gromadzenia w nim energii





Eksperymentalna przystawka pomiarowa do kart audio

Artykuł przedstawia bardzo prostą, uniwersalną dwukanałową przystawkę, przeznaczoną do wygodnego przeprowadzania pomiarów audio oraz do rozmaitych testów. Zaletą jest prostota, a z drugiej strony łatwość dostosowania do konkretnych potrzeb i eksperymentów, w szczególności do badania problemu szumów.

Przystawka pomiarowa – po co i do czego?

Rozważania projektowe

Kwestie zasilania przystawki pomiarowej

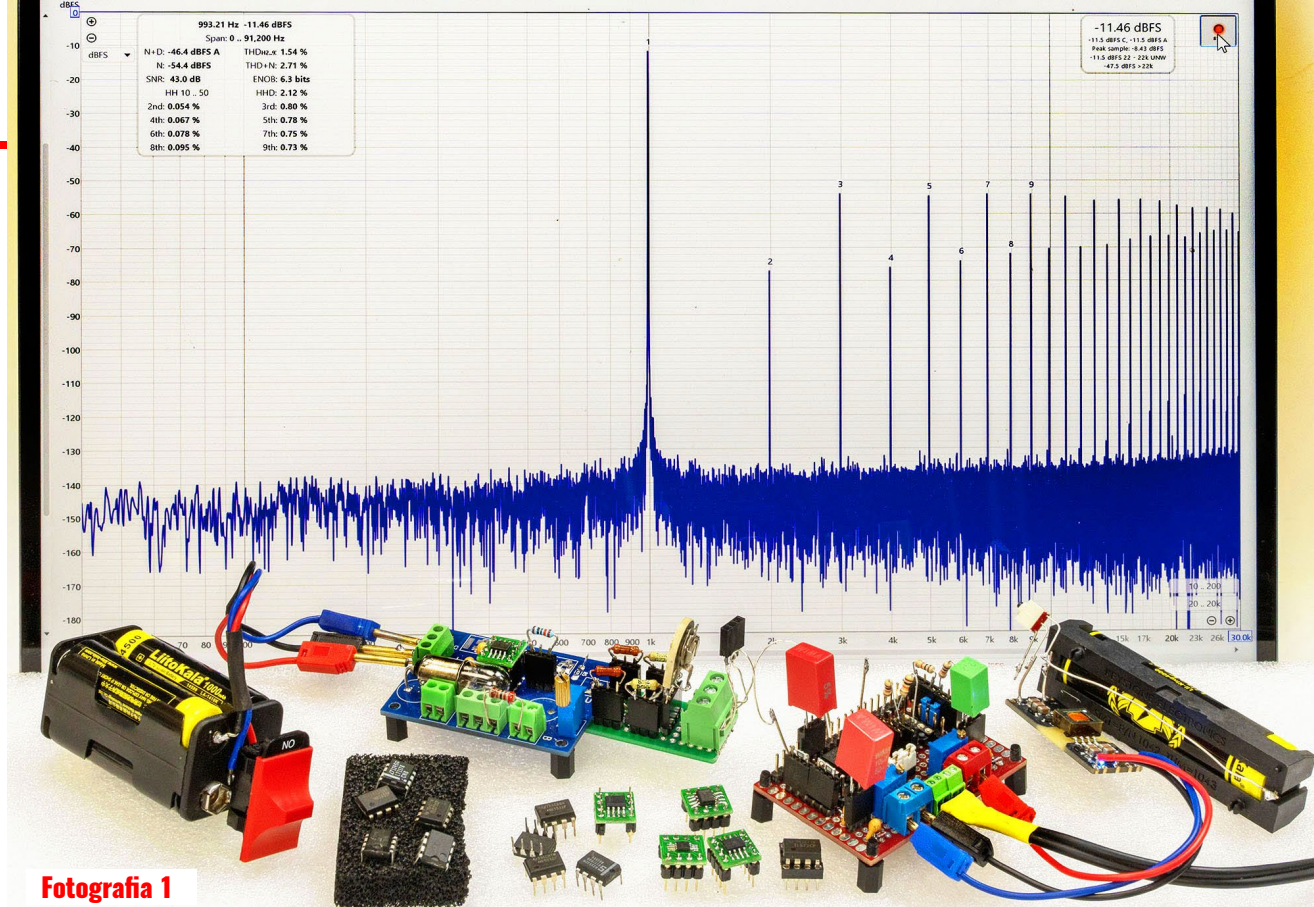
Zabezpieczenie wejścia

Dobór właściwości przystawki pomiarowej

Poziomy napięcia i wzmacnienie

Komputerowe karty audio są wykorzystywane do pomiarów, między innymi w roli oscyloskopu, ale zdecydowanie częściej w roli **analizatorów widma**. Co prawda pasmo pomiarowe jest stosunkowo wąskie, nieco mniejsze niż połowa częstotliwości próbkowania przetwornika ADC, ale za to rozdzielczość i dynamika są nieporównanie większe. Dziś powszechnie dostępne są liczne komputerowe karty audio z 24-bitowymi przetwornikami ADC o maksymalnym próbkowaniu 192 kHz. Daje to pasmo pomiarowe do około 90 kiloherców i dynamikę teoretycznie ponad 140 decybeli. Przy pomiarach audio, zwłaszcza w roli analizatorów widma, często kluczowe jest uzyskanie jak największej dynamiki.

Gdy chcemy wykorzystać kartę audio w roli analizatora widma, niezbędna jest też regulacja czułości wejściowej w bardzo szerokim zakresie, dużo szerszym, niż pozwala oryginalna karta. Pożądana jest możliwość pomiaru dużych sygnałów o amplitudach rzędu 10, a nawet 100 woltów, ale też małych sygnałów rzędu miliwoltów. I właśnie **opisywana eksperymentalna przystawka daje taką możliwość**. W praktyce trudno uzyskać dynamikę rzędu 140 decybeli i to nie tylko z uwagi na różne niedoskonałości przetwornika ADC. Ograniczenia wprowadzają też obwody analogowe doprowadzające sygnał do przetwornika. W opisywanej przystawce możliwe jest optymalizowanie dynamiki i zniekształceń.



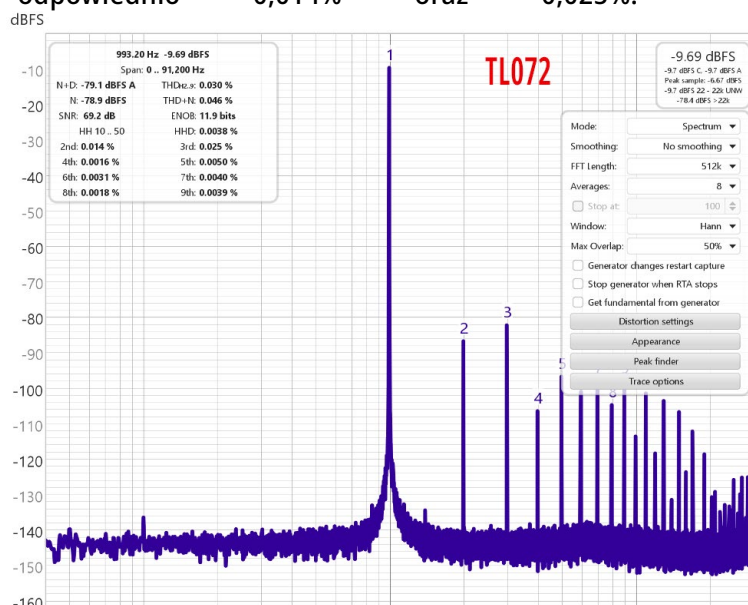
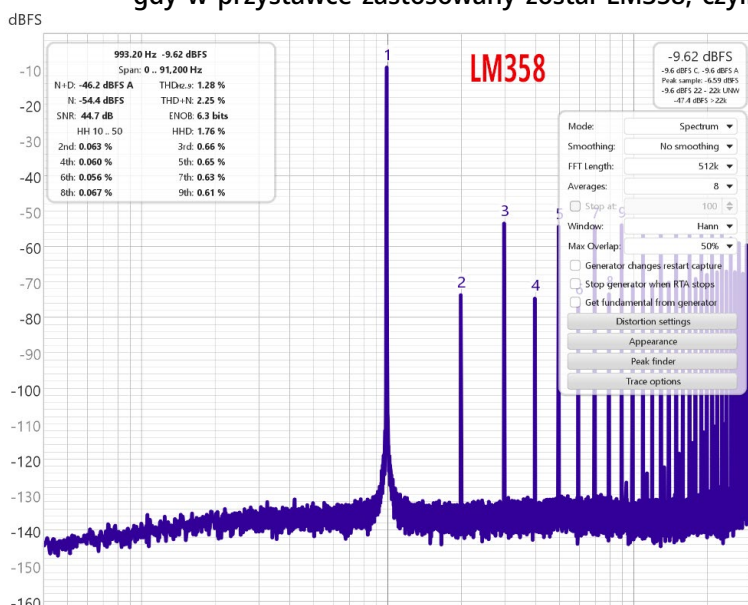
Fotografia 1

Fotografia tytułowa i fotografia 1 pokazują stanowisko pomiarowe podczas testów przystawki z kostką LM358. Podstawą jest laptop (na czas pomiarów odłączony od sieci energetycznej) z zainstalowanym programem REW. Kablem USB dołączona jest popularna i tania karta Behringer UMC202HD, do której dołączona jest opisywana przystawka. Do pomiaru zniekształceń wykorzystany jest prościutki generator przebiegu sinusoidalnego 1 kHz o znikomych zniekształceniach, a z lewej strony widać też niepodłączoną „wzorcową” kartę Cosmos E1DA.

Rysunek 2 pokazuje zawartość harmoniczną, gdy w przystawce zastosowany został LM358, czyli

najpopularniejszy wzmacniacz operacyjny. „Na oko” szumy własne są na poziomie nieco niższym niż -130 decybeli (dBFS), co można byłoby zaakceptować, ale poziom zniekształceń harmonicznnych jest nieakceptowalnie duży. Zawartość drugiej harmonicznej to 0,063%, ale trzeciej aż 0,66%. Całkowite zniekształcenia są większe niż 1 procent – potwierdza się znana prawda, że LM358 nie nadaje się do układów audio.

Rysunek 3 pokazuje szumy i zniekształcenia przystawki z kostką TL072. Szumy są tutaj mniej więcej o 10 decybeli, czyli około trzykrotnie mniejsze, a zawartość drugiej i trzeciej harmonicznej to odpowiednio 0,014% oraz 0,025%.



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



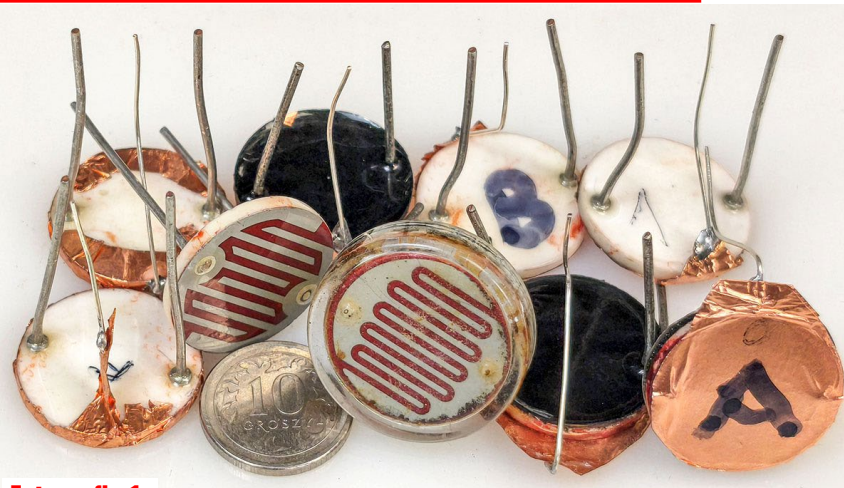
Czy można zrobić MOSFET-a w warunkach domowych?

Tranzystory stworzono mniej więcej w połowie XX wieku, czyli stosunkowo późno. Zrobienie działającego tranzystora przede wszystkim wymaga zastosowania jakiegoś materiału półprzewodnikowego. Czy dziś, w epoce półprzewodników, realizacja tranzystora polowego leży w zasięgu każdego hobbysty?

Próba realizacji IGFET-a
Pomiary stałoprądowe

Pomiary zmiennoprądowe
Uwagi końcowe

W artykule Zbuduj diodę LED! oraz w filmie YT o tym samym numerze podkreślałem radość i satysfakcję, jaką daje realizacja elementów półprzewodnikowych. Kolejną inspiracją była strona <http://sparkbangbuzz.com/cds-fet/cds-fet.htm>, gdzie można znaleźć opis... tranzystora polowego z ciekłą izolowaną bramką w postaci kropki wody. Zrealizowałem bardzo podobne modele, ale nie z kropką wody, tylko z metalową, miedzaną folią w roli izolowanej bramki. Kupiłem dwie odmiany dużych, 20-milimetrowych fotorezystorów. Część moich fotorezystorów i „MOSFET-ów” widoczna jest na fotografii 1.



Fotografia 1

FET Transistor Homemade From Cadmium Sulfide Photocell.

Updated May 10 2009

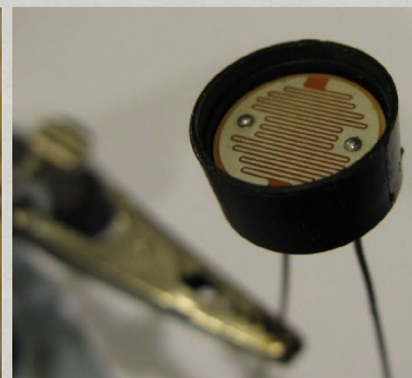
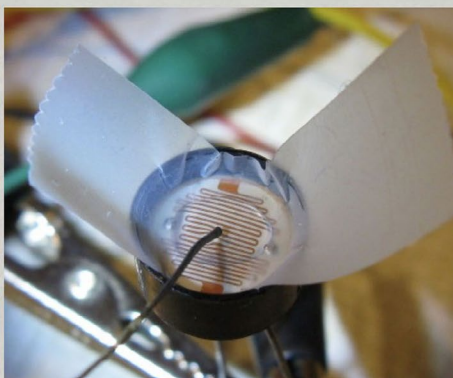
By Nyle Steiner K7NS May 7 2009.

