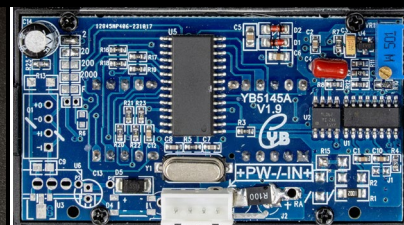
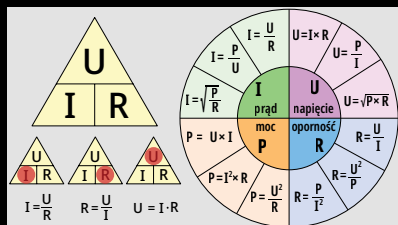
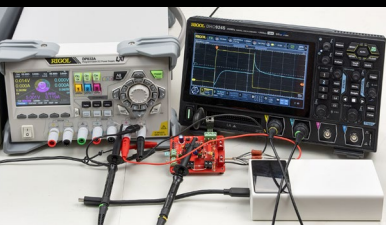
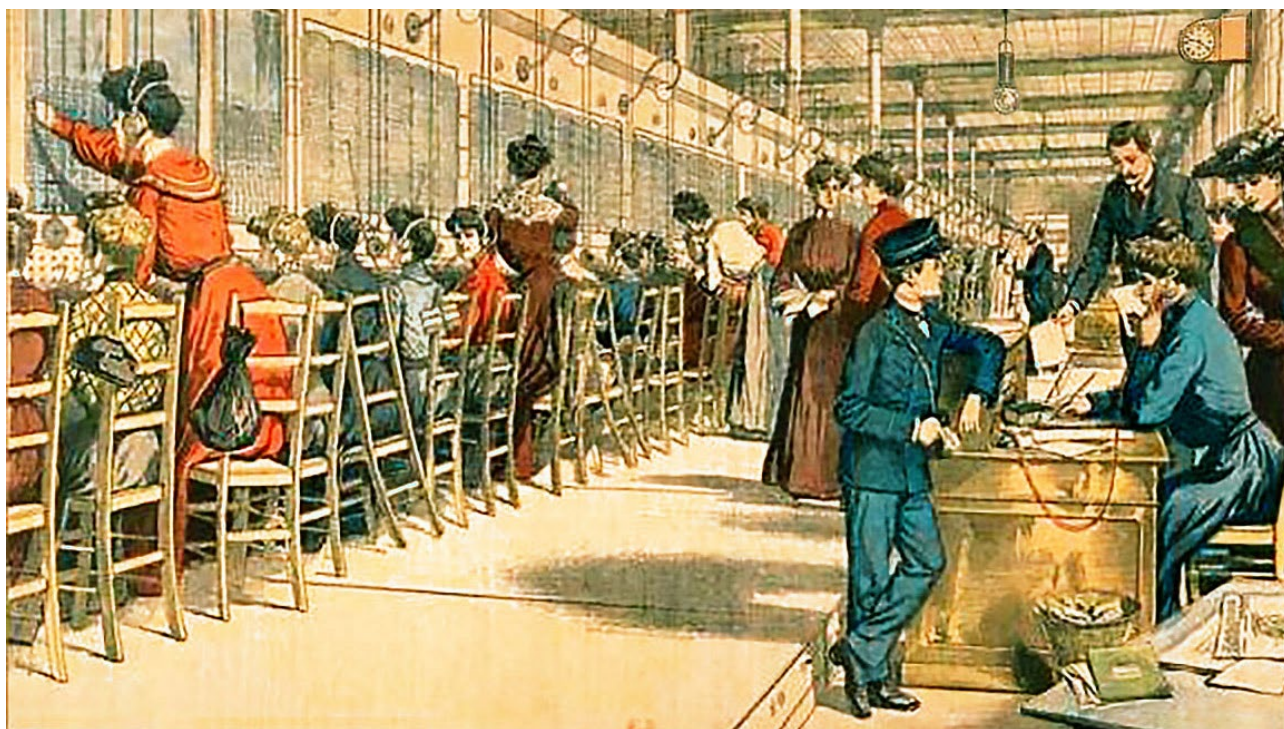




Jak działa smartfon, czyli... zaczniemy od telefonu

- Niewzruszony fundament elektroniki przewodowej • Zasilacze – niedoceniane parametry dynamiczne
 - Mały liniowy zasilacz: Pomiar napięcia i prądu • Zapisane w pamięci – historia i typy pamięci
- NanoVNA – gdzie i którą wersję kupić? • Mikroprocesorowa ośla łączka • Czy elektryczność leczy, czy szkodzi?
 - Ścieżki i przelotki na płytkach drukowanych • Moja droga do okiełznania szybkich sygnałów





Jak działa smartfon, czyli... zaczniemy od telefonu

Moja żona chętnie dowiedziałaby się, jak działa jej smartfon. Nie można tego wyjaśnić krótko. W przystępny sposób przedstawiam to w cyklu artykułów. Pierwszy omawia prapoczątki i tło historyczne, a następne – generacje telefonii 0G i 1G 2G, 3G, 4G, 5G oraz skróty GSM, UMTS, GPRS, EDGE, LTE, NR.

Smartfon? Zaczniemy od samego początku
Radiowa telekomunikacja bezprzewodowa

Pagery

Moja żona, która ma zainteresowania artystyczne a nie techniczne mówi, że chętnie dowiedziałaby się, jak działa jej smartfon. Pyta, jak możliwe jest przesyłanie na dowolne odległości rozmów, tekstów, zdjęć i filmów?

Ja to wiem i mogę wyjaśnić na dowolnym poziomie szczegółowości. Problem w tym, że dla osoby nietechnicznej próba zrozumienia, jak działa smartfon przypomina próbę wskoczenia do pędzącego

pociągu ekspresowego (**fotografia 1**). Wskoczyć jednym susem się nie da. Ale są inne sposoby. I właśnie ten artykuł jest początkiem nowej serii, stopniowo wyjaśniającej tajemnice współczesnej techniki.

Problem w tym, że współczesne smartfony to efekt 225 lat intensywnych badań elektryczności i przykład wykorzystania najróżniejszych wynalazków. Początkowo chodziło o przenośny telefon – o połączenie klasycznego telefonu przewodowego



Fotografia 1

z odbiornikiem i nadajnikiem radiowym, ale z czasem niejako obrósł dodatkowymi funkcjami i dziś jego pierwotna rola telefonu stała się dla wielu użytkowników marginalna. Bo choćby dzisiejsze komunikatory to o wiele więcej niż telefon.

Smartfon to miniaturowe kino, a raczej telewizor – powszechnie służy do wyświetlania filmów i fotografii. Ale to też kamera oraz studio fotograficzne i studio filmowe. Smartfony zawierają rozmaite czujniki, w tym lokalizator GPS, magnetometr (kompas) inklinometr, miernik przyspieszenia i obrotu.

A do tego zawiera niezliczone aplikacje o zadziwiających możliwościach.

Smartfon? Zaczijmy od samego początku

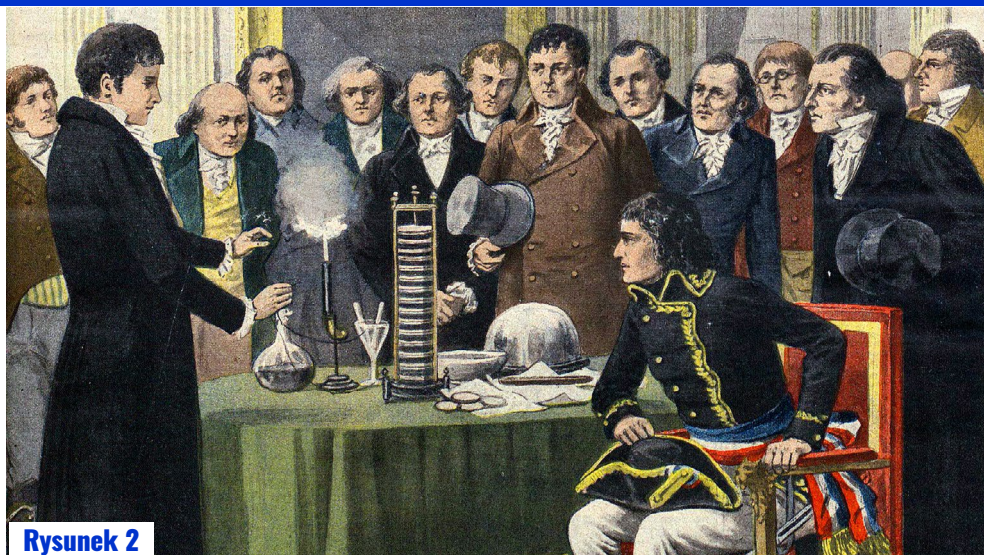
Dzisiejsze smartfony to jedne z najbardziej zaawansowanych urządzeń powszechnego użytku i mnóstwem funkcji i aplikacji. I jak wytłumaczyć działanie tego wszystkiego osobie „nietechnicznej”?

Spróbuję to zrobić powoli, małymi krokami, w sposób jak najbardziej przystępny. Dlatego także i w tym artykule „przemycam” pewne niezbędne informacje techniczne. A w jednym z następnych artykułów wytłumaczę, co to znaczy, że *wszystko to dziś realizowane jest w sposób cyfrowy, a nie analogowy*.

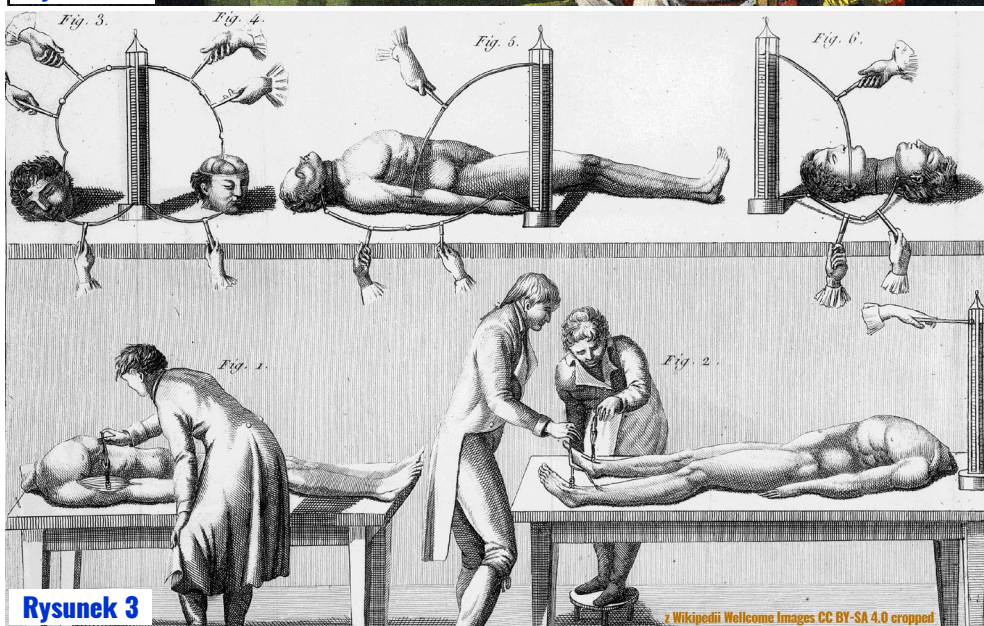
Ale odkrywanie licznych tajemnic smartfonów trzeba zacząć od funkcji telefonu. Dlatego koniecznie trzeba wrócić do jego historii.

Otóż intensywne badania elektryczności, a ściślej elektromagnetyzmu zaczęły się w roku 1800, gdy włoski hrabia Alessandro Volta przedstawił Napoleonowi (**rysunek 2**) swą nowo odkrytą baterię. Jego odkrycie stało się głośne i z początku nawet „nietechniczni” mogli nadążyć za rozwojem nauki i techniki, bo eksperymenty z elektrycznością były proste.

