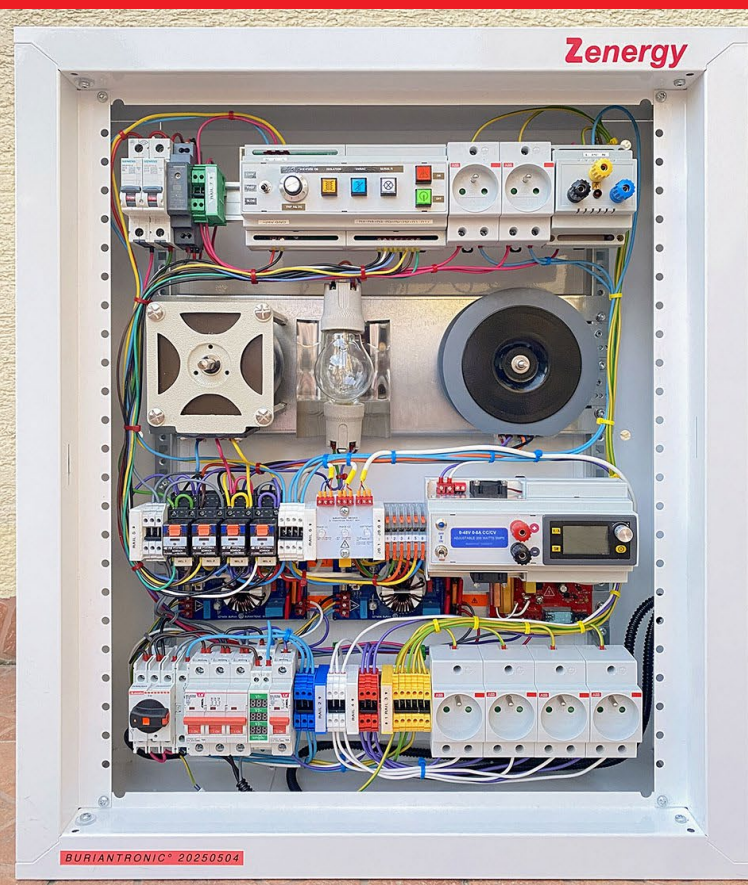
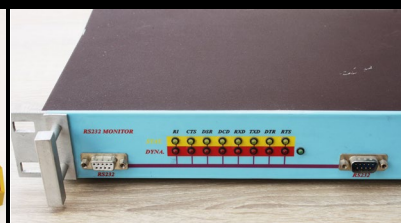
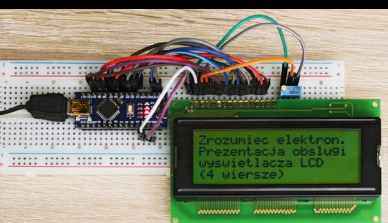
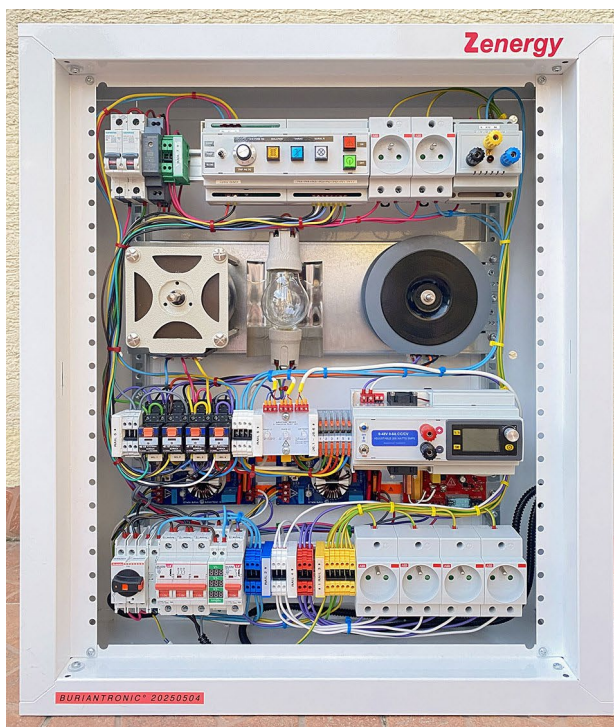




Dr Frankenstein Project – stanowisko energetyczne

- Mały liniowy zasilacz: Dobór potencjometru • Generowanie dźwięków • Amatorskie wzorce pojemności
- Czterowierszowy moduł wyświetlacza LCD • Monitor interfejsu RS232 • Mierniki LCR typu XJW01 i pokrewne
- Akumulatory kwasowe czy akumulatory zasadowe • Schematy wewnętrzne układów scalonych
- Zapisane w pamięci – historia i typy pamięci • Moja droga do okiełznania szybkich sygnałów





Dr Frankenstein Project – stanowisko energetyczne (1)

Masz dosyć zasilania byle czym i byle jak testowanych urządzeń, pracujących pod napięciem sieci 230 V? Drażni Cię kłębowisko kabli i bylejakość łączonych naprędce żarówek, chroniących badany sprzęt? Czy zawsze chciałeś mieć autotransformator lub transformator bezpieczeństwa? Jeśli tak, ten artykuł jest dla Ciebie!

Czym jest tablica warsztatowa i skąd taki pomysł?

Kilka istotnych uwag i ograniczeń.

Konstrukcja mechaniczna i wybór obudowy

Konstrukcja elektryczna

Schemat ideowy i działanie tablicy

Moduł zabezpieczenia przeciwprzepięciowego i filtra EMI

To jest pierwsza część obszernego artykułu przedstawiającego realizację warsztatowego stanowiska energetycznego, przeznaczonego do pracowni elektronika. Elektronika, który ma w pracowni mnóstwo urządzeń zasilanych z sieci, ale także przeprowadza rozmaite eksperymenty z układami i urządzeniami dołączonymi wprost do sieci energetycznej.

Czym jest tablica warsztatowa i skąd taki pomysł?

Od dobrych dwudziestu lat chodziła za mną myśl, aby uporządkować kwestie podłączania do sieci 230 V zarówno budowanych przeze mnie urządzeń, jak i posiadanego sprzętu warsztatowego (choć ambitniej brzmiałoby „laboratoryjnego”)

Ponieważ w najbliższym czasie będę przeprowadzał się do nowego domu, gdzie po raz pierwszy w życiu udało mi się wygospodarować pomieszczenie przeznaczone tylko i wyłącznie na własną pracownię, uznałem, że jest to idealny moment, by wyposażać ją w zamontowane na stałe, umocowane na ścianie nad blatem warsztatowym, stanowisko energetyczne. Pamiętam, że podobne rozwiązania widywałem w pracowniach fizycznych szkół, do których uczęszczałem. Jako że edukację szkolną zakończyłem w ubiegłym wieku, stanowiska jakie pamiętam, były wykonane w duchu minionej epoki: płyty czołowe z tekstolitu, wielkie mierniki wychyłowe, porcelanowe bezpieczniki topikowe, rozłączniki dźwigniowe – wszystko tyleż solidne, co toporne.

Ale miało swój urok! Nadawało szkolnej sali wygląd laboratorium, wspomnianego niedawno przez Piotra Góreckiego, doktora Frankensteina i przede wszystkim – niezawodnie działało.

Postanowiłem zaprojektować coś mniejszego, bardziej nowoczesnego, ale z zachowaniem kilku elementów nostalgicznych, a na własny użytek ochrzciłem całość nazwą „Dr Frankenstein Project” :)

Paradoksalnie, w sensowności projektu utwierdziła mnie historia związana bezpośrednio z jego konstrukcją. Eksperymentując z jednym z użytych w nim, zasilanym bezpośrednio z sieci modułów, używałem najpowszechniejszej techniki lutowania w pająku, płytki stykowej (tak, wiem, płytka stykowa pracująca pod napięciem sieci to zły pomysł i nie należy tego robić) oraz zwykłego kabla z wtyczką, wkładanej po omacku do listwy „wszystko zasilającej” przykręconej pod biurkiem. Każda manipulacja na układzie wymagająca odłączenia zasilania, sprowadzała się do wyciągnięcia i ponownego wetknięcia wtyczki do przedłużacza (tak, niech pierwszy bez winy rzuci kamieniem).

Za którymś razem stwierdziłem, że najwyraźniej popsuta mi się stacja lutownicza, bo nie mogę przylutować jakiegoś elementu. Już po sekundzie, z pewnym przerażeniem odkryłem, że stacja nie działa dlatego, że nie jest zasilana. Pomyliłem podobne do siebie wtyczki: lutownicę odłączyłem od sieci, zasilany układ przez cały czas był do niej podłączony, a ja dłuubałem w nim, siedząc z bosymi stopami opartymi o podłogę.

Opisywane urządzenie pozwala radykalnie ograniczyć ryzyko takich potencjalnie katastrofalnych pomyłek, ale jego możliwości są znacznie, znacznie większe. Dla pełnej funkcjonalności wymaga podłączenia do sieci trójfazowej, ale w razie jej braku można zastosować zasilanie jednofazowe, tracąc w zasadzie tylko jedną z kilku funkcji. Zawiera cztery odrębne tory zasilające o maksymalnej wydajności prądowej 3×16 A (łącznie dla wszystkich torów), przy czym są one skonfigurowane następująco.

- **Tor 1.** Służy do podłączania urządzeń wymagających zasilania dwu- lub trójfazowego 230/400 V. Jest ze wszystkich najprostszy, bo składa się tylko z wyłącznika nadprądowego 3×16 A oraz gniazda siłowego 3×16 A, zamontowanego na boku obudowy.

- **Tor 2.** Służy do zasilania sprzętu laboratoryjnego i warsztatowego. Zawiera wyłącznik nadprądowy 16 A, filtr przeciwprzepięciowy / EMI z optyczną sygnalizacją stanu, a zakończony jest czterema standardowymi gniazdami szynowymi z uziemieniem.

0–48 V i prądu w zakresie 0–8 A (pod warunkiem nie przekroczenia mocy 200 W)

- **Tor 4.** Jest sercem całego stanowiska i służy do kontrolowanego zasilania budowanych i serwisowanych urządzeń. Umożliwia separację galwaniczną od sieci, płynną regulację napięcia przemiennego przy użyciu autotransformatora (wariaka) w zakresie 0–280 V, płynne, regulowane elektronicznie ograniczenie prądu w zakresie około 0,5–10 A z możliwością zmiany charakterystyki reakcji (szybki / wolny) oraz załączenie szeregowego obciążenia zabezpieczającego w postaci żarówek, z możliwością wyboru ich mocy (40/60/100 W). Wszystkie te funkcje można wygodnie aktywować przełącznikami. Tor ten umożliwia pobór prądu nie przekraczający 10 A, natomiast załączenie funkcji separacji / regulacji napięcia ogranicza go automatycznie do 2,5 A (chwilowo do 3 A). Zakończony jest standardowym gniazdem tablicowym oraz dodatkowo zespołem izolowanych gniazd laboratoryjnych. Napięcie i prąd kontrolowane są wychyłowymi miernikami analogowymi, natomiast od strony wejścia zabezpieczony jest blokiem ochrony przeciwprzepięciowej i filtrem EMI, identycznym, jak zastosowany w torze drugim.

Całość tablicy załączana jest do sieci domowej obrotowym rozłącznikiem separacyjnym, wyposażona jest też w fabryczny woltomierz trójfazowy na szynie DIN. Nie przewidziałem osobnego wyłącznika różnicowoprądowego, ponieważ dedykowany temu urządzeniu RCD został zainstalowany na głównej tablicy rozdzielczej budynku.

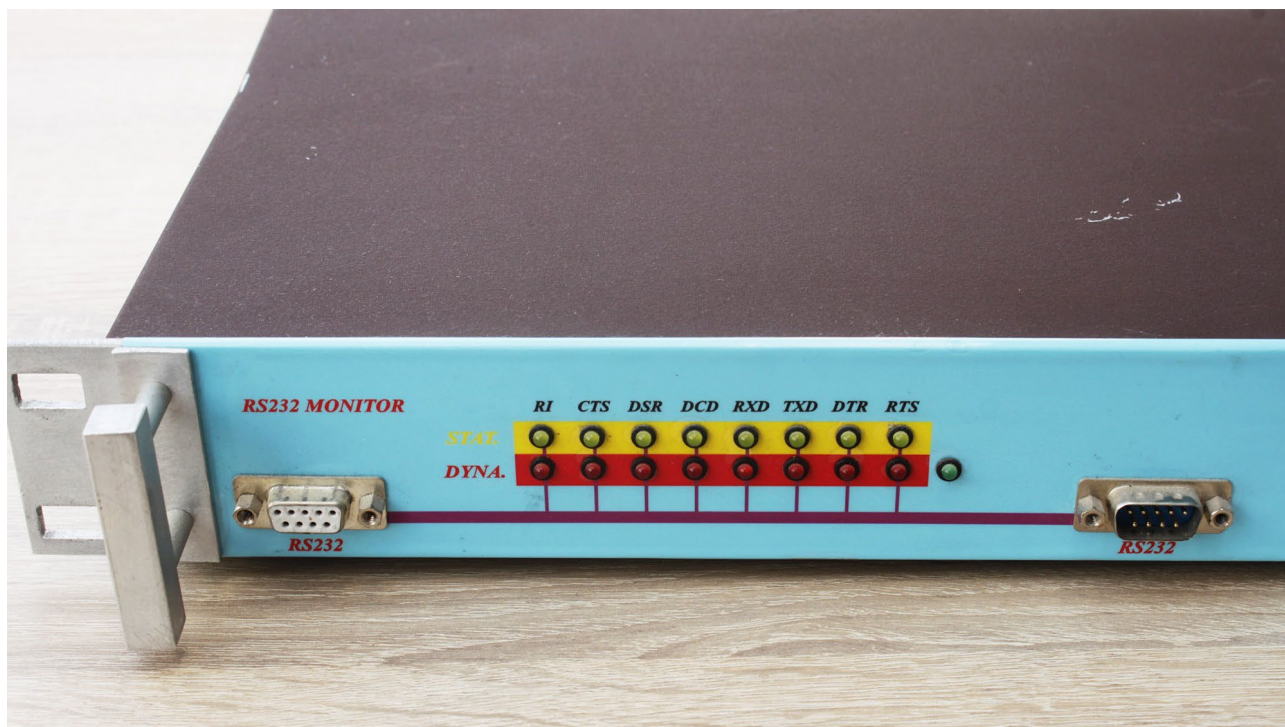
Kilka istotnych uwag i ograniczeń.

Nie spodziewam się, by szczególnie liczne grono Czytelników chciało wykonać opisane urządzenie, zrobiłem je bowiem w odpowiedzi na swoje własne potrzeby, w dodatku nieco na wyrost. Mam jednak nadzieję, że zarówno sama koncepcja budowy, jak i wykorzystane moduły, będą inspiracją do innych projektów, także tych radykalnie mniejszych i prostszych. W każdym jednak przypadku trzeba wziąć pod uwagę kilka czynników, z których najważniejsze jest bezpieczeństwo. **Ten projekt zdecydowanie nie jest dla początkujących!!!**

Nawet średnio zaawansowani elektrycy, o ile nie mieli wcześniej do czynienia ze specyfiką instalacji elektroenergetycznych, powinni podejść do zadania z należytym respektem.