CDS Photocell Made Into A FET Transistor

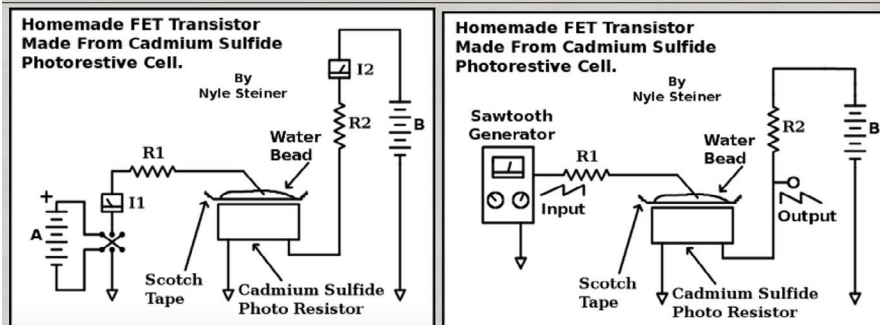
Fotorezystor to obwód dren – źródło, a izolowaną bramkę stanowi u mnie miedziana folia naklejona na „twarz” fotorezystora. Miałem obawy dotyczące foliowej bramki, bo powszechnie wiadomo, że kiedyś, jeszcze w latach 20. XX wieku powstały pomysły na „półprzewodnikową lampę elektronową”, ale nie udało się jej zrealizować z uwagi na znikomą głębokość wnikania pola elektrycznego w dość dobrze przewodzący prąd półprzewodnik. Dlatego najpierw, w roku 1948, przypadkowo pojawił się „potworek” w postaci tranzystora bipolarnego, niedługo potem złączowy FET, a MOSFET-y z izolowaną bramką wynaleziono dopiero na początku lat 70. Dlatego miałem obawy i trochę nieufnie podchodziłem do tematu.

Rysunek 2 pokazuje fragmenty wspomnianej strony. Opis jest przekonujący, a jeszcze większe wrażenie robią fotografie. Strona ta przekonuje, że realizacja tranzystora polowego w warunkach domowych jest możliwa. Postanowiłem to zbadać.

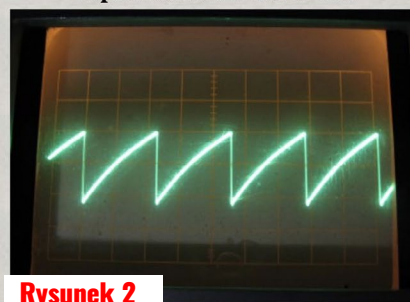
Zbadałem w miarę dokładnie i jako dowód na sukces mógłbym przedstawić **fotografię 3**, pokazującą moje stanowisko podczas eksperymentów. Niebieski przebieg na ekranie to sygnał z generatora, a sygnał żółty pobierany jest „z wyjścia tranzystora”. Wielkości przebiegów na ekranie oscyloskopu sugerują, że mamy element wzmacniający!



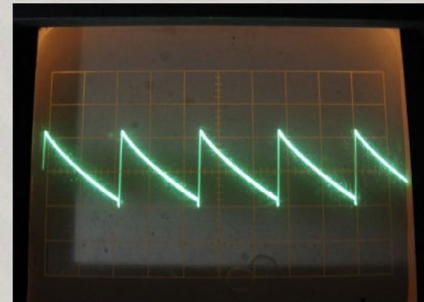
The picture above shows how transistor action was observed by improvising an insulated gate to a cadmium sulfide photo resistor. The picture was taken in normal light but the experiment had to be performed in the dark.



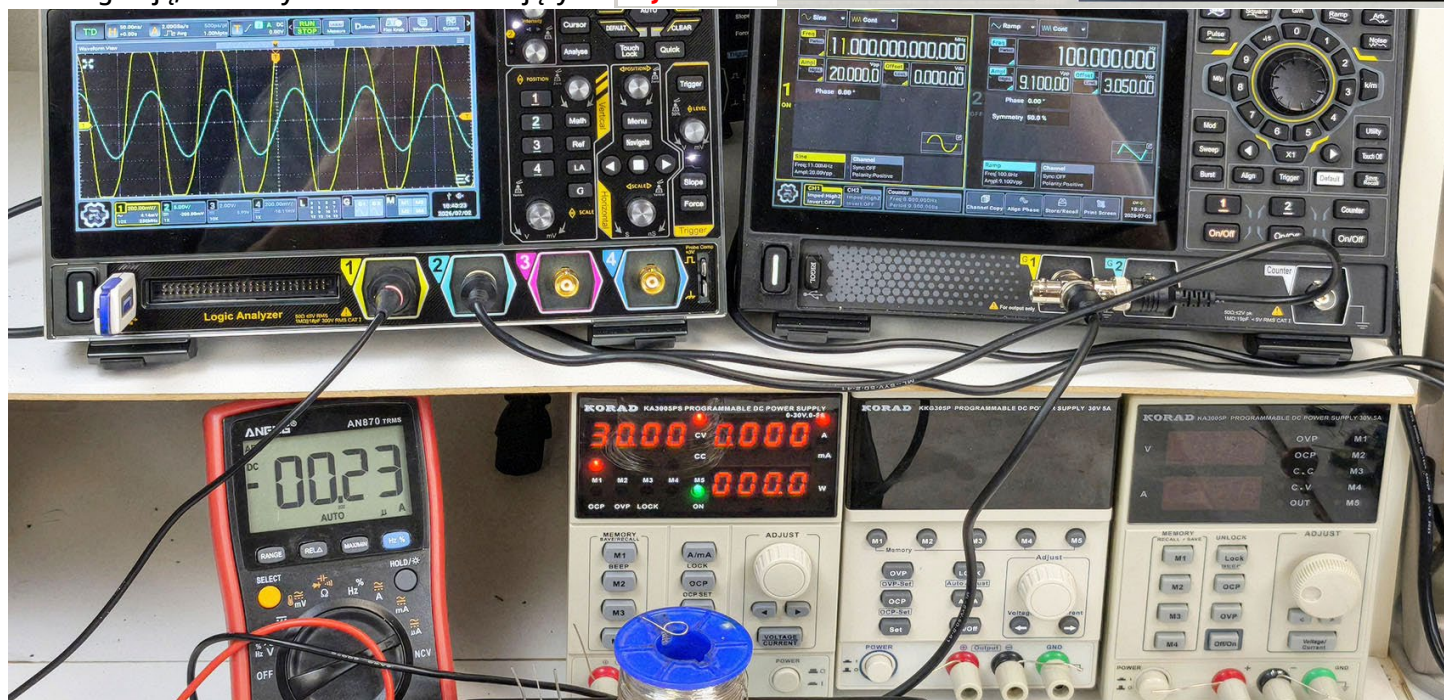
Input Sawtooth Waveform



Output Is Inverted Sawtooth Waveform



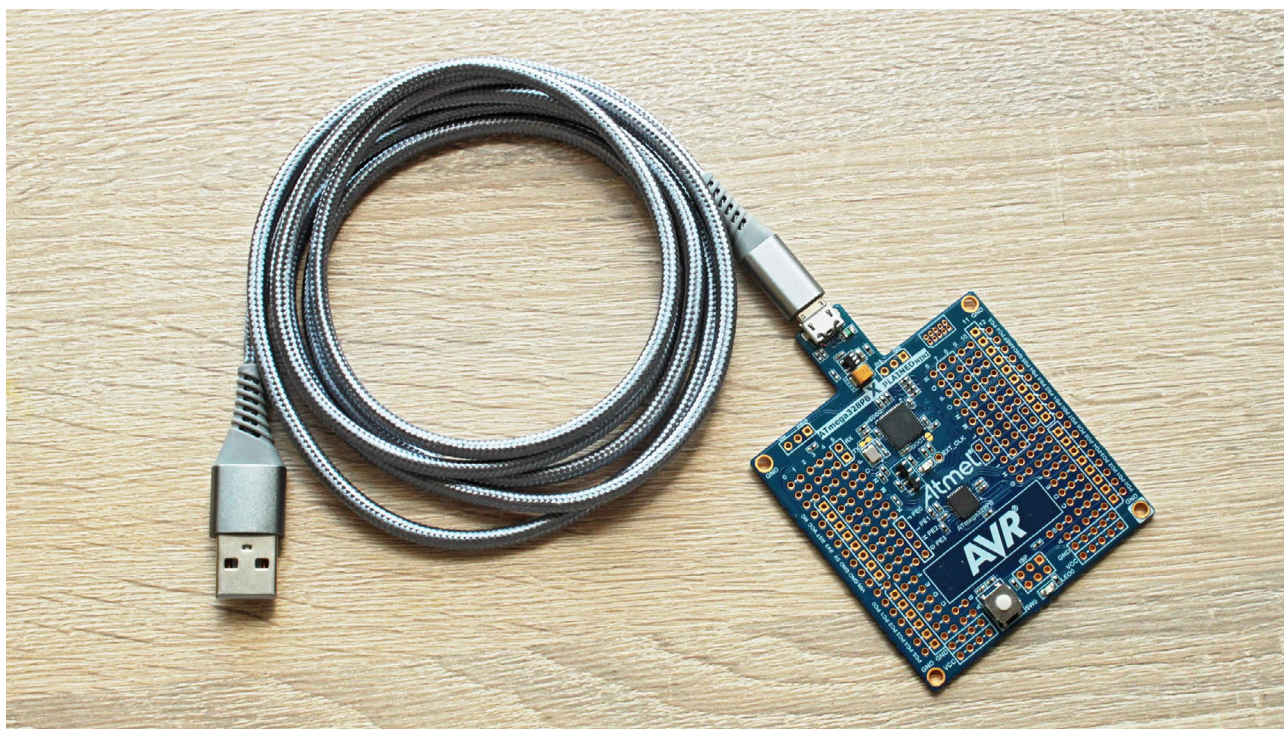
Rysunek 2



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

Fotografia 3





Programator i debugger w Xplained Mini

Firma Microchip oferuje moduł Xplained Mini, który ma wbudowany programator (mEDBG) pozwalający na umieszczenie kodu programu w pamięci Flash mikrokontrolera. Ma on jeszcze dodatkową cenną cechę – pozwala debugować program w docelowym systemie.

Krokowe uruchomienie programu
Podgląd stanu zmiennych

Debugowanie przerw
Zatrzymanie debugowania

Poszukiwanie błędów w programach jest dosyć pracochłonnym zajęciem. Z własnego doświadczenia wiem, że potrafi to zająć nawet 75% czasu poświęconego na opracowanie programu. Wypatrzenie błędów „własnymi oczami” czasem jest bardzo trudne. Zdarza się, że człowiek „zafiksuje” się nad rozwiązaniem programowym i nie dostrzeże własnych błędów. W wielu wypadkach program dla mikrokontrolera jest na tyle rozległy, że trudno jest ogarnąć całość, szczególnie gdy istotne są interakcje między zmiennymi rozrzuconymi po całym programie. Pewną pomocą może być użycie programatora/debugera wbudowanego w moduł *Xplained Mini*. Pozwala on na uruchomienie pro-

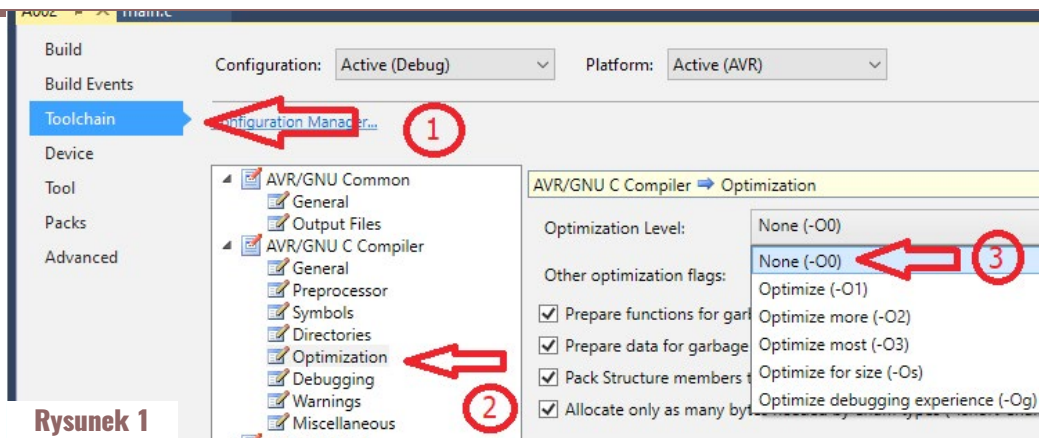
gramu w rzeczywistym układzie, w sposób krokowy i stanowi istotne wyposażenie warsztatowe. Dodatkowo można podglądać i modyfikować stan jego zmiennych.