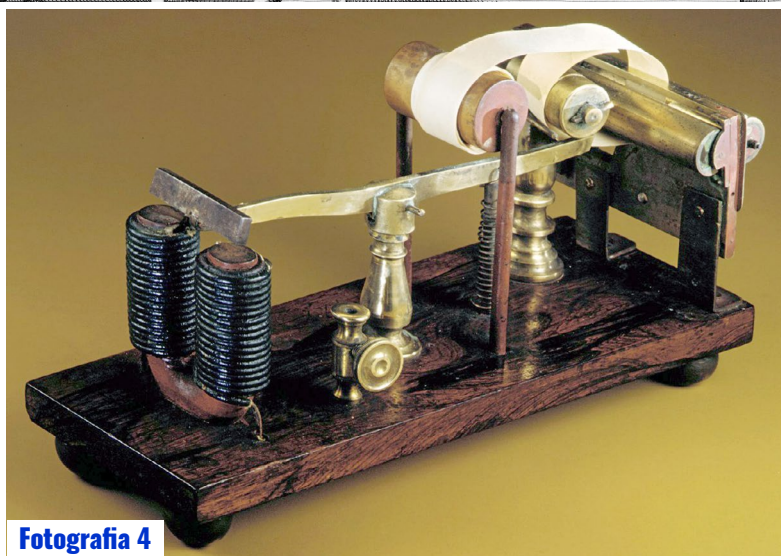
Elektryczność z początku była nie tylko wielce tajemniczą ciekawostką, zagadką, ale przypuszczano, że jest on „iskrą życia”, co wiąże się z maka-



Rysunek 2



Rysunek 3



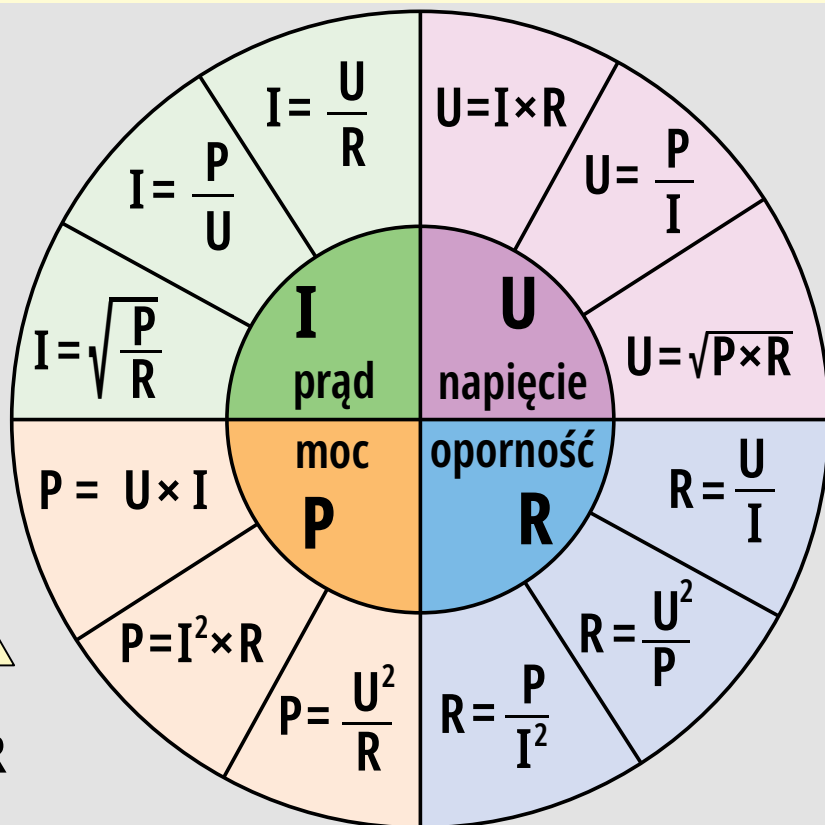
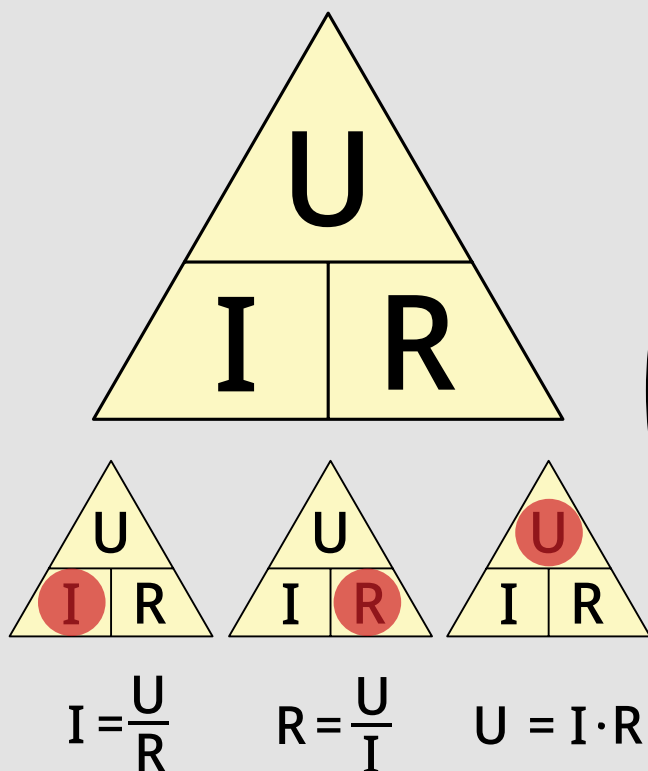
Fotografia 4

<https://piotr-gorecki.pl/h026>, <https://piotr-gorecki.pl/h025>.

Dla nas istotne jest, że pierwsze praktyczne wy-

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

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



Niewzruszony fundament elektroniki przewodowej

W innej serii artykułów uzasadniłem na pozór dziwny fakt, że energia elektr(omagnetyczna) zawsze przekazywana jest bezprzewodowo, „przez powietrze”. W tym artykule pokazuję, że gdy używane są przewody, możemy wykorzystać bardzo proste zależności dotyczące napięcia i prądu elektrycznego.

Fundament „elektroniki przewodowej” Pojęcie oporności, rezystancji

Współczesna cywilizacja opiera się na elektromagnetyzmie i elektryczności. Dziś powszechnie mamy do czynienia z **bezprzewodowym** przekazywaniem energii i informacji. Zamiast telefonów stacjonarnych, przewodowych, mamy telefony bezprzewodowe. Zamiast rozmaitych kabelków powszechnie wykorzystywane są łącza Wi-Fi i Bluetooth, nie mówiąc już o bezprzewodowych sposobach od dawna stosowanych w systemach radia oraz telewizji naziemnej i satelitarnej. Tam wykorzystuje się fale radiowe, a ściślej – fale elektromagnetyczne, inaczej mówiąc – pole elektromagnetyczne, czymkolwiek ono jest.. Na pewno przenoszona jest tam energia.

Najważniejsze wzory elektroniki przewodowej

Ewidentnymi przykładami przenoszenia energii przez pole elektromagnetyczne są ładowarka bezprzewodowa do smartfona, kuchenka indukcyjna i kuchenka mikrofalowa. Energię zawsze przenosi pole elektromagnetyczne. Okazuje się, że w przypadku wykorzystania przewodów, energię też przenosi pole elektromagnetyczne. Jednak zrozumienie szczegółów dotyczących fal i pola elektromagnetycznego nie jest łatwe. **Tak się jednak szczęśliwie składa, że w przypadku wykorzystania przewodów, choć nie przekazują one energii tylko nadają jej kierunek, to co się w nich dzieje, wiernie odzwierciedla transfer energii.**

Choć energię zawsze przenosi pole elektromagnetyczne, w przypadku wykorzystania przewodów nie musimy sobie zawracać głowy polem elektromagnetycznym i wektorem Poyntinga, ponieważ wtedy **łatwe do zmierzenia napięcie i prąd w przewodach pokazują wielkość i „okoliczności” przemian czy transportu energii**. Bezprzewodowe przekazywanie energii zawsze znajduje odzwierciedlenie w wartościach parametrów pomocniczych – prądu i napięcia elektrycznego. W tym artykule nie będziemy wnikać w to, czym jest prąd i napięcie elektryczne. Ważne jest to, że potrafimy je bardzo łatwo mierzyć: za pomocą amperomierzy i woltomierzy. A to genialnie upraszcza zagadnienie.

Fundament „elektroniki przewodowej”

Podkreślam, że poniższe informacje absolutnie nie są „całą prawdą o elektryczności”, ale na początek można uznać, że **fundamentem elektrotechniki i elektroniki „przewodowej” są dwie proste zależności i dwa reprezentujące je wzory:**

$$P = U \times I$$

$$R = U / I$$

Pierwszy wzór ($P = U \times I$) dotyczy najważniejszego: energii. Ale nie ilości energii, tylko tempa przemian czy przekazywania energii. Otóż pomnożenie pomocniczych parametrów: napięcia U oraz prądu I , daje informacje o mocy P . A moc P to właśnie tempo przemian czy przekazywania energii.

Drugi wzór ($R = U / I$) określa rezystancję R , czyli okoliczności, warunki, specyfikę przekazywania (przemian) energii elektrycznej. I właśnie stosunek U / I wyraża te „okoliczności przekazywania energii”. Szczegóły w tym i w następnych artykułach serii, ale najpierw zajmijmy się energią i mocą.

Otóż naprawdę najważniejsze w całej elektronice okazują się właśnie kwestie związane z **energiją E** i z tempem jej przekazywania/przekształcania, czyli z **mocą P**. Jednak w praktyce dużo trudniej jest określać całkowitą ilość **energii E**, a zdecydowanie łatwiej jest zmierzyć i obliczyć „aktualną” **moc**, czyli tempo przekazywania czy przemian energii, właśnie z prostego wzoru $P = U \times I$. Im większy iloczyn $U \times I$, tym w szybszym tempie energia elektryczna jest przekazywana czy też przetwarzana.

Dla bardziej wnikliwych: jeśli to **tempo (moc)**, wyrażone wzorem $P = U \times I$ jest niezmiennie w czasie, to **ilość** energii E przekazanej czy przetworzonej łatwo obliczymy mnożąc moc i czas: $E = P \times t = U \times I \times t$.

Natomiast jeżeli tempo $U \times I$ zmienia się z upływem czasu, to zamiast zwykłego mnożenia należy

Interesujący i dość istotny jest inny szczegół. Otóż o mocy P mówi się też w przypadku mechaniki. W szkole na lekcjach fizyki najpierw omawiana jest mechanika i tam mówi się o energii E , i o pracy W , wyrażanych w dżulach $[J]$. Moc P określana jest tam jako ilość pracy W wykonanej w jednostce czasu, co wyraża wzór: $P = W / t$. Jednostką mocy jest dżul na sekundę $[J/s]$, czyli wat $[W]$.

W elektronice też mamy do czynienia z energią, tylko z odmienną formą energii. Ale **energia jest tylko jedna** (energia rozumiana ogólnie jako zdolność do wykonania pracy, zdolność do zmian), natomiast **ma różne formy, postacie**. Energia elektryczna, ściślej elektromagnetyczna, podobnie jak mechaniczna, też powinna być wyrażana w dżulach $[J]$, a moc w dżulach na sekundę $[J/s]$, czyli w watach $[W]$.

Jednak w przypadku energii elektrycznej prawie nigdy nie mówimy o dżulach, tylko wyrażamy ją w watosekundach $[Ws]$, watogodzinach $[Ws]$ lub w kilowatogodzinach $[kWh]$ gdzie $1 Ws = 1 J$, $1 Wh = 3600 Ws = 3600 J$, a $1 kWh = 3600000 Ws = 3600000 J$. Można więc powiedzieć że „w elektronice wat jest ważniejszy od dżula”. Może nie tyle ważniejszy, co bardziej popularny w praktyce, bo łatwiej mierzyć tempo przemian/przekazywania energii, czyli moc P wyrażoną w watach, niż całkowitą ilość, porcję energii przekazanej/przetworzonej w określonym czasie.

