

ELEKTRONIKA

dla wszystkich

nr 12/2025 (359) • grudzień • www.elportal.pl

DIY PLUS
tylko dla prenumeratorów

PROJEKTY dla elektroników

- ▶ Zagraj we własną grę kółko i krzyżyk, część 1
- ▶ Podwójny wzmacniacz RF przeznaczony do wzmacniania sygnałów z generatorów sygnału
- ▶ Termostat z termoparą typu K
- ▶ Źródło Czasu Wi-Fi dla zegarów GPS

DIY dla wszystkich

- ▶ Zastosowanie radaru dopplerowskiego w systemie ADAS, część 4
- ▶ Urządzenie do cheating'u z użyciem ChatGPT, o małych rozmiarach

TUTORIALE

- ▶ Łączność w różnych mediach, część 2 – ziemia
- ▶ Ekscytacje Maxa: Migające diody LED i śliniacy się inżynierowie, część 27
- ▶ Chirurgia obwodowa: Zniekształcenia i obwody zniekształcające, część 1

Kółko i krzyżyk na bramkach logicznych

Podwójny wzmacniacz RF



Pomocna dłoń



automatykaB2B.pl

EP.com.pl

Największy
portal dla
elektroników
konstruktorów

eprasa.pl 030b21959



FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

przełączniki
półprzewodniki
złącza
przekładniki
radiatory
obudowy
i wiele więcej...

www.piekarz.pl

TAWOIA Glass (szkło kwarcowe)

<https://sklep.avt.pl/pl/menu/tawoia-glass-4505.html>



BESTSELLERY sklepu AVT – sklep.avt.pl

3 unikalne
serie
gniazdek
i
włączników

Rabat dla Czytelników EdW
przy zakupie podaj kod **EdW2505GW**

Kod ważny do 30.09.2025

Rabat dla Prenumeratorów EdW
przy zakupie podaj numer prenumeraty

-5%

-10%

Ceramic Loft (ceramika)

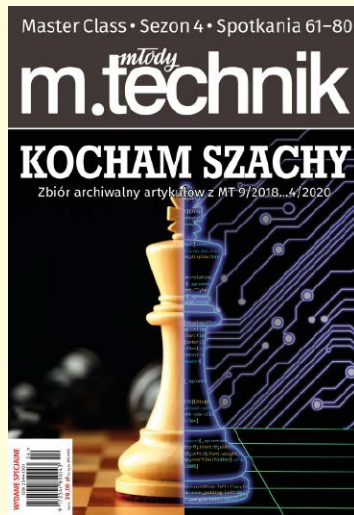
<https://sklep.avt.pl/pl/menu/seria-ceramic-loft-4190.html>



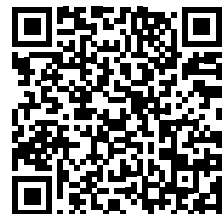
Retro PRL (bakelit)

<https://sklep.avt.pl/pl/series/retro-prl-3237.html>



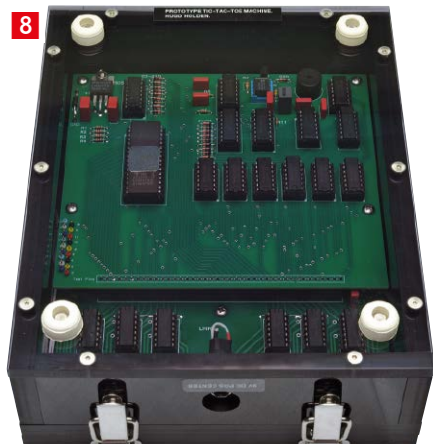


pakiet promocyjny
KOCHAM SZACHY
7 e-booków z rabatem
50%



Dla prenumeratorów – 30% rabatu!

Promocja internetowa – w formularzu zamówienia online zaznacz pole „Jestem prenumeratorem wydawnictwa AVT, kupuję ze zniżką” i podaj swój numer prenumeraty.



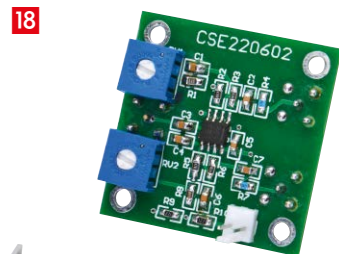
8

Projekty dla elektroników:

Zagraj we własną grę kółko i krzyżyk, część 1	8
Podwójny wzmacniacz RF przeznaczony do wzmacniania sygnałów z generatorów sygnału	18
Termostat z termoparą typu K	22
Źródło Czasu Wi-Fi dla zegarów GPS	32

Tutoriale:

Łączność w różnych mediach, część 2 – ziemia	42
Chirurgia obwodowa: Zniekształcenia i obwody zniekształcające, część 1	52
Ekscytacje Maxa: Migające diody LED i śliniacy się inżynierowie, część 27	58
Edukacja w EdW dla szkół i uczelni: Wykład 36 – Szumy w układach elektronicznych	65



18

22

DIY dla wszystkich:

Zastosowanie radaru dopplerowskiego w systemie ADAS, część 4	75
Urządzenie do cheating'u z użyciem ChatGPT, o małych rozmiarach	79

32

Elektronika dla Wszystkich – Junior:

Osiemnaste spotkanie z najmłodszymi pasjonatami elektroniki	81
Na zdjęciu na okładce Dawid, Szkoła Podstawowa nr 86, Wrocław	

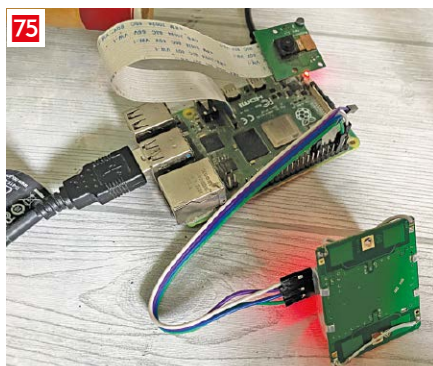
DIY PLUS tylko dla prenumeratorów zamawiających prenumeratę na www.UlubionyKiosk.pl

Bezprzewodowy nadajnik ściemniacza LED 434 MHz

– kompatybilny z Arduino	90
Odbiornik bezprzewodowy 434 MHz – ściemniacz LED	
– kompatybilny z Arduino	90

Rubryki stałe:

Od redakcji	5
Pocztą	6
Prenumerata	7



75

A za miesiąc w styczniowym EdW



Głośniki z misek sałatkowych – tanie, łatwe i efektywne

Brzmi jak żart, ale to działa! Kilka bambusowych misek z IKEI, para głośników koaksjalnych i odrobina zręczności wystarczą, by stworzyć niezwykle stylowe kolumny o zaskakująco dobrym brzmieniu. Sferyczna obudowa eliminuje dyfrakcję i niepożądane rezonanse, a wbudowany bas-refleks pozwala uzyskać czysty dźwięk już od 70 Hz. Całość można zmontować prostymi narzędziami, bez warsztatu stolarskiego. To niedrogi projekt dla majsterkowiczów, który łączy zabawę z akustyką i świetnie prezentuje się zarówno na biurku, jak i w salonie.

Stanowisko do testowania głośników i zwrotnic

Prosty, ale niezwykle funkcjonalny układ, który pozwala szybko i bezpiecznie przetestować głośniki – zarówno nisko-, jak i wysokotonowe. Dzięki wbudowanemu wzmacniaczowi, przełączanym tłumikom i zasilaniu fantomowemu można mierzyć impedancję, charakterystykę częstotliwościową i fazową, a nawet kalibrować zwrotnice. Współpracuje z darmowym oprogramowaniem REW lub Speaker Workshop, tworząc kompletne, komputerowe stanowisko pomiarowe dla każdego konstruktora kolumn i miłośnika dobrego brzmienia.

Zbuduj niepokonanego przeciwnika w kółko i krzyżyk – część 2

W drugiej części projektu klasycznej gry logicznej poznajemy sekrety algorytmu, dzięki któremu komputer nigdy nie przegra. Autor ręcznie opracował setki możliwych przebiegów rozgrywki i zamienił je na dane zapisane w pamięci EPROM. Dowiemy się, jak zoptymalizowano liczbę potrzebnych odpowiedzi maszyny oraz jak zbudować i uruchomić całość – od montażu płytek, przez wykonanie eleganckiej obudowy z czarnego akrylu, aż po wklejanie magnesów w żetony. Efekt końcowy? Stylowe, logiczne wyzwanie, które można z dumą postawić na biurku.

Dla studentów i uczniów: garść technicznych tutoriali

Dla szukających inspiracji: ciekawe projekty DIY

Dla najmłodszych: kolejny zestaw z serii AVTEDU

**W kioskach
od 29 grudnia**

Grudniowe lampki, chrzest śniegu i... szum wszechświata

Okres świąteczny to czas, gdy światło ma wyjątkowe znaczenie. Krótkie dni rozjaśniamy ciepłem lampek, lutownic i małych projektów, które przypominają, że elektronika – podobnie jak zima – potrafi zachwycać prostotą i nastrojem. Ten numer „Elektroniki dla Wszystkich” to zaproszenie do twórczego spokoju – od eksperymentów w warsztacie po rodzinne lutowanie przy choince.

Świąteczny numer rozpoczynamy od projektu lekkiego, ale niezwykle intrygującego – gry w kółko i krzyżyk, w której maszyna nigdy nie przegrywa, mimo że nie ma w niej mikrokontrolera, programu ani sztucznej inteligencji! Cała logika została zrealizowana wyłącznie na bramkach i pamięci EPROM, a ruchy gracza wykrywają czujniki Halla. To nie tylko zabawa, lecz także doskonała lekcja cyfrowego myślenia w najczystszej postaci.

Kolejnym prezentowanym projektem jest podwójny wzmacniacz RF zbudowany na układzie OPA2677. To niepozorne urządzenie o wymiarach zaledwie 38×38 mm potrafi wzmocnić sygnały do 75 MHz iysterować jednocześnie dwa urządzenia lub układy testowe. Idealny dodatek do generatorów o zbyt niskim poziomie wyjściowym – mały, lekki i gotowy do pracy od razu po zmontowaniu.

Następny projekt to termostat z termoparą typu K, zdolny mierzyć temperatury powyżej 1000°C. Oparty na układzie MAX31855, prosty w budowie i precyzyjny w działaniu. Dzięki wbudowanemu przełącznikowi może bezpośrednio sterować grzałkami, piecami czy wentylatorami – to rozwiązanie nie tylko dla hobbystów, ale i profesjonalnych majsterkowiczów.

Dla miłośników precyzji czasowej przygotowaliśmy źródło czasu NTP synchronizowane przez Wi-Fi, które może zastąpić wielokrotnie opisywany moduł GPS w zegarach. Zbudowane na Raspberry Pi Pico W, pobiera czas z Internetu i udostępnia dane NMEA oraz sygnał 1 PPS. Doskonałe rozwiązanie do synchronizacji zegarów analogowych i nie tylko – dokładne, tanie i niezależne od satelitów.

Nie zabrakło też świeżej dawki humoru i praktyki od Maxa Maxfielda, który tym razem tłumaczy działanie przetworników ADC na przykładzie sterowania głową SMAD. W typowym dla siebie stylu – z pasją, żartem i ogromem wiedzy.

Choć okres świąteczny zwykle kojarzy się z aniołkami fruwającymi po nieboskłonach, dr David Maddison ma kaprys zabrać nas w zupełnie innym kierunku. Tym razem schodzi pod ziemię – w drugiej części cyklu Łączność w różnych mediach opisuje sposoby przesyłania sygnałów przez skały, tunele i kopalnie. Od kabli szczelinowych, zwanych też promieniującymi, po kompletne systemy leaky feeder i komunikację VLF – to fascynująca wyprawa w miejsca, gdzie fale radiowe docierają z trudem lub nie docierają wcale.

A teraz coś, co łączy kolejne miesiące w subtelną dźwiękową opowieść. W tytule wstępniaka z październikowego numeru EdW pojawił się „szepc”. W listopadowym zastąpił go „szmer”. W grudniowym pora na „szum” – nie tylko w tytule, ale i na wykładzie. Jos Verstraten opisuje tym razem szumy w układach elektronicznych – od szumu termicznego po wyładowań i migotania. To materiał obowiązkowy dla każdego, kto chce zrozumieć, dlaczego nawet „idealny” rezystor wciąż „coś szepcze”.

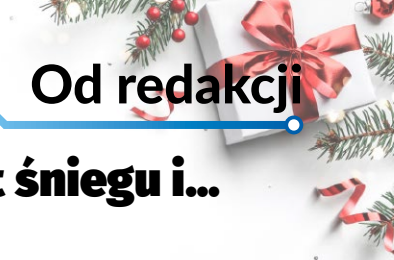
Ian Bell w inauguracyjnej części cyklu „Zniekształcenia i układy zniekształcające” wyjaśnia, czym są nieliniowości, jak powstają i jak je analizować. Zrozumienie THD, harmonicznych i intermodulacji staje się tu nie tylko przystępne, ale i inspirujące – to doskonałe wprowadzenie do przyszłych projektów efektów audio.

Ashwini Kumar Sinha w czwartej części cyklu o systemach ADAS prezentuje zastosowanie radaru Dopplera w wykrywaniu przedkości i odległości pojazdów. Prosty moduł HB100, połączony z Raspberry Pi lub Jetsonem, pozwala stworzyć własny model systemu wspomagania kierowcy – z analizą danych w Pythonie.

A na koniec – coś dla najmłodszych i tych, którzy wciąż mają w sobie dziecko: świąteczne zestawy AVTXMAS. Mikołaj, Elf Pomocnik i Bałwanek LED to nieprzypadkowo wyjątkowo proste projekty do lutowania, z diodą LED, przełącznikiem i baterią CR2032. Idealne do wspólnej zabawy z najmłodszymi przy świątecznym warsztacie, w blasku lampek choinkowych i z zapachem kalafonii zamiast piernika. Bawiąc się z najmłodszymi pamiętajcie jednak, że zestawy zawierają baterię CR2032, a tę łatwo połknąć. Niech więc zapach świątecznych ciast i innych przysmaków, nie przyćmi Waszej codziennej czujności i uważności.

Grudniowy numer EdW to podróż od szumów elektroniki po światełka Mikołajów, Elfów i Bałwanek na świątecznej choince. Niech te projekty rozświecą nie tylko Wasz warsztat, ale też oczy najmłodszych – i przypomną, że w elektronice, podobnie jak w życiu, sporo radości daje to, co zbudujemy sami, a jeszcze więcej – z naszymi Milusińskimi!