Mamy tu do czynienia nie tylko z groźnym napięciem sieci, ale z jeszcze groźniejszym napięciem

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



Monitor interfejsu RS232

Popularnym wariantem transmisji szeregowej jest interfejs RS232. Był on standardem w komputerach PC ale ustąpił miejsca USB. Jednak nadal jest powszechnie używanym w automatyce przemysłowej oraz absolutnym liderem w układach z mikrokontrolerami (szczególnie w zastosowaniach amatorskich).

Idea monitora interfejsu RS232

Schemat monitora

Opis w VHDL-u

Testy działania

Transmisja szeregowa to podstawowy kanał do komunikacji między mikrokontrolerami. W każdym ze współcześnie występujących układów tego typu znajduje się podzespół, którego zadaniem jest nadzór nad szeregowym przesyłaniem danych. Oprócz samych linii przesyłających dane, często występują dodatkowe sygnały biorące udział z komunikacji. Więcej informacji o samym interfejsie RS232 można znaleźć w artykule *Mikroprocesory i mikrokontrolery – interfejs szeregowy RS232* (ZE 4/2024).

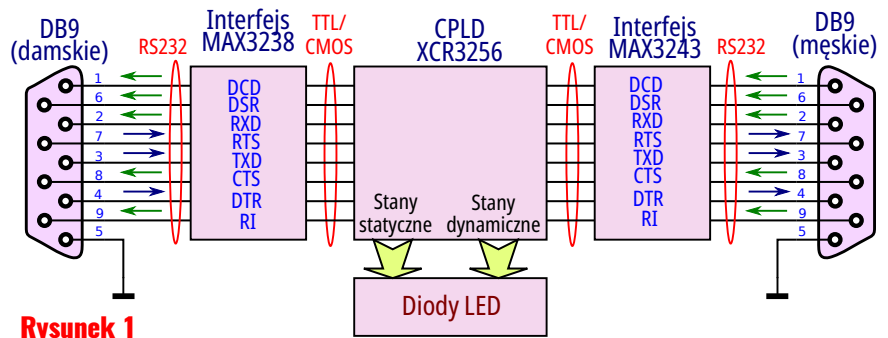
Idea monitora interfejsu RS232

Koncepcja rozwiązania monitora interfejsu RS232 sprowadza się do zbudowania urządzenia, które jest „wpięte” w kabel RS232 i sygnalizuje stany występujące na poszczególnych liniach wchodzących w skład tego interfejsu. Może ono nawet pełnić rolę „kabla” łączącego stronę DTE ze stroną DCE. Jako strony DTE i DCE niekoniecznie należy utożsamiać

komputer i modem, mogą to być dowolne urządzenia wykorzystujące transmisję szeregową. Praktycznie wszystkie współczesne mikrokontrolery oferują możliwość przesyłania danych za pośrednictwem transmisji szeregowej. Wiele z nich oprócz samych linii transmisyjnych zawiera w sobie pełny interfejs szeregowy (jako dodatkowe linie modemowe). Taki monitor jest w stanie zasygnalizować zdarzenia, jakie występują w połączeniu.

Sam monitor nie ingeruje w sygnały występujące w złączu RS232, które jest tutaj reprezentowane jako DSUB9. Po przetworzeniu stanów występujących w liniach RS232 na poziomy logiczne, są one „obserwowane” przez układ CPLD, którego zadaniem jest jedynie sygnalizowanie poprzez diody LED stanów tam występujących i po ponownej konwersji do poziomów „wydanie” dalej. Sam interfejs RS232 jest niesymetryczny (ma trzy sygnały przesyłane w jedną stronę oraz pięć sygnałów przesyłanych w drugą stronę)

i z tego powodu po jednej jest zastosowany układ MAX3238, a po drugiej MAX3243. Występujące diody LED mają za zadanie sygnalizować aktualny stan na poszczególnych stykach interfejsu. Ten zestaw jest określony jako stany statyczne (pokazują to, co jest aktualnie w danej chwili). Ze względu na to, że impulsy mogą być na tyle krótkie, że trudno jest je zaobserwować, dodany jest drugi zestaw LED-ów przeznaczonych do sygnalizacji stanów dynamicznych. Diody te są aktywowane do świecenia przy każdej jakiegokolwiek zmianie stanu na jakiegokolwiek linii i jest ono odpowiednio przedłużone, by można je było zaobserwować. Ideę pokazuje **rysunek 1**.

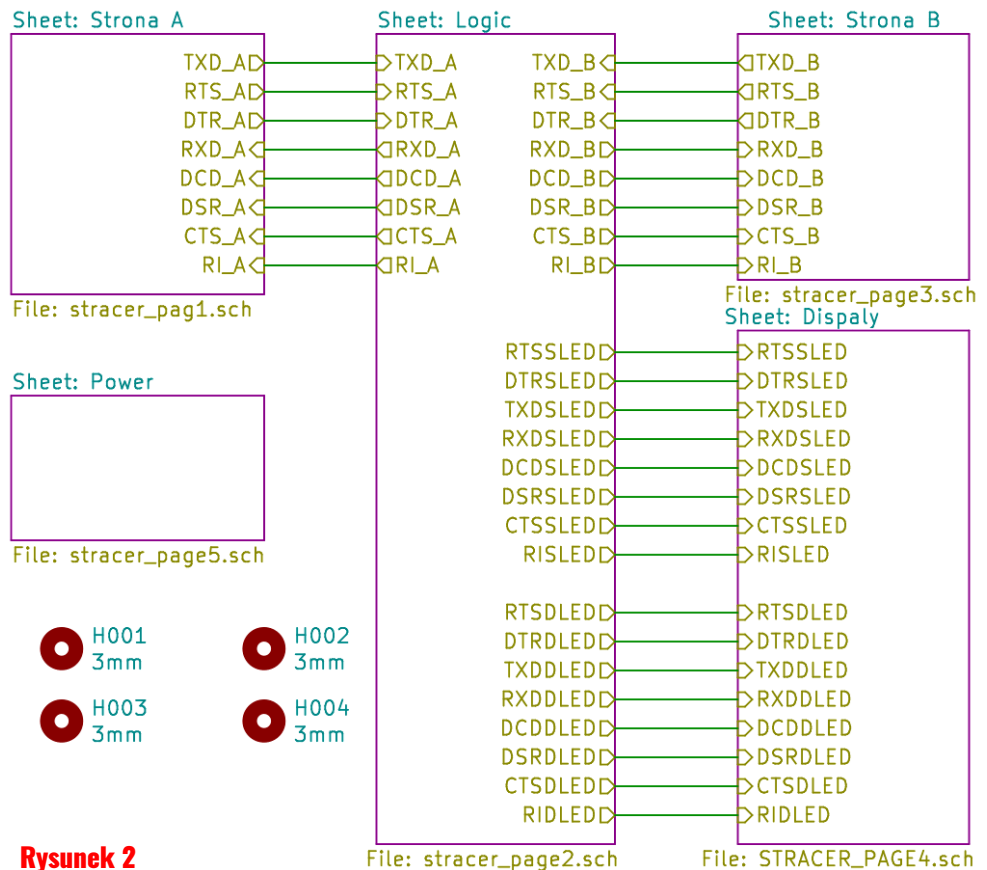


Rysunek 1

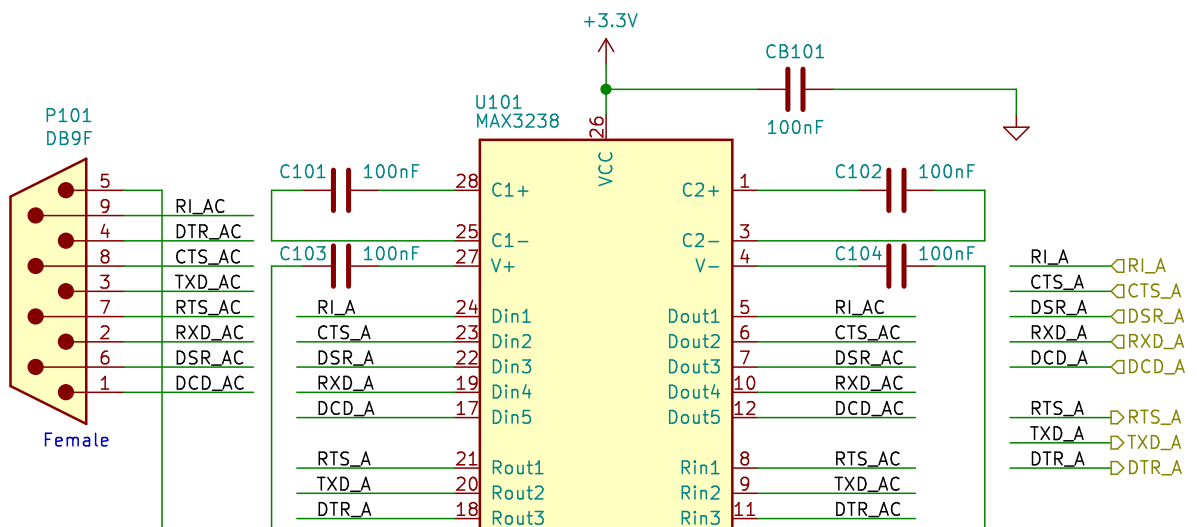
Schemat monitora

W skład monitora wchodzi kilka bloków o określonej funkcjonalności (**rysunek 2**). Głównym jego elementem jest układ logiki programowalnej CPLD z oferty Xilinx (obecnie wchodzącego w skład AMD). Ten układ wymaga zasilacza dającego napięcie 3,3 V, jak również wszystkich sygnałów logicznych o tej amplitudzie. Z tego powodu jest zastosowany układ (MAX3238) przetwarzający sygnały ze standardu RS232 na poziomy logiczne dopuszczalne dla układu CPLD (tu jest pięć linii) oraz w drugą stronę z logiki 3,3 V do wymagań RS232 (trzy linie). Pokazuje to **rysunek 3** i jest to strona „A” (do nazw styków występujących w złączu RS232 dodana jest literka „A”).

Po „przeciwnej” stronie



Rysunek 2



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.

W pełnej wersji dostępnej dla Patronów ten artykuł oczywiście ma więcej stron.



Generowanie dźwięków

Podzespół mikrokontrolera określany jako licznik/zegar ma wiele możliwych trybów pracy. Najczęściej jest używany do odmierzania interwałów czasu, ale nie jest to jego jedyna możliwość. Jak można wykorzystać informacje wynikające z „Mikroprocesorowej oślej łączki”?

Licznik/zegar jako generator fali prostokątnej Środowisko projektu

Kolejka zdarzeń uwarunkowanych czasowo

Mikrokontroler ATMEGA328 dysponuje kilkoma zegarami/licznikami. Wszystkie je można wykorzystać w dowolny sposób, poprzez odpowiednie ich zaprogramowanie (wpis do rejestrów sterujących odpowiednich danych konfiguracyjnych). Jedną z takich możliwości jest wykorzystanie go jako licznika pracującego w trybie CTC (ang. *Clear Timer on Compare*), automatycznego wyzerowania licznika po osiągnięciu przez niego określonego stanu, mówiąc najprościej do generowania sygnału o określonej częstotliwości.

Licznik/zegar jako generator fali prostokątnej

Do tego zadania warto wybrać licznik/zegar numer 1, gdyż jest to licznik 16-bitowy, przez co można uzyskać sygnały o w miarę dobrej częstotliwościowej dokładności (**rysunek 1**). Dysponuje on dwoma niezależnymi kanałami (zestaw A oraz B – w prezen-

towanym programie jest wykorzystany jedynie A). Sam licznik zlicza impulsy pochodzące z systemowego zegara z ewentualnym uwzględnieniem wstępnego podzielnika (preskalera). Stan licznika (jako liczba 16-bitowa) jest porównywany z liczbą zapisaną w odpowiednim rejestrze (również 16-bitową). W sytuacji gdy zaistnieje równość obu liczb, licznik jest wyzerowany oraz istnieje możliwość wykonania przy tej okazji dodatkowych czynności: może być sygnałem odpowiedniego przerwania (ten wariant nie jest wykorzystany) lub działań zmieniających stan odpowiedniego pinu portu (to jest naszym celem). Istotnymi rejestrami w tym przypadku są: *TCCR1A*, *TCCR1B* (konfigurujące pracę zespołu licznikowego 1) oraz *OCR1A* (przechowujący wzorec do porównania), **rysunek 2**.

W tym miejscu warto wspomnieć o możliwości przenoszenia oprogramowania pomiędzy różnymi

modelami mikrokontrolerów AVR (oczywiście na poziomie tekstu programu źródłowego). Rozpatrując migrację z mikrokontrolera ATMEGA328 do przykładowo ATMEGA32, należy zapoznać się ze szczegółami realizacji generacji fali prostokątnej przez docelowy mikrokontroler. Okazuje się, że w tym przypadku budowa i znaczenie poszczególnych bitów w rejestrach *TCCR1A* i *TCCR1B* jest zgodna. Same rejestry mają inną lokalizację w przestrzeni adresowej mikrokontrolera, ale nie stanowi to problemu, kompilacja programu dołączy właściwy plik definiujący zasoby mikrokontrolera, toteż kompilator sam „zauważy” zmianę adresacji. Również znaczenie poszczególnych bitów w rejestrze *TCCR1A* różni się o szczegóły nieistotne w tym przypadku, budowa i znaczenie bitów w rejestrze *TCCR1B* jest identyczna. Bardzo znacząca różnica występuje w lokalizacji wyjścia generowanego przebiegu. W mikrokontrolerze ATMEGA328 wyjście *OC1A* jest umiejscowione w *PORTB.1*, natomiast w mikrokontrolerze ATMEGA32 wyjście *OC1A* znajduje się w *PORTD.5*. Z tego powodu warto uelastyczyć program poprzez możliwości stworzone przez *#define*, gdyż koniecznością jest skonfigurowanie wyprowadzenia *OC1A* jako wyjścia.