To była tylko dygresja na temat dżula i wata. Ważne jest to, że podczas przekazywania czy przemian energii elektrycznej, według zależności $P = U \times I$ mamy do czynienia z bardzo różnymi sytuacjami. Najprościej biorąc, dane tempo przekazywania energii P można uzyskać przy bardzo różnych wartościach napięcia U oraz prądu I . Przykładowo moc elektryczną 100 watów, czyli tempo przekazywania energii 100 dżuli na sekundę, można uzyskać, gdy napięcie wynosi 100 woltów a prąd 1 amper, albo gdy napięcie wynosi 1 wolt, a prąd 100 amperów. Jest też nieskończenie wiele różnych innych kombinacji. Oto kilkanaście bardziej i mniej realnych przykładów:

$$100 W = 1\,000\,000 V \times 0,0001 A$$

$$100 W = 100\,000 V \times 0,001 A$$

$$100 W = 10\,000 V \times 0,01 A$$

$$100 W = 1000 V \times 0,1 A$$

$$100 W = 230 V \times 0,4348 A$$

$$100 W = 100 V \times 1 A$$

$$100 W = 20 V \times 5 A$$

$$100 W = 12 V \times 8,333 A$$

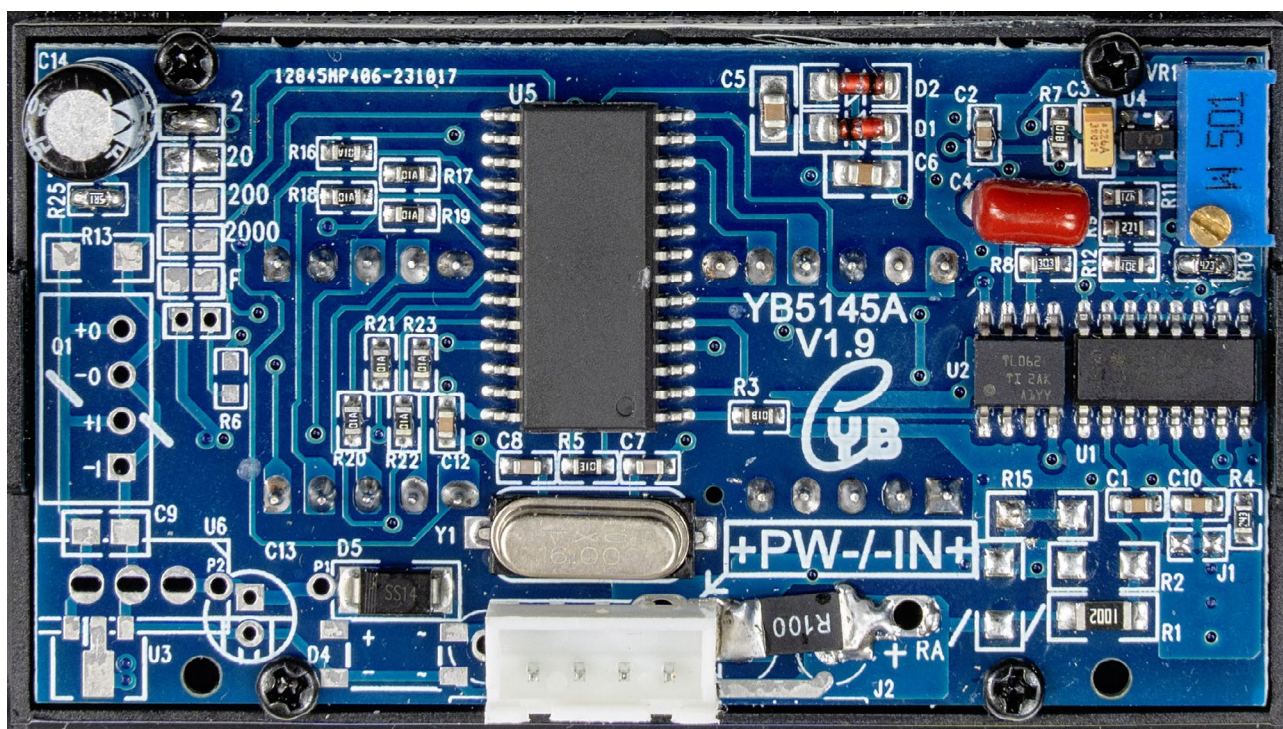
$$100 W = 10 V \times 10 A$$

$$100 W = 5 V \times 20 A$$

$$100 W = 1 V \times 100 A$$

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

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



Mały liniowy zasilacz: Pomiar napięcia i prądu

Oto siódmy artykuł serii pokazującej, jak elektronik dla swej ogromnej satysfakcji może, na bazie scalonego stabilizatora LM317 lub podobnego, zaprojektować i zrealizować niewielki zasilacz liniowy przeznaczony do swojej pracowni. Zajmiemy się kwestiami pomiaru napięcia oraz prądu wyjściowego.

Pomiar napięcia i prądu wyjściowego
Kompensacja prądu potencjometru

Problem wskazań woltomierza
Cyfrowe moduły pomiarowe

W poprzednich artykułach tej serii podałem wszystkie informacje, potrzebne do zaprojektowania i zrealizowania liniowego zasilacza regulowanego z kostką LM317 i dodatkowymi obwodami ogranicznika. Teraz, zgodnie z zapowiedzią, zajmiemy się kwestią pomiaru prądu i napięcia.

Pomiar napięcia i prądu wyjściowego

W zasilaczu ze stabilizatorem LM317 można dość prosto mierzyć napięcie i prąd wyjściowy. Wiele osób nadal darzy wielką sympatią mierniki wskazówkowe, na przykład takie, jak na **fotografii 1**.



Fotografia 1

Można je wykorzystać, tylko jest pewien problem, bo mają skalę najczęściej 0...100, a w zasilaczu maksymalne wartości zapewne będą inne.

Napięcie oczywiście mierzymy na wyjściu zasilacza, natomiast prąd wyjściowy jako spadek napięcia na boczniku. Tym bocznikiem może być rezystancja R_S . Wtedy jednak trzeba pamiętać, że przy zmianie zakresów ogranicznika prądowego (**rysunek 2**) zmieniać się będą wskazania miernika prądu. Można powiedzieć, że wtedy **będzie to miernik pokazujący procentową wartość w stosunku do prądu ograniczania na danym zakresie**.

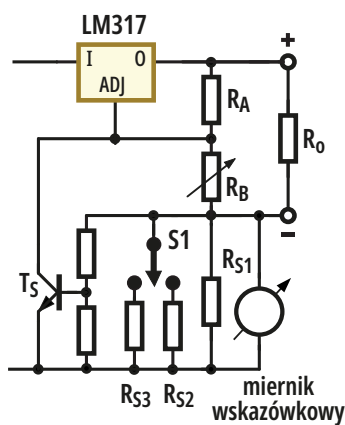
Bocznik amperomierza można włączyć w innym punkcie zasilacza, a amperomierz może być, i zapewne będzie, cyfrowy.

Rysunek 3 pokazuje przykłady. Wtedy będziemy mierzyć nie „procenty”, tylko wartość prądu w (mili)amperach. Na rysunku B celowo pokazuję dwie niezbyt praktyczne możliwości pomiaru prądu, żeby pokazać pewne istotne szczegóły. Otóż jest oczywiste, że wskazanie amperomierza byłoby wtedy zawyżone o prąd pobierany przez sam stabilizator LM317.

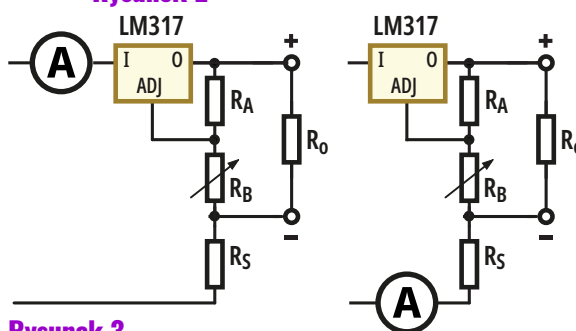
Mniej oczywiste są problemy przy na pozór najprostszym i oczywistym włączeniu amperomierza na wyjściu zasilacza. Otóż wtedy w grę wchodzi kilka kwestii, zwłaszcza gdy zachcemy wykorzystać gotowe moduły mierników cyfrowych.

Na rynku jest mnóstwo podwójnych modułów mierników napięcia i prądu. Przytłaczająca większość nie nadaje się do naszych celów, bo ma bardzo słabą dokładność przy pomiarze małych prądów.

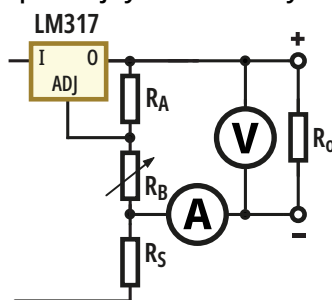
W zasilaczu należy zastosować naprawdę dobry wskaźnik cztero-cyfrowy z automatycznie przełączanym zakresem amperomierza, gdzie rozdzielczość na niższym zakresie to 0,1 miliampera. Moduły takie omawiam szerzej w oddzielnym artykule. Jeden z problemów widoczny jest



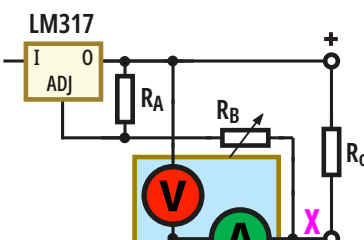
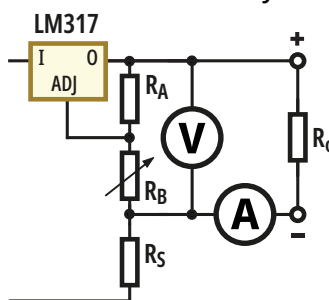
Rysunek 2



Rysunek 3



Rysunek 4



rze jest włączona wewnątrz pętli stabilizacji napięcia.

Nie zaznaczam na tych schematach rezystancji bocznika w amperomierzu, ale ona zawsze tam jest i nie możemy o niej zapomnieć. Spadek napięcia na niej wynosi od kilkudziesięciu do 200 miliwoltów.

Mierzona wartość prądu jest zawsze prawidłowa (słusznie zakładamy, że rezystancja wewnętrzna woltomierza jest duża). Jednak w obu przypadkach wzrost poboru prądu zmniejsza napięcie wyjściowe na rezystancji obciążenia R_O ze względu na nieunikniony spadek napięcia na rezystancji wewnętrznej (bocznika) amperomierza.

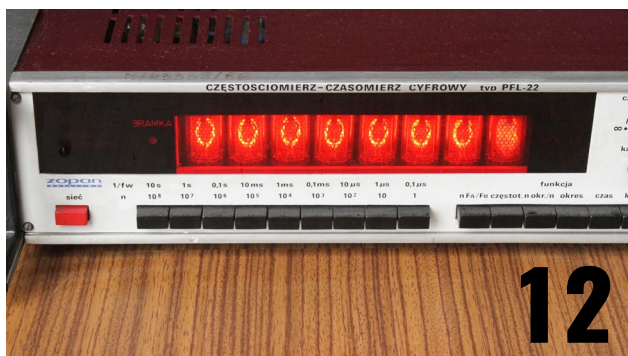
W pierwszej wersji z rysunku C woltomierz pokazuje prawidłowo wartość napięcia na obciążeniu, ale to napięcie zmniejsza się trochę ze wzrostem poboru prądu. Natomiast w drugiej konfiguracji, według rysunku C, woltomierz nie sygnalizuje problemu, bo nie pokazuje prawdziwego napięcia wyjściowego, tylko dobrze stabilizowane napięcie „przed amperomierzem”.

Praktyczny problem w tym, że jeśli chcielibyśmy wykorzystać dwufunkcyjny moduł VA, to z uwagi na jego konstrukcję jesteśmy skazani na tę drugą konfigurację z rysunku C.