Mariusz Ciszewski



Zasilacz dużej mocy

Szanowna Redakcjo!

Jestem stałym od wielu lat czytelnikiem EdW i czytam to czasopismo z wielką uwagą. Jest jedno „ale”... – to mianowicie brak jakichkolwiek konkretnych opisów układów zasilaczy liniowych i impulsowych z wykorzystaniem tranzystorów MOSFET z kanałem N lub P, które miałyby skuteczne układy regulacji prądu (zabezpieczenia przeciążeniowe, zwarcziowe i przepięciowe – zrealizowane na elementach dyskretnych, niekoniecznie SMD).

We wcześniejszych numerach EdW poświęcono wiele miejsca opisom zasilaczy liniowych z serii 78XX, 79XX, a także zbudowanych z elementów dyskretnych – z tranzystorami mocy o polaryzacji NPN lub PNP.

W moim przypadku poszukuję schematów i opisów sprawdzonych, regulowanych zasilaczy liniowych oraz impulsowych z tranzystorami MOSFET (z kanałem N lub P) o zakresach regulacji napięcia 0...38 V i prądzie regulowanym 0...60 A – czyli takich, które mogłyby zasilac silniki komutatorowe 12 V/350 W lub 24 V/400 W, z wymienionymi wyżej zabezpieczeniami.

Nie ufam schematom i opisom zasilaczy publikowanym w Internecie, bo w większości są to tylko rozważania teoretyczne, a nie praktycznie sprawdzone konstrukcje.

Moje pytanie do Redakcji: czy można liczyć na to, że EdW poświęci w przyszłości więcej miejsca zasilaczom tego typu?

W imieniu własnym oraz wielu czytelników, którzy chętnie zapoznali się z praktycznymi i sprawdzonymi układami zasilaczy z MOSFET-ami, z góry dziękuję!

Z poważaniem

Andrzej W.

Red. Drogi Czytelniku,

rzeczywiście wygląda na to, że na łamach EdW nie był dotychczas publikowany zasilacz o tak dużej mocy. 38 V przy prądzie 60 A przekłada się na moc ponad 2 kW! Taka moc z pewnością wybiega daleko poza potrzeby statystycznego Czytelnika EdW.

Prawdopodobnie najmocniejszym projektem zasilacza opublikowanym dotąd w naszym czasopiśmie (EdW 1/2002) był zestaw AVT2462 pod nazwą Zasilacz 10 A 10...20 V, którego elementem mocy był tranzystor Darlingtona BDW84C o dopuszczalnym prądzie 12 A (przy odpowiednim chłodzeniu). To dużo, ale wciąż zbyt mało jak na potrzeby Czytelnika (prąd do 60 A).

Projekt o zbliżonych parametrach opublikowała również „Elektronika Praktyczna” (EP 8/1993). Ten z kolei oparty był na układzie

L4970A – przetwornicy typu step-down (buck), która nie jest już produkowana (ma status EOL).

Jeśli do redakcji trafi godny zaprezentowania projekt zasilacza o parametrach zasugerowanych przez Czytelnika, jego opracowanie z pewnością zostanie opublikowane.

Na tę chwilę pozostaje skorzystać z rozwiązań dostępnych na rynku, choć te nie rzadko przekraczają możliwości finansowe przeciętnego hobbysty. Poniżej kilka propozycji, które wydają się spełniać wymagania mocowe Czytelnika:

- **EA-PS 9036-60 EA Elektro-Automatik** (0...36 V/0...60 A).
- **TDK-Lambda GEN40-60-1P230** (0...40 V/0...60 A).
- **Siglent SPS5042X** (0...40 V/0...60 A).

Wszystkie wymienione modele to zasilacze wysokiej klasy, oferujące nie tylko precyzyjną regulację napięcia i prądu, lecz także rozbudowane możliwości sterowania i programowania parametrów pracy – zarówno lokalnie, jak i zdalnie (przez interfejs USB, LAN lub analogowy). Taka funkcjonalność, w połączeniu z wysoką mocą wyjściową i kompleksowymi zabezpieczeniami (CV/CC, OVP, OCP, OTP), sprawia, że są to urządzenia klasy profesjonalnej, przeznaczone do stanowisk testowych, laboratoriów badawczych i automatyki przemysłowej, a nie do zastosowań amatorskich. Ich cena odzwierciedla tę kategorię – zazwyczaj wynosi od kilku do kilkunastu tysięcy złotych za sztukę, w zależności od producenta, typu interfejsów i dostępnych funkcji programowalnych.

Choć dla hobbystów mierzących się z zamiarem zbudowania prostszego odpowiednika nie mamy w swoich materiałach gotowego rozwiązania, zamiast „ryby” możemy polecić na ten moment „wędkę”:

- **Cykl „Taki zwyczajny zasilacz...”** (EdW, numery: 09...12/2019, 02/2020, 04...06/2020, 08...09/2020, łącznie 10 odcinków) szczegółowo omawia budowę stabilizatora liniowego, w którym elementem regulacyjnym jest tranzystor MOSFET. W ramach kolejnych części Czytelnik znajdzie: koncepcję realizacji i schematy, dobór i analizę MOSFET-a, omówienie problemów i modyfikacji (w tym z użyciem wzmacniaczy operacyjnych), zagadnienia zakłóceń, parametry dynamiczne oraz zestawienie testowe (przrządy i przystawki pomiarowe), a także porównanie różnych typów zasilaczy (liniowych i impulsowych). Cykl zawiera również elementy dotyczące zabezpieczeń i ograniczeń, choć nie stanowi one



głównego tematu poszczególnych części, a pojawiają się przy okazji analizy praktycznych problemów konstrukcyjnych.

- **Cykl „Tranzystory: Sterowanie MOSFET-ami”** (EdW 01/2020...05/2020) szczegółowo omawia tranzystory MOSFET – ich właściwości, dobór, procesy przełączania, straty i czasy przełączania. W cyklu poruszono również aspekty praktyczne sterowania tranzystorami, jednak bez szczególnego skupienia na topologiach high-side i low-side oraz dogłębnej klasyfikacji typów N i P.

Wspomniane artykuły opisują nie tylko gotowe schematy, ale także metodykę projektowania, dobór tranzystorów i zasady zabezpieczania układów przed zwarciem, przeciążeniem czy przegrzaniem. Łącząc te rozwiązania, można zbudować solidny, sprawdzony w praktyce zasilacz warsztatowy lub sterownik silników o znacznej mocy.

-15%
NA START
192,80 zł

-30%
po pierwszym roku
prenumeraty
158,80 zł

-40%
po drugim roku
prenumeraty
136,10 zł

-50%
po trzecim roku
nieprzerwanej prenumeraty
113,40 zł

Odkryj korzyści z **prenumeraty drukowanej** – większe oszczędności z każdym rokiem!

Rozpocznij swoją przygodę z *Elektroniką dla Wszystkich*. Decydując się teraz na roczną prenumeratę drukowaną, otrzymasz nie tylko dostęp do najnowszych wydań, ale i **znakomity start dzięki zniżce 15%** na pierwsze zamówienie!

Prenumerata to nie tylko wygoda dostępu do treści, ale także sposób na znaczące oszczędności. Dołącz do grona naszych stałych czytelników i ciesz się coraz lepszymi warunkami.

Im dłużej jesteś z nami, tym więcej oszczędzasz:

- po roku nieprzerwanej prenumeraty zapewnimy Ci **30% rabatu** na kolejny rok,
- po dwóch latach wierności zaoferujemy **40% rabatu**,
- po trzech latach lojalności osiągniesz **najwyższy poziom rabatu – 50%**!

Jak otrzymać rabat za lojalność?

Zaloguj się na swoje konto prenumeratora na www.UlubionyKiosk.pl i zamów prenumeratę, korzystając z przycisku PRZEDŁUŻ w zakładce „Prenumeraty”.

Przeglądaj wcześniej, płać mniej – postaw na **e-prenumeratę**!

Wybierz prenumeratę cyfrową PDF i ciesz się dostępem do czasopisma nawet 7 dni przed oficjalną premierą w kioskach. Oszczędzaj czas i pieniądze – skorzystaj z **rabatu 30%** na roczną e-prenumeratę w cenie 126,80 zł.

Dodatkowa oferta dla prenumeratorów wersji drukowanej: jeśli już subskrybujesz wersję papierową, możesz dokupić równolegle e-wydania w cenie 36,20 zł/rok – z **niesamowitym rabatem 80%**.

Zyskaj nieograniczony dostęp do zasobów dla pasjonatów elektroniki!

Tylko prenumeratorzy mają pełny dostęp do:

- cyfrowego archiwum *Elektroniki dla Wszystkich* na www.elportal.pl/archiwum
- projektów DIY+ na www.elportal.pl/diy

Zamów prenumeratę drukowaną lub e-prenumeratę na www.UlubionyKiosk.pl lub przez przelew na konto Wydawnictwa AVT, a po zaksięgowaniu wpłaty wyślemy Ci mailowo kod dostępu do portalu.

ARCHIWUM



Zacznij korzystać z pełnych zasobów już dziś!

W artykule opisujemy komputer do gry w kółko i krzyżyk. Nie ma on jednak procesora, sygnałów zegarowych, rejestrów ani przerzutników. Używa jedynie bramek logicznych i pamięci ROM, która zachowuje się jak złożona tablica bramek. Interfejs użytkownika jest prosty, ale zarazem elegancki: plastikowe żetony z nadrukowanymi znakami „X” i „O” są umieszczane w jednym z dziewięciu wgłębień na planszy, a komputer sygnalizuje swój ruch, zapalając diody LED w odpowiednim polu.



Zagraj we własną grę kółko i krzyżyk, część 1

W tego typu systemie jedyne opóźnienie obliczeniowe wynika z czasu propagacji bramek lub urządzeń logicznych. Ponieważ nasz komputer reaguje wyłącznie na statyczne stany logiczne, nie może się pomylić, zawiesić ani przestać działać.

Szybkość nie jest istotna dla tego zastosowania, ale niski pobór prądu już tak. W systemach taktowanych jakimś przebiegiem zegarowym pobór prądu wzrasta wraz z jego częstotliwością. Komputer statyczny wykonany w technologii CMOS jest wyjątkowo energooszczędny i doskonale nadaje się do zasilania z baterii. Pobór prądu ze źródła zasilania 9 V waha się od 75 mA do 90 mA, z czego większość jest pobierana przez świecące diody LED.

Urządzenie zawiera generator, ale nie jest on używany do żadnych obliczeń, a służy jedynie do generowania liczb losowych (RNG) wykorzystujący technikę „wirującego koła” opisaną dalej.

Zanim zabrałem się za rozwiązanie tego problemu, uznałem, że warto najpierw zastanowić się, jak właściwie w kółko i krzyżyk grają dwie osoby. Gracz rozpoczynający partię ma wyraźną przewagę nad przeciwnikiem,

dlatego w kolejnych rozgrywkach należy zaczynać grę na przemian. Dzięki temu szanse obu stron pozostaną wyrównane. Oznacza to, że komputer musi umieć zarówno wykonać pierwszy ruch, jak i pozwolić człowiekowi rozpocząć grę.

Przewidywalny gracz może rozpocząć grę w ten sam sposób za każdym razem, na przykład zaczynając na centralnym polu, aby uzyskać maksymalną przewagę. Ale takie przewidywalne zachowanie szybko staje się bardzo nudne, więc gdy komputer jest pierwszym graczem, powinien zmieniać swoją strategię otwarcia.

Kiedy człowiek wygrywa grę, prawdopodobnie ogłasza to z wielkim entuzjazmem, więc komputer potrzebuje sposobu na powiadomienie człowieka o wygranej.

Wreszcie, dwaj ludzie grający ze sobą są na tyle niedoskonalymi, że prędzej czy później jeden z nich popełnia błąd w ocenie. Pozwoliłoby to drugiemu człowiekowi, który nie popełnił żadnych błędów, nie tylko zapobiec wygranej drugiego człowieka (powodując remis), ale także pokonać drugiego gracza.

Prawidłowa i kompletna maszyna Kółko i Krzyżyk nie tylko byłaby niepokonana przez

człowieka, ale też powinna być w stanie wygrać z nim przy każdej okazji. Wymaga to analizy każdego możliwego błędu, jaki człowiek może popełnić podczas rozgrywki.

W świetle powyższych cech zdecydowałem, że sposobem na zaprojektowanie maszyny będzie początkowo stworzenie gry planszowej dla dwóch graczy. Każdy gracz mógłby umieścić dysk z etykietą „X” lub „O” na planszy. Działałaby ona dobrze nawet przy braku zasilania elektrycznego, a dwóch graczy mogłoby cieszyć się grą razem, jak podczas zwykłej gry.

Jeśli jednak jeden z graczy (ludzi) „zaginie”, a gra jest zasilana, maszyna wkracza do akcji, aby zastąpić jednego z nich. Uruchamiany jest wtedy scenariusz człowieka kontra maszyna.

Maszyna musi być w stanie wykonywać funkcje, które posiadałby „bezbłędny” gracz – człowiek: nigdy nie popełniać błędów, nigdy nie zostać pokonanym, ale także wygrywać, gdy nadarzy się okazja.

Gra została zaprojektowana jako plansza dla gracza z żetonami „X” i „O”, natomiast komputer ma jedynie „mózg” – bez oczu i rąk. Dlatego to człowiek umieszcza na planszy żetony w imieniu maszyny. Komputer wskazuje

swoje pole, zapalając diodę LED w miejscu, w którym chce umieścić swój żeton.

Zauważ, że w tym projekcie maszyna zawsze gra jako O, a człowiek jako X.

Plansza może wykryć obecność żetonu „X” lub „O” na planszy gracza. Komputer „wie”, gdzie na planszy znajduje się każdy z nich i „wie”, kiedy jest ruch maszyny, kiedy człowieka.

Gdy nadchodzi kolej maszyny, jej ruch jest obliczany w czasie krótszym niż 200 ns. Po przeanalizowaniu rozmieszczenia żetonów „X” i „O” komputer zapala diodę LED w polu, w którym chce umieścić swój żeton „O”.

Na planszy znajduje się dziewięć pól, więc w pełnej partii można wykonać dziewięć ruchów. Gdy pierwszy ruch wykonuje „X”, w grze pojawia się pięć żetonów „X” i cztery żetony „O”. Jeśli natomiast jako pierwszy zaczyna „O”, sytuacja się odwraca – pięć żetonów „O” i cztery żetony „X”.