Rozpatrzmy prosty program, generujący określoną częstotliwość na tym wyjściu, pokazany w istotnym fragmencie na **listingu 1**. Warto zauważyć, że program po skonfigurowaniu licznika 1 nic nie robi (wszystko „dzieje się” samo):

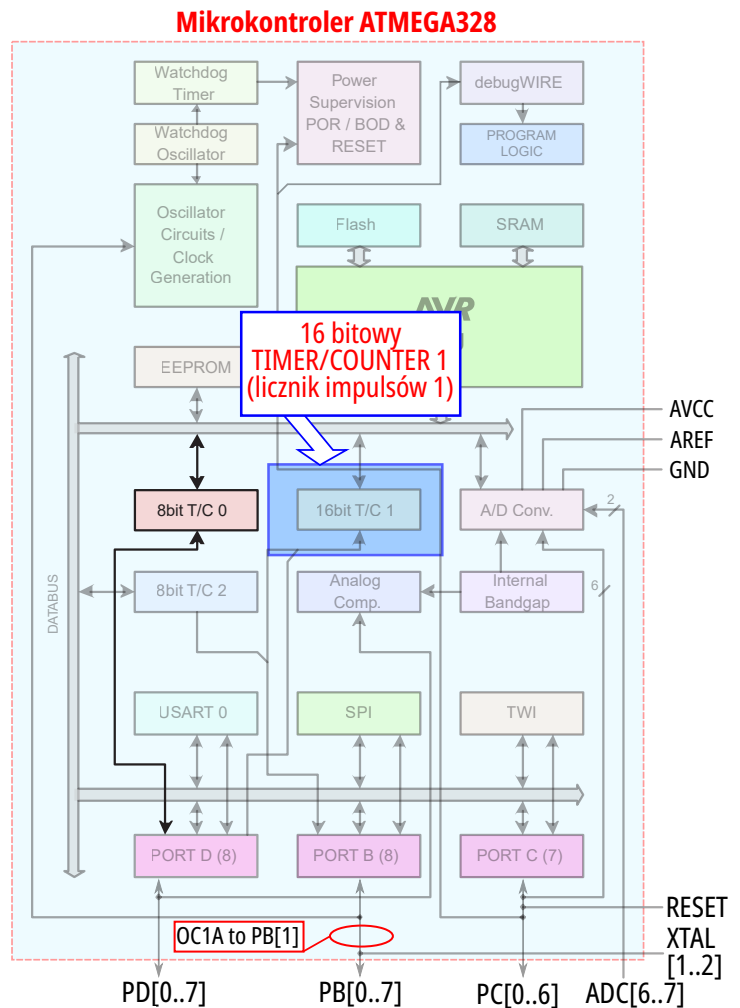
```
#define OCR1AConfigPort    DDRB
#define OCR1AConfigPin    1

#define nop() __asm__ __volatile__ ("nop")

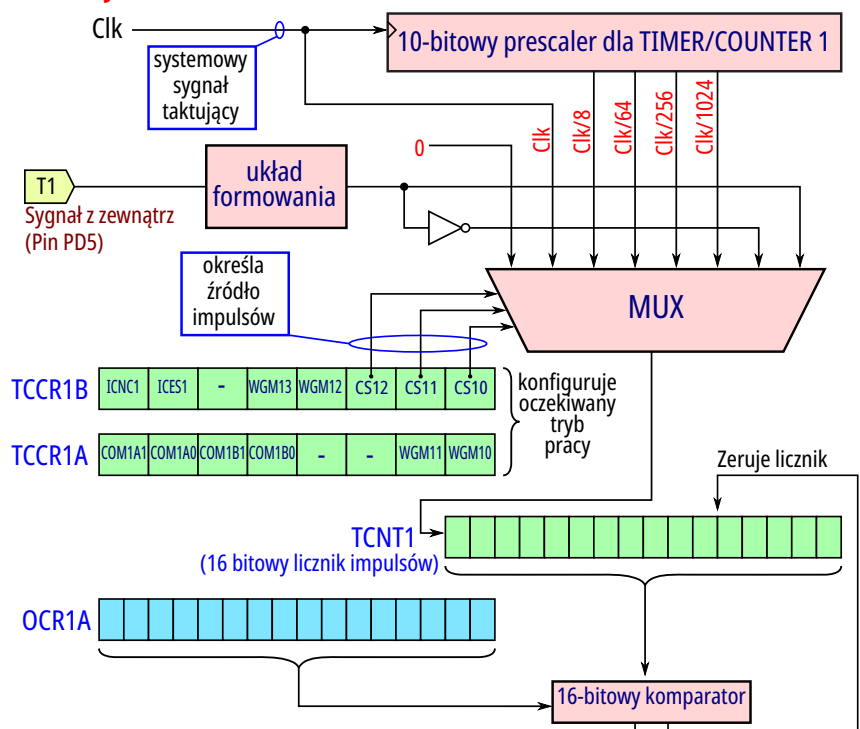
static void HardwareInit ( void )
{
    TCCR1A = ( 1 << COM1A0 ) ;
    TCCR1B = ( 1 << WGM12 ) | ( 1 << CS11 ) ;
    OCR1A = 100 ;
} /* HardwareInit */

static void EnvirInit ( void )
{
    OCR1AConfigPort |= 1 << OCR1AConfigPin ;
} /* EnvirInit */

int main ( void )
{
    HardwareInit ( ) ;
```

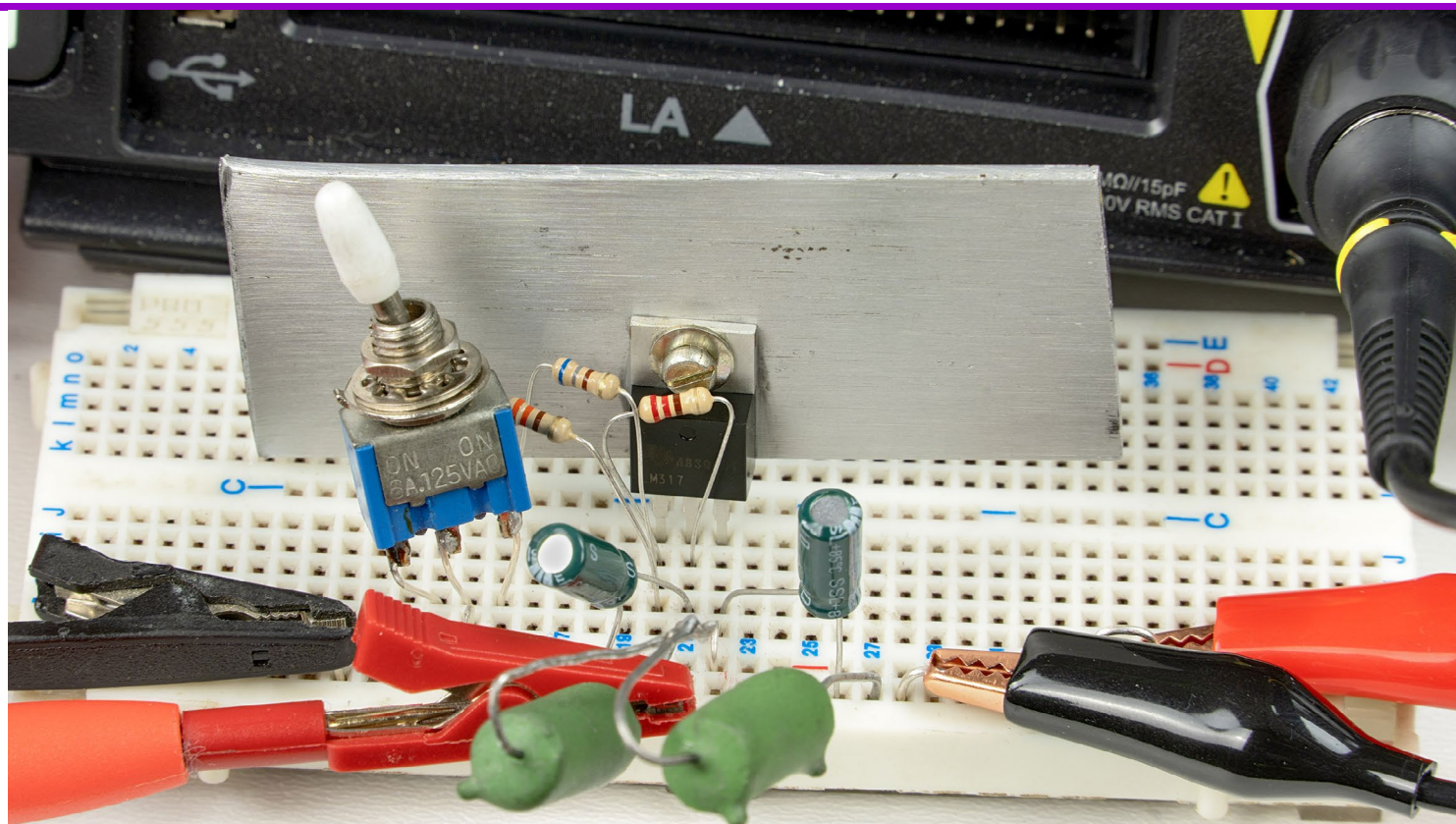


Rysunek 1



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.

W pełnej wersji dostępnej dla Patronów ten artykuł oczywiście ma więcej stron.



Mały liniowy zasilacz: Dobór potencjometru

Oto piąty artykuł serii pokazującej, jak elektronik dla swej ogromnej satysfakcji może zaprojektować i zrealizować niewielki zasilacz liniowy do swojej pracowni. Jeszcze raz podkreślam: liniowy, czyli niewytwarzający zakłóceń, które są zmorą dominujących dziś na rynku zasilaczy impulsowych

Problem trzasków i przerw

Problem prądu i mocy potencjometru

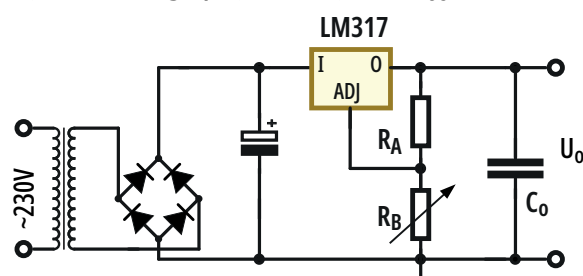
Jaki potencjometr do zasilacza regulowanego?

Minimalny prąd wstępnego obciążenia

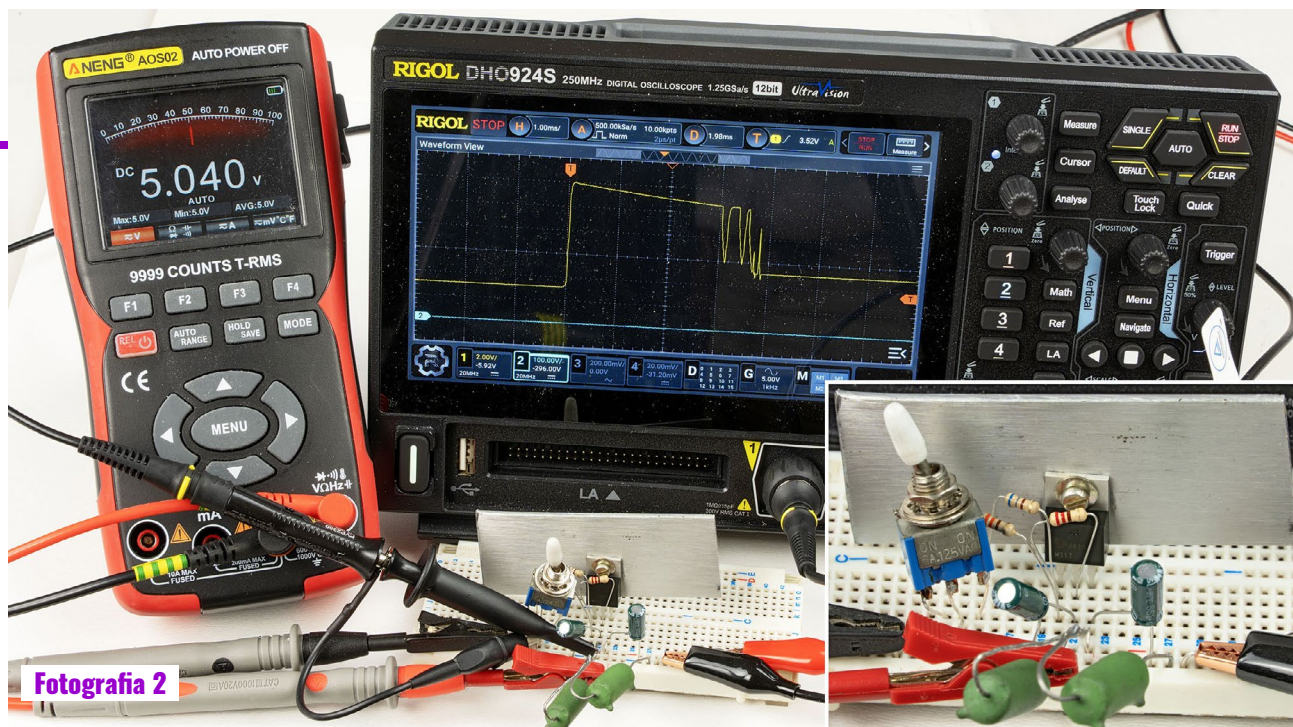
Realizacja praktyczna

We wcześniejszych artykułach tego cyklu przedstawiłem wstępnie właściwości prototypów zasilaczy pokazanych na powyższej fotografii tytułowej oraz omówiłem różne aspekty doboru transformatora, prostownika i filtru. W poprzednim artykule opisałem zasadę regulacji stabilizatora LM317. Wiemy, że zasilacz warsztatowy o płynnie regulowanym napięciu wyjściowym można zrealizować stosując potencjometr w roli R_B według **rysunku 1**. Tak, ale trzeba mieć świadomość pewnych istotnych problemów praktycznych. Jeden to problem trzasków i przerw, drugi to

problem mocy strat wydzielanych w potencjometrze. A trzeci to związek rezystancji potencjometru z kwestią minimalnego prądu obciążenia wyjścia LM317.



Rysunek 1



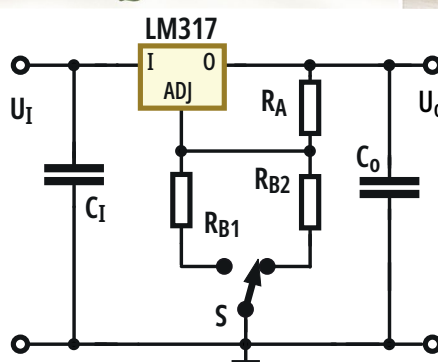
Fotografia 2

Problem trzasków i przerw

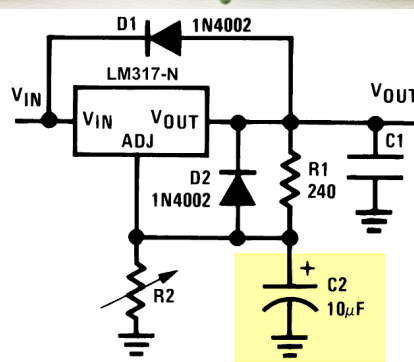
Jeżeli dolny rezystor R_B (lub potencjometr) zostanie odłączony, to napięcie wyjściowe skoczy do maksymalnej wartości, która będzie 1...2 wolty mniejsza niż napięcie na wejściu stabilizatora. Widać to na **fotografii 2**, pokazującej taki skok napięcia przy wykorzystaniu przełącznika w układzie z **rysunku 3**. W tylnej części tego impulsu występują drgania spowodowane wibracją styków przełącznika.