Warto też rozważyć wersję z **rysunku 5** z dwufunkcyjnym modułem VA. Zmieniamy miejsce dołączenia potencjometru R_B – dołączamy go „za amperomierzem” i dzięki temu prawidłowo stabilizujemy napięcie na wyjściu. Tak, ale teraz dwufunkcyjny moduł mierzy nie napięcie wyjściowe, tylko nieco większe napięcie „przed amperomierzem”, które będzie wyższe o spadek napięcia na amperomierzu. Prąd jest wtedy mierzony prawidłowo, ale wskazania napięcia są zawyżone o spadek

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

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



Wspólnie projektujemy: Częstościomierz, część 12

Stopień wejściowy
Wzmacniacz analogowy

Konwersja do sygnału cyfrowego

Pozbadaniu poszczególnych elementów mogących mieć zastosowanie w budowie toru przetwarzającego wejściowy sygnał analogowy do postaci cyfrowej, przyszedł czas na integrację części składowych w jedną całość.

Stopień wejściowy

Żaden przyrząd pomiarowy nie powinien zbyt ingerować w mierzone obwody (ideałem jest brak interakcji przyrządu z mierzonym sygnałem, stan do którego należy dążyć). Ten wymóg sprowadza się do stworzenia odpowiednich obwodów wejściowych, które cechują się wysoką impedancją. Do budowy takich stopni można wykorzystać tranzystory polowe (często określane jako tranzystory unipolarne), które pod względem własności znacząco różnią się od tranzystorów bipolarnych (nnp lub pnp). W ofercie powszechnie dostępnych tranzystorów polowych występują tranzystory JFET (skrót od Junction Field-Effect Transistor – tranzystor polowy złączowy) oraz MOSFET (od Metal-Oxide-Semiconductor Field-Effect Transistor – tranzystor polowy z izolowaną bramką). Ich działanie (JFET oraz MOSFET) jest podobne, sprowadza się do sterowania przepływem prądu przez kanał (między *drenem* a *źródłem*), poprzez pole elektryczne uzyskane w wyniku doprowadzenia napięcia do elektrody, określanej jako *bramka* (napięcie między bramką a źródłem). Ich wspólną cechą

rezystancja wejściowa tranzystorów polowych sięga setek GΩ, natomiast różnica między nimi sprowadza się do konstrukcji bramki. W tranzystorach JFET bramka jest odizolowana od kanału za pomocą złącza półprzewodnikowego spolaryzowanego zaporowo. W przypadku MOSFET, zastosowany jest rzeczywisty izolator (przykładowo dwutlenek krzemu).

W stopniu wejściowym zostały zastosowane tranzystory JFET (o symbolu J111). Są to tranzystory JFET z zubożonym kanałem N (tak jak tranzystory bipolarne dzielą się na npn oraz pnp, tranzystory polowe dzielą się na te z kanałem typu N i typu P). Tranzystor z kanałem zubożonym (w przeciwieństwie do tranzystorów z kanałem wzbogaconym) oznacza, że przy braku wysterowania wyprowadzenia bramki, tranzystor przewodzi (między drenem z źródłem płynie prąd). Aby powstrzymać przepływ prądu, należy między bramkę a źródło doprowadzić napięcie ujemne. Dla użytego J111 to napięcie odcięcia bramka – źródło wynosi -3...-10 V (jak pokazuje rysunek 1).

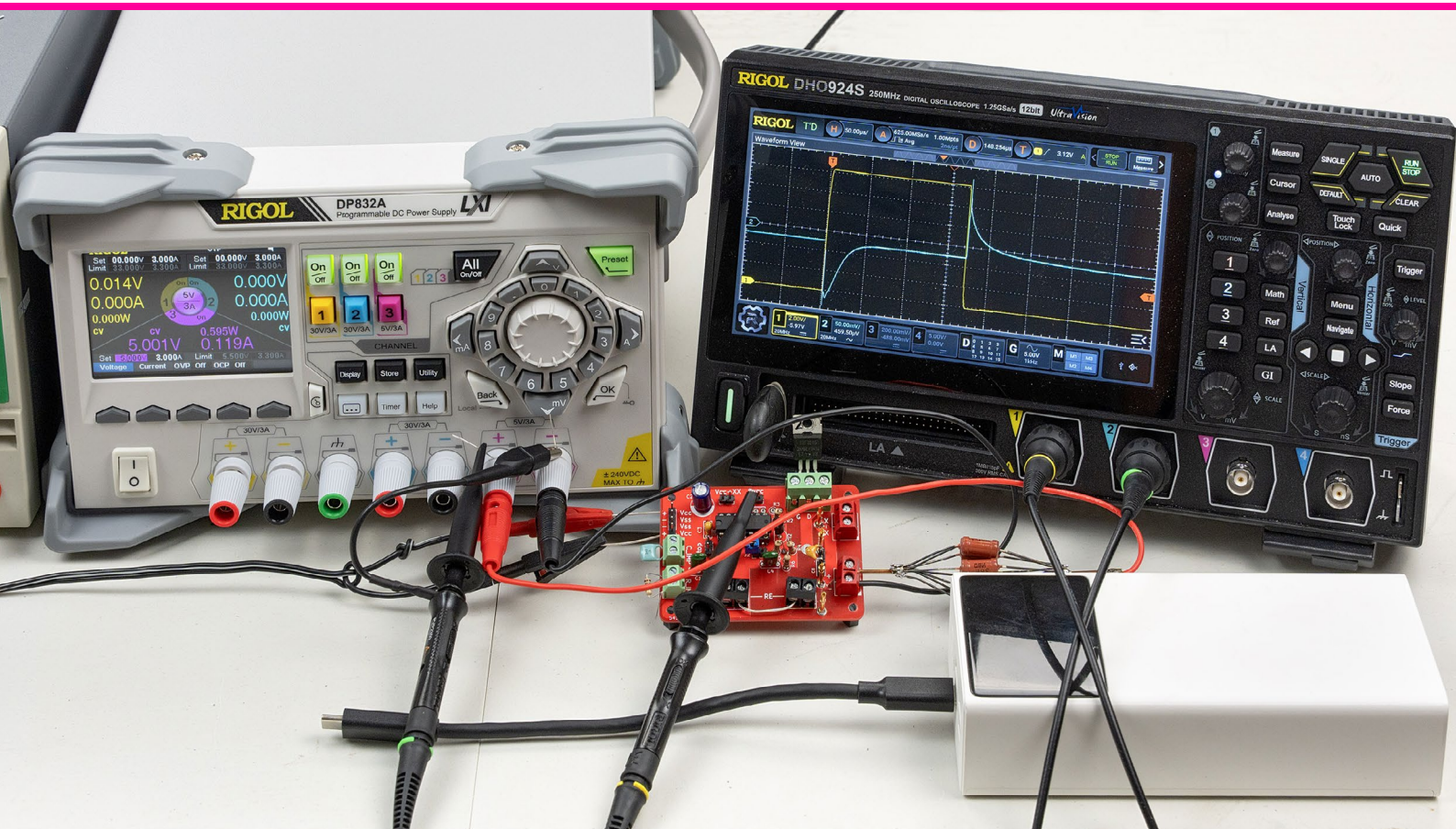
Electrical Characteristics $T_a = 25^\circ\text{C}$ unless otherwise noted

Symbol	Parameter	Test Condition	Min.	Typ.	Max.	Units
Off Characteristics						
BV _{gs}	Gate-Source Breakdown Voltage	$I_g = -1.0\mu\text{A}$, $V_{ds} = 0$	-35			V

J111 / J112 /

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

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



Zasilacze – niedoceniane parametry dynamiczne

O właściwościach zasilaczy decydują nie tylko podstawowe parametry, jak moc, zakres napięć wyjściowych oraz prąd maksymalny. W wielu zastosowaniach istotne są niedoceniane parametry dynamiczne. Artykuł pokazuje różnice parametrów dynamicznych kilku zasilaczy nazywanych laboratoryjnymi.

Testy czterech zasilaczy
Nowy Korad KA3005PS
Stary Korad KA3005P
Rigol DP832A

KUAIQU PS3010
Kwestia szumów wyjściowych
Podsumowanie

Osoby słabo zorientowane chcąc kupić zasilacz do pracowni elektronicznej zwracają uwagę głównie, a niekiedy tylko, na dwa elementarne parametry: na zakres napięć wyjściowych oraz prąd maksymalny.

Bardziej zorientowani są świadomi, że dobry zasilacz warsztatowy czy laboratoryjny powinien mieć szereg istotnych cech. Między innymi dobre właściwości dynamiczne. W poniższym artykule na przykładzie czterech zasilaczy pokażę, jak różne mogą być parametry dynamiczne zasilaczy, słusznie czy zupełnie niesłusznie nazywanych laboratoryjnymi.

Współczesny zasilacz jest skomplikowanym układem elektronicznym. Jednym z jego zadań jest utrzymanie nastawionego napięcia, niezależnie od wartości prądu obciążenia. Najprościej biorąc, wartość prądu obciążenia nie powinna zmieniać napięcia wyjściowego, o ile tylko prąd ten nie przekracza granic dopuszczalnych dla tego zasilacza.

Jeżeli weźmiemy woltomierz i zmierzmy napięcie wyjściowe bez obciążenia i po obciążeniu znacznym prądem, to zwykle cieszymy się, że napięcie praktycznie się nie zmienia – „jest sztywne jak drut”.

Owszem, takie próby wykazują, że stabilizacja napięcia jest skuteczna i że wyjściowa rezystancja statyczna ma małą wartość, rzędu kilku czy kilkunastu miliomów.

Tak, to cieszy, ale to dotyczy warunków statycznych, gdy prądy i napięcia mają ustaloną wartość. Inaczej jest w przypadku szybkich zmian prądu obciążenia. **Fotografia tytułowa** pokazuje przebieg napięcia wyjściowego dobrego profesjonalnego zasilacza (linia niebieska) przy obciążeniu go prostokątnym impulsem prądu – jak pokazuje przebieg żółty.

Lepiej widać to na **rysunku 1**. Na tym i na kolejnych takich rysunkach najważniejsza jest skala napięcia przebiegu niebieskiego, w tym przypadku 50 mV/działkę, oraz skala czasu, w tym przypadku 20 mikrosekund na działkę.

Jak widać, prostokątny impuls prądu obciążenia powoduje obniżenie się napięcia wyjściowego o około 120 miliwoltów, ale tylko na krótki czas poniżej 50 mikrosekund. Z kolei gdy prąd przestaje płynąć, następuje skokowe podwyższenie napięcia wyjściowego o około 150 miliwoltów, ale na mniej więcej 1 mikrosekundę, a potem skok napięcia zmniejsza się od 100 mV do wartości nominalnej, praktycznie w ciągu kilkudziesięciu mikrosekund.

Czy to jest dopuszczalne w dobrym, profesjonalnym zasilaczu kosztującym kilka tysięcy złotych? Czy aby przy takiej cenie napięcie wyjściowe nie powinno być „idealnie gładkie”?