Oznacza to, że w grze pozostaje albo żeton „X”, albo żeton „O”, w zależności od tego, czy grę rozpoczął – człowiek (X), czy maszyna (O). W związku z tym zapewniono dodatkowe miejsce do przechowywania nieużywanego żetonu. Miejsce to działa również jako dwubiegunowy przełącznik elektryczny do konfigurowania obwodów komputera w celu ustalenia, kto zaczyna pierwszy. Jest to nie tylko wygodne, ale także pozwala uniknąć konieczności stosowania jakichkolwiek przełączników mechanicznych.

Jeśli X (człowiek) ma rozpocząć grę jako pierwszy, na wolnym polu jest umieszczany żeton „O”, ale jeśli O ma rozpocząć grę jako pierwszy, żeton „X” jest umieszczany na dodatkowym polu. Instrukcja ta jest wygrawerowana na powierzchni planszy gracza wraz z informacją, że człowiek gra żetonami „X”, a maszyna żetonami „O”.

Gdy gra jest włączona, a na planszy nie ma początkowo żadnych żetonów, wszystkie diody LED świecą się. Jest to sygnalizacja, że nastąpi „losowe uruchomienie”. Diody LED na planszy zapalają się szybko w sekwencji, jedna po drugiej. Jest to za razem testowanie diod, podobnie jak kiedyś zwyczajowo zapalano w urządzeniach na krótko wszystkie lampki zaraz po włączeniu.

Jeśli żeton „X” zostanie umieszczony w wolnym miejscu (co oznacza, że maszyna zaczyna jako pierwsza), „obracające się koło” zatrzyma się i zablokuje pierwszy ruch komputera. Diody LED, które pozostają zapalone, wskazują losową pozycję dla pierwszego ruchu O.

Gdy pokrywa gry (górna pokrywa na zawiasach) jest zamknięta, bezpiecznie

przechowuje wszystkie żetony wewnątrz, dzięki czemu nie zgubią się ani nie oddzielią od planszy.

Prototyp gry jest zasilany z zasilacza 9 V, ale ponieważ pobór prądu jest mały, urządzenie może być zasilane z baterii 9 V.

Krótki film przedstawiający maszynę podczas pracy można obejrzeć na stronie <https://youtu.be/IE9a5ZJZCgE>.

Projekt układu

Zainspirował mnie fakt, że w 1958 roku Dick Smith zbudował maszynę do gry w kółko i krzyżyk z części pochodzących z centrali telefonicznej. Najprawdopodobniej były to wówczas części zabytkowe, a większość z nich pochodziła z centrali telefonicznej z lat 30. ubiegłego wieku.

Uznałem, że możliwe powinno być zrobienie podobnego urządzenia przy użyciu bramek logicznych. Bardzo lubię bramki logiczne z serii 74 i zwykle używam typów 74xxx (TTL) lub 74LSxxx (logika Schottky o niskim poborze mocy). Istnieją również wersje CMOS, takie jak seria 74HCTxxx. Wykonują one te same funkcje logiczne przy niższym zużyciu energii, więc wybrałem je do tego projektu. Użyłem również niebieskich

diod LED, ponieważ świecą bardzo jasno przy małym prądzie.

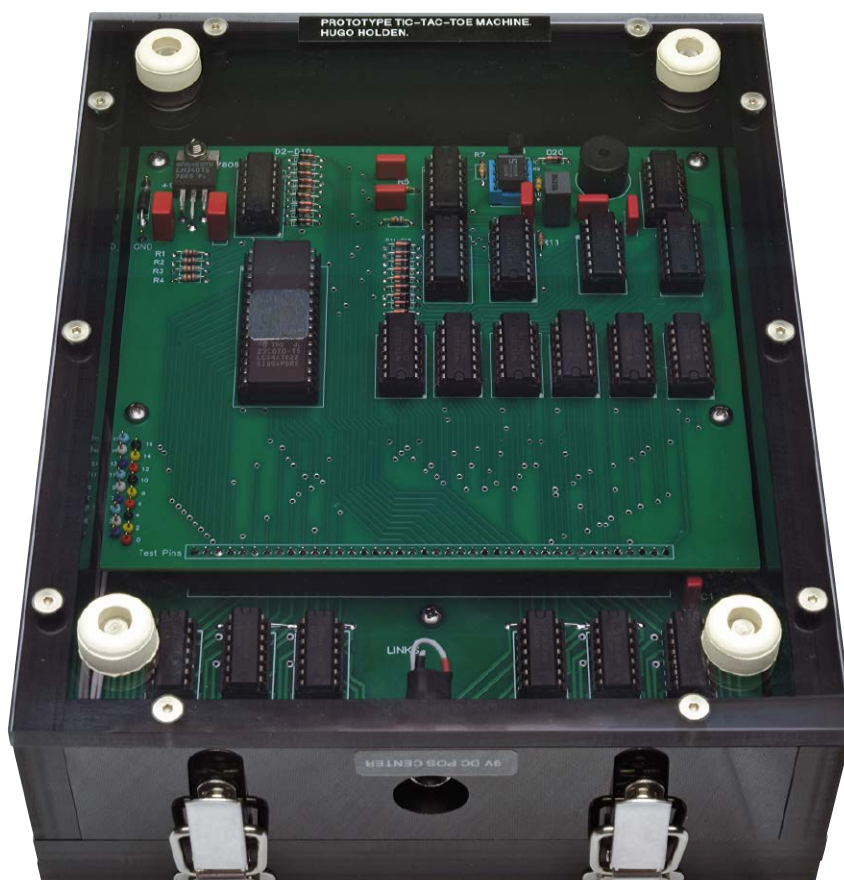
Układy komputera są rozmieszczone na dwóch płytkach drukowanych, „płytkę do gry” ze wszystkimi elementami interfejsu użytkownika (czujniki Halla, diody LED itp.) oraz „płytkę obliczeniową”, która zawiera wszystkie układy sterujące.

Aby przeanalizować wzorce rozgrywki i wykonać prawidłowy ruch, używam pamięci EPROM lub EEPROM. Potrzebują one 18 linii adresowych do przetwarzania logiki planszy gracza. Do sterowania działaniem komputera układ wykorzystuje informacje o „parzystości” wytwarzane na planszy w zależności od tego, który gracz zaczyna pierwszy.

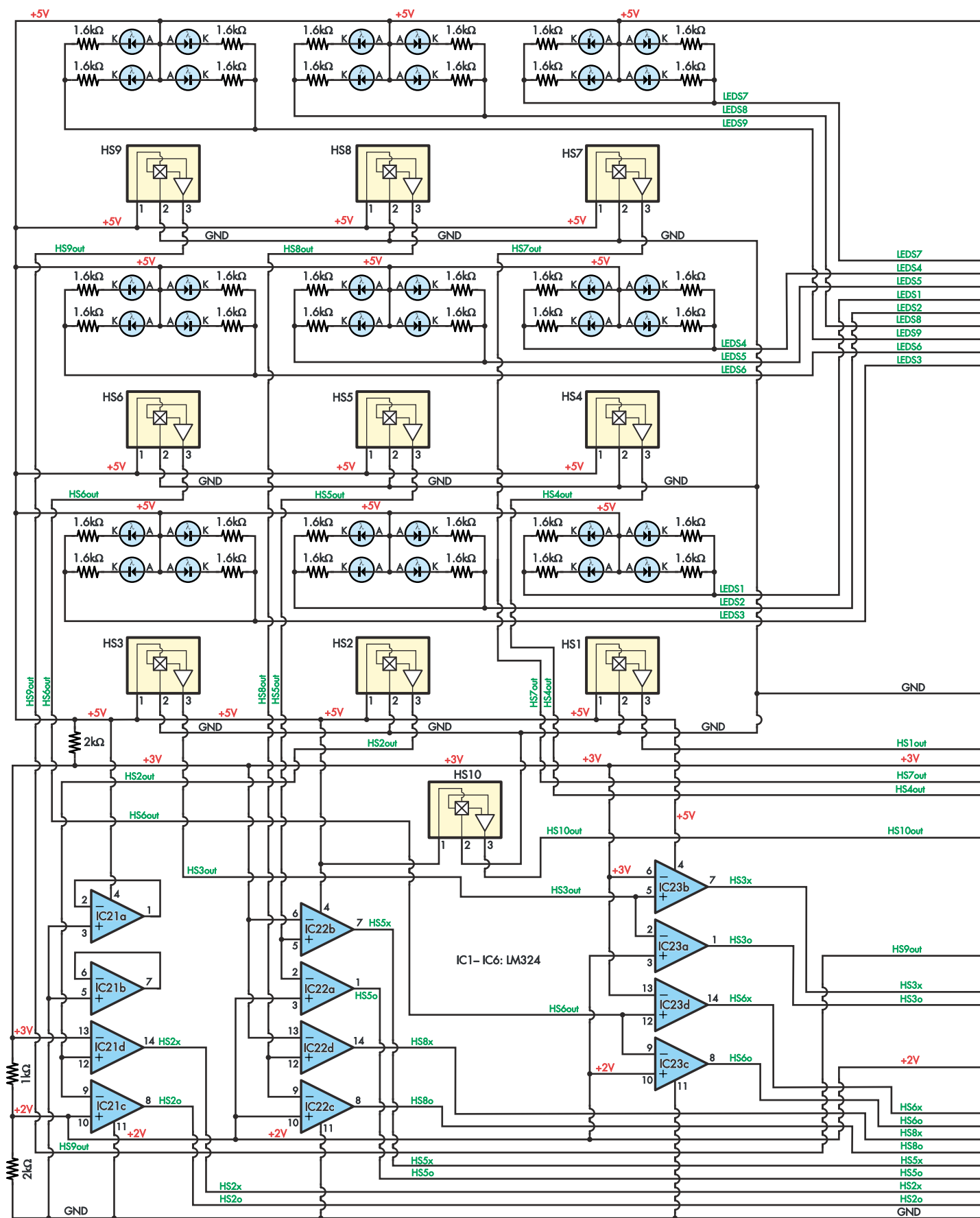
Odpowiednią pamięcią EPROM dostępną w obudowach PLCC lub DIL jest AT27C020 firmy Microchip. Kompatybilna jest też łatwo dostępna i niedroga pamięć EEPROM Winbond W27C020. W moim prototypie użyłem podobnej pamięci EPROM 27C020 kasowanej promieniami UV, ponieważ miałem ją pod ręką, wraz z odpowiednim programatorem.

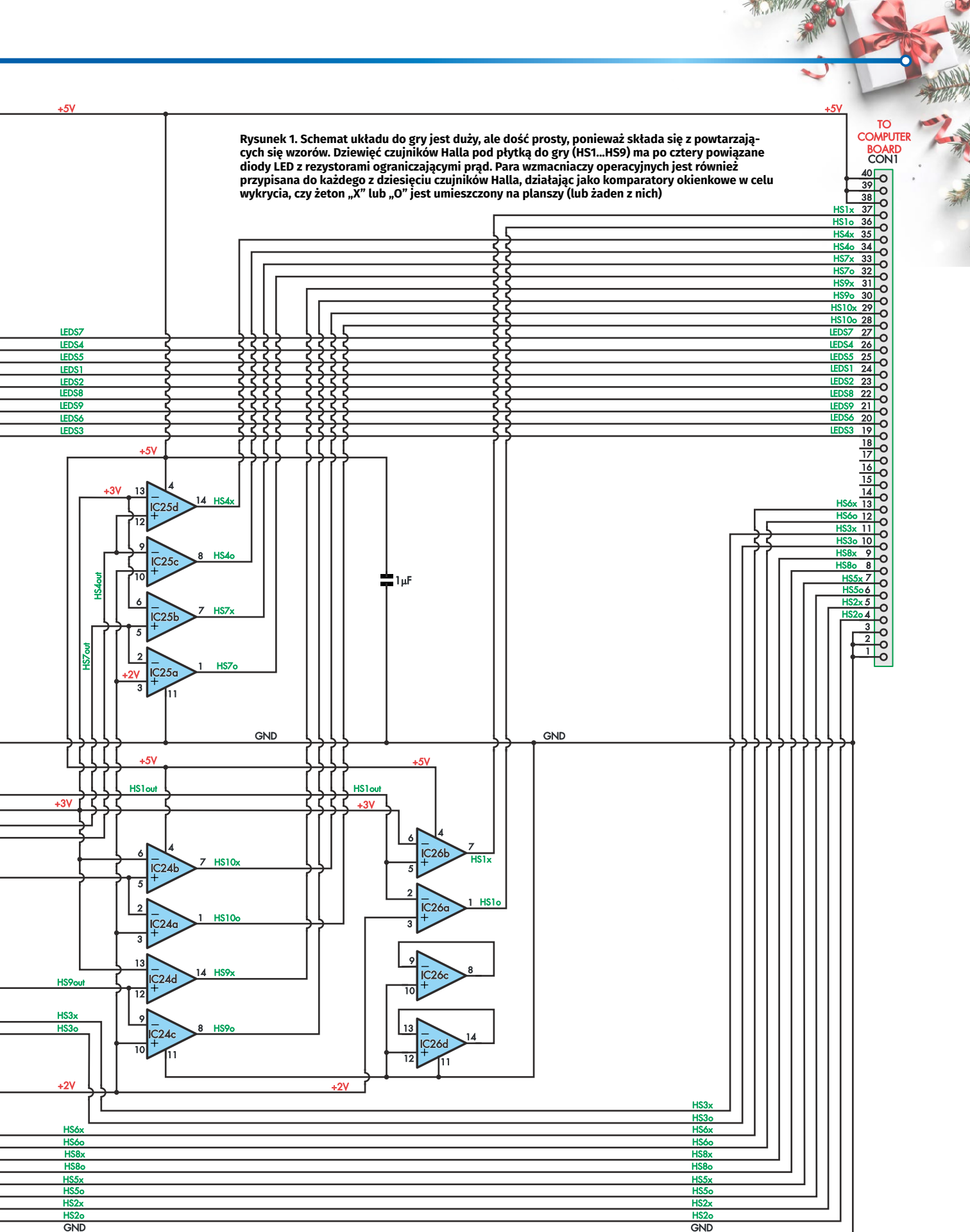
Zabawa ze słabymi magnesami

Chcąc uniknąć mechanicznych przełączników i nie mając mikrokontrolera do obsługi



Spód komputera Kółko i Krzyżyk ma przezroczystą akrylową pokrywę, która pozwala zajrzeć i zobaczyć dwie główne płytki PCB





interfejsu dotykowego, zdecydowałem się użyć magnesów ferrytowych firmy Jaycar. Są one umieszczone w elementach gracza, wraz z ratiometrycznymi (analogowymi) czujnikami Halla i służą do wprowadzania danych przez użytkownika.