Krokowe uruchomienie programu

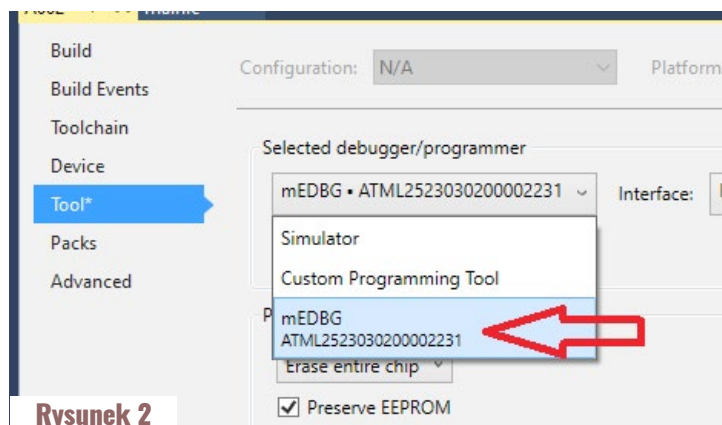
Napisałem prosty program, którego zadaniem jest żeby co pewien interwał czasu włączać i gasić diodę LED, istniejącą w płytce modułu. Program (W037) pokazuje **listing 1**, gdzie w odpowiednich pętlach program oczekuje na zmianę stanu zmiennej (*Event*), która zachodzi w funkcji obsługi przerwania, dioda LED jest przyłączona jako bit o numerze 5 (stała *LEDPin*) do PORTB (stała *LEDPort*):

```
int main ( void )
{
    SoftwareInit ( ) ;
    InitEnvir ( ) ;
    HardwareInit ( ) ;
    for ( ; ; )
    {
        for ( ; ; )
        {
            if ( Event )
            {
                Event = 0 ;
                LEDPort |= ( 1 << LEDPin ) ;
                break ;
            } /* if */ ;
            nop ( ) ;
        } /* for */ ;
        for ( ; ; )
        {
            if ( Event )
            {
                Event = 0 ;
                LEDPort &= ~ ( 1 << LEDPin ) ;
                break ;
            } /* if */ ;
            nop ( ) ;
        } /* for */ ;
    } /* for */ ;
    return ( 0 ) ;
} /* main
```

Do celów diagnostycznych warto wyłączyć opcje optymalizacji generowanego kodu. Domyślnie projekty w środowisku Atmel Studio mają włączoną optymalizację kodu. Chociaż działanie programu jest zgodne z zapisanymi instrukcjami, to niekoniecznie w pełni odpowiada programowi. Kompilator często zmienia kolejność instrukcji, czasami łączy funkcje ze sobą. Gdy funkcja jest oznaczona jako *static* (nie będzie użyta poza bieżącym modulem) i jest wywołana jeden raz, to w zależności od poziomu optymalizacji, zamiast ją wywołać może ją przenieść całkowicie w miejsce wywołania, co redukuje wielkość programu o kilka instrukcji maszynowych. Również możliwe jest, że zmienne, zamiast być lokowane w pamięci operacyjnej, znajdują się w rejestrach roboczych mikrokontrolera. Szczególnie dotyczy to zmiennych lokalnych w funkcjach (redukuje to zajętość pamięci RAM). Uruchamianie programu pod kontrolą debugera prowadzi do niespodziewanych działań:



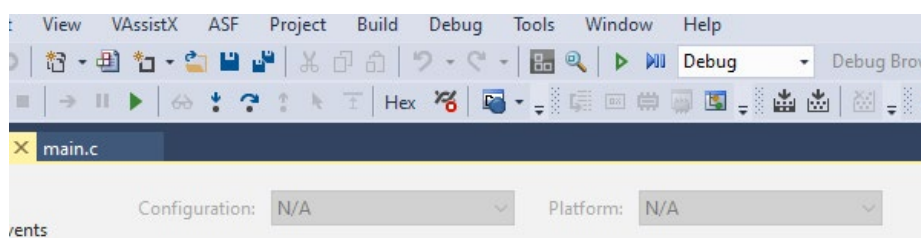
Rysunek 1



Rysunek 2

ich wartości (ponieważ nie mają one swego adresu). Na czas debugowania warto wyłączyć optymalizację kodu programu. W tym celu należy wejść w opcję *Properties - Toolchain - Optimization* i zmienić poziom na *None (-O0)*, **rysunek 1**. Po „zbadaniu sprawy” warto ponownie włączyć optymalizację kodu, by uzyskać bardziej optymalną wersję końcową a zmiany ustawień opcji warto zapisać na dysku.

Uzyskany program podlega standardowej kompilacji. Powstały kod programu należy załadować do pamięci Flash mikrokontrolera, jednak wcześniej należy wybrać odpowiednie warianty. Dołączenie modułu *Xplained* do komputera jest automatycznie rozpoznawane (**rysunek 2**), co staje się równoznaczne z tym, że może być bieżącym aktualnym urządzeniem programującym. Uruchomienie krokowe wymaga zmiany pewnej opcji (**rysunek 3**) – w okienku *Interface* należy wybrać *debugWIRE*.



Rysunek 3

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



Moduł wyświetlacza z układem TM1637

Praktycznie w każdym systemie mikroprocesorowym niezbędne są elementy wyświetlające informacje. Na rynku dostępnych jest wiele rozwiązań o różnych cechach użytkowych. Jednym z nich jest tytułowy wyświetlacz ze specjalizowanym układem scalonym TM1637.

Moduł wyświetlacza Polecenia układu TM1637

Moduł pod kontrolą mikrokontrolera

Moduł wyświetlacza z układem TM1637 to popularny, tani i łatwy w obsłudze komponent elektroniczny, służący do wyświetlania cyfr oraz kilku znaków semigrafiki (w tym liter pozwalających na prezentację danych w systemie szesnastkowym). Sam układ (TM1637) potrafi trochę więcej, niż oferuje moduł wyświetlacza. Technicznie może on sterować wyświetlaczem o maksymalnie 6 cyfrach oraz dysponuje możliwością obsługi prostej klawiatury o 16 przyciskach. Tytułowy moduł wykorzystuje układ w ograniczonym zakresie, co nie oznacza, że nie można zastosować TM1637 we własnych konstrukcjach wykorzystujących jego pełne możliwości.

Moduł wyświetlacza

Układ TM1637 autonomicznie obsługuje kilku-cyfrowy wyświetlacz 7-segmentowy LED w trybie multipleksowania. Ten sposób polega na szybkim zapalaniu i gaszeniu poszczególnych segmentów danej cyfry, przez synchroniczne włączanie zasilania odpowiedniej cyfry. Dzięki temu, że proces zachodzi z częstotliwością kilkudziesięciu Hz, ludzkie oko, ze względu na swoją bezwładność, odbiera obraz jako ciągły. Autonomiczność rozwiązania wymaga, by istniało kilka rejestrów 8-bitowych do przechowywania informacji, które segmenty w obrębie każdej cyfry mają być włączone (siedem segmentów wyświetlacza oraz sterowanie kropką lub dwukropkiem)

oraz układ cyklicznie przełączający zasilanie anod wyświetlaczy. Wszystkie te operacje realizuje układ scalony TM1637, toteż nie ma konieczności realizacji w programie jakichkolwiek operacji związanych z odświeżaniem danych. Jedyne działania jakie musi podjąć mikrokontroler, to wysłanie przez dedykowany interfejs szeregowy danych o włączeniu odpowiednich segmentów wyświetlacza (segment do świecenia włącza się jedynką logiczną). Układ sterujący nie ma wbudowanego generatora znaków, czyli program sam musi „rozłożyć” dany znak na poszczególne segmenty. Takie rozwiązanie z jednej strony zmusza do stworzenia w programie rozwiązania przekształcającego znak na sterowanie poszczególnymi segmentami, ale z drugiej strony umożliwia wyświetlanie znaków semigrafiki.

Wspomniany interfejs szeregowy w oryginalnej dokumentacji określony jest jako *two-wire serial* (**rysunek 1**), co może sugerować, że jest interfejs typu I²C (w mikrokontrolerach ATMEGA również używa się określenia *two wire serial interface*). Pomimo pewnych podobieństw do I²C nie jest to ten interfejs. Główna różnica polega na braku adresowania układów oraz odwróconej kolejności przesyłanych bitów (w I²C przesyłanie bitów zaczyna się od najbardziej znaczącego, w TM1637 od bitu najmniej znaczącego), przez co TM1637 nie jest zgodny ze standardem I²C. Brak możliwości adresowania implikuje, że na jednej magistrali szeregowej może być przyłączony jeden moduł.

Schemat modułu pokazuje **rysunek 2**, gdzie jego złącze zawiera 4 sygnały: dwa do zasilania modułu (dopuszczalna wartość napięcia zasilającego to 3,3 V...5 V, co pozwala na użycie modułu w syste-

TITAN MICRO™
ELECTRONICS

LED Drive Control Special Circuit TM1637

Features description

TM1637 is a kind of LED (light-emitting diode display) drive control special circuit with keyboard scan interface and it's internally integrated with MCU digital interface, data latch, LED high pressure drive and keyboard scan. This product is in DIP20/SOP20 package type with excellent performance and high quality, which is mainly applicable to the display drive of induction cooker, micro-wave oven and small household electrical appliance.

Function features

- Applied power CMOS technique
- The display mode (8 segments*6 bit) supports output by common anode LED.
- Keyboard scan (8*2bit), with enhanced identification circuit with anti-interference keys
- Luminance adjustment circuit (adjustable 8 duty ratio)
- Two-wire serial interface (CLK, DIO)
- Oscillating type: Built-in RC oscillator
- Built-in power-on reset circuit
- Built-in automatic blanking circuit
- Package type: DIP20/SOP20

Rysunek 1

mach o obniżonym napięciu zasilającym), sygnał zegarowy (CLK) oraz przesyłanych danych (DIO). Źródłem sygnału CLK jest zawsze mikrokontroler, natomiast w przypadku obsługi jedynie wyświetlacza (bez klawiatury), sygnałem DIO steruje mikrokontroler z wyjątkiem jednego bitu potwierdzenia ACK, który jest wymuszany przez TM1637.

Polecenia układu TM1637

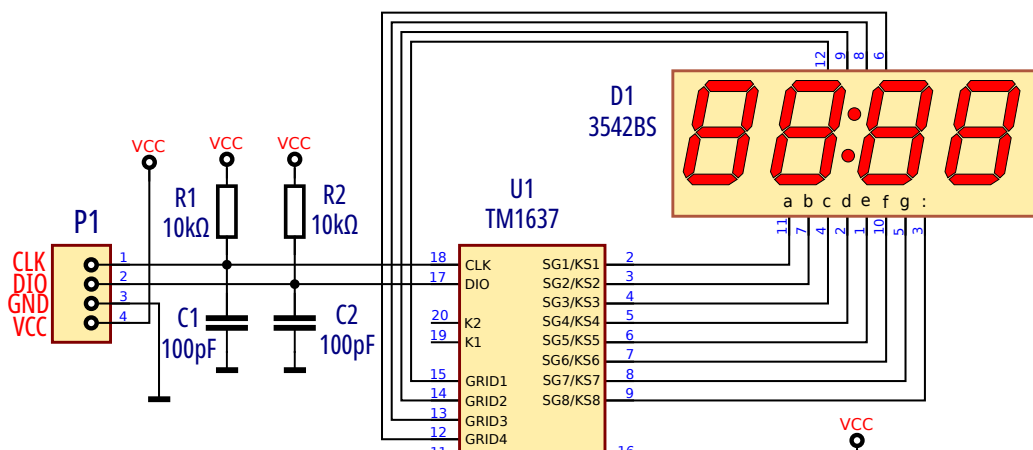
Komunikacja z układem TM1637 opiera się na przesyłaniu jednobajtowych rozkazów za pomocą interfejsu szeregowego. Polecenia te są identyfikowane przez dwa najbardziej znaczące bity wysyłanego bajtu.

Kombinacja binarna 01xxxxxx (40 hex) oznacza polecenie konfiguracyjne, które na młodszej części zawiera odpowiednie szczegóły. Możliwe są:

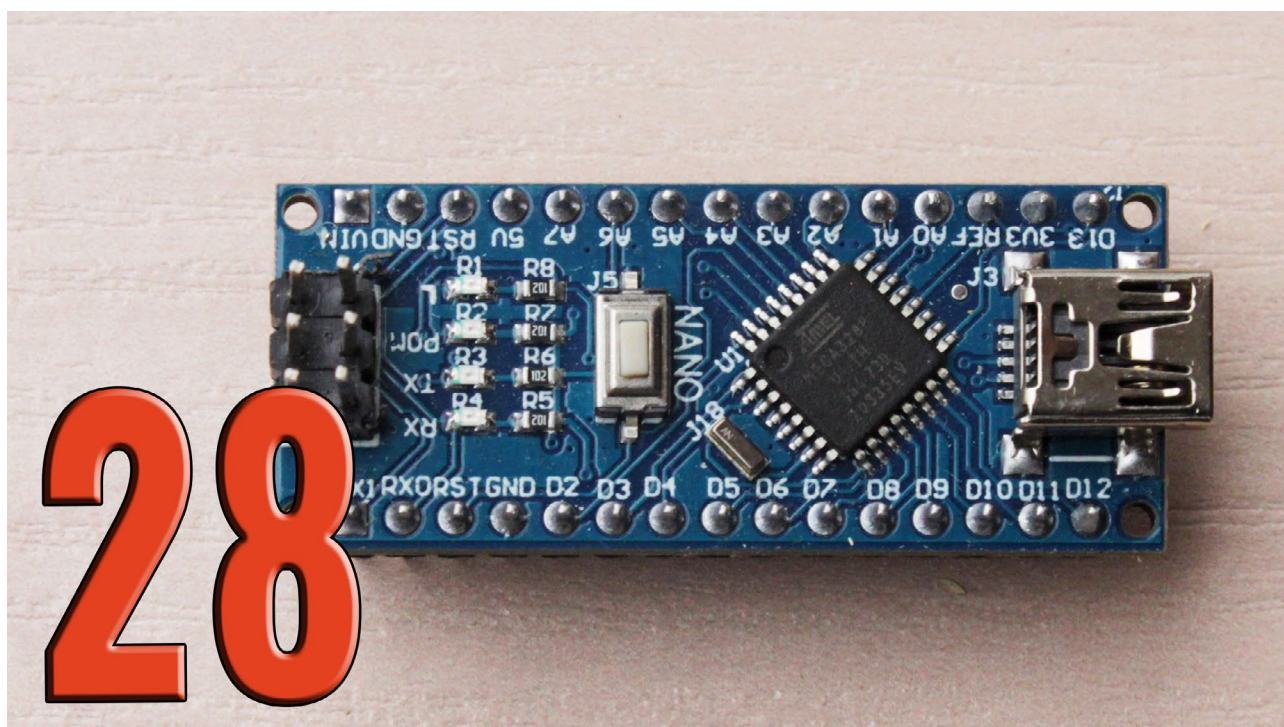
- polecenie o kodzie 40 hex oznacza automatyczną inkrementację adresu, po wysłaniu adresu pierwszej cyfry, każdy kolejny wysłany bajt danych automatycznie trafia na następną pozycję wyświetlacza,
- polecenie o kodzie 42 hex oznacza odczyt klawiatury, w przypadku modułu z fotografii tytułowej ta

możliwość nie jest wykorzystana (nie ma klawiatury),

- polecenie o kodzie 44 hex oznacza ustawienie adresu (pozycji na wyświetlaczu, gdzie ma być umieszczony znak, przeciwieństwo autoinkrementacji), każdy wyświetlany znak



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



Mikroprocesorowa ośła łączka, część 28

Dziedziczenie pozwala na realizację różnorodnych rozwiązań. Tu można wskazać na funkcjonalność typu HAL, która pozwala na adaptację obsługi sprzętu do konkretnych uwarunkowań sprzętowych. Inną możliwością jest rozszerzenie funkcjonalności istniejących komponentów.