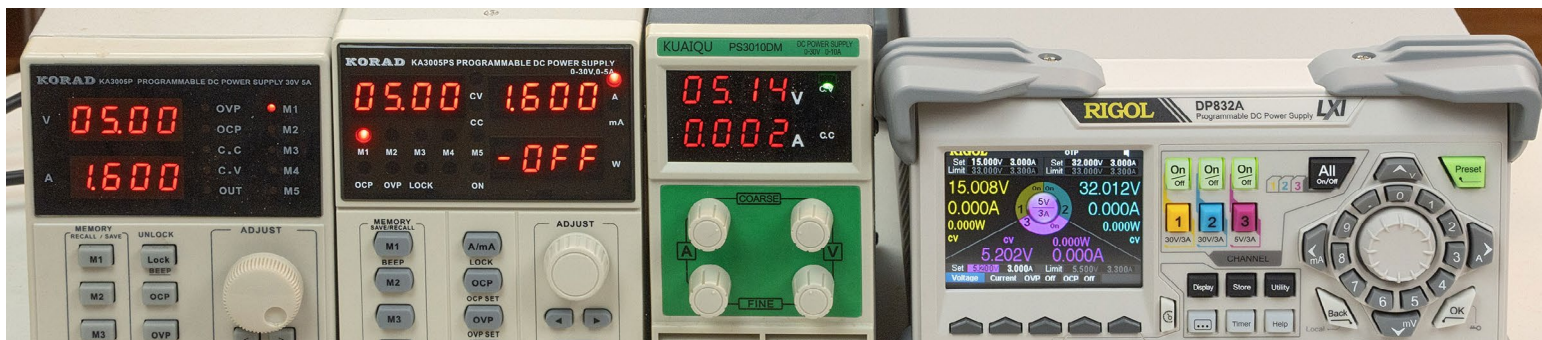
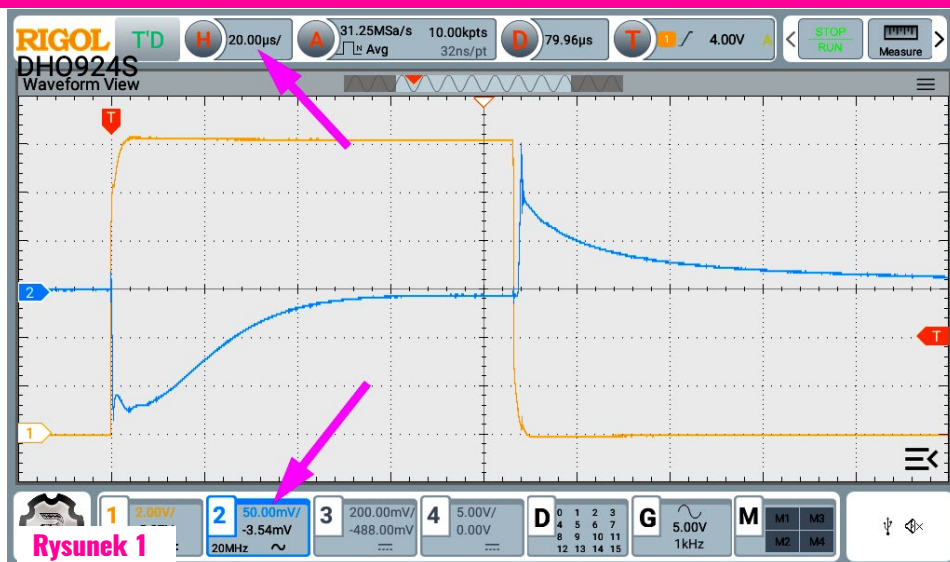
Tak „zafalowania” są nie tylko dopuszczalne, ale wręcz nieuniknione. Takie skoki przy nagłej zmianie obciążenia występują we wszystkich zasilaczach, niezależnie od ceny. Powodem jest fakt, że zasilacz zawiera pętlę (ujemnego) sprzężenia zwrotnego, w której występują elementy wprowadzające jakieś opóźnienie. Najprościej biorąc, takie skoki są nieuniknione właśnie z uwagi na nieusu-

walne opóźnienia w pętli regulacji. W lepszych zasilaczach skoki te są mniejsze i trwają krócej, w gorszych – odwrotnie. W gorszych dodatkowo mogą występować gasnące oscylacje. A w niedopracowanych takie oscylacje mogą się utrzymywać ciągle.

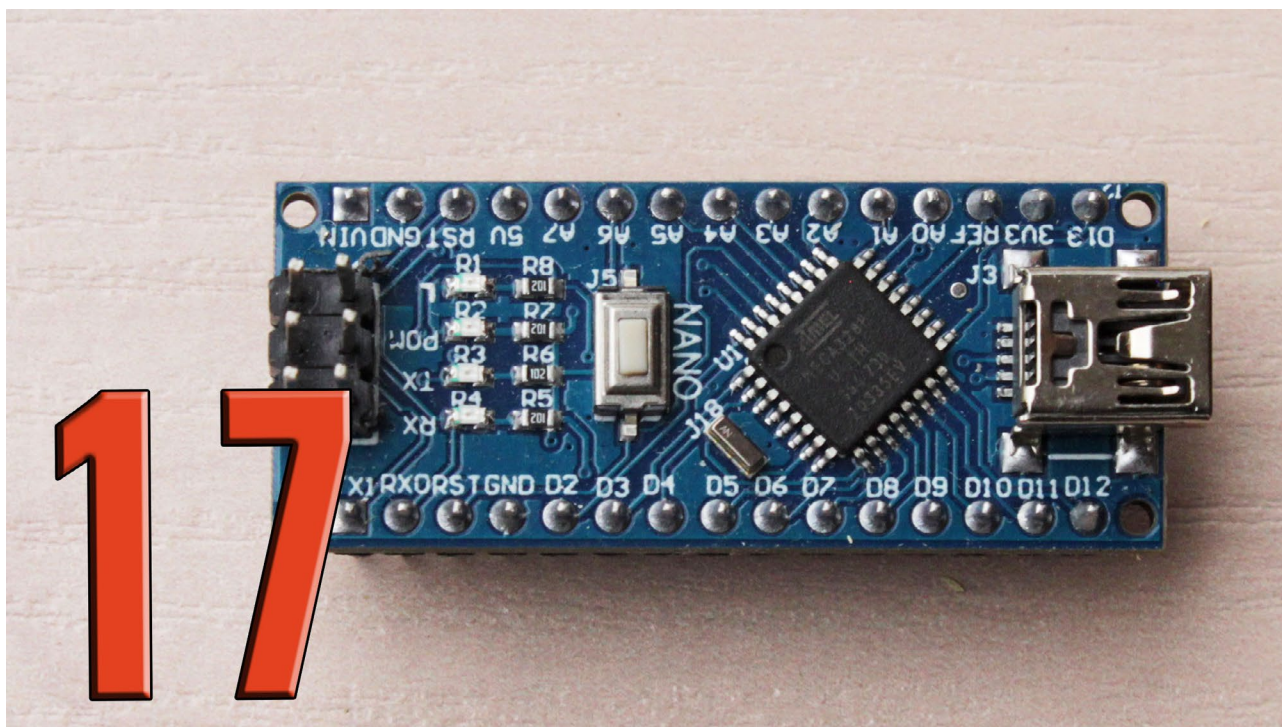
Testy czterech zasilaczy

Aby porównać, jak to wygląda w różnych zasilaczach, trzeba je obciążyć albo takim samym impulsem prądowym, albo impulsem o wartości prądu maksymalnego danego zasilacza. Ja porównałem odpowiedź zasilaczy pokazanych na **fotografii 2**. Od lewej mamy bardzo popularny, cieszący się od lat dobrą opinią „stary” Korad KA3005P. Dalej jest nowa wersja: Korad KA3005PS. Oba 30 V 5 A. Następnie tani impulsowy PS3010 (30 V 10 A), czyli najsilniejszy, 10-amperowy. A z prawej jest Rigol DP832A z dwoma kanałami 30 V 3 A i jednym 5 V 3 A.

Zanim przejdę do testów, uwaga dotycząca wygody użytkowania. Na fotografii 2 widać, że „stare Korady KA3005” i PS3010 mają zdecydowanie za nisko umieszczone zaciski wyjściowe, co jest naprawdę irytujące podczas obsługi. Wady tej nie ma „nowy Korad KA3005PS”, a tym bardziej Rigol DP832A, gdzie te zaciski są umieszczone wyżej.



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



Mikroprocesorowa ośła łączka, część 17

Tworzenie oprogramowania dla mikrokontrolerów to sztuka łączenia kilku elementów. Jednym z nich jest znajomość języka. Drugi to umiejętność tworzenia algorytmów i rozwiązań. Równie istotnym jest aspekt czysto rzemieślniczy, jak posługiwanie się narzędziem.

[Dyrektywy kompilacji warunkowej](#)
[Definiowanie symboli](#)

[Typy zmiennych](#)

Uważam, że nadszedł czas na pewne podsumowanie dotyczące samego języka programowania. W dotychczasowych częściach pewne szczegóły były przedstawione pobieżnie, niektóre odłożone na później. Są również takie, które nie wystąpiły, czas więc na uzupełnienia.

Dyrektywy kompilacji warunkowej

Przy tworzeniu programu, użycie każdego elementu musi być poprzedzone jego określeniem. Rzadko się zdarza, że program składa się z jednego pliku, a nawet w takiej sytuacji wykorzystywane są elementy biblioteczne. Filozofia języka C (jak i każdego innego) sprowadza się do ustalenia ścisłego znaczenia zaledwie kilku elementów – słów kluczowych, które mają określone znaczenie i nie można ich używać

w innym zastosowaniu. Wszystkie pozostałe detale, takie jak funkcje, które występują w programie a nie są napisane naszą ręką to elementy biblioteczne (nikt nie będzie realizował działań pozwalających przykładowo obliczyć e^x , tylko zastosuje istniejącą funkcję biblioteczną `exp ( x )`). Jednak aby z niej skorzystać kompilator musi wiedzieć, że takie coś istnieje i czym to coś jest. Z tego powodu w języku C istnieją pliki określane jako nagłówkowe (plik z rozszerzeniem `.h`), które zawierają jedynie zapowiedzi dodawanych elementów (w przypadku zmiennych oraz funkcji) lub definiują nowe typy zmiennych. Kompilując swój program (lub jego fragment) kompilator musi „przejsć” przez definicję tych elementów. Najprościej jest poinformować go, by chwilowo „przeskoczył” do innego pliku, wczytał go i następnie

powrócił do aktualnie kompilowanego. Na taki ruch pozwala dyrektywa `#include <nazwa pliku>` (jeżeli dotyczy to stałych bibliotek) lub `#include "nazwa pliku"` (jeżeli jest to nasz plik – programista również może tworzyć własne pliki o charakterze bibliotecznym, realizujące jakieś często wykorzystywane działania). Nie ma ograniczenia, by dołączany plik (przez `#include`) nie zawierał dołączeń kolejnych plików. To wnosi możliwość wielokrotnego definiowania tego samego elementu. Aby temu przeciwdziałać, język C jest wyposażony w kompilacje warunkowe. Klasyczny przykład pokazuje **listing 1** (zaczepnięty z pliku zawierającego obsługę alfanumerycznych modułów LCD):

```
#ifndef _LCDModule_
#define _LCDModule_
#include "lcdequip.h"
#endif
```

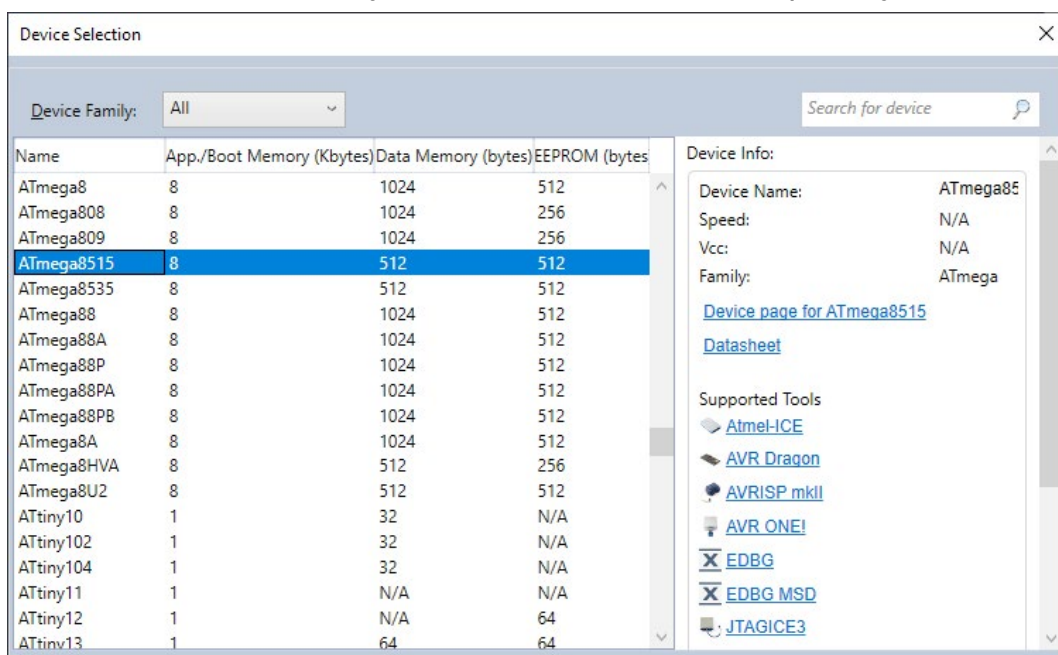
Mamy tu zastosowanie `#ifndef symbol` (jeżeli **nie jest** zdefiniowany symbol) i w przypadku braku jego definicji, kolejne zapisy są standardowo kompilowane. W przeciwnym wypadku wszystkie wiersze, aż do `#endif`, są pominięte. To nie wyczerpuje możliwości kompilacji warunkowych. Wariantem przeciwnym do powyższego jest `#ifdef symbol`. Tutaj warunek jest spełniony, jeżeli specyfikowany symbol jest zdefiniowany (przez `#define`). Kompilacje warunkowe można zastosować w dowolnym miejscu. Przydatne zastosowanie pokazuje **listing 2**. Jest mikrokontroler AT90S8515 (obecnie już przestarzały), który ma możliwość przyłączenia zewnętrznej pamięci RAM lub w jej miejsce dowolnego kontrolera (przykładowo układu transmisji szeregowej 16C450,

```
static void HardwareInit ( void )
{
    TIMSK = ( 1 << TOIE0 ) ;
    TCCR0 = ( 1 << CS01 ) ;
    TCNT0 = 0 ;
#ifdef __AVR_ATmega8515__
    MCUCR = ( 1 << SRE ) | ( 1 << SRW10 ) | ( 1 << ISC01 ) ;
    GICR = ( 1 << INT0 ) ;
#else
    MCUCR = ( 1 << SRE ) | ( 1 << SRW ) | ( 1 << ISC01 ) ;
    GIMSK = ( 1 << INT0 ) ;
#endif
} /* HardwareInit */
```