Magneszy ferrytowe (Jaycar Cat LM1616) są dostarczane w paczce po 12 sztuk, ja użyłem tylko 10, pięć dla żetonów „X” i pięć dla „O”. Wyciąłem otwór o średnicy 14,5 mm i głębokości 4,5 mm w dolnej części każdego plastikowego żetonu o średnicy 25 mm i grubości 10 mm. Następnie przykleiłem do nich magnesy używając żywicy epoksydowej utwardzającej się w czasie 24 godzin.

Celowo nie użyłem silnych magnesów neodymowych, ponieważ ich pole magnetyczne byłoby zbyt silne i mogłoby zakłócać działanie pobliskich elementów. Mają one również zwyczaj przeskakiwania do siebie nawzajem lub do najbliższego obiektu magnetycznego i magnesowania go, co sprzyja przypadkowemu wymazywaniu nośników magnetycznych. Chociaż użyłem słabych magnesów, ich pole magnetyczne jest wystarczające do działania czujników Halla.

Schemat układu planszy do gry przedstawiono na **rysunku 1**. Siatka 3×3 czujników Halla oznaczonych HS1...HS9 i cztery diody LED powiązane z każdym z nich tworzą planszę do gry. Dziesiąty czujnik Halla, HS10, służy do określania, który gracz zaczyna pierwszy, jak opisano wcześniej. Sześć poczwórnych wzmacniaczy operacyjnych LM324 (w sumie 24 pojedyncze wzmacniacze, cztery nieużywane) odpowiednio przetwarza napięcia wyjście czujników Halla.

Bez przyłożonego pola magnetycznego, napięcie wyjściowe DC każdego czujnika Halla mieści się w zakresie od 70 mV od 2,5 V. Elementy gracza X i O mają magnesy przyklejone wewnątrz nich w przeciwnych orientacjach. Umieszczenie żetonu „O” powoduje, że sygnał wyjściowy z czujnika Halla jest niski (poniżej 2 V), natomiast umieszczenie żetonu „X” na czujniku powoduje, że sygnał wyjściowy jest wysoki (powyżej 3 V).

Wzmacniacze operacyjne są używane jako komparatory do generowania stanu logicznego 1 (wysokiego poziomu) na wejściu Data X lub O, gdy na płytce umieszczony jest odpowiednio element gracza X lub O. Układy LM324 są przydatne do tego zadania, ponieważ ich zakres napięć wyjściowych, gdy są zasilane z 5 V, jest idealnie dopasowany do poziomów logicznych TTL.

Ogólny schemat każdego czujnika pokazano na **rysunku 2**. Jest on powielony 10 razy na planszy. Łańcuch rezystancyjny 2 kΩ/1 kΩ/2 kΩ na zasilaniu 5 V wytwarza poziomy odniesienia 3 V i 2 V.

Dwa wzmacniacze operacyjne LM324 są używane jako komparator okienkowy. Wyjście jednego z nich przechodzi w stan wysoki, gdy czujnik Halla wytwarza napięcie wyższe niż 3 V, a wyjście drugiego przechodzi w stan wysoki, gdy czujnik wytwarza napięcie mniejsze od 2 V.

Jak pokazano na **rysunku 3**, cztery diody LED są rozmieszczone wokół każdego czujnika z dwóch powodów. Jednym z nich jest to, że wynikająca z tego symetria nadaje przyjemny wygląd, drugi powód to redundancja. Jeśli jedna dioda LED ulegnie awarii

(co może się zdarzyć), gra nadal będzie działać normalnie. Każda dioda LED w grupie 4 ma anodę podłączoną do +5 V z oddzielnym rezystorem katodowym.

Rezystory 1,6 kΩ ograniczające prąd powodują przepływ prądu o natężeniu około 1 mA przez każdą diodę, gdy są włączone, ponieważ jednak są to diody niebieskie, uzyskujemy dość dużą jasność świecenia nawet przy tak małych prądach.

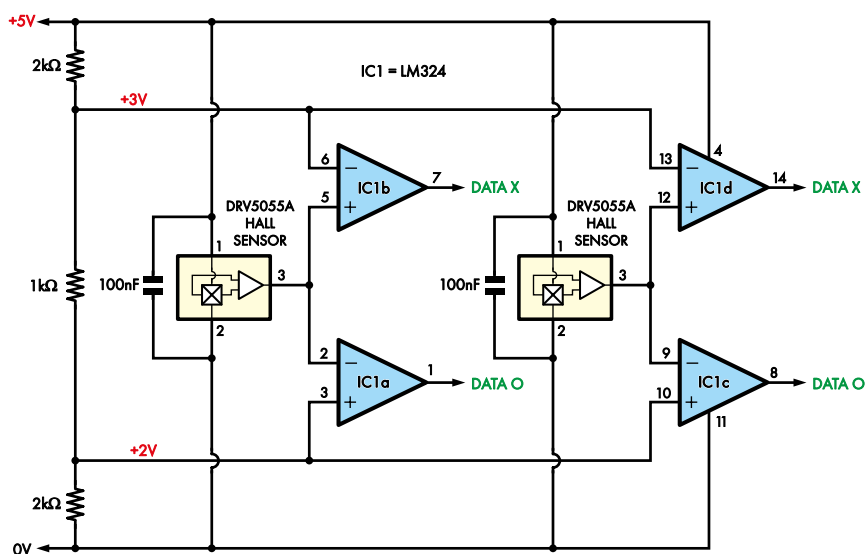
Producenci czujników Halla zalecają stosowanie kondensatorów filtrujących na ich pinach zasilających, więc dodałem je jako kondensatory ceramiczne do montażu powierzchniowego. Mimo to czujniki działały dobrze bez nich, prawdopodobnie dlatego, że nasz komputer jest statyczny, więc na liniach zasilających nie występują zakłócenia.

Płyta z grą łączy się z płytą obliczeniową za pomocą 40-stykowego złącza SIL (gniazdo żeńskie na płycie z grą i wtyk męski na płycie obliczeniowej), co pozwala uniknąć dużej ilości przewodów.

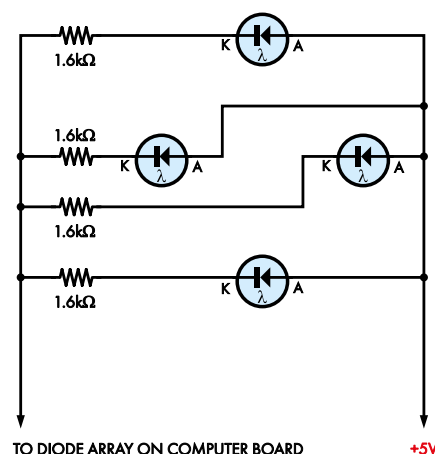
Opcje czujników Halla

„Dostępnych” jest kilka różnych wersji czujnika Halla. Użyłem cudzysłów, ponieważ z ich dostępnością bywa różnie. Wersja A1 jest najbardziej czuła i właśnie jej użyłem. Niestety, w chwili pisania tego tekstu czujników takich nie było w żadnym magazynie, choć mam nadzieję, że zmieni się to wkrótce po publikacji ([oryginalny artykuł był wydany w styczniu 2023 r. – przyp. red.](#)).

Wersja A2 jest łatwiejsza do zdobycia, ale ma o połowę mniejszą czułość niż wersja A1. Na szczęście można to dość łatwo



Rysunek 2. Para czujników Halla jak na rysunku 1 w izolacji. Pokazano dwa, ponieważ współdzielą one jeden poczwórny wzmacniacz operacyjny. To, czy linia DATA X lub DATA O osiągnie stan wysoki, gdy magnes zostanie umieszczony w pobliżu czujnika, zależy od orientacji tego magnesu, a tym samym polaryzacji pola magnetycznego. Kondensatory bocznikujące 100 nF byłyby niepotrzebne w prototypie, więc nie są pokazane na rysunku 1



Rysunek 3. Na obwodzie każdego wgłębienia pod żeton na planszy rozmieszczone są cztery diody LED. Uzyskuję się tym samym przyjemną symetrię. Każda z nich ma własny rezystor ograniczający prąd diody LED do około 1 mA. Niebieskie diody LED świecą wystarczająco jasno przy takim prądzie, nie zużywając dużo energii, nawet jeśli świecą się jednocześnie wszystkie zamontowane na planszy (36)

skompensować, zmieniając pojedynczy rezystor 1 kΩ w ciągu 2 kΩ/1 kΩ/2 kΩ na 510 Ω. Spowoduje to, że progi komparatora okienkowego będą wynosić 2,25 V i 2,75 V zamiast 2 V i 3 V, kompensując zmniejszone zakres napięcia wyjściowego z czujników Halla.

Płyta główna

Na **rysunku 5** przedstawiono schemat płytki komputera, w której zastosowano wszystkie elementy przewlekane. W projekcie nie są używane kondensatory elektrolityczne, a jedynie foliowe. Dzięki temu układ powinien działać przez wiele lat.

Na płycie znajdują się punkty testowe dla wszystkich linii adresowych EEPROM (od A0 do A17), chociaż nie jest to pokazane na schemacie. Przydały się one do potwierdzenia poprawnego działania płytki podczas programowania.

Plansza gracza musi rozpoznawać każdy możliwy wzór lub kombinację żetonów „X” człowieka i „O” maszyny. Musi ona rozróżniać, w każdej z dziewięciu lokacji, czy umieszczono tam żeton „X”, „O”, czy też nie umieszczono tam żadnego żetonu.

Aby przekonwertować te informacje na format binarny odpowiedni do oceny komputerowej, inny element umieszczony w każdej lokalizacji generuje liczbę binarną, jak pokazano na niebieskich i czarnych liczbach na rysunku 4. Wartości te są sumowane, aby uzyskać pojedynczą 18-bitową liczbę reprezentującą stan planszy w danym momencie, która jest używana jako adres EEPROM do wyszukiwania następnego ruchu.

Na przykład, jeśli żeton „O” zostanie umieszczony na środku (kwadrat 5), zostanie

wygenerowana wartość dziesiętna 8192. Następnie, jeśli żeton „X” zostanie umieszczony na planszy w miejscu 07, zostanie wygenerowana wartość dziesiętna 64. Z uwagi na powyższe dla tego prostego stanu dwóch żetonów na planszy, wygenerowany adres to $64 + 8192 = 8256$.

To właśnie w tej lokalizacji w pamięci EEPROM przechowywany jest następny ruch gracza O. Innymi słowy, pamięć EEPROM generuje wartość (odczytując pamięć pod tym adresem), w celu zapalenia odpowiedniej diody LED w miejscu, w którym ma być umieszczony następny żeton „O” maszyny.

Brak elementów gracza na planszy generuje adres równy zero, a pamięć EEPROM ma wartość 255 (0xFF szesnastkowo, 11111111 binarnie) zapisaną w tej lokalizacji. Cztery dolne bity będące jedynekami oznaczają, że dekodery 74HCT42 IC2 nie generuje żadnego stanu wyjściowego, ponieważ jest to nieprawidłowy kod. W tym stanie dekodery nie zapala żadnej diody LED. Zamiast tego wartość ta uruchamia początkowy proces randomizacji ruchu.

Następnie, jeśli żeton „X” gracza zostanie umieszczony na wolnym polu, losowo wybrane pole pozostanie podświetlone dla pierwszego żetonu O, który zostanie umieszczony, zgodnie z licznikiem, na którym zatrzymał się IC7 i czterobitową wartością przedstawioną IC6.

Gdy figury znajdują się na planszy i nadchodzi kolej maszyny na wykonanie ruchu, elektronika gry przechodzi w „tryb obliczeniowy”. Wyjścia pamięci EEPROM

są aktywowane przez linię sterującą /COMPUTE, która łączy się z pinem 22 układu IC1. Wyjścia EEPROM są aktywowane, co skutkuje zapaleniem się odpowiednich diod LED pokazujących, gdzie maszyna chce umieścić swój żeton.

Przez resztę czasu, gdy linia sterująca /COMPUTE jest w stanie wysokim, wyjścia EEPROM przechodzą w stan wysokiej impedancji (obwód otwarty) i są podciągane przez cztery rezystory 10 kΩ do napięcia zasilającego. Oznacza to, że dekodery kodu BCD na kod dziesiętny, IC2 (74HCT42), otrzymuje same jedynki (z powodu podciągania), więc nie świecą się żadne diody LED.