Problem można zredukować stosując zalecany przez producenta (**rysunek 4**) dodatkowy kondensator filtrujący między masą a końcówką ADJ. **Fotografia 5** przedstawia wielkość takiego skoku w układzie z fotografii 2 z kondensatorami o różnej wartości, podczas przełączania z napięcia 3 V na 5 V.

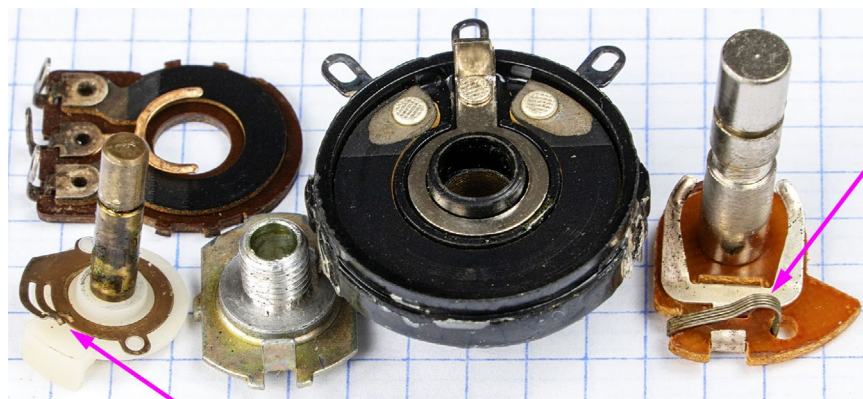
I takie skoki napięcia mogą wystąpić nie tylko w przypadku zastosowania przełącznika, ale też w przypadku użycia potencjometru z wyjątkowo silnie zabrudzo-



Rysunek 3

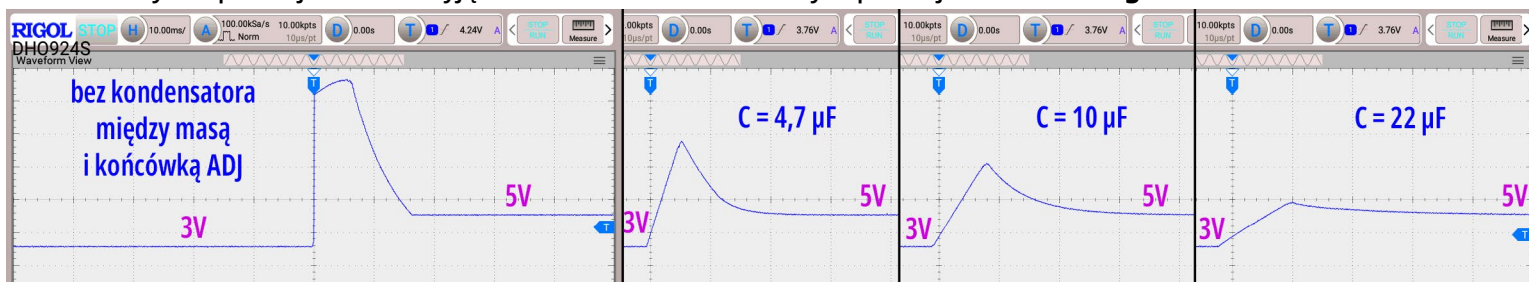


Rysunek 4



Fotografia 6

nym stykiem suwaka. Suwaki potencjometrów mają podwójne lub wielokrotne styki, żeby minimalizować trzaski, co widać na dwóch przykładach starych krajowych potencjometrów – **fotografia 6**.



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



Mierniki LCR typu XJW01 i pokrewne

W poniższym artykule przedstawiam kolejną, bardzo interesującą historię trzeciej rodziny amatorskich mierników pojemności i indukcyjności. Są to przyrządy niewątpliwie amatorskie, natomiast wykorzystują znaną od lat metodę pomiarową, powszechnie stosowaną w profesjonalnych miernikach LCR (RLC).

Dobra koncepcja, słabe rozwiązanie
Automatically balanced bridge method

Od projektu hobbysty do oferty rynkowej

Historia amatorskich mierników XJW01 ma kilka-naście lat. Zaczyna się bowiem w roku 2011, kiedy to pewien chiński hobbysta, nauczyciel pracujący w szkole średniej, przedstawił swoją konstrukcję miernika LCR, czy jak kto woli miernika RLC. Według deklaracji dokładność pomiarów sięgała 0,3%.

Jego opracowanie wzbudziło duże zainteresowanie wśród chińskich elektroników amatorów. Co ciekawe, głównie wśród tych niezamożnych i mniej zaawansowanych. Dla nich bardzo kusząca była perspektywa samodzielnej realizacji niedrogiego przyrządu o znakomitej dokładności 0,3%.

Wątek na forum **crystalradio.cn** zaczyna się w grudniu 2011 i dziś ma 385 dość długich stron:

<http://www.crystalradio.cn/thread-231933-1-1.html>

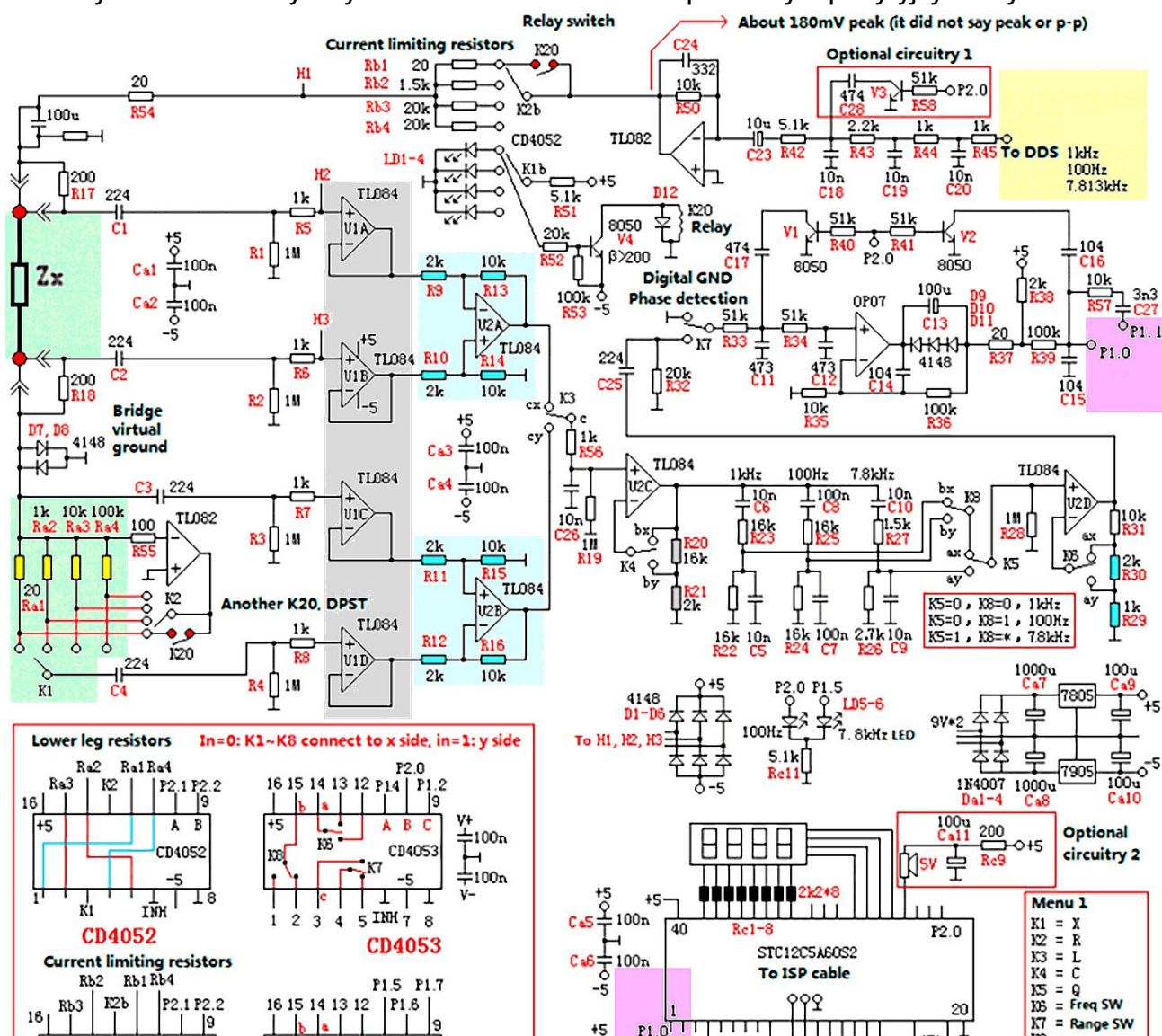
Wpisów jest tam nieprzeliczone mnóstwo, a włączenie automatycznego tłumaczenia na polski pozwala poznać zaskakujące szczegóły dotyczące tego projektu. Ściszej biorąc, dotyczące pierwszej wersji tego projektu, gdzie częstotliwości pomiarowe wynoszą 100 Hz, 1 kHz i 7,813 kHz. W roku 2013 pojawił się drugi wątek, dotyczący ulepszonej wersji o maksymalnej częstotliwości pomiarowej 100kHz: <http://www.crystalradio.cn/thread-373788-1-1.html>

Autorem opracowania jest użytkownik nazywający się **XJW01**, stąd później oznaczenie przyrządu. Nieco zmodyfikowany pierwotny schemat pokazany jest na **rysunku 1**. Podstawą jest skromniutki dziś STC12C5A60S2 o architekturze 8051, ze swoim standardowym, 10-bitowym przetwornikiem ADC, oczywiście typu SAR. Szczegóły w karcie katalogowej: <https://www.stcmicro.com/datasheet/STC12C5A60S2-en.pdf>

Przyrząd działa na zasadzie pomiaru impedancji, a impedancja dotyczy przebiegów sinusoidalnych. Na rysunku 1 żółtymi podkładkami wyróżniłem punkty oznaczone DDS, ale nie ma tu prawdziwego generatora DDS, a sinusoida jest uzyskiwana przez filtrowanie przebiegu impulsowego z procesora. Przebieg sinusoidalny jest podawany na trochę dziwny (pół)mostek, wyróżniony zielonymi podkładkami, zawierający mierzoną impedancję Z_x oraz powiedzmy że wzorcowe rezystory $Ra1...Ra4$.

W układzie realizowany jest czteropunktowy pomiar Kelvina. Napięcia na Z_x i RaX są mierzone różnicowo. Szara podkładka wyróżnia bufory wejściowe, a niebieskie podkładki – dwa wzmacniacze różnicowe. Tak zmierzone przebiegi podawane są na obwody pomocnicze formujące sygnał (w tym na detektor fazy) i finalnie trafiają na dwa wejścia procesora (P1.0, P1.1), co jest zaznaczone fioletową podkładką.

Jak widać, w układzie wykorzystane jest wiele kłuczy analogowych CD405x. Warto też zwrócić uwagę, że oprócz niezbyt mocnego, taniego procesora i jednego, trochę lepszego wzmacniacza OP07, w układzie zastosowane są głównie wzmacniacze operacyjne TL08x, bardzo tanie, niezbyt dobre, uważane dziś za przestarzałe. Nic dziwnego, bo jest to typowa konstrukcja amatorska sprzed lat, wykorzystująca popularne elementy, a Autor miał kłopoty z doborem potrzebnych precyzyjnych rezystorów.



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.

W pełnej wersji dostępnej dla Patronów ten artykuł oczywiście ma więcej stron.

Range LED

CMOS switch wiring

C21 C22

K1, K2 = CD4052 K1b: range LED indicator switch
K3, K4, K5 = CD4053 K2b: current limiting resistor switch
K6, K7, K8 = CD4053 K1, K2: range switch

100Hz, 1kHz, 7.813kHz 数字电桥

矿石收音机

XJW01, 2011.10 设计于莆田

www.crystalradio.cn



Amatorskie wzorce pojemności

W czterech artykułach omówiłem trzy rodziny tanich mierników pojemności i indukcyjności, dostępnych na chińskich portalach. Wszystkie są konstrukcjami amatorskimi, a ich chińskie realizacje są zawodne. Do ich sprawdzenia potrzebne są jakieś wzorce, w szczególności wzorce pojemności.

Wzorce pojemności DIY – na ile dokładne?
Najtańsze styrofleksowe „wzorce pojemności”
Jakie wartości wzorców styrofleksowych?

Kondensatory innych wytwórców
„Wzorcowe” kondensatory ceramiczne
Inne kondensatory o wąskiej tolerancji

Artykuł Jaki miernik RLC kupić? A jakiego nie kupować? zawierał odpowiedź, na jakie tanie mierniki pojemności i indukcyjności warto zwrócić uwagę. Potem w artykułach M024, M025 oraz M026 omówiłem trzy grupy przyrządów. Wszystkie są opracowane przez amatorów. Zaskakują interesującymi rozwiązaniami, ale mają też rozmaite wady. Właśnie z uwagi na ich wady i dużą niepewność wskazań, każdy szanujący się hobbysta powinien posiadać jakieś kondensatory wzorcowe, które pozwolą sprawdzić ich działanie i dokładność. Wbrew pozorom, nabycie niedrogich kondensatorów o wąskiej tolerancji nie jest ani trudne, ani kosztowne.

Wzorce pojemności DIY – na ile dokładne?

W laboratoriach wykorzystywane są wzorce pojemności – stabilne kondensatory o niezmiennych parametrach – przykład na **fotografii 1**.



Fotografia 1



Fotografia 2

Laboratoryjne wzorce pojemności są bardzo kosztowne. Nie są przeznaczone dla amatorów. Na ile jednak bardzo dokładny pomiar i dobór pojemności jest potrzebny współczesnym hobbystom? Dawniej, w epoce analogowej, dość często były potrzebne kondensatory o dokładnie określonej pojemności, albo jak najbardziej stabilne, albo mające dokładnie określony temperaturowy współczynnik pojemności. Dziś zdecydowana większość hobbystów, w tym także radioamatorów, bardzo rzadko potrzebuje stabilnych kondensatorów o ściśle określonej wartości. Jednak także współczesnym elektronikom przydają się czasem nie tyle prawdziwe wzorce pojemności, co stabilne kondensatory o w miarę dokładnie określonej pojemności. Zwykle nie spełniają one wysokich wymagań pomiarowych, a są wykorzystywane głównie dla zaspokojenia ciekawości, służąc do sprawdzania i kalibrowania różnych mierników pojemności, zarówno multimetrów, jak i innych przyrządów, z których wiele ma kiepską dokładność.