Listing 2

wejścia `INT0`). Wymaga to nadania określonej wartości dla odpowiednich rejestrów, by aktywować obsługę zewnętrznej pamięci. Jego następca ATMEGA8515 również ma takie możliwości, ale konfigurowanie jej w szczegółach jest odmienne (inne bity i inne rejestry). Można stworzyć program, który zrealizuje to poprawnie w obu przypadkach, jak pokazuje listing 2. Przy tworzeniu projektu dla określonego mikrokontrolera specyfikowany jest jego model (przykładowo jak na **rysunku 1**). To automatycznie definiuje symbol określający zastosowany mikrokontroler, który można wykorzystać we własnym programie.

Powyższy przykład pochodzi ze środowiska AVR Studio – poprzednika oprogramowania narzędziowego Atmel Studio, który już nie wspiera mikrokontrolera AT90S8515, ale zasada pozostaje nadal ważna.



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

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

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, odpowiesz](#).



Czy elektryczność leczy, czy szkodzi?

(...) Kolega odkrył że ma napięcie elektryczne skóry na poziomie 450 mV (...), a jego znajomy (...) ma pomiary na poziomie 10–30 mV. Ja (...) miałem około 370 mV (...) Czy to możliwe, że ofiary broni „elektromagnetycznej” mają wyższe napięcie elektryczne ciała? Czy robił Pan kiedyś takie pomiary? (...)

To niestety nie jest żart! Pytanie dotyczy rzekomej „broni energii sterowanej”, wykorzystującej fale elektromagnetyczne do szkodenia człowiekowi i jest częścią e-maila, który trafił do mojej skrzynki:

Wymieniłem z Autorem pytania kilka e-maili, zadałem kilka, a raczej kilkanaście pytań i okazało się, że chodzi o jedną z teorii spiskowych, w którą wierzy grupka ludzi. Ludzi, którzy najwyraźniej nie mają pojęcia o elektronice, elektromagnetyzmie ani o elektryczności. Wszystko wskazuje, że ich wiedza techniczna jest bardzo bliska zeru, natomiast osoby te szczerze i niewzruszenie wierzą w teorie, które ktoś im podsunął. Nie tylko sami wierzą, ale także poświęcają środki i czas, żeby rozpowszechniać swoją wiarę i szukać nowych „wyznawców”.

Czy słyszał Pan o broni energii skierowanej?



Do kontakt@piotr-gorecki.pl

Na tę wiadomość udzielono odpowiedzi lub przesłano ją dalej.

Dzień dobry,

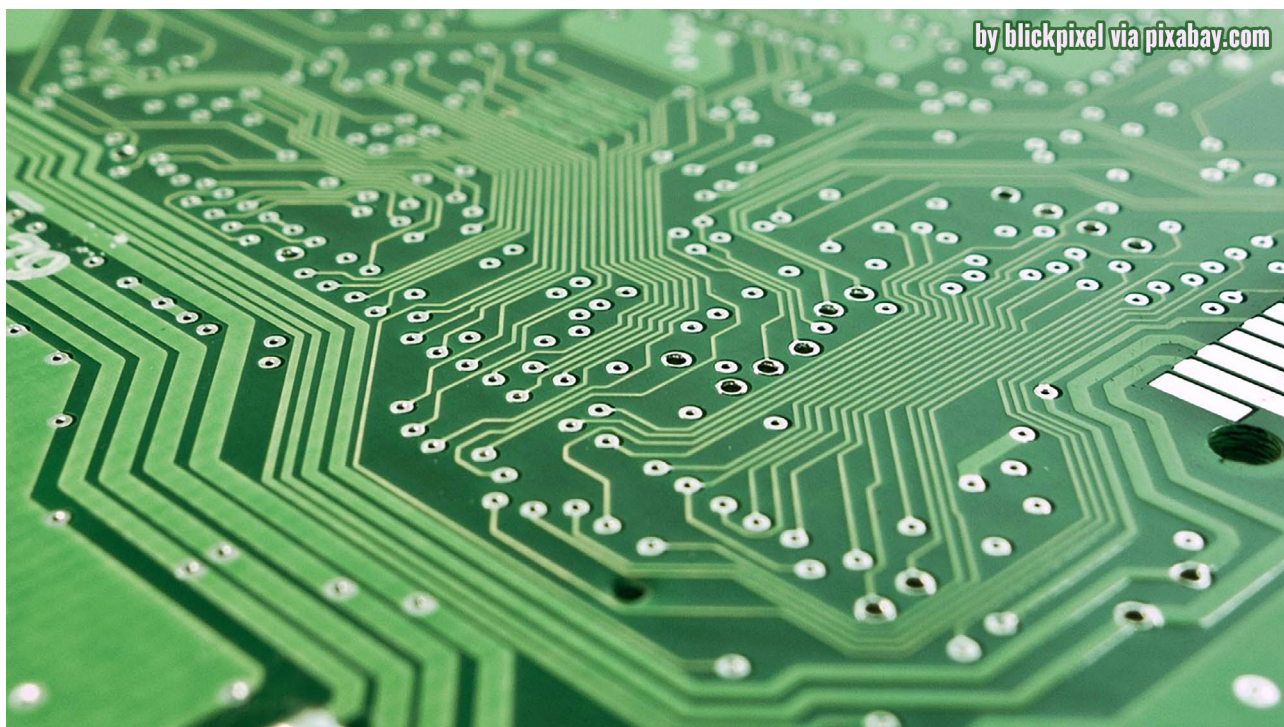
Mam pytanie związane z nielegalnymi eksperymentami z bronią energii skierowanej, a konkretnie z falami em. Kolega odkrył że ma napięcie elektryczne skóry na poziomie 450mV (twierdzi że jest ofiarą tej broni), a jego znajomy, który nie jest ofiarą ma pomiary na poziomie 10-30mV. Ja sobie kiedyś mierzyłem i miałem około 370mV ale nie miałem pojęcia że to dużo bo nikomu innemu nie mierzyłem.. Czy to możliwe że ofiary broni "elektromagnetycznej" mają wyższe napięcie elektryczne ciała? Czy robił Pan kiedyś takie pomiary i jeśli tak to jakie one były? Według podręcznika do biofizyki napięcie generowane przez ludzkie ciało nie przekracza 71mV (91mV). Te pytania i odpowiedzi na nie są ważne dla

Ps. Oglądałem Pana filmy o oscyloskopach i miernikach cyfrowych. Mam zamiar wykryć bron energii skierowanej przy pomocy oscyloskopu i cewki indukcyjnej lub hallotronow ze wzmacniaczem sygnału operacyjnego. Ten mail absolutnie nie jest żartem. Traktujemy nasze problemy bardzo poważnie. Będę wdzięczny za wszelką pomoc Łączę wyrazy szacunku

Rysunek 1

← Odpowiedz ← Odpowiedz wszystkim → Prześlij dalej ...

sob. 09.08.2025 08:17



Ścieżki i przelotki na płytkach drukowanych

W czasopiśmie ZE 2/2025 w „Listach Czytelników” poruszona została ważna kwestia szerokości ścieżek płytek drukowanych. Ten artykuł pokazuje, że często niestety lekceważone parametry ścieżek i przelotek mogą mieć kluczowe znaczenie dla prawidłowego działania urządzenia elektronicznego.

Obliczenia

Przepalenie ścieżki

Zwiększanie obciążalności prądowej ścieżek

Ścieżki na warstwach wewnętrznych

Przelotki

W czasie projektowania płytek drukowanych należy uwzględnić wiele parametrów. Także dotyczących szerokości ścieżek. Ścieżki i przelotki mają wpływ nie tylko na poprawne działanie zaprojektowanego urządzenia elektronicznego, ale także na wielkość płytki drukowanej. Wielkość płytki drukowanej z kolei przekłada się na koszt jej wykonania. We współczesnych skomplikowanych urządzeniach elektronicznych nieraz trudno jest pogodzić te wszystkie czynniki. Projekt płytki drukowanej często powstaje z uwzględnieniem wielu kompromisów.

Stopniowanie szerokości ścieżek. Jednym z parametrów płytki drukowanej jest szerokość ścieżek. W wielu prostych projektach parametr ten może nie

mieć większego znaczenia. Jednak zbyt mała szerokość ścieżek na płycie drukowanej w najlepszym przypadku spowoduje nadmierne ich nagrzewanie. Zwiększa to rezystancję takich ścieżek, co może spowodować zmianę parametrów pracy elementów na takiej płycie, a nawet usterkę. W skrajnym wypadku może dojść do przepalenia się ścieżki na płycie, a to może być powodem dalszych uszkodzeń. Dlatego należy wyrobić sobie nawyk stopniowania szerokości ścieżek płytek drukowanych i stosować go w praktyce.

Odstępy między ścieżkami. Z szerokością ścieżek jest związany kolejny parametr, jakim są odstępy ścieżek między sobą, polami lutowniczymi i elementami na płycie drukowanej.

Również ten parametr w prostych projektach zwykle nie ma większego znaczenia. W przypadku ścieżek między którymi występuje wysokie napięcie, konieczne jest już zastosowanie odpowiednich odległości, a nawet szczelin i wycięć w płytce drukowanej. Dodatkowe obostrzenia i wymagania dotyczą izolacji galwanicznej od sieci energetycznej, to jednak oddzielny, szeroki temat. Przy niskich napięciach odległości między ścieżkami w prostych projektach mogą nie mieć większego znaczenia. W urządzeniach pracujących z dużą częstotliwością przetwarzania sygnałów stosuje się tak zwane pary różnicowe. Wówczas odległość ścieżek pary różnicowej ma znaczenie, pomimo że napięcia są małe.