Początkowy randomizator

Wyobraź sobie obracające się koło z cyframi od 1 do 9 kręcące się tak szybko,

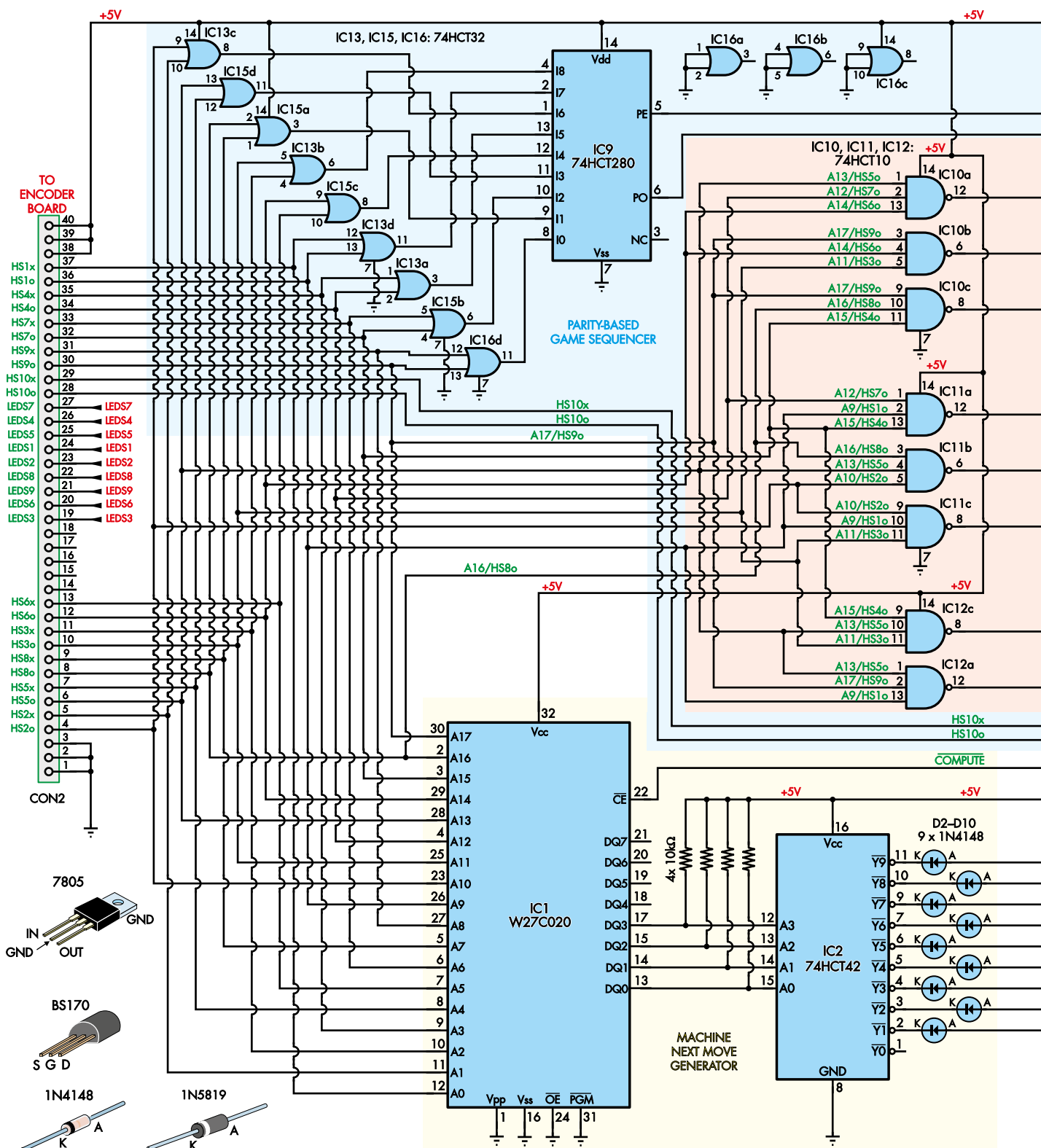
64 32768 07	128 65536 08	256 131072 09
8 4096 04	16 8192 05	32 16384 06
1 512 01	2 1024 02	4 2048 03

- Address for X array
- Address for O array
- ROM output byte value and square identity number

Rysunek 4. Każdej możliwej pozycji gry przypisana jest inna wartość potęgi dwójki w zależności od tego, czy znajduje się tam żeton „X” czy „O”. Z 18 możliwymi wartościami, liczby te mogą być użyte do zaadresowania 218 = 256 KB komórek pamięci EPROM lub EEPROM. Liczby zapisane pod tymi adresami mówią maszynie, jaki ruch ma wykonać w następnej kolejności (1..9, aby umieścić żeton na danym polu lub 0/255, aby nie wykonać żadnego ruchu)

Kółko i Krzyżyk używa zestawu 10 magnetycznych elementów do gry, przy czym jeden element określa, który gracz jest pierwszy



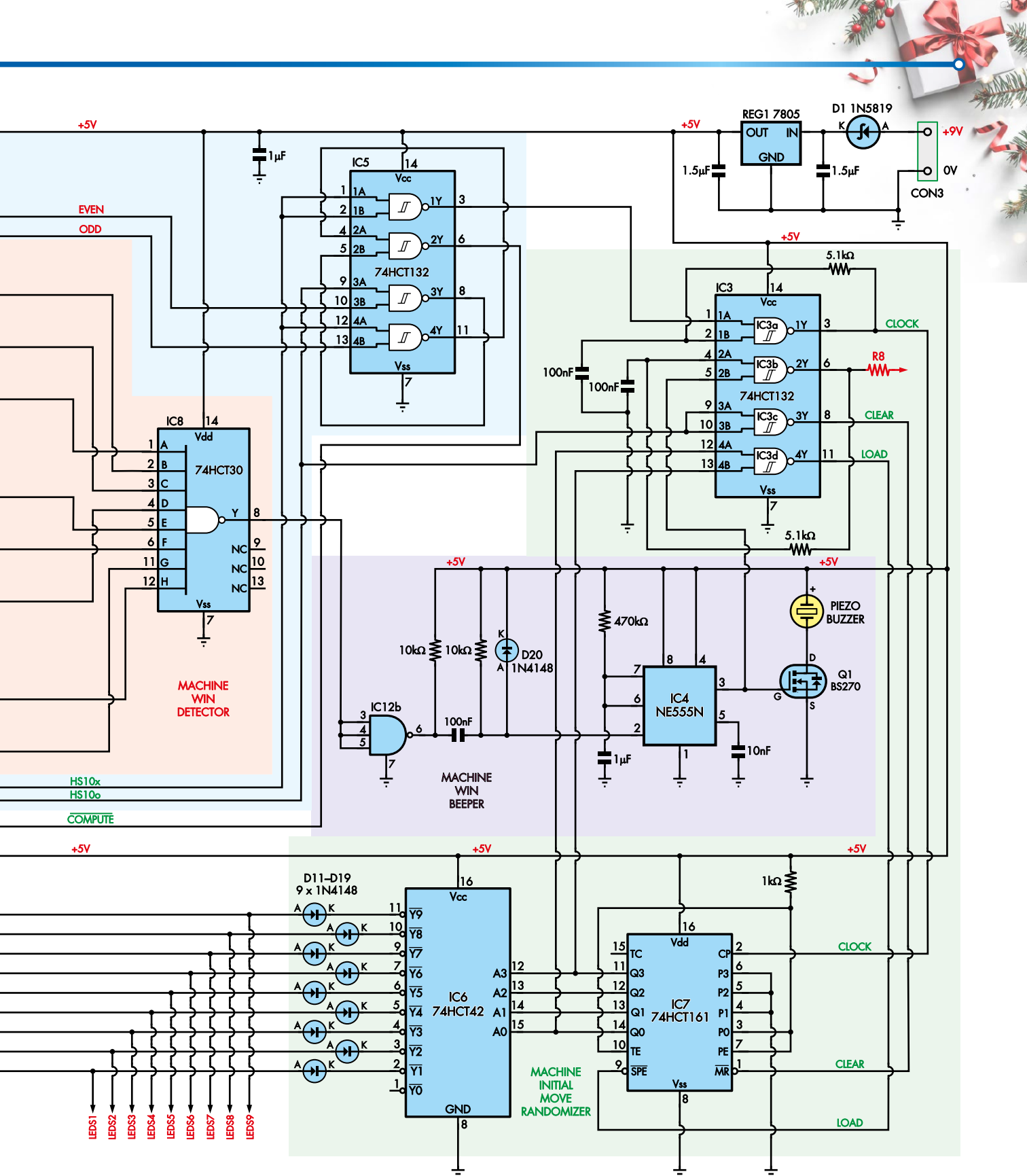


Rysunek 5. Płytki obliczeniowa zawiera układ logiczny, który decyduje o tym, który ruch należy wykonać w następnej kolejności i czy gracz wygrał. Każda sekcja opisana w tekście jest oznaczona innym kolorem, co pomoże w zrozumieniu całego schematu. Wszystkie połączenia między układami z rysunku 1 i rysunku 5 są wykonane za pomocą 40-pinowych złączy SIL CON1 i CON2

że postrzegamy to jako rozmyty obraz. Jeśli rzucisz w nie rzutką lub gwałtownie zatrzymasz hamulcem, jedna z liczb zostanie wybrana w pozornie losowy sposób.

Podobne zadanie przypada układowi licznikowi 74HCT161 (IC7). Jest on taktowany z częstotliwością około 3 kHz przez generator utworzony przez układ IC3a z rezystorem

sprężenia zwrotnego 5,1 kΩ i kondensatorem 100 nF. Gdy IC7 osiągnie wartość zliczania równą dziewięć, pin wyjściowy 11 IC3d przechodzi w stan niski, co ładuje



licznik wartością jeden, więc przy następnym impulsie zegarowym licznik przechodzi do stanu jeden.

Gdy żeton „X” zostanie umieszczony na wolnym polu na planszy (oznaczającym, że maszyna ma wystartować jako pierwsza), wyjście generatora – pin 3 układu IC3 zostaje

zablokowane. W związku z tym na planszy pozostaje liczba od 1 do 9, zapalając diody LED na danym kwadracie na planszy.

Zauważ, że przypisałem wyjściom EEPROM pokazanym na rysunku 4 wartości od 1 do 9, a nie od 0 do 8. Wynika to z faktu, że randomizator O-start jest generowany przez licznik

binarny 74HCT161 (IC7), który musi wykonać dziewięć zliczeń, od 1 do 9, i cyklicznie przechodzić do stanu 1. Ale potrzebuje również innego unikatowego stanu zerowego, gdy jest zerowany i nie zapala żadnych diod LED.

Czterobitowe wyjście IC7 (Q0...Q4) trafia do dekodera IC6 4 linie na 10 linii, a ten z kolei

zapala odpowiednie diody LED za pomocą diod D11...D19, więc nie koliduje to z kontrolą diod LED przez IC2. Jeśli IC6 otrzyma wartość wejściową 0, jego linie wyjściowe od 1 do 10 pozostają w stanie wysokim i żadne diody LED nie są przez niego zapalane.

Można by się zastanawiać, dlaczego nie użyłem pinów PE i TE układu IC7 do wstrzymania zliczania, zamiast zatrzymywania zegara. Powodem jest to, że wejścia te nie powinny być przełączane, gdy impuls zegarowy jest niski i nie ma synchronizacji między momentem zatrzymania zegara przez człowieka umieszczającego „X” na wolnym miejscu a stanem impulsu zegarowego w tym momencie.

Podczas gdy poszczególne diody LED zapalone przez wyjście randomizatora po jego zatrzymaniu pozostają włączone do końca gry, pierwszy żeton „O” jest umieszczany nad nimi, więc nie są one widoczne.

Jeśli człowiek (X) zaczyna jako pierwszy, żeton gracza O jest umieszczany na wolnym miejscu. Powoduje to, że pin 8 układu IC3 przechodzi w stan niski, zerując licznik 74HCT161, więc randomizator nie zapala żadnych diod LED. Więcej diod LED zapala się dopiero po umieszczeniu pierwszego żetonu „X”, aby zapewnić następny ruch maszyny.

Sekwencer gier wykorzystujący parzystość

Komputer powinien przejść w „tryb obliczeniowy” tylko wtedy, gdy nadejdzie jego kolej na wykonanie ruchu (umieszczenie żetonu „O”). Wzorce danych wygenerowane przez płytkę kodera, tuż po umieszczeniu wybranego „O”, nie mają znaczenia dla komputera, ponieważ człowiek nie umieścił jeszcze swojego następnego „X”.

Pytanie „kiedy pozwolić komputerowi na obliczenia” ma dwie odpowiedzi, w zależności od tego, kto pierwszy rozpocznie grę.

W przypadku, gdy maszyna zaczyna jako pierwsza, początkowo na planszy znajduje się jeden żeton „O” (wybrany początkowo przez randomizator), a następnie żeton „X” jest umieszczany przez człowieka. Teraz na planszy znajdują się dwa żetony, czyli liczba parzysta, i nadszedł czas na następny ruch maszyny. Po wykonaniu ruchu przez maszynę całkowita liczba żetonów wynosi trzy, co jest liczbą nieparzystą, co oznacza, że teraz nie jest ruch maszyny. Jednakże, jeśli człowiek zaczyna pierwszy i umieszcza swój żeton „X”, to ponieważ jeden jest liczbą nieparzystą, więc jest to czas na obliczenia w celu określenia ruchu maszyny. Po umieszczeniu „O” liczba żetonów staje się parzysta, w tej sytuacji komputer czeka, ponieważ jest to ponownie ruch X. Gracz X gra ponownie, co powoduje,

że liczba staje się nieparzysta, wprowadzając komputer w tryb obliczeniowy.

Zdałem sobie sprawę, że mogę rozwiązać ten problem za pomocą układu kontroli parzystości, z selektorem danych na jego dwóch wyjściach. Jego zadaniem będzie sterowanie linią /COMPUTE poprzez wybieranie nieparzystego lub parzystego wyjścia układu. Dziewięć sekcji z trzech poczwórnych 2-wejściowych bramek OR (IC13, IC15 i IC16) łączy linie „X” i „O” z planszy do gry

w dziewięć sygnałów „obecny element”, które następnie trafiają do układu parzystości IC9 (74HCT280).