Analizator leksykalny

Proste zastosowanie analizatora leksykalnego

Rozbudowa analizatora leksykalnego

Analizator syntaktyczny

Układ praktyczny

W poprzedniej części została opisana możliwość adaptacji obsługi odpowiedniego sprzętu elektronicznego do określonego środowiska. Warto oddzielić logikę sterowania dołączanych komponentów od rzeczywistego rozwiązania samego przyłączenia. Teraz zostanie przedstawiona koncepcja rozbudowy funkcjonalności komponentu o nowe możliwości. Tłem do rozważań jest analiza składniowa. Wiele urządzeń jest sterowanych poleceniami tekstowymi przesyłanymi z zewnątrz. To może być kanał transmisji szeregowej, poprzez który dostarczane są polecenia. Bardziej nowoczesne jest zastosowanie sieci ethernetowej. Różnica związana jest jedynie z „nośnikiem danych”. Wspólną częścią jest analizator tych poleceń wbu-

dowany w mikrokontroler. Składa się on z kilku warstw, gdzie podstawową jest analiza leksykalna.

Analizator leksykalny

Analiza leksykalna (ang. lexical analysis) to proces wstępnego przetwarzania tekstu, polegający na podziale ciągu znaków na mniejsze, logiczne jednostki zwane tokenami (np. wyrazy, liczby, znaki przestankowe). Proces ten obejmuje też odrzucenie elementów bez znaczenia, jak przykładowo spacje. Wyodrębniając podstawowe elementy (tokeny), analizator sprawdza poprawność ich zapisu i może zgłosić błąd formatu zapisu. Przykładowo jako napis (ciąg znaków zapisany między znakami apostrofów lub cudzysłowu) musi mieć znak zamknięcia

łańcucha, identyczny jak znak otwierający (apostrof lub cudzysłów). Podobnie wymagane jest, by liczba z ułamkiem dziesiętnym po znaku oddzielającym część całkowitą od ułamkowej (w tej roli występuje znak kropki), zawierała cyfry.

Do analizy leksykalnej zdefiniowana jest klasa *TLexicalAnalyzer* zawarta w (*LexicalModule.h* oraz *LexicalModule.cpp*). Postać komponentu pokazuje listing 1.

```
class TLexicalAnalyzer {
public :
    TLexicalAnalyzer ( ) ;
    void AnalError ( uint16_t ErrorNo ) ;
    void AnalCardError ( uint16_t ErrorCode ,
                        uint32_t LCValue ) ;
    void AnalStrError ( uint16_t ErrorCode ,
                      uint8_t * StrValue ) ;
    uint8_t ReadSymbol ( ) ;
    uint32_t ReadCardinal ( ) { return CardInpValue ; } ;
    float ReadReal ( ) { return RealInpValue ; } ;
    uint8_t * ReadString ( ) { return StrInpValue ; } ;
    void SetRealEnable ( uint8_t Enable ) { RealEnable = Enable ; } ;
    void OpenLineAnalyser ( uint8_t * LineBuffer ) ;
    void CheckEndOfLine ( uint16_t ErrorCode ) ;
    uint8_t ErrorsPresent ( ) { return ErrorPresent ; } ;
    void PrintAllErrors ( ) ;
    virtual void StartPrintErrors ( ) ;
    virtual void StopPrintErrors ( ) ;
    virtual void PrintError ( uint16_t ErrorCodes ,
                          uint8_t * ErrorSpelling ) ;
    virtual void PrintLine ( uint8_t * Buffer ) ;
public :
    SpellingType      InpSpelling ;
private :
    uint8_t GetNextChar ( ) ;
    void PutSpellCh ( uint8_t PCh ) ;
    void ReadString ( uint8_t Terminator ) ;
private :
    LineBuffType      SourceLine ;
    uint8_t           RealEnable ;
    uint8_t *         CurrChar ;
    uint32_t          CardInpValue ;
    float             RealInpValue ;
    uint8_t           LineChar ;
    uint8_t           LineChar2 ;
    uint16_t          SpeelCt ;
    InpStringType      StrInpValue ;
    uint16_t          ErrCodes [ MaxErr ] ;
    SpellingType       ErrSpell [ MaxErr ] ;
    uint8_t           ErrIndex ;
```

Publicznie dostępne elementy pełnią następujące funkcje:

- *TLexicalAnalyzer* – konstruktor obiektu, inicjuje wartości początkowe zmiennych obiektu, między innymi „włącza” rozpoznawanie liczb ułamkowych (aby umożliwić analizę adresów IP jako czterech liczb rozdzielonych znakiem kropki, należy „wyłączyć” analizę liczb ułamkowych),

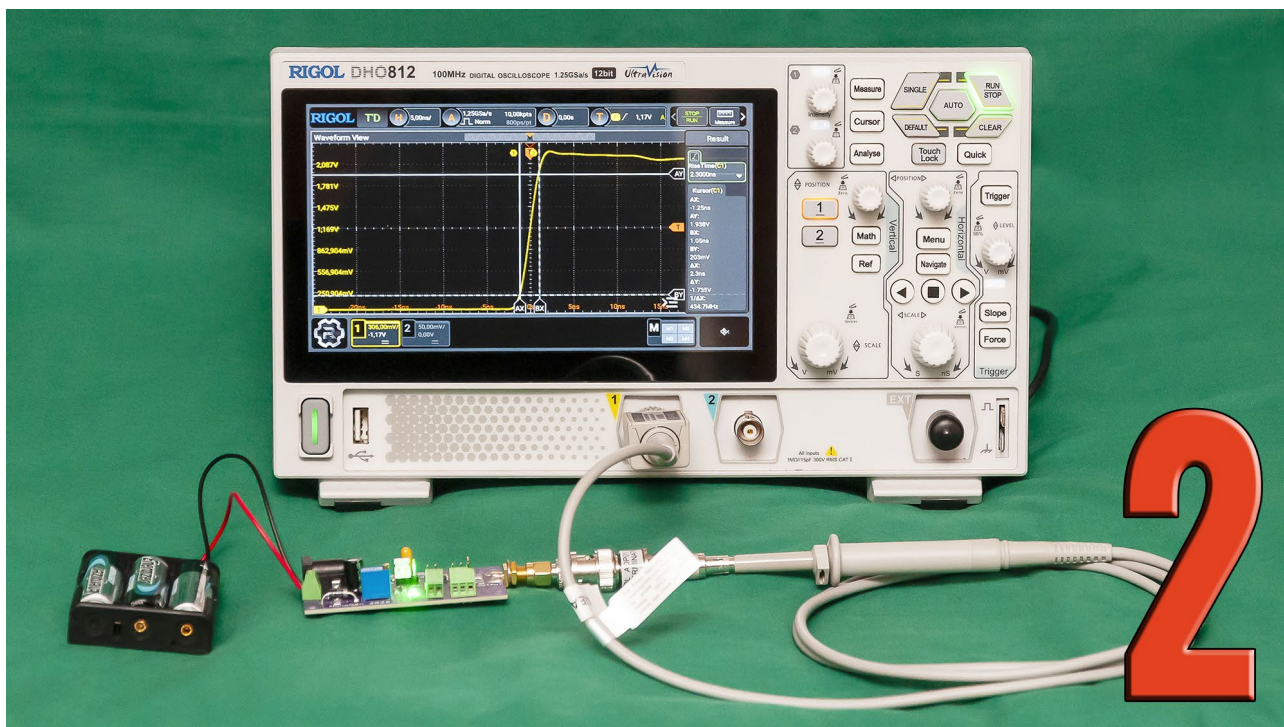
- *OpenLineAnalyser* – metoda przyłączająca do

analizatora bufor zawierający tekst do analizy, podany parametr wywołania jest tablicą znakową w rozumieniu języka C/C++ (zakończony znakiem o kodzie 0),

- *SetRealEnable* – metoda do włączania/wyłączania rozpoznawania liczb ułamkowych (aby analiza adresu IP zapisanego w konwencjonalnym formacie nie kolidowała z rozpoznawaniem liczb ułamkowych należy ją wyłączyć),

- *ReadSymbol* – podstawowa metoda do wczytania kolejnego tokena, jego treść zawsze jest zawarta w *InpSpelling*, przy zgłaszaniu błędów treść tokena jest zapamiętana, by umożliwić lokalizację błędnych zapisów, wynikiem funkcji jest liczba 8-bitowa pozwalająca na rozpoznanie wczytanego tokena, przykładowo *LexicNoSy* oznacza, że wczytany token nie jest rozpoznany, *LexicIdentSy* sygnalizuje, że został wczytany wyraz, *LexicEndLineSy* informuje, że w analizie został

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



Impulsowy test oscyloskopów i sond pomiarowych, część 2

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.

Referencyjny pomiar oscyloskopu Rigol DH0812
Architektura stanowiska pomiarowego
Konfiguracja wejściowa oscyloskopu
Analiza wpływu terminacji

Referencyjny czas narastania oscyloskopu
Metodologia podłączania sond w pomiarach nanosekundowych
Podsumowanie

Referencyjny pomiar oscyloskopu Rigol DH0812

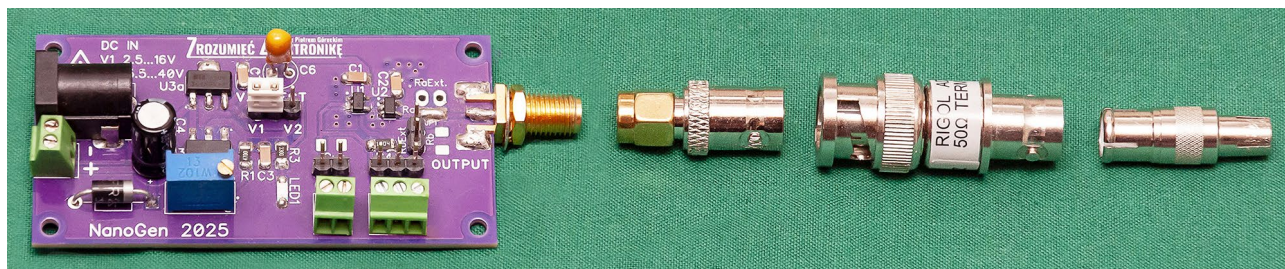
W pierwszej części cyklu przedstawiłem podstawy metody impulsowej oraz omówiłem zależność między czasem narastania a efektywnym pasmem przenoszenia toru pomiarowego. Kolejnym krokiem było wyznaczenie referencyjnej odpowiedzi impulsowej oscyloskopu Rigol DH0812, która miała stanowić punkt odniesienia dla wszystkich dalszych pomiarów sond.

Celem eksperymentu nie było jedynie określenie czasu narastania samego oscyloskopu. Równie istotne było sprawdzenie wpływu terminacji linii

transmisyjnej na obserwowany przebieg oraz pokazanie, jak łatwo można uzyskać pozornie lepsze wyniki wskutek nieprawidłowej konfiguracji pomiarowej.

Architektura stanowiska pomiarowego

Głównym źródłem sygnału testowego był nanogenerator impulsowy, opisany w poprzedniej części cyklu. Generator wytwarza impulsy o bardzo krótkim czasie narastania, dzięki czemu może pełnić funkcję szerokopasmowego źródła pobudzającego podczas badań odpowiedzi impulsowej aparatury pomiarowej.



Fotografia 7

Przy czasach narastania liczonych w pojedynczych nanosekundach wszystkie połączenia należy traktować jak elementy linii transmisyjnych. Dodatkowe złącza, przejściówki oraz przewody wprowadzają pasożytnicze indukcyjności i pojemności, które mogą powodować odbicia sygnału, lokalne rezonanse oraz zniekształcenie odpowiedzi impulsowej.

W początkowej konfiguracji stanowiska pomiarowego stosowałem bardziej rozbudowany łańcuch połączeń zawierający przejściówkę SMA-BNC (gniazdo) oraz dodatkowy łącznik typu „beczka” BNC. Analiza przebiegów wykazała zwiększony poziom oscylacji przejściowych oraz wydłużony czas ich wygaszania. Każde dodatkowe złącze stanowi bowiem lokalną nieciągłość impedancji, która przy bardzo szybkich zboczach może stać się źródłem odbić i rezonansów pasożytniczych.

W kolejnej iteracji stanowiska zastąpiłem przejściówkę SMA-BNC (gniazdo) przejściówką SMA-BNC (wtyk), co umożliwiło całkowite wyeliminowanie dodatkowego łącznika BNC (**fotografia 7**).

Pozwoliło to ograniczyć liczbę nieciągłości impedancyjnych oraz wpływ pasożytniczych elementów RLC w torze pomiarowym. W rezultacie odpowiedź impulsowa stała się bardziej stabilna, a amplituda oscylacji przejściowych uległa zmniejszeniu.

Przyjęta przeze mnie metodologia obejmowała dwa etapy:

- wyznaczenie referencyjnej odpowiedzi impulsowej samego oscyloskopu,
- pomiary z udziałem sond oraz porównanie ich wpływu na odpowiedź impulsową całego toru pomiarowego.

Konfiguracja wejściowa oscyloskopu

Kanały analogowe oscyloskopu Rigol DH0812 charakteryzują się impedancją wejściową 1 MΩ, równolegle z pojemnością około 15 pF. Oscyloskop nie ma przełączanej terminacji 50 Ω, dlatego podczas pomiarów szerokopasmowych konieczne jest zastosowanie zewnętrznego terminatora. W trakcie badań wykonałem pomiary w dwóch konfiguracjach połączeń. Pozwoliło to zobrazować jeden z najczęściej popełnianych błędów podczas pomiarów bardzo szybkich sygnałów.