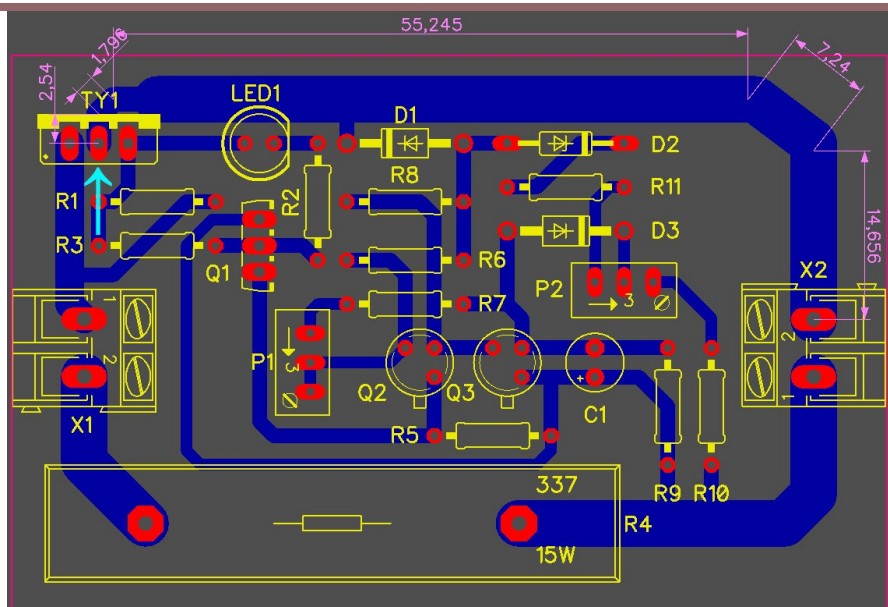
Wielkości i odległości związane są też z impedancją falową, co dotyczy bardzo szybkich sygnałów i układów w.c.z. To też oddzielne, szerokie tematy.

Klasy sieci i reguły projektowe. Większość współczesnych programów dla elektroników ma funkcję umożliwiającą zdefiniowanie tak zwanych **klas sieci** lub **klas połączeń**. Czyli zdefiniowanie już na etapie edycji schematu parametrów ścieżek płytki drukowanej dla poszczególnych sieci połączeń. Sieciom tym możemy nadawać wybrane nazwy, na przykład GND. Następnie należy zdefiniować szerokość ścieżki na płycie dla danej sieci połączeń, jej odstępy od innych ścieżek, jak również przypisać im odpowiednie rozmiary przelotek. Taką sieć możemy przypisać do określonej warstwy płytki drukowanej, na przykład dolnej. Później na etapie ręcznego lub automatycznego trasowania ścieżek program do edycji płytek drukowanych stosuje zdefiniowane klasy sieci. Jest też możliwość zdefiniowania tak zwanych reguł projektowych, które również są uwzględniane podczas trasowania połączeń. Programy do edycji płytek nie są jednak idealne i może się zdarzyć, że najczęściej z powodu dużego zagęszczenia elementów i ścieżek na płycie

niektóre połączenia nie zostaną wytrasowane.

Obliczenia

Jak w takim razie obliczyć wymaganą szerokość dla



Rysunek 1

prądy o dużej wartości? Tutaj z pomocą przychodzą twórcy programów EDA, dołączając do swoich programów kalkulatory. Dostępne są również kalkulatory szerokości ścieżek i innych paramentów płytek drukowanych działające online za pośrednictwem przeglądarki internetowej. Umożliwiają one obliczenie w prosty sposób wymaganej szerokości dla ścieżek płytki drukowanej. Do obliczeń przykładowej ścieżki płytki drukowanej przyjąłem standardową grubość warstwy miedzi laminatu – 35 μm , długość ścieżki – 82 mm i maksymalny prąd przez nią płynący do 10 A, przy wzroście temperatury o 20°C. Jest to ścieżka na płycie drukowanej z naniesionymi wymiarami, widoczna na **rysunku 1**. Jej łączna długość między otworami pól lutowaniczych to około 81,5 mm. Jest to płytka do automatycznego odłączania ładowania akumulatora samochodowego.

KiCad. W programie KiCad kalkulator ten można uruchomić bezpośrednio z Menadżera projektu, klikając na ikonę **Podręczny kalkulator elektronika**. Uruchomiony kalkulator możemy zobaczyć na **rysunku 2**. Po lewej stronie okna kalkulatora mamy rozwijane drzewo parametrów, które możemy obliczyć – między innymi szerokość ścieżki.

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 2

Istotnym elementem każdego urządzenia zawierającego mikroprocesor lub mikrokontroler jest pamięć operacyjna, której zadaniem jest przechowywanie bieżących, ciągle zmieniających się informacji. Ten rodzaj elementów na przestrzeni lat rozwijał się bardzo dynamicznie.

Koncepcja budowy pamięci RAM „Starożytne” rozwiązania pamięci RAM

Konstrukcje urządzeń bazujących na mikroprocesorach lub mikrokontrolerach muszą mieć możliwość zapisywania oraz odczytywania gdzieś informacji wynikających z procesu przetwarzania danych. W poprzedniej części zostały opisane pamięci typu ROM w różnych ich odmianach technologicznych. Ich charakterystyczną cechą jest to, że są przeznaczone do przechowywania stałych informacji, jak przykładowo kod programu mikroprocesora. Choć nowoczesne rozwiązania tych pamięci umożliwiają zapis nowych danych, nie nadają się one do budowy pamięci operacyjnych. Potrzebny jest tu nowy rodzaj określany jako RAM (ang. **R**andom-**A**ccess **M**emory – pamięć o dostępie swobodnym), pozwalający sięgnąć bezpośrednio do dowolnej komórki.

Półprzewodnikowe pamięci RAM Pamięci dwuportowe

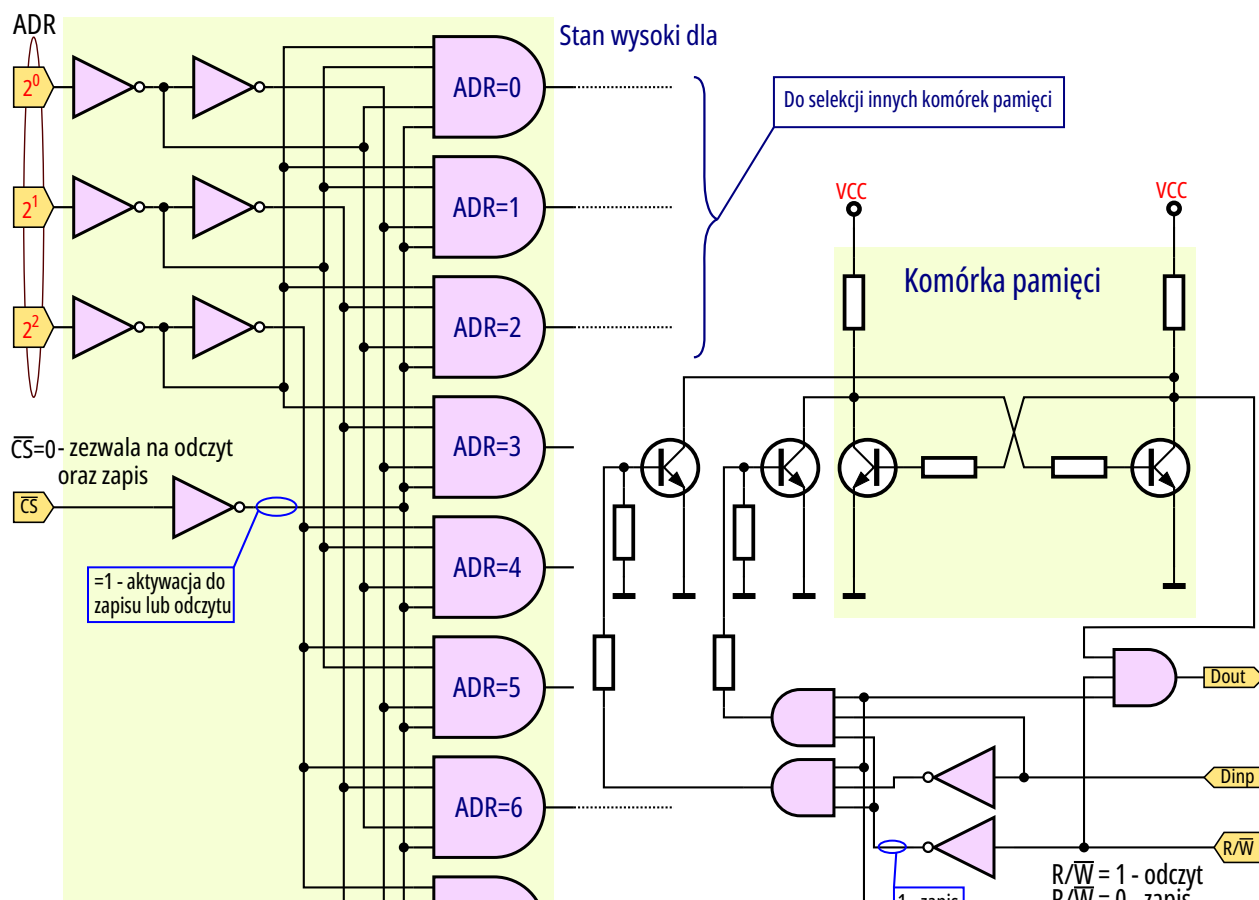
Koncepcja budowy pamięci RAM

Sama pamięć RAM jest zbiorem określonej liczby przerzutników o dwóch stanach stabilnych (bistabilnych). Ich liczba stanowi pojemność pamięci. Aby wszystko w prosty sposób współgrało ze sobą, wspomniane przerzutniki muszą być wyposażone w dodatkowe podzespoły. Tu nieodzownym elementem jest dekodery adresowy, którego zadaniem jest wyselekcjonowanie pojedynczej komórki. W stosunku do pamięci stałych (ROM), gdzie występowała szyna adresowa (pozwalała wskazać komórkę pamięci w realizowanej operacji), szyny danych („wydającej” zapisane dane) i sygnału aktywującego, tu dochodzi szyna danych „dostarczająca” dane do zapisu (początkowo dane wejściowe były rozdzielone od danych wyjściowych)

oraz sygnał określający wykonywaną operację: zapis albo odczyt. Ideę ilustruje poglądowy **rysunek 1**. Wykorzystany jest tutaj przerzutnik bistabilny (o dwóch stanach stabilnych). Takie rozwiązanie komórki można określić jako statyczne. Jeżeli nie nastąpi zapis nowej wartości, dotychczas zapisana może być przechowywana w nieskończoność (pomijając zanik zasilania). Nie jest to jedyne możliwe rozwiązanie. Innym wariantem jest ładowanie i rozładowanie kondensatora. Jeżeli kondensator ma ładunek elektryczny jest to interpretowane jako jeden stan logiczny, w przypadku kondensatora rozładowanego mamy przeciwny stan logiczny. Ponieważ kondensatory mają tendencję do rozładowania się (dodatkowo sprawdzenie naładowania również uszczupla jego ładunek), taka komórka nie może przechowywać danych w nieskończoność. Po upływie odpowiedniego czasu dane muszą być odświeżone. Różne warianty rozwiązania samej komórki przechowującej zapisany stan doprowadziły do rozdziału pamięci RAM na dwie podgrupy: pamięci statyczne (S-RAM – **S**tatic **R**andom-**A**ccess **M**emory) oraz dynamiczne (D-RAM – **D**ynamic **R**andom-**A**ccess **M**emory). Bieżący artykuł jest poświęcony pamięciom statycznym, natomiast o pamięciach

dynamicznych będę pisał w kolejnych częściach, gdyż ich obsługa znacząco odbiega od obsługi pamięci statycznych.