Nabycie bardzo dobrych, stabilnych **rezystorów wzorcowych** jest dziś łatwe (choć kosztowne, bo jeden naprawdę dobry rezystor kosztuje kilkaset złotych). Można jednak okazynie kupić tanie rezystory o tolerancji 0,1% i o współczynniku cieplnym lepszym niż 15 ppm/°C. Zdecydowanie trudniej w warunkach amatorskich zdobyć **kondensatory** o dobrych parametrach, które mogłyby służyć jako amatorskie

Fotografia tytułowa oraz **fotografia 2** pokazują przykłady takich amatorskich wzorców. Fabryczna tolerancja użytych w nich kondensatorów to co najmniej $\pm 1\%$, a jeżeli ktoś ma dostęp do precyzyjnego laboratoryjnego miernika pojemności, może zmierzyć rzeczywistą pojemność swoich wzorców, by potem, przy własnych pomiarach, uzyskać dokładność nawet 0,1%, zależnie od kondensatora i zakresu temperatur, w jakich będą dokonywane pomiary.

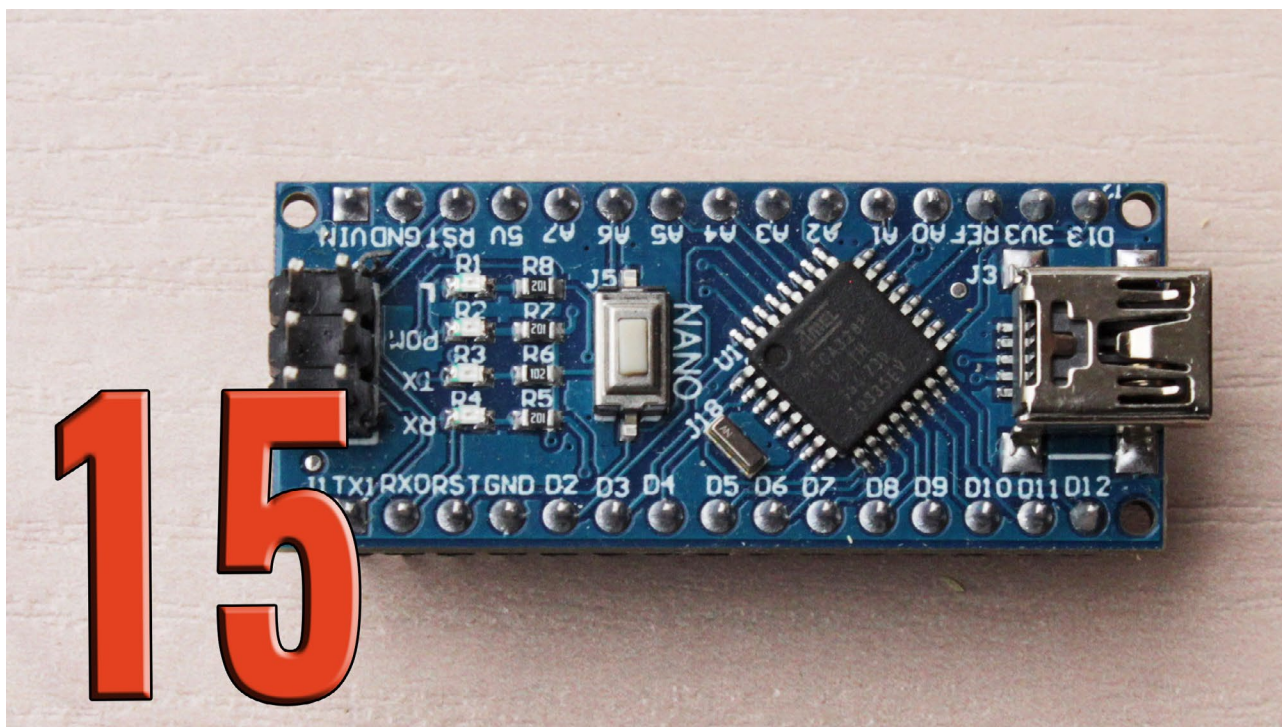
Najtańsze styrofleksowe „wzorce pojemności”

W przypadku dokładnych pomiarów napięcia czy rezystancji, błąd (niepewność) rzędu 0,5% byłby nieakceptowalny, natomiast przy pomiarach pojemności dokładność 0,5% należy uznać za zupełnie wystarczającą. Na pewno w przypadku hobbysty, ale nie tylko.

Taką tolerancję $\pm 0,5\%$ mają stare **kondensatory polistyrenowe, czyli styrofleksowe**, produkowane dawniej między innymi przez krajowy MIFLEX.

MIFLEX produkował głównie kondensatory styrofleksowe KSF-020 o słabej tolerancji, zawierające na obudowie literkę J ($\pm 5\%$), K ($\pm 10\%$) lub M ($\pm 20\%$), przeznaczone do sprzętu powszechnego użytku. Takie kondensatory mają dobre parametry elektryczne, ale dla nas są nieprzydatne, bo ich tolerancja jest zbyt szeroka. Co teraz dla nas ważne, produkowane były też kondensatory styrofleksowe do sprzętu

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



Mikroprocesorowa ośła łączka, część 15

Proste klawiatury są czasami za proste (mają zbyt mało przycisków) i wymagają rozbudowy. Wiąże się ona często z większą komplikacją sprzętową i często z bardziej rozbudowanym programem obsługi. Zróbmy więc coś ciekawego: przejście od prostej realizacji algorytmu do bardziej złożonego.

[Działanie klawiatur matrycowych](#)
[Popularne klawiatury](#)

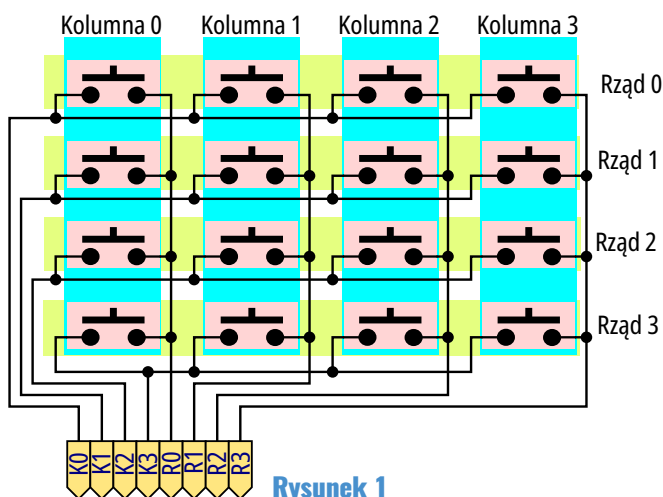
[Obsługa w programie](#)
[Elastyczność rozwiązania](#)

W tej części cyklu poświęconego mikrokontrolerom przedstawię „technikę” projektową, gdzie metodą kolejnych przybliżeń oprogramowanie będzie ulegało ewolucji. Pewne elementy, które będą wykorzystane są już znane, jak obsługa klawiatury w przerwaniach od upływu czasu. Drugim elementem jest obsługa kolejek FIFO, która doskonale radzi sobie z buforowaniem danych uzyskanych z obsługi klawiatury. Jednak najważniejszym celem jest uzyskanie takiej wersji, która jest sama w sobie niezależna od środowiska.

Działanie klawiatur matrycowych

To rozwiązanie pozwala znacząco zwiększyć liczbę obsługiwanych przycisków przy jak najmniejszej liczbie pinów niezbędnych do jej obsługi. Zbiór przycisków należy połączyć zgodnie z kon-

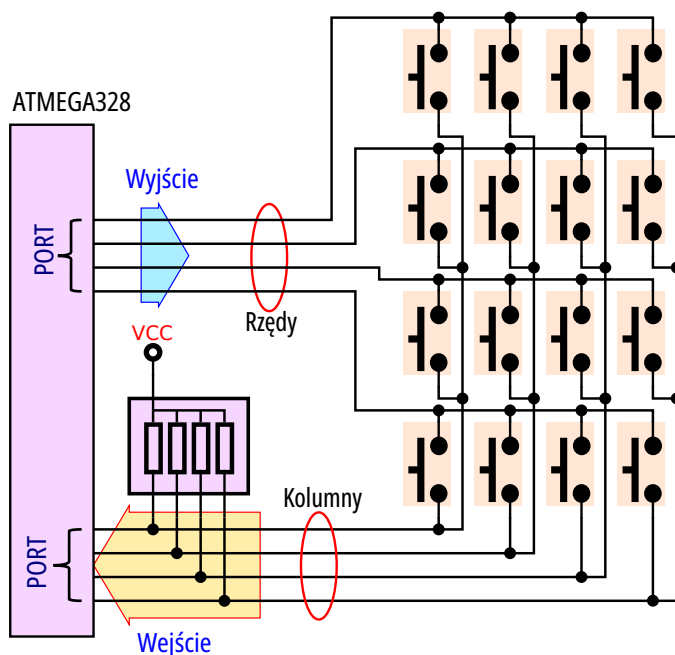
cepcją pokazaną na **rysunku 1** – można tu wyróżnić rzędy oraz kolumny tak utworzonej matrycy.



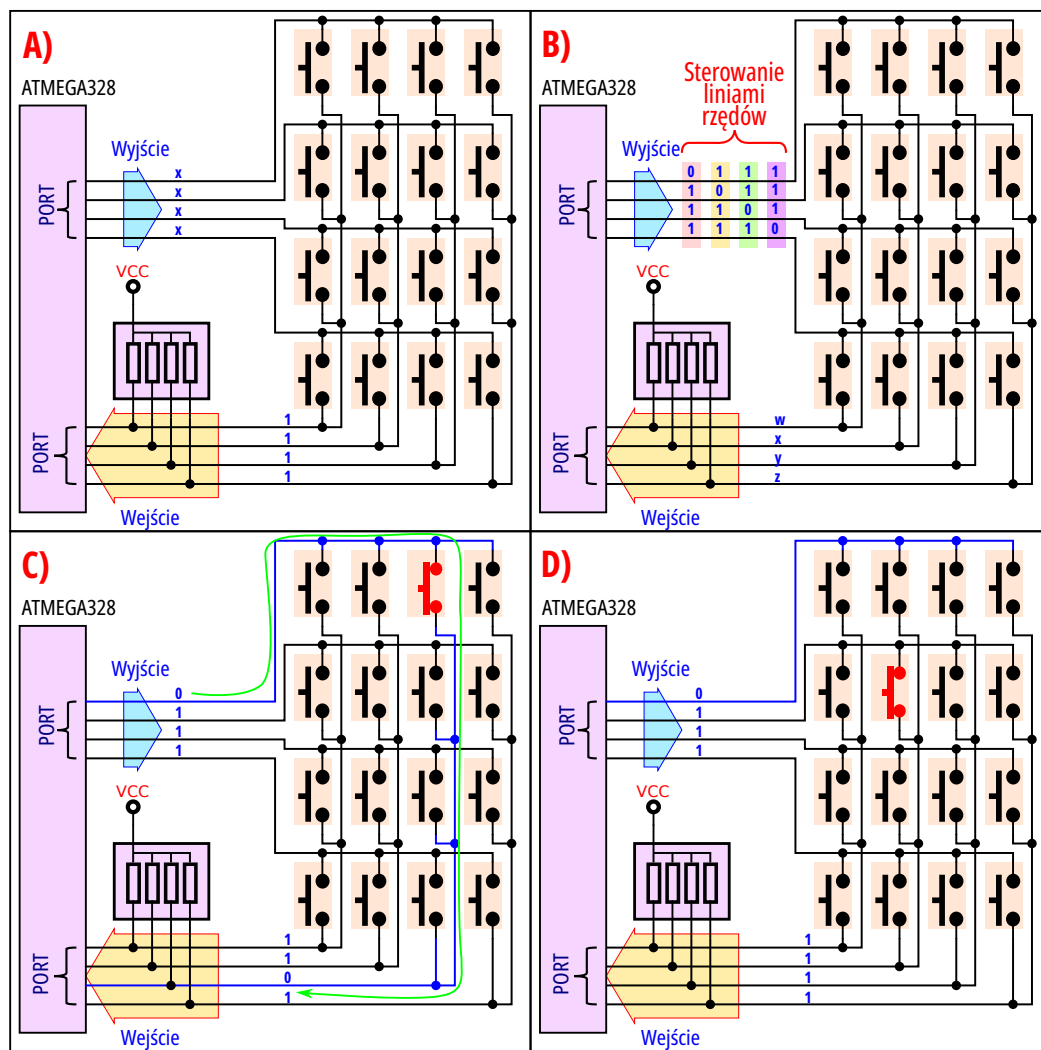
Rysunek 1

Naciśnięcie jakiegokolwiek przycisku łączy jedną linię rzędu z jedną linią kolumny. Obsługa tego typu klawiatur polega na przyłączeniu do portu wyjściowego wszystkich linii będących rzędami oraz do portu wejściowego wszystkich linii będących kolumnami klawiatury. Łatwo zauważyć, że zagadnienie jest symetryczne, toteż znaczenie rzędów oraz kolumn jest umowne (w dalszych rozważaniach mikrokontroler będzie sterował liniami rzędów oraz odczytywał linie kolumn), jak pokazuje **rysunek 2**. Istotnym elementem jest zespół rezystorów „podciągających” linie wejściowe do jedynki logicznej w porcie wejściowym. Obecnie mikrokontrolery są produkowane na bazie tranzystorów CMOS, więc nie mogą występować wejścia „wiszące w powietrzu”, gdyż daje to nieokreślony stan logiczny (gdy żaden przycisk nie jest naciśnięty jedynym elementem determinującym stabilny stan logiczny jest drabinka rezystorów).

Obsługa takiej klawiatury sprowadza się do odpowiedniego sterowania liniami rzędów. Jeżeli nie jest naciśnięty żaden przycisk, to bez względu na stan występujący na liniach rzędów, mikrokontroler odczyta z linii kolumn jako same jedynki (**rysunek 3a**). By wykryć naciśnięcie jakiegokolwiek przycisku należy na liniach rzędów „wyprodukować” jedno zero na tle samych jedynek logicznych oraz wczytać stan kolumn. Tej operacji cyklicznie musi zostać poddana każda linia rzędów (**rysunek 3b**). W przypadku, gdy jest naciśnięty przycisk łączący linię rzędu, na której panuje zero logiczne, to ten sygnał zostanie przekazany do odpowiedniej linii kolumny. Odczytując stan wszystkich kolumn na odpowiedniej



Rysunek 2



Rysunek 3

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



Elektroenergetyka – wytwarzanie energii, część 3

Wytwarzanie energii to nie tylko popularne elektrownie ciepłone, wodne czy wiatrak, ale również kilka mniej znanych technologii. W przyszłości mogą one odegrać dużą rolę w energetyce, jednak w polskim miksie energetycznym raczej nie zdobędą popularności.

Nietypowe źródła energii

Geotermia

Energia morza

Energetyka maretermiczna

Elektrownie słoneczne

Nietypowe źródła energii

Spalanie paliw kopalnych, wiatr czy woda to znane od dawna źródła energii, które są podstawą do budowy elektrowni. Poza takimi źródłami wykorzystuje się na przykład ciepło wnętrza Ziemi, a energię wód morskich próbuje się użyć dwójako: jako energię pływów morskich – elektrownie maremotoryczne – oraz jako różnicę temperatur wody – elektrownie maretermiczne.

W Polsce poza próbami zastosowania geotermii (bardziej w formie ciepłowni niż elektrowni) nie po-

opłacalnością inwestycji jako stosunku kosztów budowy do mocy.