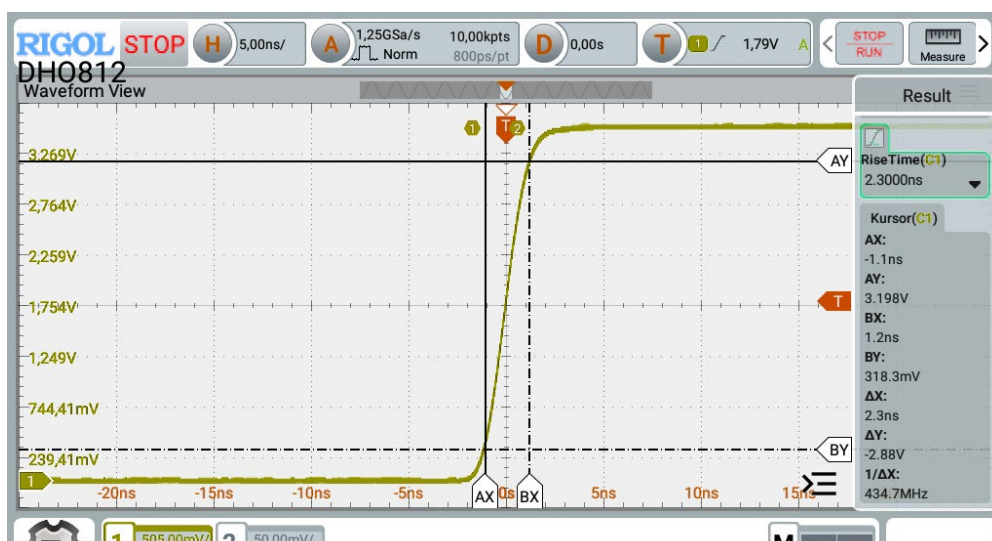
Połączenie A – bezpośrednie (błędne metrologicznie)

Łańcuch połączeń miał postać:

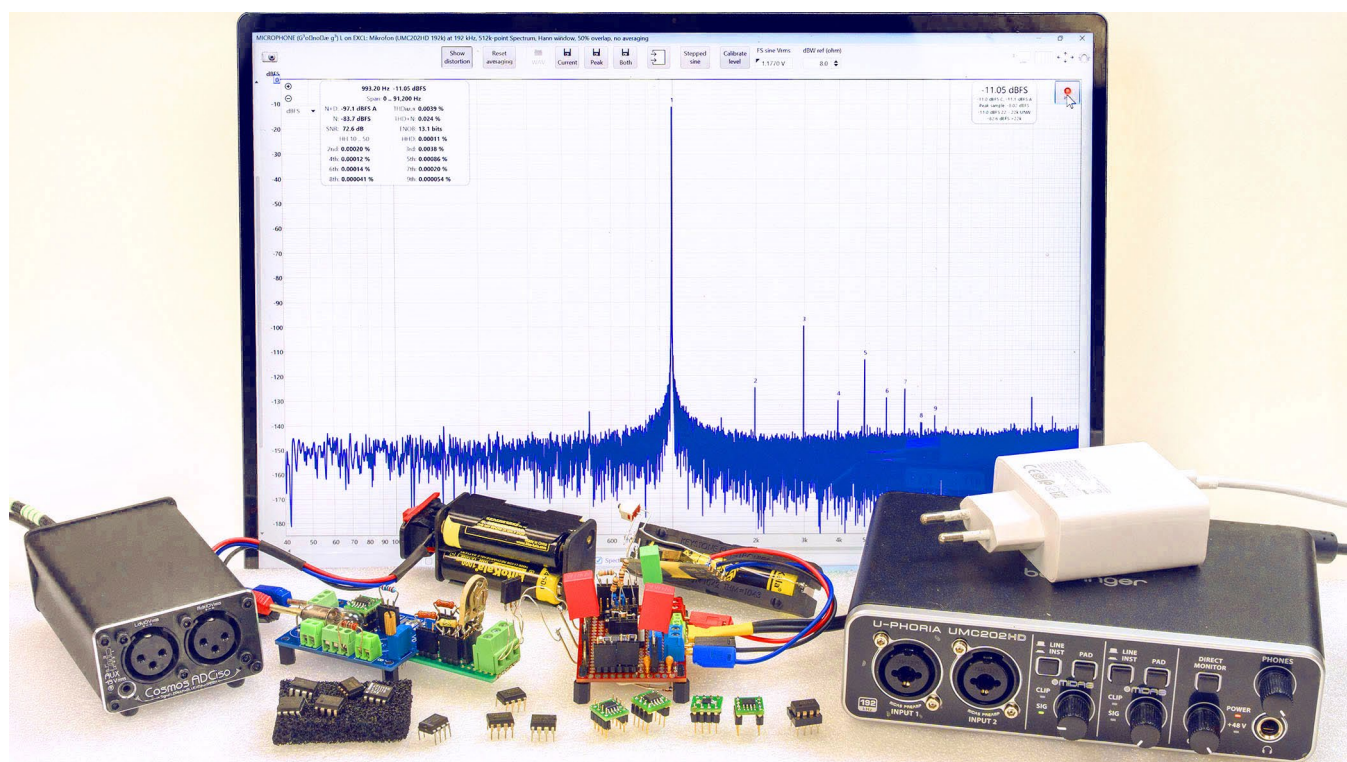
[Wyjście SMA generatora] → [adapter SMA-BNC] → [wejście oscyloskopu].

W tej konfiguracji linia transmisyjna została zakończona wejściem o impedancji około 1 MΩ. Dla

sygnałów o bardzo dużej szybkości narastania oznacza to praktycznie rozwarcie linii transmisyjnej. Sygnał docierający do wejścia oscyloskopu ulegał niemal całkowitemu odbiciu. Fala odbita powracała do punktu pomiarowego i nakładała się na falę padającą, powodując lokalny wzrost amplitudy oraz pozorne



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



Generatory z mostkiem Wiena (1)

W 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.

Komu naprawdę są potrzebne takie generatory?

Jak mierzyć skrajnie małe zniekształcenia?

Analogowe generatory „sinusa”

Generator z mostkiem Wiena - podstawy

Wartości elementów określających częstotliwość

Diodowa stabilizacja amplitudy

Wykorzystanie termistora

Stabilizacja za pomocą żarówki

Dziś dominuje technika cyfrowa, ale **nadal najlepsze generatory sinusoidalne realizowane są w sposób analogowy**. Dawniej w laboratoriach, pracowniach i w warsztatach elektronicznych używano generatorów pracujących na różnych zasadach. Dziś dominują generatory wykorzystujące bezpośrednią syntezę cyfrową (DDS). Faktem jest, że ich możliwości są coraz lepsze, a ceny coraz niższe. Jednak w niektórych zastosowaniach nadal potrzebne są inne generatory. Na przykład impulsów o bardzo stromych zboczach. Innym przykładem są generatory jak najbardziej czystego „sinusa”, czyli przebie-

gu sinusoidalnego o jak najmniejszych zniekształceniach harmonicznym (THD). Najmniejszych, czyli jakich? Dobre generatory DDS dają „sinus” o zawartości harmonicznym rzędu -60dB, czyli 0,1%. Jeśli to nie wystarczy, potrzebny jest inny generator. Dobre komputerowe karty audio mogą wytwarzać przebiegi dużo czystsze, ale są kosztowne i mają też swoje wady. Dlatego **nadal w niektórych zastosowaniach potrzebne są analogowe generatory „czystego sinusa” o zawartości harmonicznym rzędu jednego „pipiema”**: 1 ppm = 0,0001% czyli -120 decybeli, a nawet mniej.

Komu naprawdę są potrzebne takie generatory?

To interesujące, że nadal jest sporo osób, którym taki generator o ultraniskich zniekształceniach THD potrzebny jest „dla micia”. Chcą oni taki generator samodzielnie zbudować i mieć tylko dla własnej ogromnej satysfakcji. Właśnie samodzielnie zbudować, a nie kupić, bo tu kluczowe znaczenie ma walka o jak najmniejsze zniekształcenia. Dość łatwo się przekonać, że poziom zniekształceń wyznacza głównie użyty element aktywny, czyli w praktyce wzmacniacz operacyjny. To może być interesująca „sztuka dla sztuki” i forma specyficznych zawodów sportowych. Ale niekoniecznie!

Są osoby, które znają się na elektronice i chcą sprawdzać „skrajne” parametry różnych układów, a jeszcze bardziej elementów elektronicznych. W szczególności sprawdzają liniowość, czyli badać jakie zniekształcenia mogą wprowadzać różne odmiany wzmacniaczy operacyjnych, kondensatorów i rezystorów przy różnych częstotliwościach i amplitudach sygnału. Tak, nie tylko zniekształcenia wzmacniaczy, ale też zniekształcenia wprowadzane choćby przez kondensatory i rezystory, które też idealne nie są i nie do końca spełniają prawo Ohma!

„Bardzo czysty sinus” potrzebny jest miłośnikom techniki audio do sprawdzania parametrów sprzętu. To wcale nie są fanaberie pseudoaudiofilów, tylko interesujące, pożyteczne i wartościowe testy, do których potrzebne jest źródło sygnału sinusoidalnego o naprawdę znakomitej czystości, czyli o skrajnie małych zniekształceniach.

Faktem jest, że „czyste” przebiegi można dziś wytworzyć „cyfrowo”, choćby za pomocą komputerowych kart dźwiękowych. Tak, ale dobry analogowy generator „sinusa” jest potrzebny między innymi

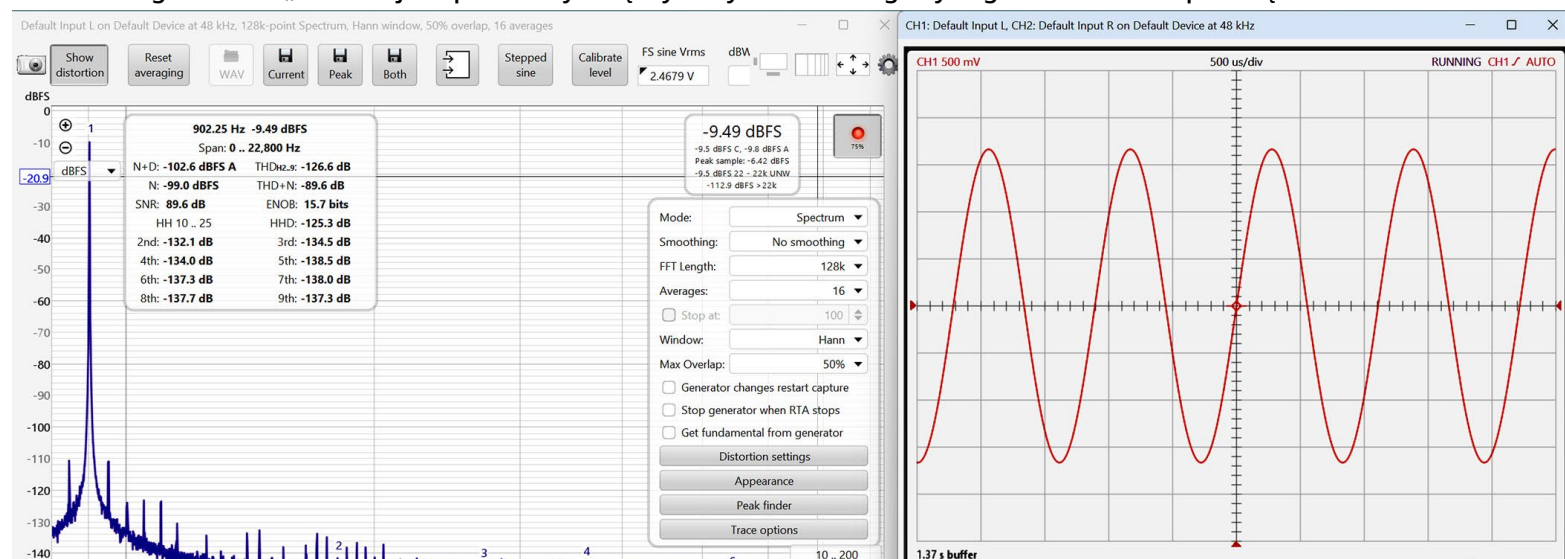
właśnie do sprawdzania parametrów takich kart, w szczególności zawartych w nich przetworników ADC. Dlatego naprawdę warto zainteresować się analogowymi generatorami „czystego sinusa”, w szczególności wersjami z mostkiem Wiena.

Oczywiście pojawia się podstawowe pytanie...

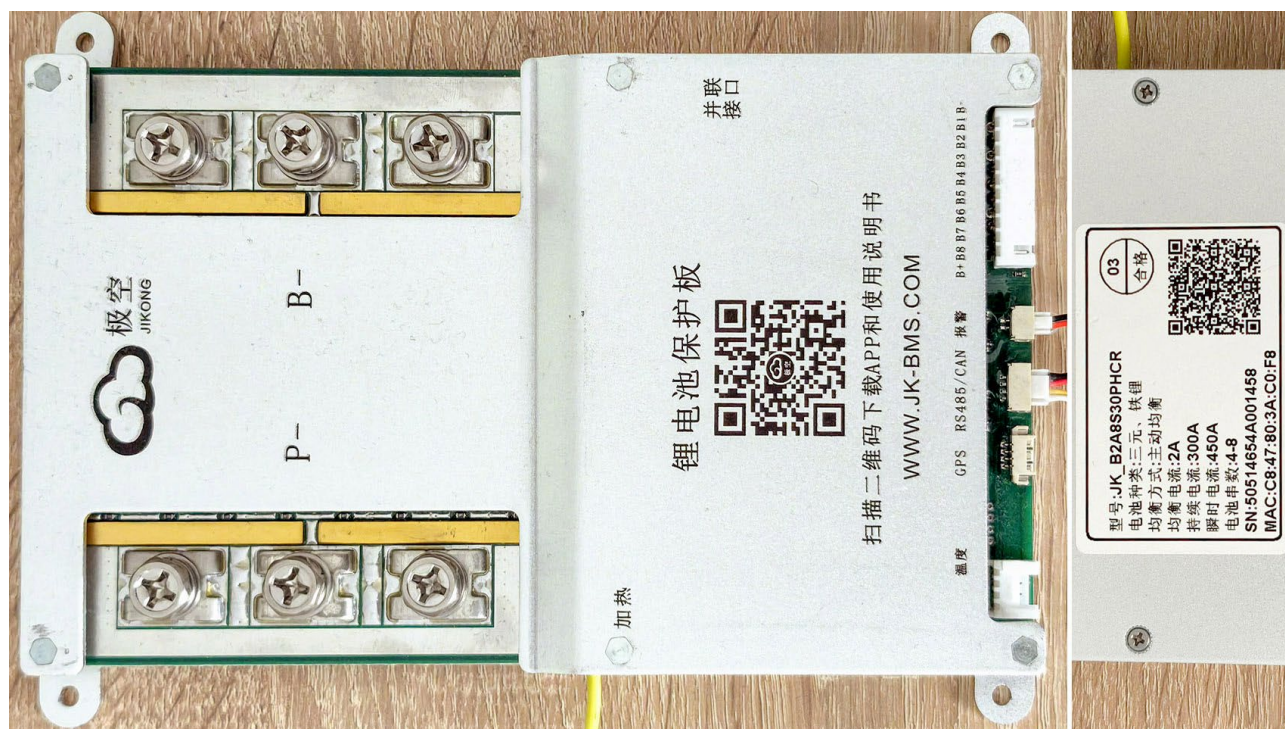
Jak mierzyć skrajnie małe zniekształcenia?