Aby określić komórkę w pamięci, którą/do której chcemy odczytać/zapisać dane, należy na szynie adresowej określić wartość adresu (wypracować odpowiednią kombinację bitową tworzącą adres). W stosunku do pamięci ROM (którą można jedynie odczytać), tu dodatkowo możliwa jest operacja zapisu. Jest ona identyfikowana stanem sygnału R/W. Zero logiczne oznacza operację zapisu oraz, w przeciwieństwie, jedynka logiczna oznacza czynność odczytu (odnosząc się do współczesnych rozwiązań, gdyż w przeszłości wcale tak nie musiało być). Realizując zapis do pamięci, na wejściową szynę danych należy podać zapisywaną informację. Wystawienie tych wszystkich sygnałów nie oznacza wykonania wymaganej operacji. Układ musi być jeszcze aktywowany do działania: na wejściu typu *chip select* należy podać stan zera logicznego (istnieją pamięci S-RAM, gdzie występują dodatkowe sygnały typu *chip select* ze stanem aktywnym jako jedynka logiczna, ale są one w kategorii dodatkowych, natomiast „główny” *chip select* ma stan aktywny jako zero logiczne). To uruchamia całą lawinę zdarzeń prowadzącą do

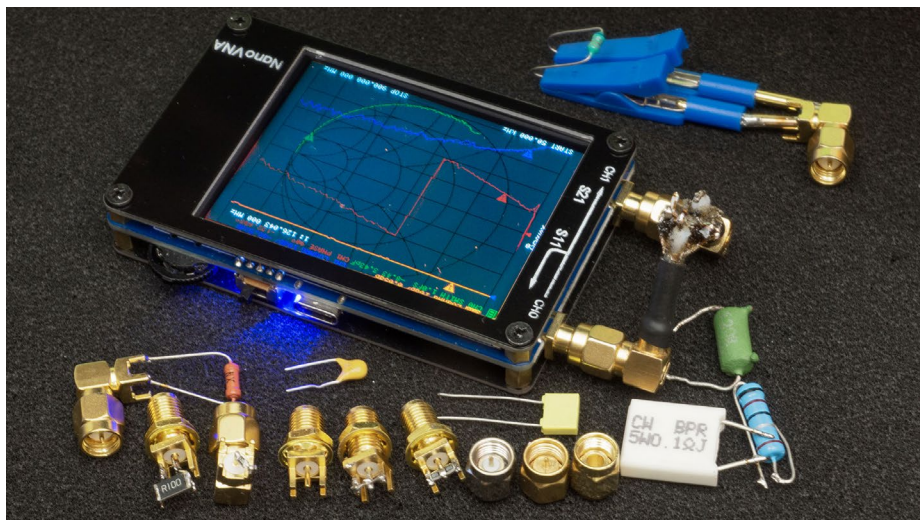


Uwaga! To jest demonstracyjny (niepełny) egzemplarz czasopisma ZE.

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

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](#).



NanoVNA – gdzie i którą wersję kupić?

Jeden z Patronów napisał: (...) Nawiązując jeszcze do zakupu nanoVNA: czy mógłby Pan polecić konkretną wersję i dystrybutora? Trochę zgłupiałem, przeglądając różne oferty. Niektóre egzemplarze mają nieekranowaną część w.cz., a wyniki testów sugerują, że wpływa to na osiągi analizatora. Pozdrawiam

Dostępnych jest kilka wersji, które omówię w oddzielnych artykułach. Kto nie chce wgłębiać się w szczegóły, niech po prostu kupi w jednym z tych źródeł wersję o akceptowalnej dla niego cenie.

Jeśli ktoś jeszcze nie wykorzystywał nanoVNA, na początek, do zapoznania się ze sprzętem naprawdę

TLDR (too long; didn't read), czyli najkrócej dla tych niecierpliwych, którzy chcą kupić, a nie czytać: zajrzyj do godnego uwagi [chińskiego sklepu internetowego Zeenko](#), gdzie można znaleźć oferty różnych wersji nanoVNA w różnych cenach – **rysunek 1**. Ten sam sprzedawca

Zeenko Store

99.3% Pozytywna recenzja | 2.4K obserwatorzy

+ OBSERWUJ

Skontaktuj się teraz

Q Szukaj w tym sklepie

10KHz ~ 1,5GHz NanoVNA-H 2,8-calowy analizator sieci z ekranem... ★★★★★ 208 sprzedano 141,59zł	Originalny wyświetlacz Zeenko NanoVNA-H4 4.0 "wersja 4.3..." ★★★★★ 223 sprzedano 279,19zł 279,38zł	Originalny Hugen 50 kHz ~ 6,3 GHz tinyVNA LiteVNA 64 4... ★★★★★ 154 sprzedano 565,59zł 565,97zł	LibreVNA 100kHz - 6 GHz w pełni 2-portowy analizator sieci... ★★★★★ 52 sprzedano 2360,39zł 2362,63zł	Originalny analizator częstotliwości fazowych Zeenko ... ★★★★★ 41 sprzedano 190,19zł 190,21zł
SMA attenuator 30dB DC-6GHz 2W Max 3dB / 6dB / 10dB / 20dB / 30dB	SAA2 tinyVNA 50k-3GHz 内置电池 金属外壳			

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

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

Schemat ideowy testera

Nieprawidłowo wykonany fizyczny model testera

Testy właściwości dynamicznych buforów cyfrowych

Wnioski końcowe

Artykuł ten ze względu na swoją obszerność został podzielony na trzy części. W pierwszym odcinku cyklu przedstawiłem moje wyniki testów dotyczących statycznych właściwości buforów cyfrowych z różnych rodzin TTL. Testy te rzuciły światło na pewne ograniczenia i problemy kompatybilności połączeń różnych rodzin układów cyfrowych.

W drugiej części skupiłem się na szczegółowym wyjaśnieniu terminacji szeregowej, która z pewnością będzie wieść prym w moim komputerze TTL. Przedstawiłem wady i zalety tego jedno-

stronnego dopasowania. Mam wielką nadzieję, że niejednemu Czytelnikowi wykład ten rozwiązał wiele wątpliwości w związku z czym terminacja ta zyskała przy bliższym poznaniu.

Poniżej część trzecia, ostatnia, w której przedstawię wyniki testów właściwości dynamicznych szybkich buforów cyfrowych. Dzięki analizie tych testów przeprowadzonych przy użyciu płytki, która całkowicie odbiega od prawidłowych zasad projektowania dla szybkich sygnałów, przekonam się, jakie dodatkowe problemy stwarzają płytki PCB, które są niewłaściwie zaprojektowane.

Schemat ideowy testera

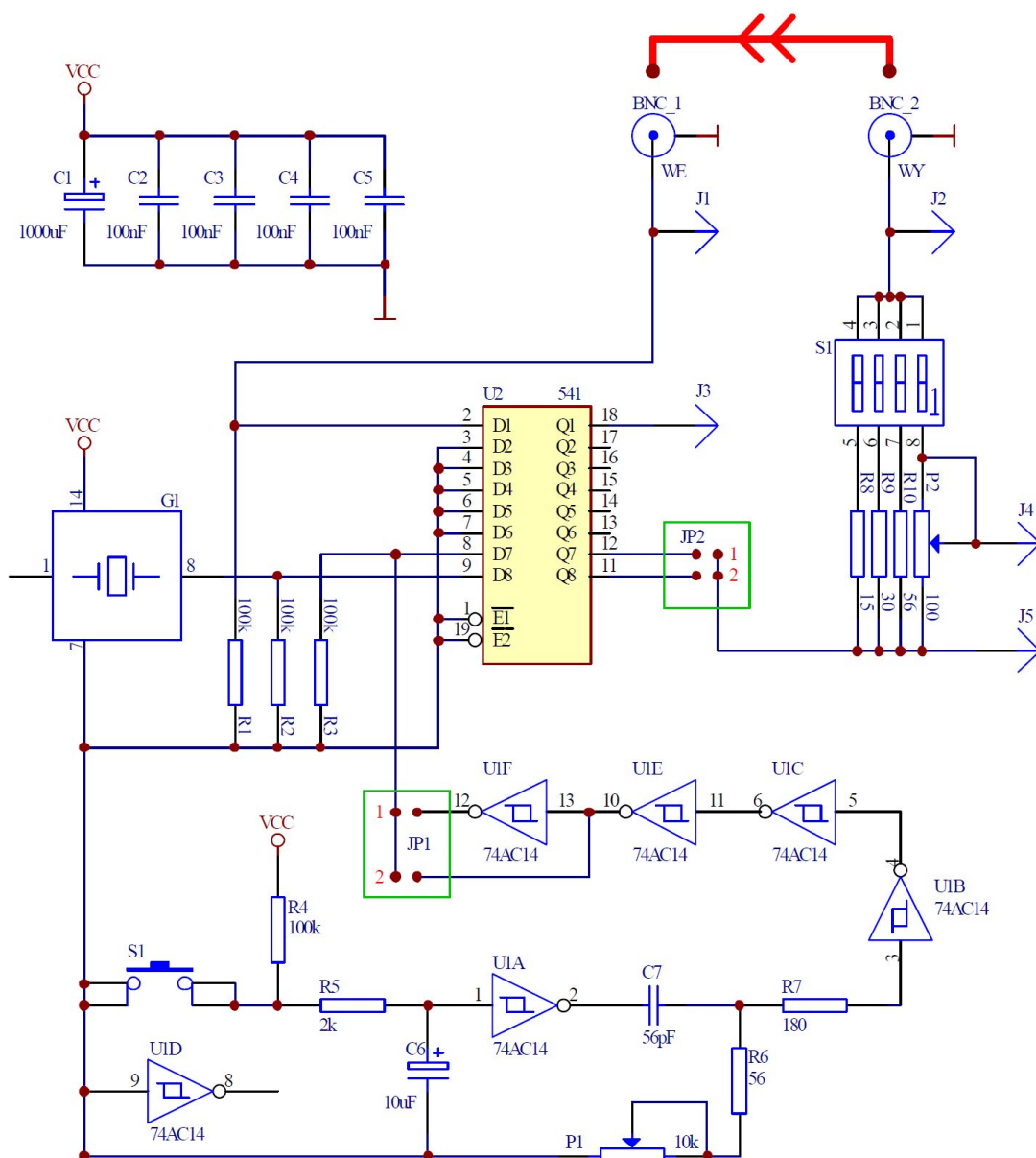
Aby przeprowadzić testy właściwości dynamicznych, najpierw musiałem zaprojektować i zbudować urządzenie według schematu pokazanego na **rysunku 1**.

Obwód wykorzystujący układ 74AC14 służy do generowania, krótkich impulsów prostokątnych o stromych zboczach. Przyciskiem **S1** wywołujemy krótki impuls szpilkowy. Przycisk ten jest podłączony do obwodu eliminującego drgania styków składających się z elementów **R4**, **R5** oraz **C6**. Rezystor **R5** dodatkowo ogranicza prąd ładujący kondensator **C6**. Stała czasowa obwodu **R4+R5C6** jest asymetryczna. Wynosi w przybliżeniu 20 ms lub 1 s. Obwód ten okazał się niezbędny, bo zdarzały się nieprawidłowe sytuacje, a ja chciałem mieć wygenerowany tylko i wyłącznie pojedynczy impuls.

Obwód z inwerterami Schmitta **U1A**, **U1B** oraz z kondensatorem **C7**, helitrimem **P1** i rezystorami **R6** i **R7** tworzy klasyczny układ różniczkujący. Gdy na wyjściu bramki **U1A** zmieni się stan z logicznego zera na logiczną jedynkę, na wejściu bramki **U1B** powstanie krótki dodatni impuls, którego zbocza, szczególnie zbocze opadające, oczywiście nie przypominają typowego impulsu prostokątnego. Dlatego zastosowane bramki są z przerzutnikami Schmitta – dzięki tej właściwości na wyjściu inwertera **U1B** uzyskamy typowy impuls prostokątny o stromych zboczach, w tym wypadku

impulsu uzależniony jest od stałej czasowej obwodu **R6+P1C7**. Rezystor **R7** zabezpiecza obwody wejściowe bramki **U1B** przed nadmiernym prądem spowodowanym gwałtownym ładowaniem/rozładowaniem kondensatora **C7**. Pozostałe inwertery **U1C**, **U1E** oraz **U1F** wykorzystywane są do transportowania tego impulsu na złącze **JP1**.

Złącze te podaje sygnał prosty lub zanegowany na jedno z wejść testowanego ośmiokrotnego bufora **U2**. Dzięki temu możemy testować bufor danych proste (541) lub zanegowane (540). Zawsze na wyjściu takiego bufora domyślnym stanem logicznym będzie stan niski. Analizując schemat z rysunku 1 możemy również zauważyć, że do kolejnego z wejść ośmiokrotnego bufora jest podłączony generator kwarcowy.



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

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

ZROZUMIEĆ **ELEKTRONIKĘ** z Piotrem Góreckim

ZE 10/2025

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>**

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.