Geotermia

Islandzkie gejzery są źródłem pomysłu na budowę elektrowni czy elektrociepłowni geotermalnej. Wydobywająca się spod powierzchni gruntu woda pod ciśnieniem i temperaturą sięgającą 80 stopni Celsjusza stanowi źródło ciepła – jak w typowej elektrowni ciepłej kocioł. W tym celu wykonuje się dwa odwierty do warstwy o odpowiednich pa-

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



Akumulatory kwasowe czy akumulatory zasadowe

Dziś absolutnie najpopularniejsze są różne odmiany akumulatorów litowych i jak na razie wygląda na to, że to do nich będzie należeć przyszłość. Zanim jednak zaczniemy szczegółowo omawiać ich właściwości, koniecznie należy poznać zalety i wady ich poprzedników: akumulatorów kwasowych i zasadowych.

Dlaczego nie akumulatory zasadowe?

Akumulatory kwasowo-ołowiowe

Zasada działania akumulatorów kwasowych

Zgodnie z zapowiedzią z poprzedniego artykułu tego cyklu, zajmiemy się teraz głównie akumulatorami kwasowymi, które pod pewnymi istotnymi względami są podobne do akumulatorów litowych. Dla uzyskania jeszcze szerszego, pełniejszego obrazu i dla porównania, na początku tego artykułu krótko omówię też kluczowe właściwości akumulatorów zasadowych. Takie informacje będą dobrym i wręcz niezbędnym fundamentem do gruntownego poznania kilku odmian akumulatorów litowych.

Dlatego nie lekceważ informacji zawartych w tym i następnym artykule cyklu, ponieważ jest to ważny punkt odniesienia i porównania.

Akumulatory kwasowe otwarte – mokre

Akumulatory VRLA, SLA, żelowe i AGM

Praca cykliczna czy buforowa?

Dlaczego nie akumulatory zasadowe?

Wprawdzie akumulatory zasadowe są coraz rzadziej wykorzystywane, jednak naprawdę warto znać ich rodzaje i właściwości. Mało kto pamięta dziś o akumulatorach niklo-żelazowych (NiFe), niklo-cynkowych (NiZn), srebrno-kadmowych (AgCd), srebrno-cynkowych (AgZn). Spośród akumulatorów zasadowych dziś wykorzystywane są tylko niklo-kadmowe (NiCd) oraz niklo-wodorkowe (NiMH). Akumulatory zasadowe NiCd i NiMH mają bardzo ważne zalety. Po pierwsze, nie boją się całkowitego rozładowania. W przeciwieństwie do kwasowo-ołowiowych i litowych, ich całkowite wyładowanie

„do zera” nie jest groźne. Podobnie umiarkowane przeładowanie nie grozi poważniejszymi skutkami. Nadmiar energii, której nie może przyjąć akumulator zasadowy wydziela się w postaci ciepła i nie jest to destrukcyjne, byle akumulator nie został nagrzwany do temperatury, która uszkodzi jego elementy. Inną poważną zaletą akumulatorów zasadowych jest ich trwałość przy pracy w takich trudnych warunkach.

Dlatego akumulatory zasadowe, a konkretnie ogniwa NiMH, a do niedawna też NiCd, są powszechnie używane w różnego rodzaju mniejszych i większych lampkach solarnych. Przykład takiej solarnej lampki ogrodowej z akumulatorem widoczny jest na **fotografii 1**.

Akumulatory zasadowe są również wykorzystywane w pewnych zastosowaniach niszowych, a nie tak dawno stosowano je też jako główne źródło energii w samochodach elektrycznych (Toyota), ale zniknęły z rynku głównie z uwagi na spór patentowy.

Akumulatory zasadowe bardzo rzadko pracowały w trybie buforowym jako baterie rezerwowe, będąc cały czas podładowywane małym prądem (rzędu 0,01 C...0,02 C, czyli liczbowo 1...2% pojemności). Zasadniczo taki prąd konserwujący z zapasem pokrywa straty samowyładowania i zapewnia nieustanną gotowość do pracy. Jednak pracujący buforowo akumulator rozładowany w czasie przerwy zasilania, należałoby po takim rozładowaniu naładować prądem dużo większym, bowiem naładowanie małym prądem konserwującym (0,01 C) trwałoby bardzo długo, a przy zużytych, starych akumulatorach pełne naładowanie mogłoby okazać się niemożliwe.

Dawniej „paluszki” NiCd oraz NiMH (o napięciu nominalnym 1,2 V) ceniono za znakomitą wydajność prądową, dużo większą, niż ogniwa jednorazowych – zwykłych manganowych i alkalicznych (o napięciu nominalnym 1,5 V). Dlaczego więc, skądinąd dobre akumulatory zasadowe, szybko znikają z rynku?

Przede wszystkim akumulatory zasadowe (i kwasowe też) mają **gęstość energii mniejszą** od litowych. Drugą poważną wadą jest to, że w akumulatorach NiCd i NiMH występuje tak zwany **efekt pamięciowy**. Otóż jeżeli akumulator nie jest rozładowywany w pełni, tylko częściowo, to niejako zapamiętuje ten fakt i jego użyteczna pojemność się zmniejsza. Aby uniknąć efektu pamięciowego wystarczy co kilka (5...6) cykli niepełnego rozładowania przeprowadzić cykl konserwujący, polegający na pełnym naładowaniu i pełnym rozładowaniu kontrolnym, przy czym pełne rozładowanie to nie rozładowanie „do zera” tylko do napięcia około 1 V



Fotografia 1

spowodowana innymi przyczynami, w tym słabą jakością ogniw, a całą winę zrzuca się na efekt pamięciowy, co na pewno jest wygodne dla producentów.

Istotną wadą zarówno akumulatorów NiCd, jak i większości NiMH jest duże samorozładowanie, zwłaszcza w podwyższonych temperaturach – już po kilku tygodniach przechowywania akumulator może się okazać pusty. To poważna wada.

Ale trzeba podkreślić, że od roku 2005 wiodąca firma Sanyo oferuje ulepszone akumulatorki NiMH, które mają bardzo małe samorozładowanie. Zaletą jest też ogromna trwałość, nawet ponad 2000 cykli pracy, płaska krzywa rozładowania oraz potwierdzona przez użytkowników stabilność parametrów, co ważne – także pojemności. Dziś Sanyo jest częścią koncernu Panasonic, a ulepszone akumulatorki są sprzedawane pod marką **eneloop** – **fotografia 2**.

Akumulatorki zasadowe znikają z rynku, ale odmiana **eneloop** nadal jest atrakcyjna.



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.

W pełnej wersji dostępnej dla Patronów ten artykuł oczywiście ma więcej stron.



Zapisane w pamięci – historia i typy pamięci, część 1

Pamięć – element bez którego nie może działać żaden komputer – na przestrzeni ostatnich kilkudziesięciu lat uległ wielkim przeobrażeniom. To co mamy dzisiaj, jest efektem prób, badań i eksperymentów – olbrzymiej pracy inżynierskiej.

Koncepcja pamięci ROM

„Starożytne” rozwiązania pamięci ROM

Półprzewodnikowe pamięci ROM

Pamięci EPROM

Pamięci EEPROM

Pamięci FLASH

Powstanie komputerów wniosło zapotrzebowanie na elementy o pewnej specyfice działania: zapisywania informacji oraz odczytywania wcześniej zapisanych danych. Pozwalało to na przechowywanie informacji o różnorodnym zastosowaniu. Wyobraźmy sobie jakiś układ, który dokonuje określonych pomiarów (choćby temperatury) i gromadzi je w sobie, aby później wykonywać odpowiednie analizy. Operacja gromadzenia to nic innego jak zapis bieżących danych w kolejnych komórkach pamięci. Operacja przetwarzania, jak przykładowo znalezienie najmniejszej i największej wartości, to nic innego, jak „przejrzenie” zapi-

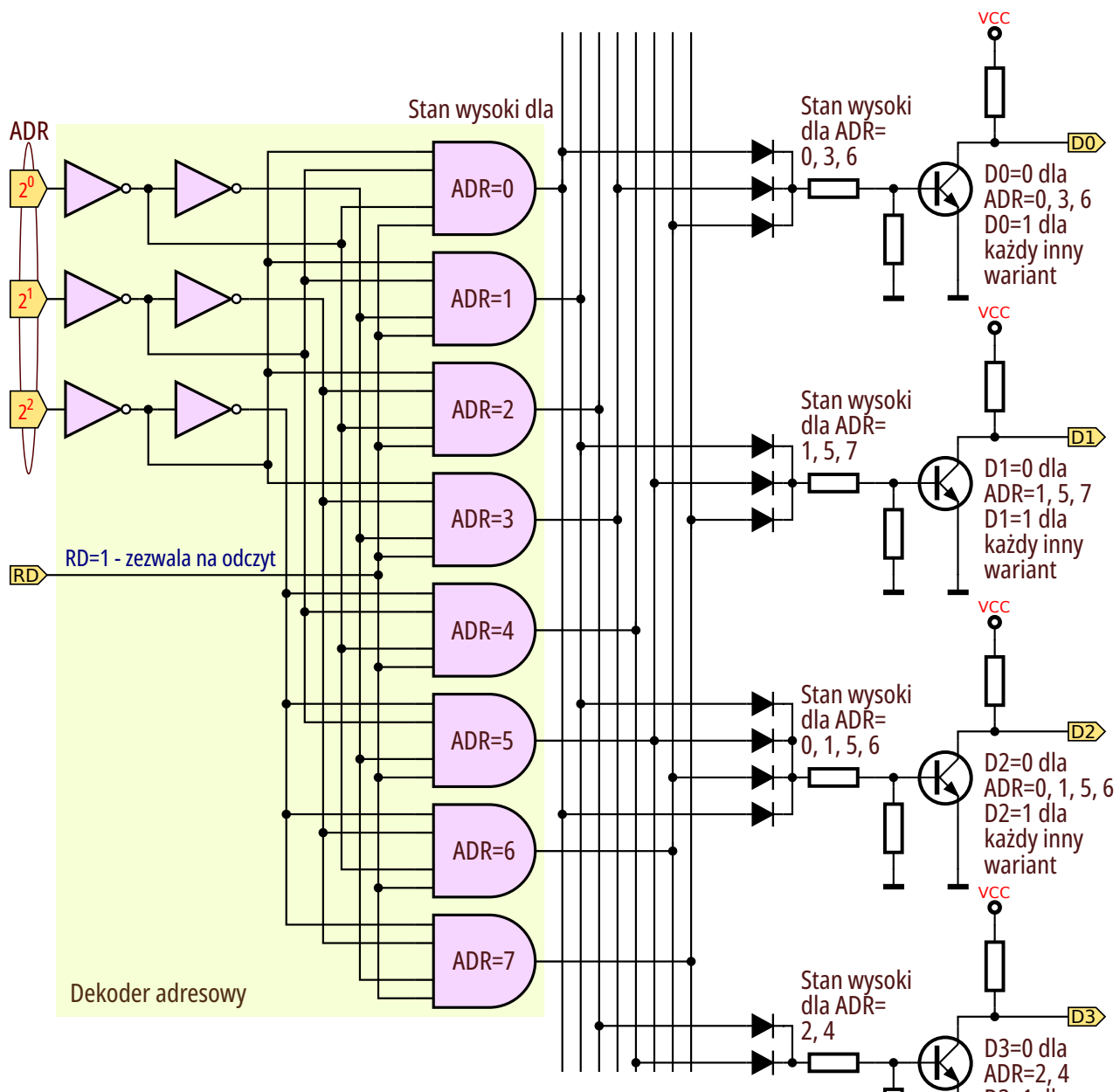
sanego stanu i znalezienie takiego elementu, który spełnia narzucone kryteria (wartości mniejszej lub większej). Takie poszukiwanie również wymaga przechowywania danych roboczych jak i wyników. Taki typ pamięci jest określany jako RAM (ang. **R**andom-**A**ccess **M**emory – pamięć o dostępie swobodnym) i pozwala bezpośrednio sięgnąć do dowolnej komórki.

W rzeczywistości występuje jeszcze jeden typ pamięci, który jest określany skrótem ROM (ang. **R**ead-**O**nly **M**emory – pamięć tylko do odczytu). Wyobraźmy sobie włączenie komputera do pracy. Zawiera on w sobie mikroprocesor,

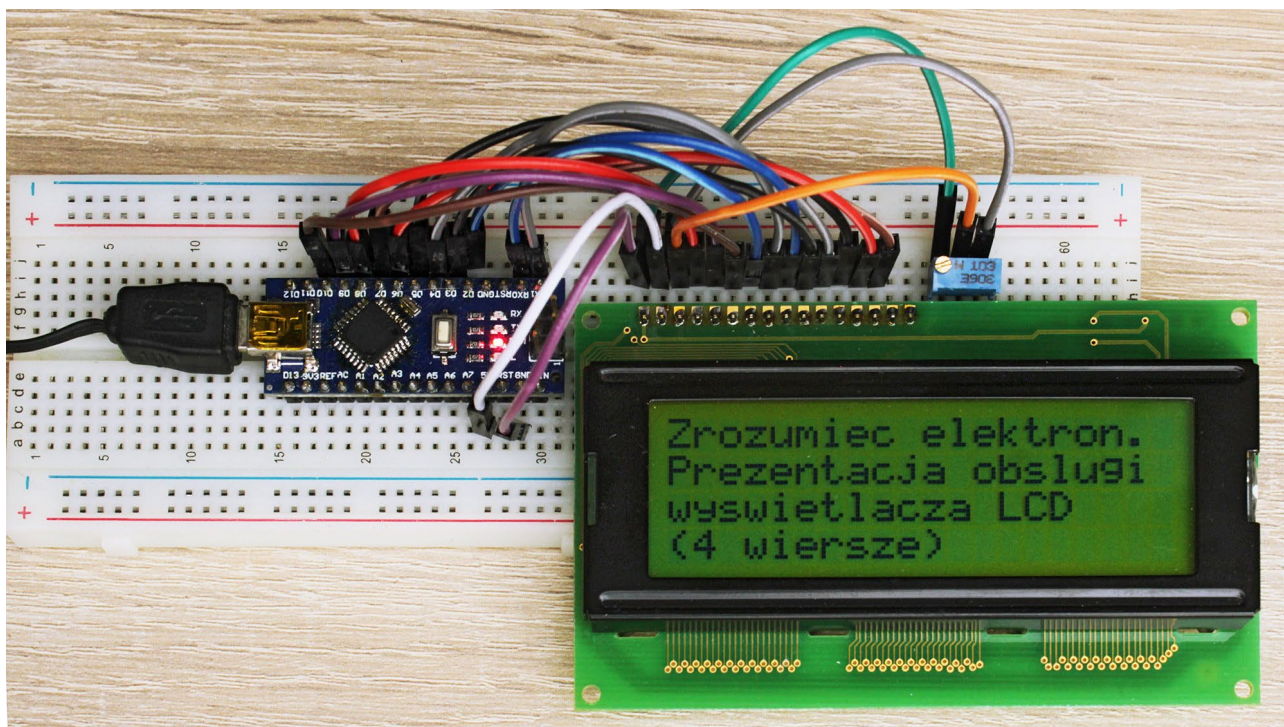
który od sygnału reset startuje z realizacją programu, który przykładowo wczyta z dysku i łączy do pamięci system operacyjny. Tu można zadać sensowne pytanie: jak to zrobić, skoro „goły” mikroprocesor nie robi nic, jak wczytać z dysku dane nie „umiejąc” obsługiwać dysku. Tu z pomocą przychodzi właśnie koncepcja pamięci typu ROM. Każdy komputer ma taką pamięć, gdzie jest umieszczony jego program rozruchowy (jako BIOS – ang. **B**asic **I**nput/**O**utput **S**ystem – podstawowy system wejścia – wyjścia). Każda pamięć ROM jest pamięcią nieulotną, czyli taką, która po wyłączeniu zasilania nie traci swojej zawartości i po ponownym włączeniu natychmiast jest gotowa do pracy, „wszystko dobrze pamięta”.