Jego wyjścia EVEN i ODD trafiają do poczwórnej 2-wejściowej bramki NAND IC5 wraz z sygnałami HS10x i HS10o, które wskazują, który gracz rozpoczął jako pierwszy. Umożliwia to układowi IC5 wygenerowanie sygnału /COMPUTE, który pochodzi z wyjścia na nóżce 6 układu IC5. Podsumowując, zależy to od tego, czy na planszy znajduje

Wykaz elementów:

- 1 zestaw części do wykonania obudowy (patrz poniżej)
- 1 dwustronna płytka drukowana z kodem 08111221, 138 mm × 166 mm („plansza gry”)
- 1 dwustronna płytka drukowana z kodem 08111222, 138 × 124 mm (płytką obliczeniową)
- 1 40-stykowe gniazdo szpilkowe, raster 2,54 mm (CON1)
- 1 40-stykowe złącze, raster 2,54 mm (CON2)
- 1 brzęczyk piezoelektryczny lub sygnalizator dźwiękowy, raster wyprowadzeń 7,5 mm
- 1 32-pinowa podstawka DIL IC (opcjonalne dla IC1)
- 3 16-pinowe podstawki DIL IC (opcjonalne dla IC2, IC6 i IC7)
- 16 14-pinowych podstawek DIL IC (opcjonalne dla IC3, IC5, IC8...13, IC15, IC16 i IC21...26)
- 1 8-stykowa podstawka DIL IC (opcjonalna dla IC4)
- 4 tuleje dystansowe 10 mm z gwintem M3
- 8 śrub z łbem walcowym M3 × 6 mm
- 1 drut miedziany ocynowany o średnicy 0,7 mm lub drut dzwonkowy o długości 100 mm

Półprzewodniki

- 1 W27C020 szybka pamięć EEPROM o niskim poborze mocy 256 KB lub jej odpowiednik, DIP-32, zaprogramowana danymi z pliku 0811122A.bin (IC1)
- 2 dekodery BCD na dziesiętny 74HCT42, DIP-16 (IC2, IC6)
- 2 poczwórne 2-wejściowe bramki NAND 74HCT132, DIP-14 (IC3, IC5)
- 1 układ czasowy 555, DIP-8 (IC4)
- 1 74HCT161 4-bitowy licznik synchroniczny z możliwością ustawiania, DIP-16 (IC7)
- 1 pojedyncza 8-wejściowa bramka NAND 74HCT30, DIP-14 (IC8)
- 1 74HCT280 9-bitowy generator parzystości, DIP-14 (IC9)
- 3 potrójne 3-wejściowe bramki NAND 74HCT10, DIP-14 (IC10...IC12)
- 3 74HCT32 poczwórne 2-wejściowe bramki OR, DIP-14 (IC13, IC15, IC16)
- 6 poczwórnych wzmacniaczy operacyjnych LM324 z pojedynczym zasilaniem, DIP-14 (IC21-IC26)
- 10 DRV5055A1QLPG liniowe czujniki hallotronowe, TO-92 (HS1...HS10)
- 1 7805 lub LM2940CT-5 (patrz tekst) stabilizator liniowy 5 V 1 A (REG1)
- 1 BS270 lub 2N7000 MOSFET N-kanalowy, TO-92 (Q1)
- 36 niebieskich diod LED 3 mm (LED1...LED36) [Jaycar ZD0134]
- 1 1N5819 dioda Schottky'ego 40 V 1 A (D1)
- 19 1N4148 małe diody sygnałowe 75 V 250 mA (D2...D20)

Kondensatory

- 2 1,5 µF 50 V MKT lub ceramiczny wielowarstwowy
- 3 1 µF 50 V MKT lub ceramiczny wielowarstwowy
- 3 100 nF 50 V MKT lub ceramiczny wielowarstwowy
- 1 10 nF 63 V MKT

Rezystory (wszystkie 1/4 W 1% osiowe)

- 1 470 kΩ 6 10 kΩ 2 5,1 kΩ
- 36 1,6 kΩ 2 2 kΩ 2 1 kΩ

Obudowa

- 1 obrobiona i wygrawerowana pokrywa wykonana z akrylu o grubości 10 mm, 160 mm × 200 mm
- 1 obrobiony i wygrawerowany panel górny (akryl o grubości 10 mm), 160 mm × 200 mm
- 1 arkusz 160 mm × 200 mm przydymionego, półprzezroczystego akrylu o grubości 3 mm (na podstawę)
- 2 arkusze 200 mm × 40 mm akrylu o grubości 10 mm (panele boczne)
- 2 arkusze 140 mm × 40 mm akrylu o grubości 10 mm (panele przedni i tylny)
- 1 zawias 150 mm
- 2 małe zatrzaski/zapięcia pokrywy
- 4 przykręcane gumowe nożki
- 20 śrub maszynowych z łbem stożkowym i gniazdem sześciokątnym o długości 10 mm i gwintem 4-40 UNC (do mocowania podstawy i panelu górnego)
- 10 śrub z łbem sześciokątnym 4-40 UNC o długości 10 mm (do zatrzasków i nożek)
- 8 wkrętów maszynowych M2 z łbem stożkowym i gniazdem sześciokątnym o długości 10 mm (do mocowania zawiasu)
- 8 metalowych wkładek gwintowanych M2 o długości 10 mm i średnicy 4 mm (do zawiasu)
- 4 nakrętki sześciokątne 4-40UNC (do mocowania nożek)
- 1 gniazdo baryłkowe do montażu w obudowie LUB
- 1 uchwyt baterii 9 V lub
- 1 uchwyt na 6 ogniw AAA (CON3)
- 1 200 mm długości lekkiego drutu osemkowego

Elementy do gry

- 10 czarnych plastikowych żetonów o średnicy 25 mm i grubości 10 mm
- 10 słabych magnesów ferrytowych [Jaycar LM1616]



W żetonach do gry są zastosowane tylko słabe magnesy

się parzysta czy nieparzysta liczba żetonów i kto zaczął pierwszy.

Niejednoznaczność stanu gry

Należy zauważyć, że w przypadku, gdy gracz O (maszyna) gra jako pierwszy i na płytce nie ma żadnych żetonów, linia /COMPUTE jest w stanie niskim, umożliwiając obliczenia. Jednakże, gdy na płytce nie ma żadnych elementów, wszystkie linie adresowe są wyzerowane, a jak wspomniano wcześniej, dane pod adresem 0 w pamięci EEPROM mają wartość szesnastkową 0xFF. Gdy 74HCT42 ma wszystkie cztery bity w stanie wysokim, nie zaświeci żadnych diod LED na płytce.

X STARTED FIRST

		O4
O2	X3	
X1		

Human X's next move,
no calculation required

O STARTED FIRST

		O1
O3	X2	
X4		

Machine O's next move,
calculation is required

Note identical binary addresses

Rysunek 6. Graficzne wyjaśnienie, dlaczego ten sam stan planszy może zostać osiągnięty w dwóch różnych grach, z których jedną rozpoczyna człowiek (X1, u góry), a drugą maszyna (O1, u dołu). Liczby wskazują kolejność ruchów, natomiast znaki „X” i „O” pokazują, który gracz umieszcza żeton na którym polu

Należy zauważyć, że ROM został zaprogramowany tak, aby generować prawidłowe wartości wyjściowe tylko dla prawidłowych kombinacji wejść, odpowiadających rzeczywistym układom żetonów na planszy. Skoro tak, to można zapytać, po co w ogóle potrzebny jest sekwencer gry – skoro błędne adresy lub stany i tak spowodowałyby wystawienie wartości 0xFF, a więc żadna dioda LED by się nie zapaliła.

Układy sekwencera nie byłyby potrzebne, gdyby gra była zaprojektowana tak, by jeden z graczy (człowiek lub maszyna) zawsze zaczynał jako pierwszy.

Na rysunku 6 pokazane są dwa możliwe identyczne wzorce, które można osiągnąć poprzez różne sekwencje gry. Dlatego generują one ten sam adres dla pamięci EEPROM. W jednym przypadku X zaczął pierwszy, natomiast w drugim przypadku pierwszy zaczął O. W następnym ruchu piątym umieszczonym żetonem może być „X” lub „O”, w zależności od tego, kto rozpoczął grę.

Właśnie dlatego sekwencer gry z układem parzystości był wymagany, ponieważ w tych dwóch przypadkach daje inny stan dla linii /COMPUTE, mimo że adres EEPROM jest identyczny.

Dźwięki w grze

Gdy maszyna wygrywa z człowiekiem, emituje sygnał dźwiękowy. W układzie znajduje się generator sterujący brzęczykiem piezoelektrycznym (oparty na układzie czasowym 555 IC4). Jednak w moim prototypie użyłem brzęczyka z wbudowanym generatorem, element14 Cat 107-2397, więc generator został pominięty.

Do gracza-człowieka nie są przypisane żadne funkcje sygnału dźwiękowego, ponieważ może on sam wydawać dźwięki, jeśli chce. Mogą, zwłaszcza gdy maszyna nie tylko nie pozwala im wygrać, ale stresuje ich jeszcze bardziej, pokonując ich przy najmniejszym błędzie lub braku uwagi. Ponieważ człowiek nigdy nie może pokonać maszyny, automatyczny komunikat o wygranej człowieka nie jest wymagany.

W związku z tym konieczne jest jedynie sprawdzenie danych żetonów „O” z planszy (a nie danych „X”) pod kątem ośmiu możliwych konfiguracji ułożenia elementów O, które wskazują na wygraną.

Opcje zasilania

Nie byłem pewien, czy pobór prądu w układzie może przekroczyć 100 mA, więc zbudowałem go ze stabilizatorem 7805. Jak się okazało, prąd zmienia się w zakresie od 75 do 90 mA. Mój

zasilany jest z zasilacza wtyczkowego Jaycar 9 V 0,5 A. Jak się okazuje, zastosowanie mniejszego stabilizatora 78L05 w obudowie TO-92 byłoby możliwe.

Gdy rozlega się sygnał dźwiękowy buzzera piezoelektrycznego prąd na krótko osiąga wartość ponad 100 mA, ale układ 78L05 powinien wytrzymać takie krótkotrwałe przeciążenie.

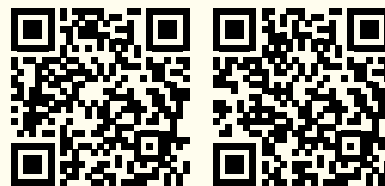
Ze względu na niski pobór mocy, gra może być zasilana baterią alkaliczną 9 V o pojemności około 500 mAh lub baterią litowo-jonową 9 V o pojemności 1200 mAh. Wewnątrz obudowy jest jednak dużo miejsca na sześć ogniw AAA w pojemniku, co znacznie wydłużyłoby czas pracy w porównaniu ze standardową baterią 9 V.

W przypadku zasilania baterijnego rozsądnie byłoby pozostawić w układzie diodę 1N5819 (D1) (aby zapobiec błędom związanym z odwrótną polaryzacją). Lepszy od 7805 byłby stabilizator 5 V o niskim spadku napięcia (LDO). Zastosowanie takiego układu pozwoliłoby maksymalnie wykorzystać żywotność baterii. Na przykład można użyć LM2940CT-5.0, który jest kompatybilny pod względem wyprowadzeń i przy prądzie 100 mA wymaga napięcia wejściowego wyższego tylko o 110 mV względem wyjściowego. Dla układu 7805 różnica tych napięć wynosi około 1,5 V.

W następnym odcinku

W drugiej i ostatniej części artykułu wyjaśnię strategię rozgrywki i sposób, w jaki wygenerowałem dane rozgrywki dla pamięci EEPROM. Omówię też proces montażu płytki drukowanej, konstrukcję obudowy i połączenie wszystkiego w działającą grę. ■

dr Hugo Holden

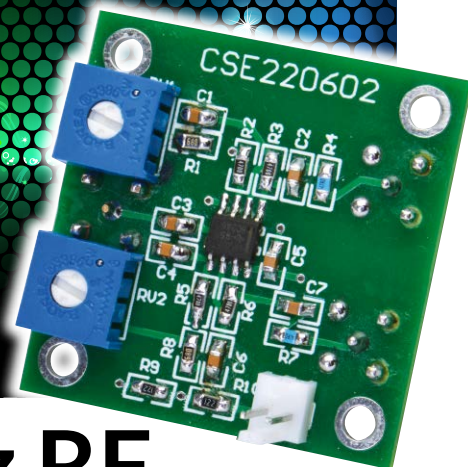


Materiały dodatkowe dostępne są na stronie Silicon Chip:
<https://www.siliconchip.com.au/Shop/8/6644>
<https://www.siliconchip.com.au/Shop/8/6645>

Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022.
www.siliconchip.com.au

Opisany w artykule niewielki wzmacniacz RF ma dwa wyjścia z indywidualnie wybieranymi wzmocnieniami. Dzięki temu można go dołączać do generatora sygnału w celu zapewnienia wyższego poziomu wyjściowego, większej mocy sygnału lub możliwości rozgałęziania sygnału do wielu urządzeń.



Podwójny wzmacniacz RF przeznaczony do wzmacniania sygnałów z generatorów sygnału

Duża część generatorów sygnału nie zapewnia wystarczająco wysokiego poziomu wyjściowego dla niektórych zastosowań. Na przedstawionej w artykule niewielkiej płytce został zamontowany szybki podwójny wzmacniacz operacyjny OPA2677 używany do wzmacniania sygnałów z zakresu 100 kHz...75 MHz z mocą od około 0 dBm (1 mW, 225 mV/-13 dBV na 50 Ω) do około 18 dBm (63 mW, 1,78 V/5 dBV na 50 Ω).

Układ OPA2677 odznacza się bardzo dobrymi parametrami technicznymi. Może pracować z napięciami od 3,3 V do 12 V, ma wyjścia typu rail to rail, wysoką zdolność sterowania i parametr GBW wynoszący 200 MHz (iloczyn pasma i wzmocnienia). Ale to, co go wyróżnia, to szybkość narastania wynosząca 1800 V/μs, co oznacza, że może zapewnić duże zmiany napięcia wyjściowego dla sygnałów o wysokiej częstotliwości.

Układ OPA2677 to podwójny wzmacniacz operacyjny i dlatego mój projekt zapewnia

dwa wyjścia dla jednego sygnału wejściowego. Poszczególne rezystory sprzężenia zwrotnego i potencjometr ustawiają wzmocnienie dla każdego wyjścia.

Maksymalne wzmocnienie wynosi $1 + (470 \Omega / 68 \Omega) = 7,9 \text{ V/V}$ przy jednoobrotowym potencjometrze montażowym 1 kΩ ustawionym na minimum. Najniższe wzmocnienie wynosi $1 + (470 \Omega / 1068 \Omega) = 1,44 \text{ V/V}$ przy ustawieniu potencjometru na maksimum. Impedancja wyjściowa wynosi 50 Ω i bezpiecznie wysteruje obciążenie 50 Ω.

Napięcie zasilania powinno idealnie mieścić się w zakresie 9 V...12 V. Można użyć zasilania 5 V DC, ale wzmocnione sygnały będą na wyjściu wzmacniacza operacyjnego ograniczone do 5 V_{pp} i 2,5 V_{pp} na obciążeniu 50 Ω lub 884 mV RMS (13,9 dBm/24 mW). Maksymalna moc wyjściowa przy zasilaniu 12 V wynosi około 25 dBm, jak pokazano w ramce specyfikacji.

Wzmacniacz jest użyteczny w zakresie od 100 kHz do 75 MHz, chociaż poprzekroczeniu

50 MHz maksymalny poziom wyjściowy zaczyna spadać. W tabeli 1 zamieszczono wyniki punktowych pomiarów dla kilku częstotliwości przy użyciu mojego generatora sygnału jako wejścia. Zmiany na wyjściu zależą w pewnym stopniu od zmian poziomu wyjściowego generatora sygnału.

Układ OPA2677 nie jest tani, kosztuje około 9 USD w sklepie Digi-Key, Mouser lub element14, ale kupiłem pięć z AliExpress za 14,50 USD. Mimo to, nawet jeśli zapłacisz 9 USD, całkowity koszt budowy tego wzmacniacza jest niewielki. Zobacz ramkę na końcu artykułu w której znajduje się krótki opis zestawu.