Nieliniowość objawia się przede wszystkim pojawieniem w sygnale składników o częstotliwościach harmonicznych. Jeżeli są to przebiegi z zakresu częstotliwości akustycznych, a nie sygnały w.cz., to akurat dziś możliwy jest pomiar tych harmonicznych także w warunkach amatorskich i nie trzeba mieć horrendalnie drogiego sprzętu pomiarowego kosztującego tysiące dolarów. Jesteśmy w dobrej sytuacji, ponieważ dostępne są niedrogie, komputerowe karty dźwiękowe, zawierające 24-, a nawet 32-bitowe przetworniki ADC. Takie karty komputerowe mogą z powodzeniem pracować w roli sprzętu pomiarowego, a konkretnie jako analizatory widma o dynamice rzędu 120, a nawet więcej decybeli. I właśnie na **fotografii tytułowej** przedstawiony jest pomiar zawartości harmonicznych zaskakująco prostego generatora przebiegu sinusoidalnego, za pomocą kosztującej około 300 złotych karty Behringer UMC202HD (i prostej przystawki).

Natomiast **rysunek 1** pokazuje pomiar zawartości widmowej sygnału innego egzemplarza takiego generatora, ale za pomocą dużo lepszej karty z 32-bitowym przetwornikiem ADC. Generator opisany jest w artykule **Eksperymentalny generator przebiegu sinusoidalnego**. Aby jednak taki generator zrealizować, warto mieć trochę więcej wiedzy na temat analogowych generatorów. Naprawdę warto!



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



BMS w akumulatorze – co to i po co?

Akumulatory litowe zostały wstępnie omówione w artykule [Akumulatory wczoraj i dziś \(2\)](#), w numerze 7/2025. A w trzech ostatnich artykułach tej serii zacząłem omawiać przyczyny oraz rozmaite sposoby zabezpieczania pojedynczych ogniw litowych. W tym artykule omawiam układy nazywane BMS.

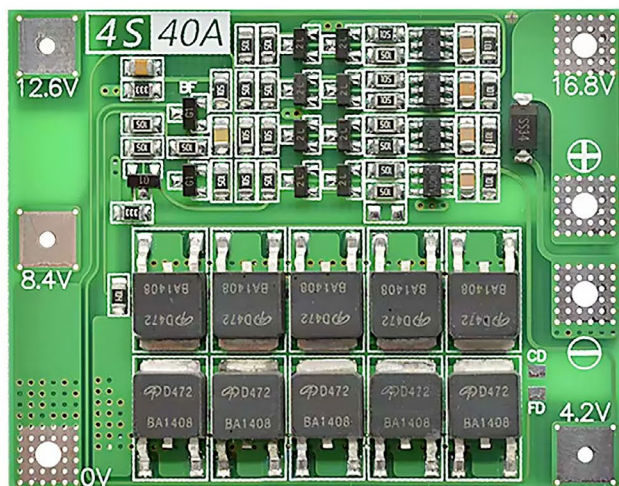
[Czy BMS i balanser to to samo?](#)
[Rozmaite rozwiązania BMS-ów](#)
[BMS-y z układami DW01](#)

[Zabezpieczenie nadprądowe w BMS-ach](#)
[BMS-y ze specjalizowanymi układami scalonymi](#)

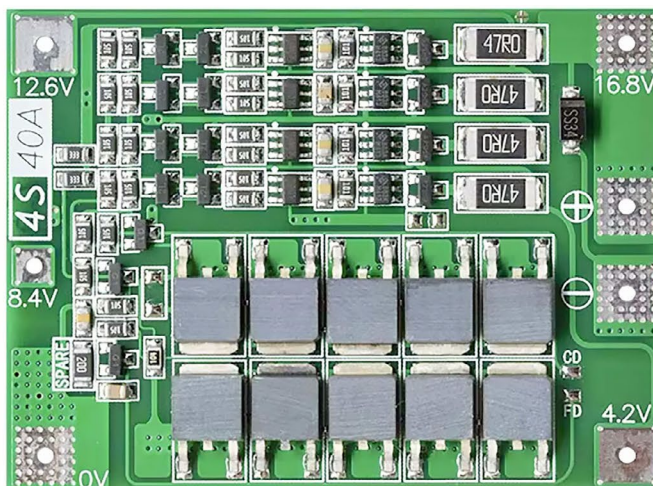
W dwóch poprzednich artykułach tego cyklu wyczerpująco omówiłem sposoby zabezpieczenia pojedynczych akumulatorów litowych, których napięcie nominalne zawiera się w granicach 3,2...3,8 V. Dla wielu zastosowań jest to napięcie małe, dlatego powszechnie wykorzystuje się też pakiety ogniw litowych połączonych w szereg (*serial connected*) lub równolegle (*parallel connected*). Zestaw trzech połączonych w szereg ogniw, oznaczony jest 3S, ma trzy razy wyższe napięcie, a pojemność w amperogodzinach taką samą, jak akumulator pojedynczy. Natomiast pakiet sześciu akumulatorów, w którym po dwa połączone są równolegle, a potem te trzy grupy szeregowo, ma oznaczenie 3S2P – ma napię-

cie trzy razy wyższe, a pojemność dwa razy większą niż pojedyncze ogniwo.

Bardzo popularne są kwasowe akumulatory 12-woltowe, składające się z sześciu ogniw – cel. Z akumulatorów litowych budowane są ich przybliżone odpowiedniki. Mianowicie zestaw 3S, czyli zestaw trzech ogniw Li-Ion ma napięcie nominalne $3 \times 3,7 \text{ V} = 11,1 \text{ V}$. Po pełnym naładowaniu napięcie wynosi $3 \times 4,2 \text{ V} = 12,6 \text{ V}$. Jeżeli chodzi o fosfatowe LiFePO₄, to zestaw 4 S, czyli cztery połączone szeregowo ogniwa mają napięcie nominalne 12,8 V, a po pełnym naładowaniu napięcie 14,6 V – nie tylko pod względem napięcia zestawu 4S z ogniwami LiFePO₄ są podobne do akumulatorów kwasowych.



Fotografia 1 „Poczwórny” BMS bez balansera



„Poczwórny” z balanserem

Oprócz zestawów 12-woltowych, popularne są zestawy o wyższym i dużo wyższym napięciu kilkunastu woltów, przeznaczone do rozmaitych pojazdów elektrycznych, np. do deskorolek, hulajnóg, rowerów i skuterów. Natomiast w samochodach elektrycznych są montowane zestawy ogniw Li-Ion, dające napięcie rzędu kilkuset woltów.

I wracamy do najważniejszego: **akumulatory litowe są delikatne, łatwo je uszkodzić**. Przedstawione w poprzednich artykułach serii sposoby ochrony **pojedynczych** ogniw okazały się rozbudowane i skomplikowane, natomiast w przypadku **zestawów ogniw połączonych szeregowo**, zadanie jest jeszcze trudniejsze, bo trzeba chronić nie tyle cały zestaw, co **indywidualnie trzeba chronić wszystkie składowe ogniwa połączone szeregowo**.

Wszystko to związane jest z faktem, że te szeregowo połączone akumulatory nie mają identycznej pojemności. Występują różnice między ogniwami i problem narasta z czasem. Problem, którego nie ma przy ochronie pojedynczego akumulatora. To doprowadza nas najpierw do tematu BMS-ów, a potem także do balanserów.

Czy BMS i balanser to to samo?

Mnóstwo osób nie ma jasności co do określeń „beemes” (BMS – Battery Management System) oraz balanser (*Balancer*). Na pewno nie są to synonimy!

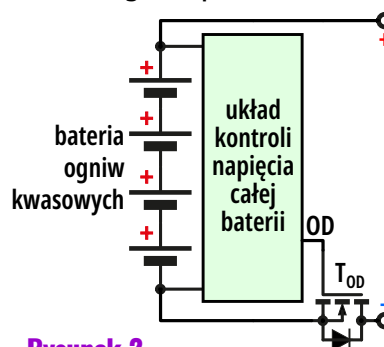
Określenie **BMS** jest szersze i obejmuje wszelkie systemy zarządzania baterią ogniw. Natomiast **balanser** do urządzenia lub blok, który ma coś balansować, czyli równoważyć, wyrównywać. W założeniu ma „wyrównywać pojemność” poszczególnych ogniw składowych, co będę omawiał później.

W każdym razie **prostsze systemy BMS nie zawierają żadnego balansera**, a jedynie obwody ochrony

Na **fotografii 1** pokazane są dwa stosunkowo proste BMS-y: jeden z balanserem, drugi bez. Bardziej rozbudowane „beemensy” mają więcej funkcji ochronnych, a tym zabezpieczenia termiczne i zawsze, mniej lub bardziej skuteczne, obwody balansera.

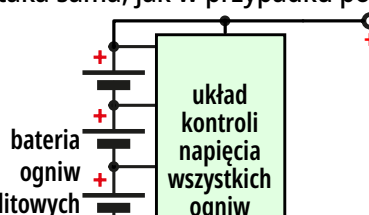
Różne rozwiązania BMS-ów

Otóż dla ochrony przed rozładowaniem, także w przypadku akumulatorów kwasowych, niekiedy stosowano wyłączniki według idei z mocno uproszczonego **rysunku 2**, gdzie przekaźnik, lub lepiej pojedynczy MOSFET odłączał akumulator, gdy jego napięcie obniżyło się do określonej granicy (dla 12-woltowych kwasowych około 10 V). W przypadku ogniw litowych komplikacja polega na tym, że **akumulatory litowe są bardzo delikatne, poszczególne ogniwa zestawu mają trochę inne parametry i dlatego nie wystarczy pomiar sumarycznego napięcia całego zestawu. Niezbędna jest kontrola napięcia każdego z ogniw**.



Rysunek 2

Podstawowa idea ochrony baterii ogniw litowych jest więc dokładnie taka sama, jak w przypadku pojedynczego ogniw: nadmierne rozładowanie dowolnego z ogniw składowych ma powodować odłą-



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



Zakup przekaźników „rtęciowych”

W ZE 5/2026, w artykule Przekaźniki ze stykami zwilżanymi rtęcią opisałem te bardzo interesujące elementy, których nie wolno stosować w sprzęcie powszechnego użytku, ale które nadal można wykorzystywać do celów specjalnych, m.in. „na potrzeby własne”. Poniżej garść informacji o możliwościach zakupu.

Problem z napięciem nominalnym

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

Dobra wiadomość jest taka, że na rynku nadal jest sporo przekaźników – kontaktronów rtęciowych. **W chwili pisania artykułu niektóre można było kupić w śmiesznie niskich cenach, kilku złotych.** Szczegóły za chwilę. Jednak ogólnie biorąc, na aukcjach i znanych portalach przekaźniki rtęciowe są bardzo drogie. Owszem, warto na przykład w wyszukiwarce e-Bay wpisać: **mercury relay**,

Wyszukiwanie przekaźników rtęciowych

i mniej dolarów, głównie od dostawców z krajów „mniej godnych zaufania”, na przykład z Bułgarii: <https://www.ebay.pl/itm/175660794496>.

Przekaźniki – kontaktrony rtęciowe są też dostępne na krajowych platformach. Ja niedawno kupiłem cztery przekaźniki HG1002T po umiarkowanej cenie około 25 złotych za sztukę. A mniej więcej dwa lata temu kupowałem 12-woltowe przekaźniki HGRM51111 po około 4 złote za sztukę. Widać je na powyższej fotografii wstępnej

Do niedawna można było kupić (używane) podobne

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



Elektroenergetyka – przesył i dystrybucja, część 4

W tym odcinku zajmiemy się obiektami, które każdy z nas mija na co dzień. W terenach miejskich są one może bardziej ukryte niż na wsiach, ale dzięki nim każdy z nas ma w swoim domu energię elektryczną. Temat na dziś to stacje transformatorowe średniego napięcia.

Stacje transformatorowe SN/nn
Stacje słupowe
Stacje wewnętrzne

Rozdzielnia średniego napięcia
Rozdzielnia niskiego napięcia
Projektowanie stacji

Stacje transformatorowe SN/nn

Stacje transformatorowe tego typu przekształcają energię elektryczną przesyłaną na napięciu średnim, na napięcie niskie, które jest przesyłane do większości odbiorców. Stacje te mogą być w wykonaniu dwojakim – jako wewnętrzne i napowietrzne. Stacje napowietrzne to stacje słupowe, zaś wewnętrzne to rozwiązania kontenerowe lub w formie murowanego budynku. Budynek może być wieżowy lub parterowy.