Koncepcja pamięci ROM

Ideę pamięci ROM ilustruje **rysunek 1**, gdzie pokazane jest przykładowe jej rozwiązanie o organizacji 8 komórek po 4 bity. Bazując na kombinacji stanów na szynie ADR, wypracowane są stany logicznej jedynki na poszczególnych wyjściach bramek. Wysoki stan z dekodera adresowego poprzez odpowiednią kombinację diod (jest dioda/nie ma diody) włącza do stanu nasycenia tranzystor wyjściowy. Można to interpretować jako logiczne zero na danym bicie dla określonej wartości adresu – w każdym innym przypadku jest jedynka logiczna. Zaprogramowanie takiej pamięci, czyli ustalenie jakie stany mają być na wyjściach dla każdej kombinacji adresowej, oznaczało wlutowanie lub nie diody łączącej tranzystor z wyjściem dekodera.



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



Czterowierszowy moduł wyświetlacza LCD

Gotowe alfanumeryczne moduły LCD stały się obecnie bardzo popularne w zastosowaniach. Standardowy moduł oferuje maksymalnie dwa wiersze po 40 znaków w wierszu. Czasami w zastosowaniach przydałoby się jednak więcej. Okazuje się, że jest to możliwe.

Środowisko do eksperymentów
Wyświetlacz czterowierszowy

Ulepszona wersja programu

Maksymalnie dwuwierszowy moduł wyświetlacza LCD, bazujący na HD44780 (lub kompatybilnym) zdomowił się nawet w amatorskich konstrukcjach. Jego obsługa nie jest skomplikowana, istnieje dość precyzyjna dokumentacja do samych modułów (tu znajdziemy istotne informacje dotyczące tego, jak go przyłączyć do budowanej konstrukcji) oraz zastosowanego kontrolera (co przekłada się na informacje dotyczące jego programowania). Naturalnym jest, że apetyty użytkowników stają się większe, toteż przemysł zaoferował moduły LCD czterowierszowe. Przyglądając się dokładnie specyfikacji takich modułów można dostrzec dwa warianty rozwiązania:

- moduł czterowierszowy z dwudziestoma znakami w wierszu (być może istnieją „krótsze” wiersze, ale takich nie napotkałem),

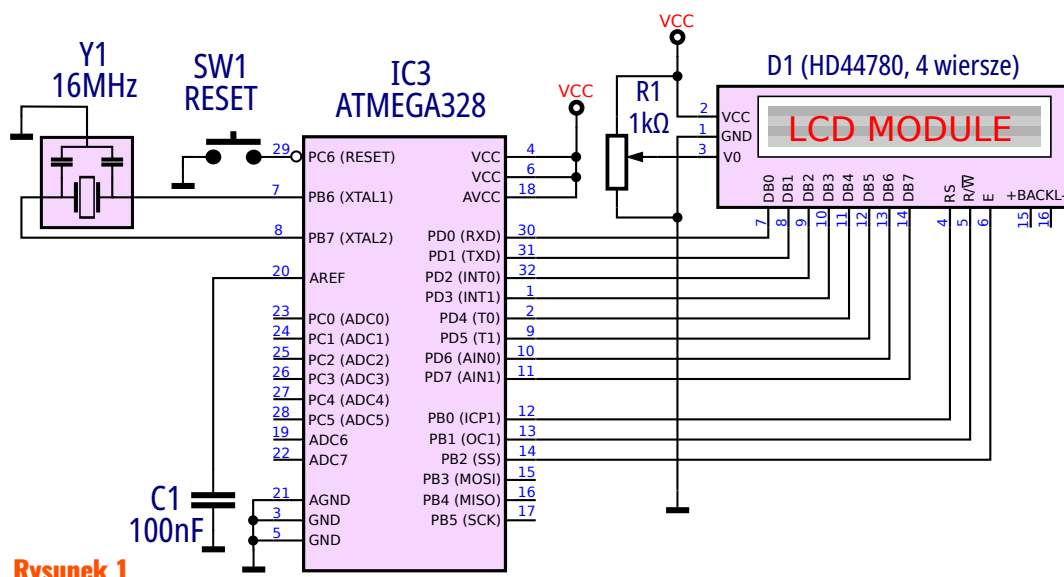
- moduł czterowierszowy z czterdziestoma znakami w wierszu (ewentualnie z pośrednią długością wiersza między 20 a 40 znaków).

Te dwa warianty znacząco różnią się między sobą, natomiast ten artykuł jest poświęcony wariantowi pierwszemu.

Środowisko do eksperymentów

Alfanumerycznym wyświetlaczom był poświęcony artykuł w ramach *Mikroprocesorowej oślej łąączki*, cz. 11 (ZE 04/2025), więc tam należy szukać

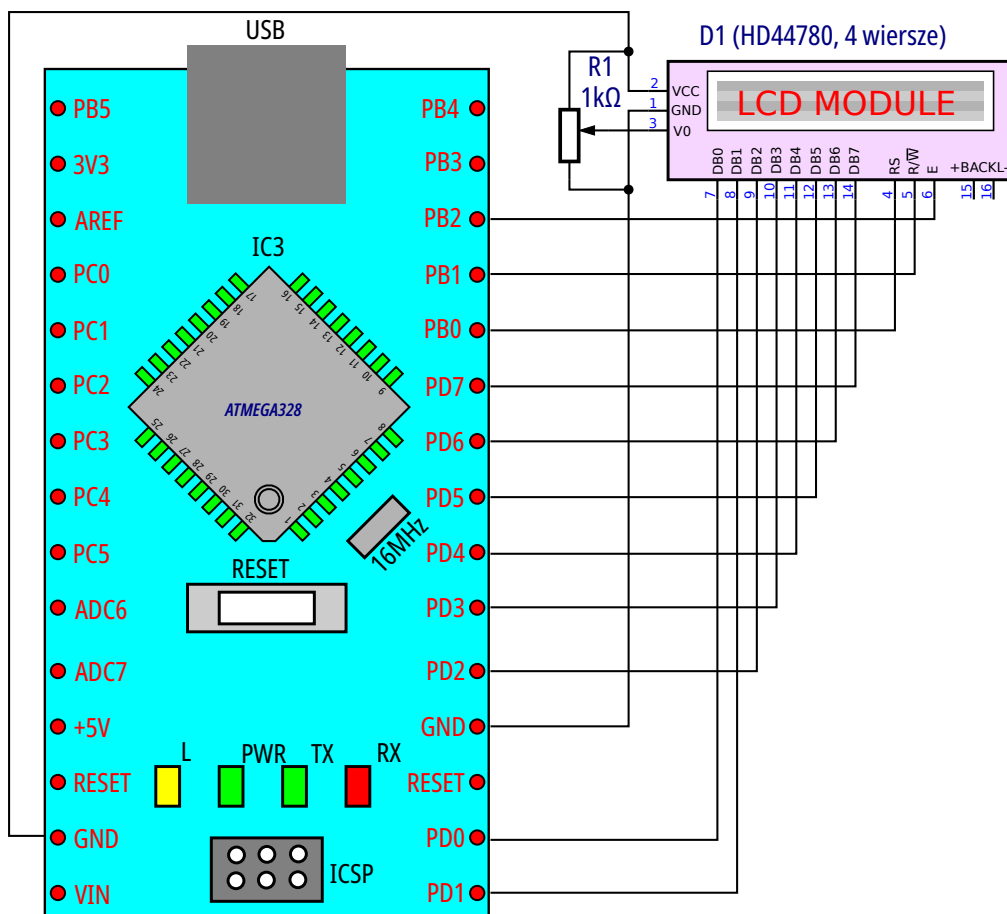
podstawowych informacji związanych z jego użyciem. W tym artykule jest on użyty w konfiguracji z 8-bitowym interfejsem, jak pokazuje **rysunek 1** (lub ewentualnie **rysunek 2**). Oczywiście możliwa jest również obsługa w wariancie 4-bitowej szyny danych. Zastosowany został moduł o oznaczeniu EW20400 (fotografia 3).



Wyświetlacz czterowierszowy

Wyświetlacz czterowierszowy z maksymalnie 20 znakami w wierszu jest „ciekawą” modyfikacją klasycznego rozwiązania – wyświetlacza dwuwierszowego. Układ HD44780 (lub kompatybilny) może obsługiwać maksymalnie 2 wiersze po 40 znaków w jednym wierszu. Rozpatrzmy ten właśnie przypadek. Kontroler zajmuje się takim sterowaniem poszczególnych „pikseli” znaków by finalnie uzyskać oczekiwany efekt „optyczny”. W gruncie rzeczy kontrolera nie interesuje (a nawet więcej – nie jest tego świadomy), gdzie zlokalizowany jest dany piksel. Jeżeli podzielić każdy 40-znakowy wiersz na dwa wiersze 20-znakowe, to uzyskamy wyświetlacz czterowierszowy. Jedyna występująca różnica w stosunku do wariantu dwuwierszowego dotyczy samego „szkła”, gdyż ono jest odpowiedzialne za fizyczne położenie każdego piksela.

Można by się spodziewać, że pierw-



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.

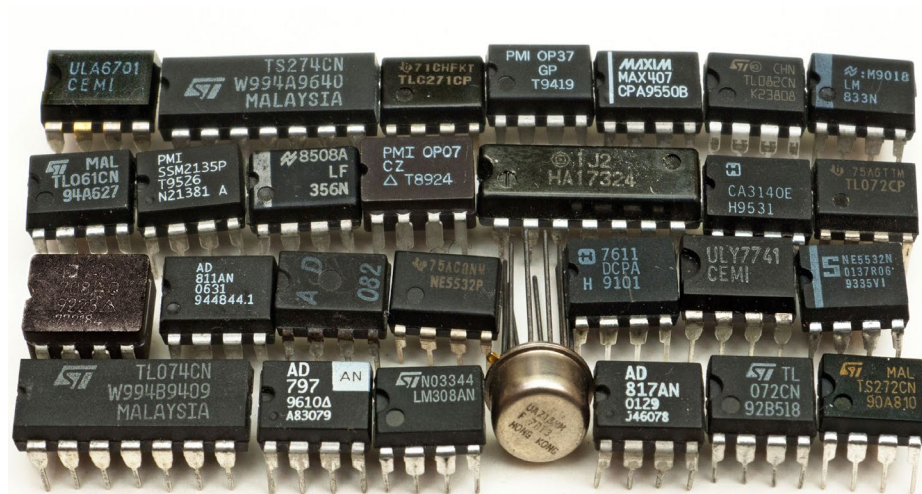
W pełnej wersji dostępnej dla Patronów ten artykuł oczywiście ma więcej stron.



Rysunek 3

Skrzynka pytań i odpowiedzi

W tej rubryce przedstawiane są odpowiedzi na wybrane pytania dotyczące elektroniki, zawarte w komentarzach do postów i filmów, nadsyłane przez Patronów i Mecenatów oraz innych Czytelników za pomocą kanałów podanych na stronie: [Zapytaj, odpowiedz](#).



Schematy wewnętrzne układów scalonych

[Gdzie znaleźć] pomoc w zrozumieniu bardziej skomplikowanych układów, np. LM393 lub LM358? Nie widzę, jak poprzez praktykę zrozumieć tego typu układy, w których jest kilkanaście tranzystorów. Chyba że spróbować zbudować podobny układ na płytce stykowej :) Może jest jakaś literatura (...)?

Czytelnik chce znaleźć schematy wewnętrzne oraz opisy działania różnych układów scalonych, by w ten sposób uczyć się elektroniki.

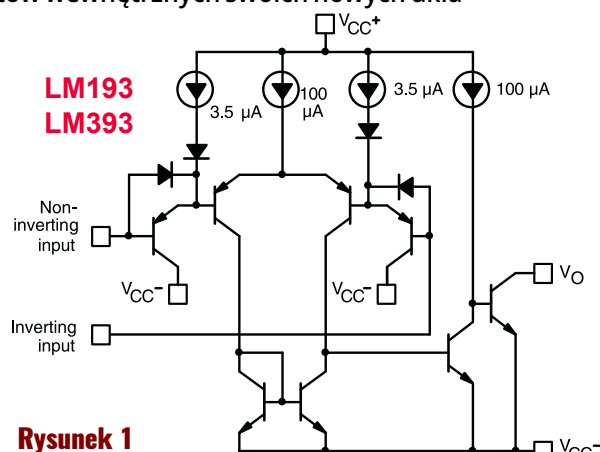
Kwestia jest interesująca. Wskazówki co do szukania schematów wewnętrznych podam w dalszej części artykułu. Ale najpierw omówię kilka poważnych problemów i ograniczeń, o których trzeba wiedzieć.

Otóż jak najbardziej warto znać budowę wewnętrzną niektórych klasycznych układów scalonych, jednak nie jest to główna ani najlepsza droga poznawania elektroniki. Wprost przeciwnie – jest to ślepa uliczka i nie warto nastawiać się na próby szczegółowej analizy innych, nowszych układów scalonych, ani wzmacniaczy operacyjnych, ani innych układów analogowych czy cyfrowych.

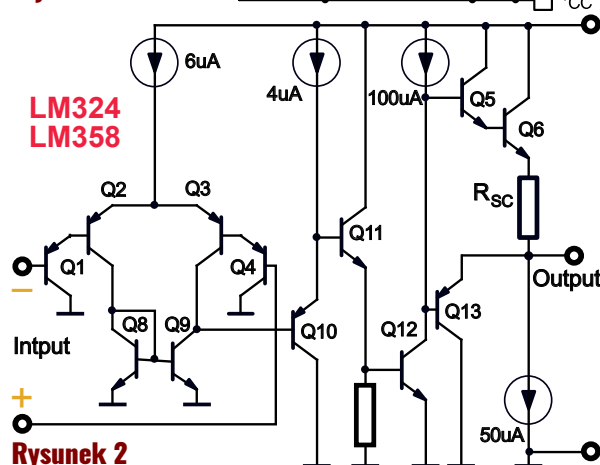
Wiele lat temu „bardziej skomplikowane układy” zawierały kilkanaście tranzystorów. Dziś zawierają ich zdecydowanie więcej, ale nie tu jest główny kłopot. Najpierw podam w punktach kluczowe problemy.