Opis schematu

Pełny schemat układu pokazano na rysunku 1. Zastosowano w nim sprzężenie zmiennoprądowe sygnału wejściowego podawanego przez złącze SMA CON1 do obu połówek podwójnego wzmacniacza operacyjnego IC1. Użyto do tego kondensatorów o pojemności 100 nF. Sygnały te są podciągane do połowy napięcia zasilającego VCC (np. 2,5 V dla zasilania 5 V lub 6 V dla zasilania 12 V) za pomocą rezystorów 470 Ω.

Kondensatory sprzęgające i rezystory podciągające tworzą filtry górnoprzepustowe o częstotliwości granicznej 3,4 kHz ($1 / [2 \pi F \times 100 \text{ nF} \times 470 \Omega]$), więc nie będą tłumić sygnałów w określonym zakresie częstotliwości roboczych, od 100 kHz do 75 MHz.

Cechy i specyfikacja

- Zakres częstotliwości pracy: od 100 kHz do 75 MHz
- Liczba wejść: 1
- Liczba wyjść: 2, indywidualnie regulowane wzmocnienie
- Zakres wzmocnienia: 1,44× (3 dB) do 7,9× (18 dB)
- Maksymalny poziom wyjściowy:
 - 25,6 dBm przy 30 MHz (360 mW na 50 Ω, 12,5 dBV, 4,25 V RMS)
 - 23,2 dBm przy 50 MHz (207 mW na 50 Ω, 10 dBV, 3,2 V RMS)
 - 13,5 dBm przy 70 MHz (22 mW na 50 Ω, 0,51 dBV, 1,06 V RMS)
- Zasilanie: 9 V...12 V DC przy 20 mA...25 mA (lub 5 V DC ze zredukowanymi maksymalnymi poziomami wyjściowymi)

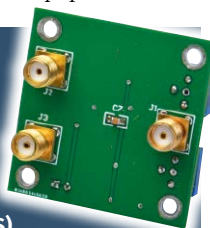


Tabela 1. Częstotliwość a maksymalny poziom wyjściowy przy zasilaniu 12 V DC

Częstotliwość	Wyjście (p-p)	Wyjście (RMS)	Wyjście (dBm)	Wyjście (dBV)
1 MHz	9,5 V	3,36 V	23,5	10,5
10 MHz	8,4 V	2,97 V	22,5	9,5
20 MHz	10,0 V	3,54 V	24,0	11,0
30 MHz	12,0 V	4,24 V	25,6	12,5
40 MHz	9,6 V	3,39 V	23,6	10,6
50 MHz	9,1 V	3,22 V	23,2	10,2
60 MHz	5,6 V	1,98 V	18,9	5,9
70 MHz	3,0 V	1,06 V	13,5	0,51

Sygnały są doprowadzone do nieodwracających wejść, więc wzmacniacze nie odwracają fazy sygnału. Wyjścia wzmacniaczy operacyjnych (piny 1 i 7) są połączone z powrotem do wejść odwracających (piny 2 i 6) przez rezystory 470 Ω , które z potencjometrami montażowymi i ich szeregowymi rezystorami 68 Ω tworzą dzielniki napięcia VR1/VR2.

Kondensatory 100 nF w obwodzie sprzężenia zwrotnego zmniejszają wzmocnienie DC tych wzmacniaczy do 1 $\times$, dzięki czemu wejściowe napięcia offsetu (do 5,3 mV) nie są wzmacniane. Częstotliwość graniczna utworzonego filtra górnoprzepustowego jest zbliżona do częstotliwości obwodów wejściowych, ponieważ wartości elementów są takie same.

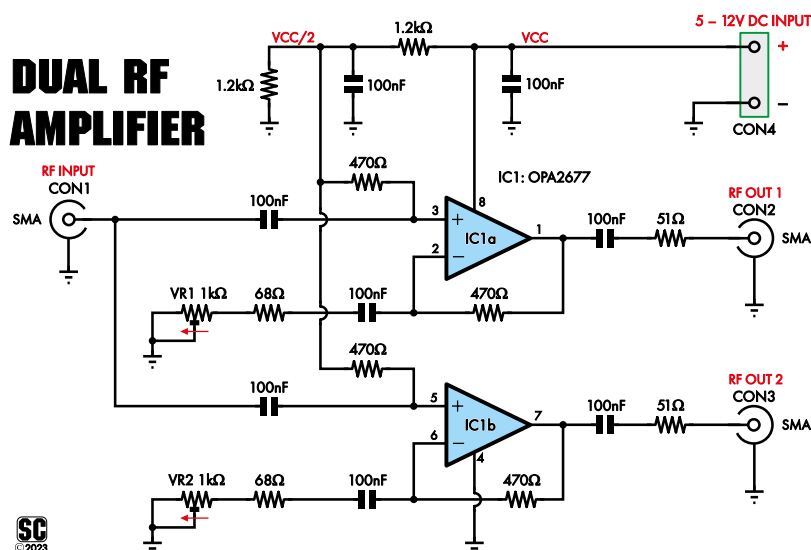
Jak wspomniano wcześniej, wzmacniacze operacyjne mają bardzo szerokie pasma wzmocnienia (GBW) i szybkość narastania napięcia wyjściowego, więc nadają się świetnie do układów wysokich częstotliwości. Ponieważ szerokość pasma wzmocnienia jest stała, maksymalna częstotliwość

sygnału spada wraz ze wzrostem wzmocnienia. Na przykład, przy GBW równym 200 MHz, możliwe jest wzmocnienie 4 $\times$ przy 50 MHz lub około 3 $\times$ przy 70 MHz.

Wyjścia dwóch wzmacniaczy operacyjnych są podłączone do złączy SMA za pośrednictwem kondensatorów 100 nF. Dzięki temu została wyeliminowana polaryzacja VCC/2 DC. Rezystory 51 Ω zapewniają dopasowania impedancji. Można je zmienić na 75 Ω , jeśli konieczna jest współpraca z urządzeniami o impedancji 75 Ω .

Szyna VCC/2 jest utworzona przez prosty dzielnik napięcia 1,2 k Ω /1,2 k Ω z kondensatorem 100 nF połączonym między złączem a masą w celu wyeliminowania tętnienia zasilania i utrzymania niskiej impedancji źródła przy wyższych częstotliwościach. Wzmacniacz operacyjny IC1 ma również obowiązkowy kondensator blokujący zasilanie o pojemności 100 nF.

Należy pamiętać, że wejście CON1 nie jest wyposażone w rezystor terminujący.



Rysunek 1. Podwójny wzmacniacz RF to prosta implementacja podwójnego wzmacniacza operacyjnego OPA2677 o wysokiej wydajności. Zastosowano sprzężenie zwrotnopiętrowe na wejściach i wyjściach, dzięki czemu sygnały mogą być spolaryzowane z offsetem równym połowie napięcia zasilającego za pomocą dwóch rezystorów i kondensatora. Potencjometry montażowe VR1 i VR2 regulują współczynnik sprzężenia zwrotnego, a tym samym wzmocnienie każdego wzmacniacza

W razie potrzeby do zacisków gniazda SMA można dodać rezystor o rozmiarze M2012/0805 (51 Ω lub 75 Ω).

Budowa

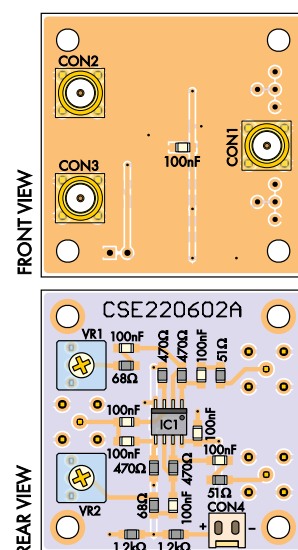
Konstrukcja jest stosunkowo prosta, ponieważ układ wykorzystuje zaledwie kilkanaście komponentów. Podwójny wzmacniacz RF jest zmontowany na dwustronnej płytce drukowanej oznaczonej kodem CSE220602A o wymiarach 38 mm $\times$ 38 mm. Zapoznaj się ze schematami montażowymi PCB (**rysunki 2 i 3**), co ułatwi Ci montaż.

Rozpocznij od wlutowania układu IC1 umieszczonego po stronie elementów. Określ lokalizację pinu 1 – poszukaj kropki lub wgłębienia w jednym rogu lub, jeśli to niemożliwe, ściętej krawędzi po stronie pinu 1. Umieść go z pinem 1 w prawym górnym rogu, z płytką drukowaną zorientowaną jak pokazano na rysunku 2.

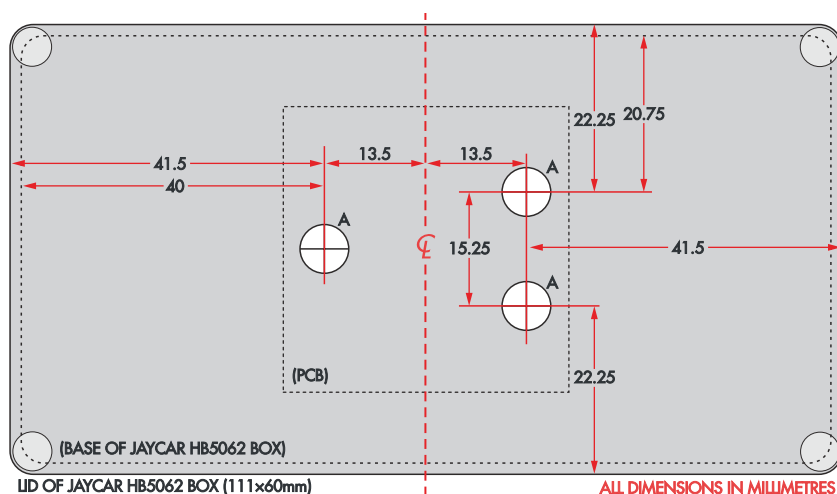
Nalóż topnik na pady, a następnie przylutuj jeden pin odrobiną lutowia i sprawdź wyrównanie pozostałych pinów. Jeśli rozmieszczenie jest poprawne, przylutuj pin znajdujący się po przekątnej. W przeciwnym razie podgrzej oryginalne połączenie lutowane i delikatnie skoryguj położenie elementu, aż znajdzie się na swoim miejscu.

Następnie przylutuj pozostałe piny, popraw pierwszy i wyczyść mostki lutownicze, które mogły powstać między pinami za pomocą kolejnej kropli topnika i plecionki lutowniczej.

Wyczyść pozostałości topnika z płytki alkoholem lub środkiem do czyszczenia topnika i sprawdź połączenia lutowane, aby mieć pewność, że wszystkie są wykonane poprawnie.



Rysunki 2 i 3. Większość elementów jest przeznaczona do montażu SMD. Montuje się je z tyłu, natomiast jeden kondensator i trzy złącza SMA znajdują się z przodu. Strona płytki z gniazdami RF jest pokryta płaszczyzną masy



Rysunek 4. Prawie każda metalowa obudowa będzie odpowiednia, ale ta zaproponowana jest bardzo kompaktowa. Pokrywa jest większa niż podstawa, więc jeśli używasz jej jako szablonu, przycinaj ją do odpowiedniego konturu. Centralny obszar można wyciąć i przenieść do dowolnej innej obudowy. Otwór z boku na gniazdo zasilania nie został tutaj pokazany, można go umieścić w dowolnym miejscu

Następnie przystąp do montażu elementów pasywnych, które nie mają biegunowości. Lutuj je podobnie jak poprzednie – najpierw przylutuj jeden koniec, wycentruj element i po krótkiej chwili, gdy lut zastygnie, przylutuj drugi koniec.

Rezystory będą oznaczone kodami wskazującymi ich wartości (np. 122 lub 1201 dla 1,2 kΩ), natomiast kondensatory nie będą oznaczone, ale wszystkie mają tę samą wartość (100 nF). Gdy wszystkie elementy SMD zostaną zamontowane na górnej warstwie, odwróć płytke i przylutuj pojedynczy kondensator po drugiej stronie.

Pozostaje jeszcze tylko sześć elementów przewlekanych: dwa potencjometry montażowe, gniazdo zasilania i trzy gniazda SMA. Najlepiej jest zamontować gniazda SMA w następnej kolejności, aby mieć dobry dostęp do ich pinów. Wciśnij je całkowicie i przylutuj wszystkie pięć pinów, pamiętając, że ze względu na ich masę termiczną lutowanie czterech zewnętrznych pinów może wymagać dostarczenia większej ilości ciepła lub topnika.

Na koniec zamontuj dwa potencjometry i gniazdo. Użyj jednoobrotowych potencjometrów montażowych, ponieważ wieloobrotowe mają prawdopodobnie zbyt dużą indukcyjność. Zasilający przewód ośmiokowy możesz przylutować bezpośrednio do płytki, ale spolaryzowane złącze jest wygodniejsze. Jego dokładna orientacja nie ma znaczenia, o ile podczas podłączania przestrzegasz oznaczeń + i –.

Obudowa

Płytką jest niewielka, zmieści się więc w większości obudów. Ze względu na ekranowanie RF preferowana jest obudowa metalowa. Warto zwrócić uwagę na to, że odlewane obudowy o wymiarach 51 mm × 51 mm sprzedawane przez Jaycar i Altronics są zbyt małe, mogą nie pomieścić płytki drukowanej.

Na **rysunku 4** pokazano pozycje otworów, które należy wywiercić w pokrywie lub podstawie, a następnie można zamontować płytkę za pomocą nakrętek złącza SMA.

Wywierć otwór z boku obudowy, w celu dopasowania gniazda baryłkowego do montażu w obudowie i podłącz je do CON4. Płytką nie ma zabezpieczenia przed odwrotną polaryzacją, dlatego upewnij się, że przewód dodatni (zwykle pin na tylnej ścianie gniazda baryłkowego) jest podłączony do strony + CON4.

Nie ma potrzeby stosowania wyłącznika zasilania, ponieważ można po prostu odłączyć wtyczkę zasilania. Jeśli jednak chcesz dodać przełącznik zasilania, wszystko co musisz zrobić, to wywiercić otwór w dogodnym miejscu, zamontować przełącznik i podłączyć go szeregowo z dodatnim przewodem z gniazda baryłkowego do CON4.

Jeśli chcesz dodać zabezpieczenie przed odwrotną polaryzacją, przylutuj diodę Schottky'ego 1N5819 do gniazda zasilania typu baryłkowego tak, aby jej anoda była połączona z dodatnim stykiem gniazda, a katoda z przewodem prowadzącym do płytki lub przełącznika. Spowoduje to niewielki spadek napięcia zasilania – około 0,3 V – co może nieznacznie zmniejszyć maksymalny poziom sygnału wyjściowego.