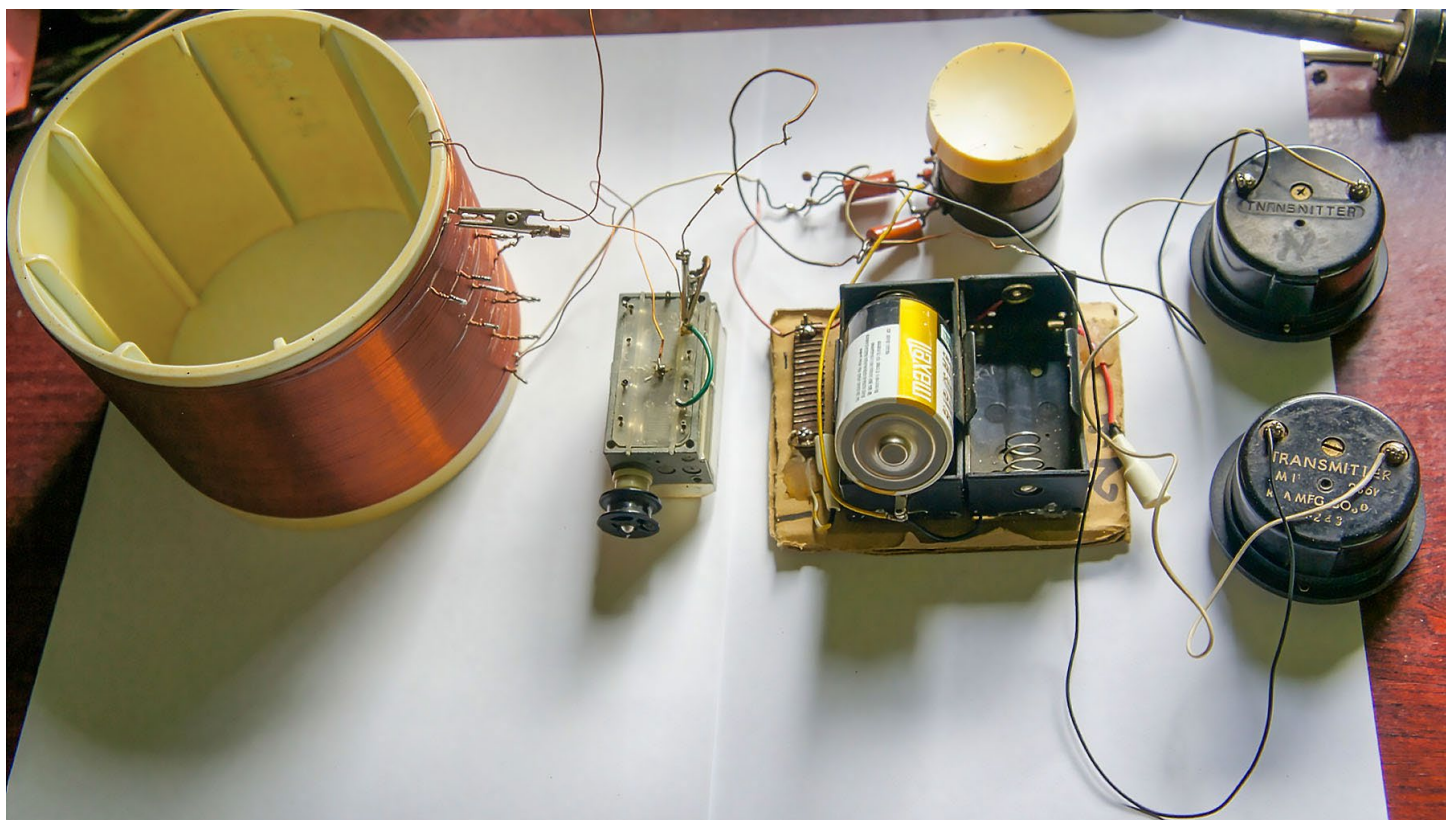
Stacje słupowe

Stacje słupowe budowane w czasach PRL, to zazwyczaj



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





Odbiornik detektorowy i nie tylko detektorowy

W artykule przedstawione są praktyczne, bardzo pożyteczne informacje na temat odbiorników detektorowych, a co jeszcze bardziej interesujące – są informacje o wykorzystaniu diody tunelowej w roli jednocześnie i detektora, i wzmacniacza. Jest też wiele konkretnych, praktycznych porad z życia wziętych.

Trochę mojej historii

Antena

Wykorzystanie diody tunelowej

Słuchawki

Moje działania i plany

Nazywam się **Krzysztof Kleczyński**, jestem inżynierem elektronikiem po Politechnice Warszawskiej. Obecnie na emeryturze – mam 78 lat. Mieszkam koło Rawy Mazowieckiej, a do anteny programu pierwszego Polskiego Radia pod Solcem Kujawskim mam w prostej linii około 210 kilometrów.

Trochę mojej historii

Elektroniką zainteresowałem się już w szkole podstawowej. Był rok chyba 1959 albo 1960. i w tej mojej dziurze, Rawie nie można było dostać części radiowych. Dioda germanowa była marzeniem, a kryształki były dawno wyrzucone na śmieci. Do-

stępny był tylko kondensator zmienny z demobilu. W końcu wyprosiłem od przewijacza silników drut nawojowy i zrobiłem cewkę. Jakiś kryształek gdzieś załatwiłem i słuchawki. Mieszkalem w domu, gdzie nie można było założyć anteny, więc chciałem się podłączyć do sieci przez kondensator i oczywiście był ogień, dym i na tym się skończyło. Potem były tranzystory (1962 rok), więc odbiornik kryształkowy był już „be”. Po wielu latach przeniosłem się z Rawy na wieś, wybudowałem dom i było już miejsce na prawdziwą antenę. Ma ona około trzydzieści metrów długości i jest rozciągnięta od wschodu do zachodu, tak że odbieram najlepiej z południa i północy.

Wysokość zamocowania to około 7 metrów. Uziemienie zakopane głęboko, dwa wiadra ocynkowane, doprowadzenie 4 mm średnicy, w izolacji. Ważne by miejsce połączenia drutu z wiadrem okleić lepka plasteliną aby nie dopuścić tam wilgoci. Cewka radia detektorowego nawinięta jest na dziesięciocentymetrowym odcinku rury PCV o średnicy jakieś 130 mm. Cewka nawinięta jest drutem w emalii 0,6 mm. Marginesy na początku i końcu po 10 mm. Zrobiłem 8 odczepów, równo rozmieszczonych, aby można było eksperymentować z dopasowaniem anteny i diody.

Jako detektor – dioda germanowa, kryształek, dioda krzemowa – wszystko gra, nawet stare kombinerki i żyłka z ołówkiem zaostrzonym – to nie jest krytyczne.

Teraz najlepsze: ja z racji wieku jestem już porządnie głuchy i zwykłe słuchawki są dla mnie za słabe – na słuchawkach takich 2 kiloomy, normalnych, coś tam słyhać, ale słabo. Swego czasu trafiłem w necie na stronę **Crystal Radio Stay Tuned**. To była dosłownie encyklopedia na temat odbiorników detektorowych – od A do Z. Były tam opisane słuchawki, tzw. Sound Power, zbudowane podobnie jak stary głośnik kotwiczny albo głośnik od tzw. hitlerka – taki z papierowym koszem. Takie słuchawki są o wiele, wiele czulsze od normalnych, ale u nas nie były produkowane.

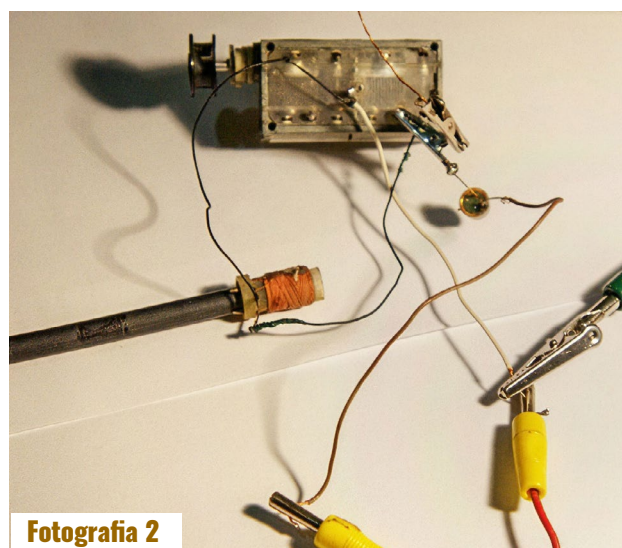
Antena

Fotografia 1 przedstawia moją antenę rozciągniętą od nowego domu do starego. Linka jest w izolacji, więc nawet gdy jakaś gałąź jej dotknie, to nie ma to większego wpływu na jakość odbioru. Antena musi być, ale nie musi być koniecznie taka. Kiedyś miałem antenę **DWZO**, czyli Drut Wyrzucony Za Okno – było tego około piętnaście metrów drutu nawojowego 0,5 mm w emalii, rozciągniętego od okna w domu do gruszki na podwórku i też działało zupełnie dobrze. Uziemienie do odbiornika detektorowego musi być dobre, np. w blokach może być wykorzystany do tego kaloryfer. Wszystko zależy od odległości od nadajnika i chęci eksperymentowania. Kiedyś, przed laty, mój nieżyjący już wujek z Łodzi opowiedział mi, że on, przed wojną, jako młody chłopak zakładał na uszy słuchawki radiowe i dotykał jedną końcówką przewodu słuchawkowego do zlewu w kuchni i słyszał w słuchawkach muzykę. Zlew był emaliowany ale w miejscu gdzie on dotykał, emalia była uszkodzona i metal był zardzewiały albo pokryty takim ciemnym nalotem.

Wtedy nie za bardzo mu wierzyłem ale teraz wierzę

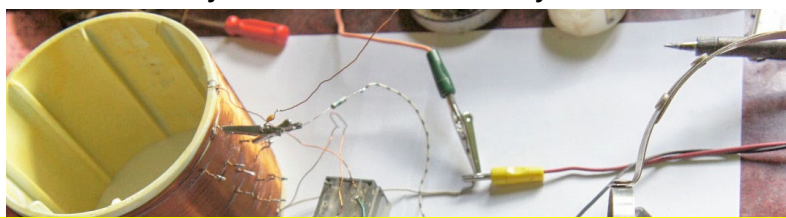


Fotografia 1



Fotografia 2

Fotografia 2 przedstawia odbiornik z cewką z anteny ferrytowej od radia tranzystorowego. Mamy tutaj dodatkową możliwość strojenia rdzeniem ferrytowym cewki. Cewka jest długofalowa i ma indukcyjność około dwóch milihenrów. Na **fotografii 3** cewka jest tradycyjna, ma 138 zwojów drutu w emalii o średnicy 0,5 mm.



Fotografia 3

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



Co to jest ładowanie kondensatora?

Podstawowe zależności w kondensatorze

(Nie)podobieństwo kondensatora i akumulatora

Źródłem właściwości kondensatora jest (nadal bardzo tajemnicze) pole elektryczne. Nie tylko początkujący mają kłopot ze zrozumieniem szczegółów dotyczących pola elektrycznego, jego wektorowego natężenia, jego skalarnego potencjału, konsekwencji występowania i ładunków dodatnich, i ujemnych, dlatego ***na początkowych etapach edukacji powszechnie wykorzystuje się modele, w szczególności modele hydrauliczne***. Jak pisałem we wcześniejszych artykułach, wykorzystuje się co najmniej dwie analogie hydrauliczne – wodne: ciśnieniową (poziomą) oraz wysokościową (pionową). Obie są bardzo mocno niedoskonałe i obie mogą wprowadzić w błąd.

Tak, ale w miarę przystępnie pokazują niektóre ważne aspekty zagadnienia, dlatego pomimo wad warto je wykorzystać. W obu odpowiednikiem napięcia elektrycznego jest ciśnienie wody, a odpowiednikiem prądu elektrycznego – przepływ wody.

W literaturze można znaleźć modele hydrauliczne kondensatora, gdzie kluczową rolę odgrywa elastyczna membrana. Przykład pustej w środku kuli przedzielonej gumową membraną, z **rysunku 1** pochodzi z interesującej strony, którą prowadzi W. Beaty (<http://amasci.com/emotor/cap1.html>). Ten model prawidłowo pokazuje, że przez kondensator nie może płynąć prąd stały. Pokazuje, jak należy rozumieć ładowanie i rozładowanie. Wielkość kuli i jej pojemność słusznie kojarzy się z pojemnością kondensatora. Oprócz tych zalet, model ten ma wiele poważnych ograniczeń, wad i może wprowadzić w błąd.

Pokrewny model pokazany jest na **rysunku 2**. Kondensator przedstawiony jest jako tłok poruszający się wewnątrz cylindra. Ma zalety poprzedniego, a obecność sprężyn wskazuje na magazynowanie w nich energii.

Dużo gorszy hydrauliczny model kondensatora (jako animowany GIF) można łatwo znaleźć w Wikipedii – **rysunek 3**. To elastyczna membrana zamontowana w rurze. Podstawowy problem, że tu nigdzie nie widać pojemności. To bardzo kiepski model.

Żaden z tych modeli nie ilustruje w przystępny sposób najważniejszych kwestii: związku prądu, napięcia i energii. Do tego zdecydowanie lepsza jest „wysokościowa” analogia kondensatora, która oczywiście niestety, też ma liczne wady.

W analogii „wysokościowej” mamy dobry pod pewnymi względami model kondensatora w postaci pojemnika na wodę o kształcie pionowej rury, wysokiego walca – **rysunek 4**. Model ten też ma bardzo poważne wady: nie widać tu, iż przez kondensator nie może płynąć prąd stały. U dołu zbiornika narysowałem dwa dopływy/odpływy, ale wystarczyłby jeden.

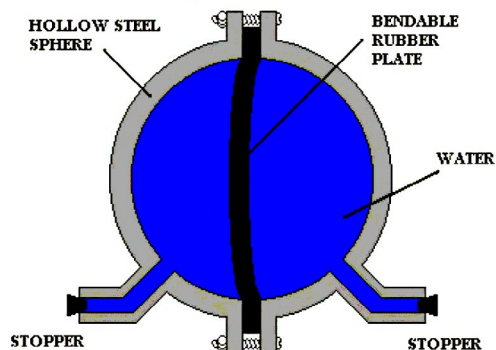
W tej analogii „wysokościowej” napięcie między okładkami kondensatora, czyli „ciśnienie elektryczne”, reprezentowane jest oczywiście przez ciśnienie wody na dole zbiornika, wyznaczone przez wysokość słupa wody. Czyli napięcie na okładkach kondensatora reprezentuje tu wysokość wody w rurze.