1. Tylko w przypadku najstarszych układów scalonych dostępne są w miarę pełne schematy wewnętrzne. Niestety, z reguły producenci nie udostępniają dziś schematów wewnętrznych swoich nowych układów scalonych.

2. Wiele kluczowych właściwości zależy nie od wnętrza układu scalonego, tylko od właściwości obwodów dołączonych do najważniejszych końcówek układu scalonego. Nie trzeba więc znać



Rysunek 1



Rysunek 2

„szczegółów wewnętrznych”, a jedynie pożądana, i to nie zawsze, może być wiedza o budowie obwodów wejściowych i wyjściowych. W szczególności dotyczy to wzmacniaczy operacyjnych: o właściwościach decyduje głównie budowa stopnia wejściowego.

3. W praktyce elektronikowi wcale nie jest potrzebna wiedza o budowie wewnętrznej, dużo ważniejsze okazują się parametry. I na tym powinien skoncentrować się praktyk. Owszem, parametry układu scalonego zależą od budowy wewnętrznej, ale np. przy wyborze wzmacniacza operacyjnego do bardziej wymagającego zastosowania decydują katalogowe parametry, a nie szczegóły budowy wewnętrznej. Elektronik powinien bardzo dużo uwagi i wysiłku poświęcić zdobywaniu gruntownej wiedzy o parametrach elementów elektronicznych, a nie o budowie wewnętrznej. Bo parametry katalogowe tak naprawdę mówią o niedoskonałości elementów, a to w praktyce jest najważniejsze. Ja te kwestie omawiałem szeroko w serii artykułów dotyczącej parametrów różnych podzespołów. Oto [link do skróconej wersji pierwszego artykułu tej serii](#).

4. Tylko w przypadku najstarszych układów scalonych ich budowę można wiernie przedstawić za pomocą najpopularniejszych klasycznych symboli, używanych do elementów dyskretnych (tranzystorów bipolarnych i rozmaitych tranzystorów polowych). Układy scalone często zawierają specyficzne struktury półprzewodnikowe, które nie mają dobrych odpowiedników w postaci elementów dyskretnych. Wśród popularnych symboli elementów nie ma takich, które wiernie przedstawiałyby budowę scalonych struktur i ich właściwości. Analiza działania bardziej skomplikowanych struktur scalonych nie może opierać się na najpopularniejszych symbolach elementów graficznych, które dotyczą elementów dyskretnych: pojedynczych klasycznych tranzystorów i różnych odmian diod.

5. Wiedza teoretyczna o budowie wewnętrznej jest potrzebna nielicznym naukowcom, którzy projektują układy scalone. Ale to są szerokie, trudne zagadnienia. Natomiast elektronik praktyk powinien po pierwsze, umieć wykorzystać dostępne układy scalone, a po drugie, umieć walczyć z problemami, jakie sprawiają w realnych układach. I tutaj wiedza o budowie we-

wania urządzeń z układami scalonymi to zupełnie inna dziedzina, gdzie przede wszystkim trzeba rozumieć sens parametrów użytych układów scalonych. A budowa wewnętrzna ma niewielkie znaczenie, bo liczą się inne kwestie.

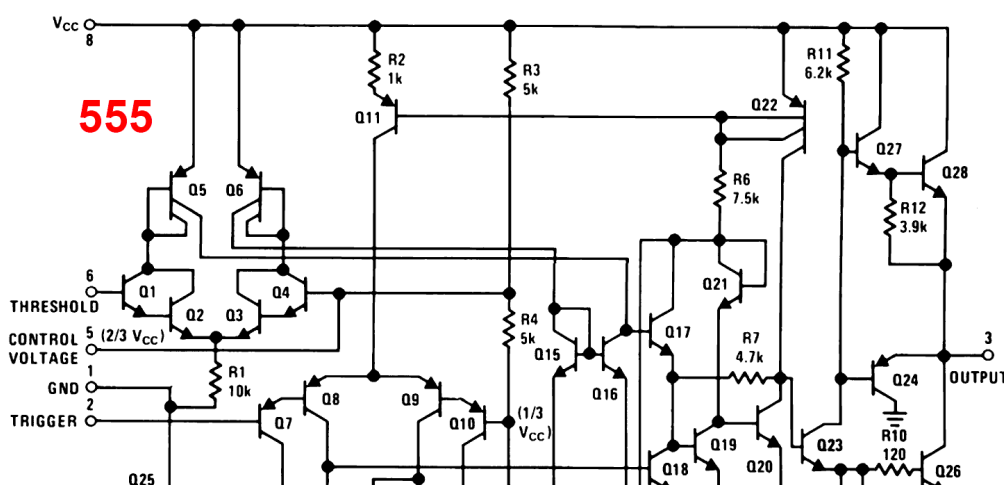
Na przykład w wielu zastosowaniach wzmacniaczy operacyjnych najważniejsze okazują się dokładność, precyzja, stabilność. A wszystko to zależy od parametrów użytych w układzie elementów. Należy skupić się na parametrach.

W wielu zastosowaniach układy scalone, w szczególności wzmacniacze operacyjne, stabilizatory liniowe i impulsowe, sprawiają poważne przykre niespodzianki. Często zamiast pełnić swoje właściwe funkcje zamieniają się w niepożądane generatory lub są na granicy samowzbudzenia, co nazywamy „dzwonieniem”. I dla elektronika najważniejsza jest właśnie wiedza o przyczynach samowzbudzenia i „dzwonienia”. A to praktycznie nie ma nic wspólnego z wiedzą na temat budowy wewnętrznej.

Choćby z wymienionych pięciu powodów nie zachęcam do koncentrowania się na szczegółach budowy wewnętrznej układów scalonych.

Gościwie zachęcam natomiast do działań praktycznych, bo wtedy można się najwięcej nauczyć. Wróć do tego, a na razie omówię pewne schematy. Omówię niektóre wzmacniacze operacyjne, ale najpierw przedstawię inny znamieny przykład.

Otóż na **rysunku 3** mamy schemat wewnętrzny jednego z najbardziej popularnych układów scalonych, a mianowicie kultowej kostki 555, którą opracował Hans R. Camenzind w roku 1971 i która została wypuszczona na rynek przez Signetics. Ze schematu po pierwsze nie wynika, jak działa taki układ, ale co ważniejsze, w ogóle nie ma śladu informacji, dlaczego znakomita jest stabilność czasu wytwarzanych impulsów lub przebiegów.



Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.

W pełnej wersji dostępnej dla Patronów ten artykuł oczywiście ma więcej stron.



Moja droga do okiełznania szybkich sygnałów

Cykl ten opowiada o mojej coraz większej świadomości dotyczącej zjawisk falowych, a szczególnie dopasowania falowego. Zjawiska te nabierają coraz bardziej znaczenia w moim projekcie komputera TTL, ze względu na to, że będzie on zbudowany z wykorzystaniem szybkich układów cyfrowych.

Kalkulator TTL, czyli nie jest aż tak źle...
Komputer TTL – tu już nie ma taryf ulgowych...

Testy właściwości statycznych szybkich buforów cyfrowych

Materiał ten ze względu na swoją obszerność został podzielony na trzy części. W tym artykule przedstawię moje wyniki testów dotyczących statycznych właściwości buforów cyfrowych z różnych rodzin TTL. Testy te rzucą światło na pewne ograniczenia i problem kompatybilności połączeń różnych rodzin układów cyfrowych.

W drugim artykule skupię się na szczegółowym wyjaśnieniu terminacji szeregowej, która będzie wieść prym w moim komputerze TTL. Przedstawię wady i zalety tego jednostronnego dopasowania. Mam wielką nadzieję, że dla wielu Czytelników terminacja ta zyska przy bliższym poznaniu.

W trzecim artykule przedstawię wyniki testów

właściwości dynamicznych szybkich buforów cyfrowych. Omówiona tam analiza wyników pozwala wyciągnąć istotne wnioski, które z pewnością zmaterializuję w fizycznym projekcie mojego komputera. Testy te zostaną przeprowadzone przy użyciu płytki, w której nie zastosowano prawidłowych zasad projektowania dla szybkich sygnałów. Dzięki temu będzie można się przekonać, jakie dodatkowe problemy stwarzają płytki PCB, które są nieprawidłowo zaprojektowane.

Kalkulator TTL, czyli nie jest aż tak źle...

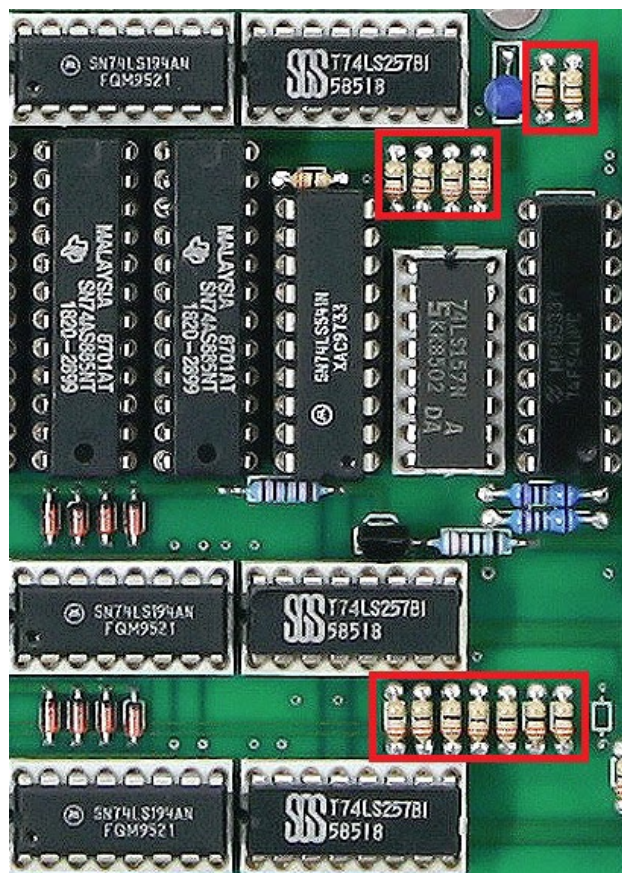
Gdy projektowałem mój kalkulator TTL wiedziałem, że problemy z kondycjonowaniem sygnałów

nie będą istotne nawet na tak ogromnej, ośmiowarstwowej płycie PCB o wymiarach 720 mm × 500 mm. Falowa natura energii elektromagnetycznej nie dawała się tak jeszcze we znaki. Ale i tam, projektując tysiące połączeń między układami scalonymi przeważnie z serii **74LS**, trzeba było być świadomym wielu problemów, gdy połączenia zaczęły przekraczać 40...50 cm. Kilka układów scalonych, szczególnie buforów, pochodziło z serii **74F** przede wszystkim z powodu większej obciążalności prądowej ich wyjść. Tu dystans realnej ścieżki, która łączyła ten szybki układ z innymi obwodami musiał być możliwie jak najkrótszy. W niektórych obwodach udało się skrócić te długości połączeń układowych, w innych było to całkowicie awykonalne.

W przypadku kalkulatora zostało zastosowane tylko kilka zaleceń projektowych, które opierały się, a w zasadzie jedynie ocierały się o reguły **HIGH-SPEED**. Najczęstszym sposobem radzenia sobie z potencjalnymi problemami, które swe źródło mają w zjawiskach falowych (szczególnie gdy chodzi o niedopasowanie impedancyjne) było wstawienie szeregowego rezystora na wyjściu bramki, bufora czy innego funkora logicznego. Prezentuje to **fotografia 1**, na której w czerwonych ramkach zaznaczone są takie rezystory. Jednak tak naprawdę żadna ścieżka nie była w moim projekcie typową, prawdziwą linią długą. Czas narastania/opadania zbocza impulsu prostokątnego był dłuższy niż czas propagacji sygnału na ścieżce połączeniowej, przez którą ten impuls podążał. Jednak trzeba zdawać sobie sprawę, że falowa natura przekazywanej energii elektrycznej daje o sobie znać zawsze, nawet przy prądzie stałym, z tego względu, że energia w postaci fali elektromagnetycznej zawsze przekazywana jest bezprzewodowo.

Nie ma ostrej granicy, a tak naprawdę wyraźnej korelacji między czasem narastania/opadania zbocza sygnału prostokątnego a długością ścieżki, która jest przewodnicą tego sygnału od źródła do obciążenia. Gdybyśmy rozpatrywali przebieg sinusoidalny, to teoretycznie problemy związane z niedopasowaniem impedancyjnym dają o sobie już w jakimś stopniu znać, gdy długość fali λ tego przebiegu jest mniejsza niż 1/10 długości ścieżki.

Jednak w technice cyfrowej nie występują przebiegi sinusoidalne tylko przebiegi prostokątne. Jak wiadomo przebiegi prostokątne zbudowane są z wielu sinusoidalnych harmonicznych. Takim książkowym, wzorcowym przykładem jest przebieg prostokątny o wypełnieniu równym 50%. Gdy rozmontujemy taki



Fotografia 1

tego amplituda każdej z tych harmonicznych jest o $1/n$ razy mniejsza od przebiegu podstawowego (pierwszej harmonicznej), gdzie n to numer nieparzystej harmonicznej. I tu zaczyna się problem z wyznaczaniem częstotliwości takiego przebiegu, bo jak wyznaczyć realną częstotliwość przebiegu, który składa się z wielu składników sinusoidalnych? Przyjęło się, a w zasadzie najczęściej stosuje się termin, który nosi miano częstotliwości efektywnej przebiegu prostokątnego. Wzór ten korzysta z trzeciej harmonicznej $F_{eff} = 0,35/t_r$, gdzie t_r oznacza czas narastania/opadania zbocza impulsu prostokątnego. Dzięki temu wzorowi możemy w przybliżony sposób określić „częstotliwość” przebiegu prostokątnego, a wyznaczając częstotliwość automatycznie możemy określić również długość fali – jednak zawsze trzeba mieć świadomość, że to tylko pewne ułatwienie, przybliżenie.

Wracając do szeregowych rezystorów na wyjściu różnych funklorów logicznych w moim kalkulatorze – ich rezystancja waha się od 15 Ω do 33 Ω . Jednak funkcja tego rezystora nie jest bezpośrednio związana z terminacją szeregową, o której szerzej opowiem w następnym artykule. Funkcja tych rezystorów związana jest z łagodzeniem ostrych zboczy impulsu.

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

ZROZUMIEĆ **E**LEKTRONIKĘ z Piotrem Góreckim

ZE 8/2025

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>
oraz konto buycoffee.to: [buycoffee.to/ piotr-gorecki](https://buycoffee.to/piotr-gorecki)**

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

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

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

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

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