Pojemność i zdolność magazynowania energii

Zacznijmy od „oczywistej oczywistości”: sam fakt, że mamy kondensator o jakiejś pojemności wcale

amasci.com/emotor/cap1.html

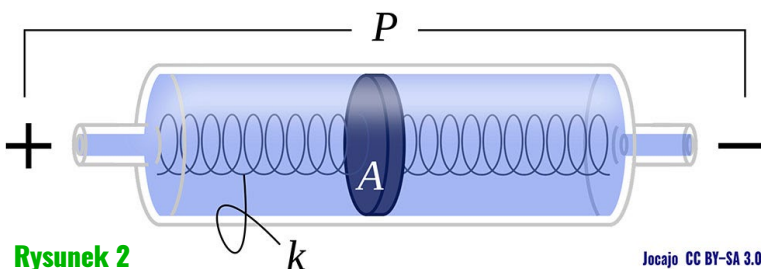
I have beefs with textbook explanations of capacitors too. "Capacitors store charge?" No! Flat out wrong! Wait and hear me out, I'm not insane. ;)



CAPACITOR
(water analogy)

When we "charge" a conventional metal-plate capacitor, the power supply pushes electrons into

Rysunek 1

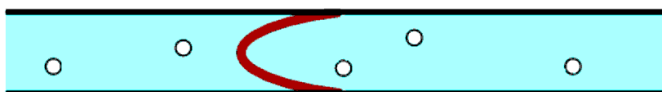


Rysunek 2

Jocajo CC BY-SA 3.0

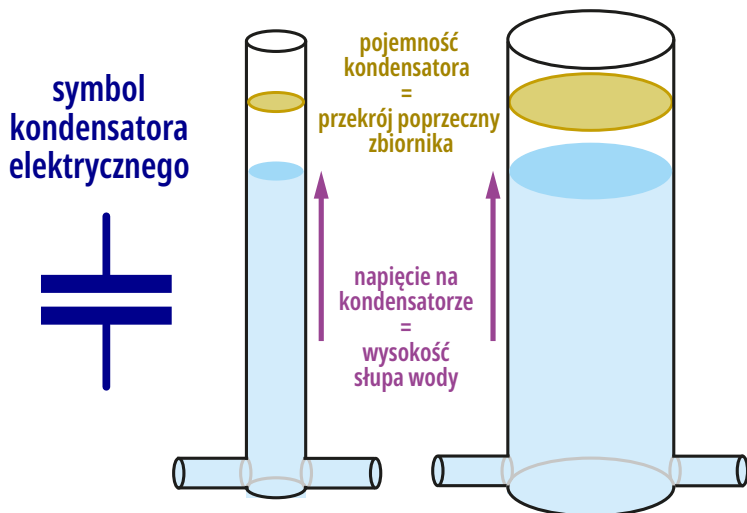
<https://commons.wikimedia.org/wiki/File:CapacitorHydraulicAnalogyAnimation.gif>

File:CapacitorHydraulicAnalogyAnimation.gif

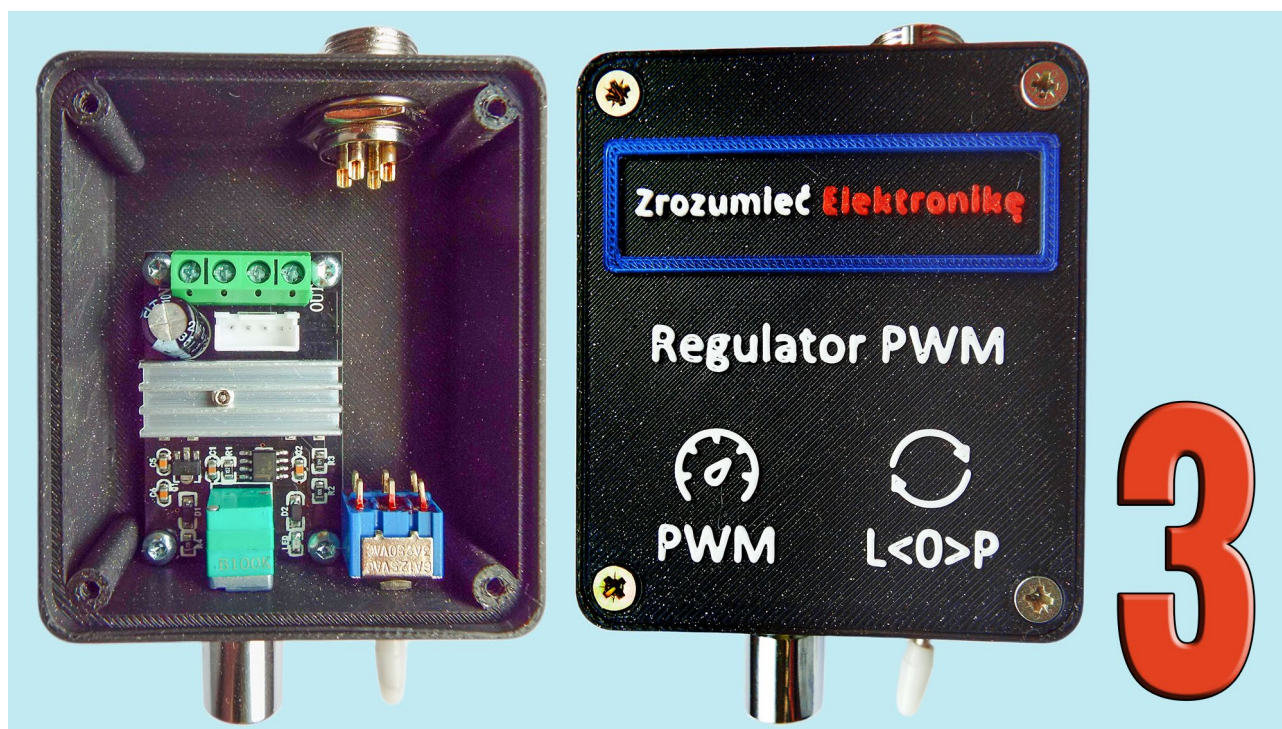


MIME type: image/gif, looped, 8 frames, 3.0 s)

Rysunek 3



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



Druk 3D – wykorzystanie kreatorów, część 3

W dwóch poprzednich częściach artykułu zostało szczegółowo przedstawione, jak za pomocą kreatora samodzielnie zrealizować obudowę do konkretnej płytki drukowanej. Trzecia część pokazuje jak utworzyć do niej pokrywę z własnymi napisami, ikonami i elementami graficznymi.

Płyta czołowa z drukarki 3D

Podział obudowy

Tworzenie, dodawanie i kolorowanie logo

Dodawanie napisów

Ikony i dodawanie ikon

Dodanie tekstu

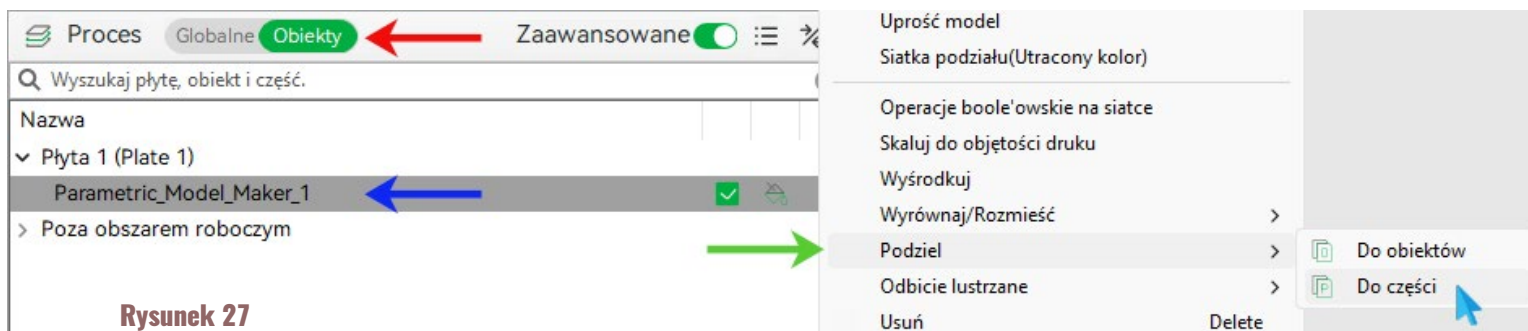
Podsumowanie

Płyta czołowa z drukarki 3D

W poprzednich częściach opisu projektowania obudowy za pomocą kreatora, poruszyłem zagadnienie jakości druku napisów na pokrywie zaprojektowanej obudowy. Na fotografii 22 z drugiej części artykułu widać było, że napisy z kreatora obudów są słabej jakości. Napisy z programu Bambu Studio, widoczne na fotografii 23, też nie wyglądają dużo lepiej. Postanowiłem więc zaprojektować zupełnie nową płytę czołową do opisywanej obudowy. Płytę tę możemy zobaczyć na fotografii tytułowej oraz na fotografiach 24 i 26 z drugiej części artykułu.

Podział obudowy

Chcąc utworzyć indywidualny projekt pokrywy do naszej obudowy, musimy zaprojektować w serwisie <https://makerworld.com> i wyeksportować wybraną obudowę bez napisów na jej pokrywie. Projekt obudowy wczytujemy do programu Bambu Studio narzędziem z ikoną sześciangu ze znacznikiem + w obwódce (Ctrl+I), wskazaną żółtą strzałką na rysunku 32 (w dalszej treści artykułu). Aby móc edytować naszą obudowę trzeba w panelu **Proces** włączyć przełącznik **Obiekty**, wskazany czerwoną strzałką na **rysunku 27** (na następnej stronie).



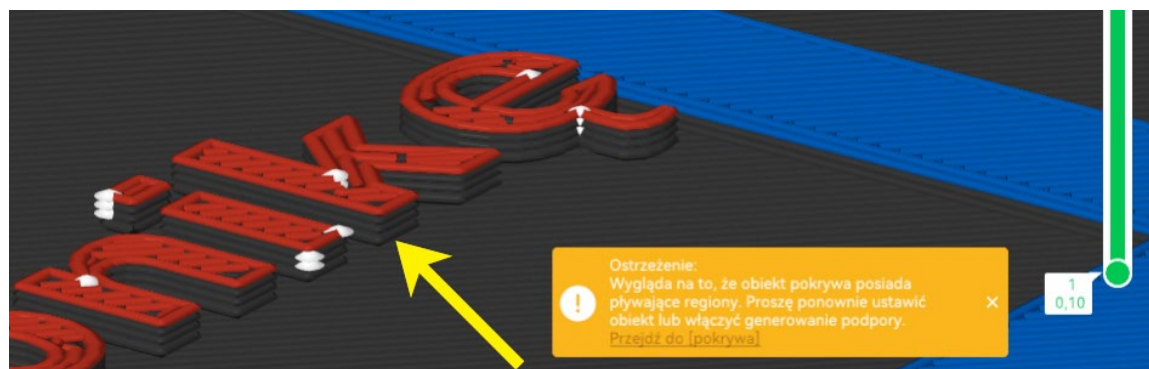
Rysunek 27

Po wczytaniu okazuje się, że nasz projekt to jeden obiekt o nazwie **Parametric_Model_Maker_1**, jak wskazuje to niebieska strzałka na rysunku 27. My natomiast potrzebujemy oddzielić pokrywę obudowy do edycji jako odrębny obiekt. Klikamy prawym klawiszem myszki na nasz projekt i z menu kontekstowego myszki wybieramy opcję **Podziel** → **Do części**, jak wskazuje zielona strzałka. Po wykonaniu tej operacji, jak widzimy na **rysunku 28**, nasz projekt został podzielony na dwa obiekty o nazwach **Parametric_Model_Maker_1_1** i **Parametric_Model_Maker_1_2**, widocznych na **rysunku 28**. Czerwona strzałka pokazuje korpus naszej obudowy, a zielona jej pokrywę. Klikając na te obiekty myszką możemy zaobserwować na podglądzie obudowy, że osobno zaznaczana jest obudowa i osobno jej pokrywa. Usuwany zaznaczony korpus obudowy klawiszem **Del** i zapisujemy (**Ctrl+Shift+S**) projekt pokrywy w nowym pliku, wybierając z menu **Plik** → **Zapisz projekt jako...** Korzystając opcji **Cofnij** (**Ctrl+Z**) możemy cofnąć usunięcie obudowy i usunąć pokrywę, tym razem zapisując w osobnym pliku korpus obudowy. W ten sposób mamy rozdzielone i zapisane w odrębnych plikach obudowę i jej pokrywę. Po usunięciu korpusu obudowy należy zaznaczyć **Parametric_Model_Maker_1_2** i wcisnąć klawisz **F2** z klawiatury. Wówczas będziemy mogli zmienić domyślną nazwę **Parametric_Model_Maker_1_2** na odpowiadającą danemu obiektowi. Ja zmieniłem u siebie tę nazwę na **pokrywa obudowy**. Na dalszym etapie projektu, nadając poszczególne obiektem własne nazwy można łatwiej się

Rysunek 28

Tworzenie, dodawanie i kolorowanie logo

Pierwsze próby utworzenia logo za pomocą narzędzi wbudowanych w Bambu Studio zakończyły się u mnie niepowodzeniem. Stworzyłem sześciang, który miał być ramką i umieściłem go na górnej płaszczyźnie pokrywy. Zrobiłem drugi sześciang, jako tak zwaną część ujemną, która utworzyła wgłębienie w pierwszym sześcianie i wraz z pierwszym sześcianiem miała tworzyć ramkę logo. Wewnątrz powstałej ramki dodałem tekst „Zrozumieć Elektronikę”. Pomimo wizualnie poprawnego projektu, Bambu Studio wyświetlało ostrzeżenie „**pływające regiony**”, widoczne na **rysunku 29**. Prawdopodobnie część ujemna, mająca utworzyć wgłębienie w ramce logo, powodowała, że później dodany napis był traktowany przez Bambu Studio jako zawieszony w powietrzu. Dodatkowo, po pokolorowaniu tekstu, na podglądzie wydruku, pokryta kolorem była tylko wierzchnia warstwa napisu, co wskazuje żółta strzałka z **rysunku 29**.



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

ZROZUMIEĆ **E**LEKTRONIKĘ z Piotrem Góreckim

ZE 9/2026

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.**