W zależności od potrzeb i zastosowania można wywiercić kilka małych otworów w obudowie, aby można było włożyć cienki śrubokręt do regulacji potencjometrów VR1 i VR2 po jej zamknięciu.

Zamiast tego można po prostu ustawić inne stałe wzmocnienie za pomocą obu potencjometrów, a następnie użyć tego, które odpowiada Twoim potrzebom w danym momencie.

Przed przykręceniem pokrywy należy odłączyć wtyczkę CON4 od płytki, podłączyć

Wykaz elementów:

- 1 dwustronna płytka drukowana oznaczona kodem CSE220602A, 38 mm × 38 mm
- 1 odlewana aluminiowa obudowa, wystarczająco duża, aby zmieścić płytkę [np. Jaycar HB5062, 111 mm × 60 mm × 30 mm],
- 1 źródło zasilania 9 V...12 V DC 50 mA
- 1 OPA2677IDDA podwójny wzmacniacz operacyjny o wysokiej wydajności, SOIC-8 [element14, Mouser, Digi-Key],
- 2 jednoobrotowe górne potencjometry montażowe typu 3362P o rezystancji 1 kΩ (VR1, VR2)
- 8 kondensatorów ceramicznych SMD 100 nF 50 V X7R, rozmiar M2012/0805
- 3 pionowe gniazda żeńskie SMA (CON1...CON3)
- 12-stykowe spolaryzowane złącze z dopasowaną wtyczką i stykami (CON4)
- 1 gniazdo zasilania montowane w obudowie dopasowane do wtyczki
- 1 krótki odcinek kabla typu 8
- 1 przełącznik SPDT do montażu w obudowie (opcjonalny; przełącznik zasilania)
- 1 dioda Schottky'ego 1N5819 (opcjonalnie; patrz tekst)

Rezystory

2 1,2 kΩ	4 470 Ω	2 68 Ω	2 51 Ω
----------	---------	--------	--------

zasilacz do gniazda baryłkowego i za pomocą multimetru sprawdzić, czy polaryzacja zasilania na wtyczce jest prawidłowa. Następnie podłącz urządzenie i doprowadź sygnał do gniazda wejściowego.

Sprawdź, za pomocą oscyloskopu, miernika poziomu sygnału lub miernika częstotliwości, w zależności od tego czym dysponujesz, czy na wyjściach pojawia się wzmocniony sygnał.

Obsługa

Obsługa urządzenia jest bardzo prosta. Wystarczy włączyć zasilanie, podać sygnał, wyregulować, jeśli to konieczne, poziom za pomocą potencjometrów VR1 lub VR2 i odebrać sygnał wyjściowy z odpowiedniego gniazda. Poziom/wzmocnienie sygnału CON2 reguluje się za pomocą VR1, a poziom/wzmocnienie sygnału CON3 reguluje się za pomocą VR2.

Należy pamiętać, że VR1 i VR2 są podłączone w taki sposób, że obrót w kierunku przeciwnym do ruchu wskazówek zegara zwiększa wzmocnienie, a obrót w kierunku zgodnym z ruchem wskazówek zegara zmniejsza je. ■

Charles Kosina



Materiały dodatkowe dostępne są na stronie: <https://www.siliconchip.com.au/Shop/8/6744>

Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. www.siliconchip.com.au



FN-SWM10

Zgrzewarka do ogni – spawarka punktowa z kolorowym wyświetlaczem i funkcją powerbank FNIRSI SWM10



FN-DPOS-350P

Dwukanałowy oscyloskop 350 MHz, FNIRSI DPOS350P



FN-2C53T

Dwukanałowy oscyloskop z multimetrem i generatorem 50 MHz FNIRSI 2C53T

BESTSELLERY sklepu AVT – sklep.avt.pl

Mierniki Testery FNIRSI

Rabat dla Czytelników EdW
przy zakupie podaj kod **EdW2505FN**

Kod ważny do 30.09.2025

-3%

Rabat dla Prenumeratorów EdW
przy zakupie podaj numer prenumeraty

-6%



FN-LCR-ST1

Miernik pęsetowy, tester elementów FNIRSI LCR-ST1



FN-LCR-P1

Tester elementów FNIRSI LCR-P1



FN-HRM10

Tester rezystancji wewnętrznej akumulatorów FNIRSI HRM-10



FN-G1200

Mikroskop cyfrowy G1200 z wyświetlaczem 7 cali, powiększenie x1200, tryb foto/video



FN-DWS200-F245

Stacja lutownicza 200 W z kolbą F245, FNIRSI DWS200



FN-1014D

Oscyloskop dwukanałowy 100 MHz, Generator sygnału DDS, FNIRSI 1014D

Termostat z termoparą typu K

Za pomocą opisanego w artykule termometru można łatwo mierzyć temperaturę w bardzo szerokim zakresie, i w zależności od wyniku pomiaru odpowiednio sterować zewnętrznym urządzeniem. Jako element pomiarowy wykorzystano termoparę typu K. Urządzenie jest wyposażone w przełącznik umożliwiający bezpośrednie sterowanie obciążeniem, dzięki czemu może pełnić funkcję termostatu pracującego w trybie ogrzewania lub chłodzenia.

Termometr/termostat z termoparą typu K (nazywany dalej jako termometr lub termostat) może mierzyć temperaturę w bardzo szerokim zakresie. Został w nim zastosowany przełącznik, który może sterować zasilaniem elementu grzejącego lub sprężarki lodówki.

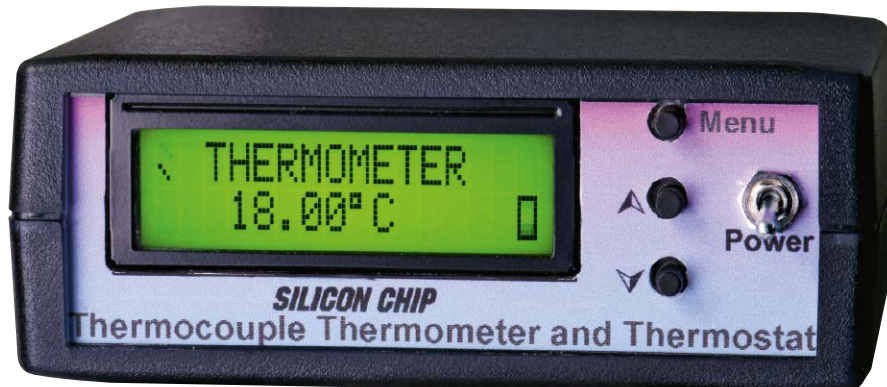
Duża część multimetrów cyfrowych może mierzyć temperaturę za pomocą termopary, zwykle nie mogą one jednak sterować urządzeniami w celu ogrzewania lub chłodzenia.

W przypadku ogrzewania, zasilanie może być włączane, gdy temperatura jest niższa od zadanej i wyłączane, gdy osiągnie zadaną wartość. Alternatywnie, zasilanie jest włączane w celu chłodzenia, gdy temperatura jest wyższa od zadanej i wyłączane, gdy spadnie poniżej progu.

Nasz termometr ma regulowaną histerzę zapobiegającą szybkiemu przełączaniu przełącznika w pobliżu progu regulacji. Histeresa wprowadza pewną różnicę między temperaturami, w których przełącznik będzie włączał i wyłączał sterowane urządzenie. Regulacja możliwa jest w zakresie od 0°C do 60°C. Zwykle histeresa wynosi tylko około 1°C...2°C.

Temperatura jest wyświetlana na dwuwierszowym, 16-znakowym wyświetlaczu LCD. Urządzenie może wprawdzie wyświetlać temperaturę od -270°C do +1800°C, rzeczywisty zakres zależy jednak od użytej sondy. Niektóre sondy typu K działają w zakresie od -50°C do +250°C, inne od -50°C do +900°C, a jeszcze inne od -40°C do +1200°C. Są też takie sondy, które działają tylko powyżej 0°C.

Sondy z termoparą mogą być izolowane lub nieizolowane. Izolowane sondy nie



mają połączenia elektrycznego z termoparą, więc sonda może dotykać materiału, który jest uziemiony lub znajduje się pod pewnym napięciem stałym bez generowania błędnych odczytów.

Nieizolowane sondy nie powinny być używane w miejscach, w których występuje różnica potencjałów między masą termometru a sondą. Nasz termometr pokaże w takim przypadku błąd (zwarcie do masy lub zwarcie do zasilania).

Termometr jest umieszczony w niewielkiej obudowie z elementami sterującymi z przodu służącymi do załączania bądź wyłączania zasilania, wyboru widoku wyświetlacza i regulacji ustawień.

Z tyłu obudowy znajdują się gniazda wejścia zasilania 12 V DC i termopary typu K. Znajduje się tam również przepust kablowy do wprowadzenia przewodów do obudowy i podłączenia ich do styków przełącznika termostatu za pomocą zaciśków śrubowych. Dostępne są styki wspólne (C), normalnie otwarte (NO) i normalnie zamknięte (NC).

Zasada działania termopary typu K

Termopara typu K tworzona jest poprzez styk dwóch różnych przewodów. W tym typie sondy jeden przewód jest wykonany ze stopu chromu i niklu (chromelu), drugi natomiast jest wykonany

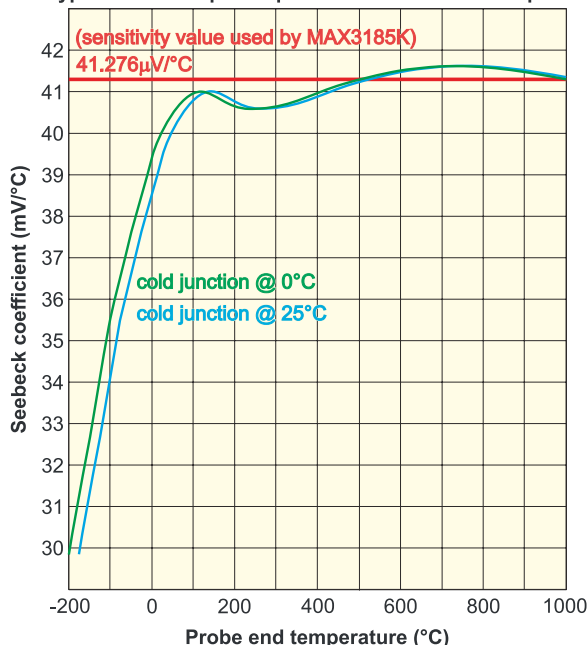
ze stopu aluminium, manganu, krzemu i niklu (o nazwie alumel).

Te dwa przewody stykają się ze sobą tylko na końcu sondy. Pozostałe końce przewodów łączą się z dwupinową wtyczką termometru.

Zasada działania termopary opiera się na zjawisku polegającym na tym, że na styku połączenia dwóch różnych metali wytwarza się napięcie zależne od temperatury. Termopara typu K ma nominalną czułość 41,276 $\mu\text{V}/^\circ\text{C}$.

Jednak korzystanie z tej jednej wartości ma swoje ograniczenia – czułość nie jest stała, ale zmienia się wraz z temperaturą. Na przykład

K-Type Thermocouple output coefficient versus temperature



Rysunek 1. Często uważa się, że termopara typu K ma liniową czułość 41,276 $\mu\text{V}/^\circ\text{C}$ (współczynnik Seebecka), ale w rzeczywistości zmienia się ona w pokazany sposób. Uzyskanie dokładnych odczytów wymaga uwzględnienia tej zmienności, zwłaszcza w niższych temperaturach

termopara typu K ma czułość $35,54 \mu\text{V}/^\circ\text{C}$ w temperaturze -100°C i $41,61 \mu\text{V}/^\circ\text{C}$ w temperaturze $+750^\circ\text{C}$. Jeśli przyjmie się stałą wartość czułości, jej zmienność wprowadzi błędy odczytu temperatury.

Czułość termopary typu K w zależności od temperatury pokazano na **rysunku 1**. Zmiana napięcia wyjściowego przypadająca na jeden $^\circ\text{C}$ nazywana jest współczynnikiem Seebecka. Odnosi się ona do zmiany napięcia spowodowanej różnicą temperatur między sondą a wtyczką termopary.

Typowy wykres przedstawia współczynnik Seebecka z końcówką wtykową termopary w temperaturze 0°C .

Współczynnik jest dość spójny w przedziale od 75°C do 1000°C , ale gwałtownie spada dla temperatur w obszarze ujemnym. Gdyby w naszym termometrze zastosowano wartość czułości $41,276 \mu\text{V}/^\circ\text{C}$, odczyty byłyby naprawdę dokładne tylko w trzech temperaturach: 0°C , 500°C i 1000°C .

Utrzymywanie temperatury wtyczki termopary (można to sobie utożsamiać z połączeniem drutu, z którego jest wykonana termopara i miedzianym przewodem pomiarowym) na poziomie 0°C nie jest zbyt wygodne. Ponadto temperatura końcówki wtyczki może zmieniać się wraz z temperaturą otoczenia. Sterownik termopary mierzy jej temperaturę i wykorzystuje ten odczyt do kompensacji odczytów na końcu sondy. Procedura ta nazywa się „kompensacją zimnego złącza” (koniec wtyczki jest zdefiniowany jako zimne złącze).

Pomimo nazwy, końcówka wtyczki niekoniecznie jest zimniejsza niż sonda, może być też cieplejsza.

Na **rysunku 1** umieściliśmy dodatkową krzywą, gdy zimne złącze ma temperaturę 25°C . Daje to wyobrażenie o przesunięciu wykresu przy różnych temperaturach zimnego złącza.

Jeśli temperatura zimnego złącza wynosi 25°C , a końcówka sondy termopary wskazuje 0°C , termopara w rzeczywistości mierzy -25°C . W tym miejscu wartość współczynnika Seebecka gwałtownie spada wraz ze spadkiem temperatury mierzonej przez termoparę. Uzyskanie dokładnych odczytów w tej części krzywej jest trudne.

W naszym termometrze został zastosowany układ scalony MAX31855 firmy Maxim. Zapewnia on cyfrowe wyjście danych odczytu termopary dostosowane do kompensacji zimnego złącza. Sam układ scalony mierzy temperaturę zimnego złącza. Daje to odczyt z dokładnością do $\pm 2^\circ\text{C}$ w zakresie od -200°C do $+700^\circ\text{C}$ (bez uwzględnienia błędów związanych z samą termoparą). Jednak

ta dokładność nie uwzględnia zmienności odczytów spowodowanej zmianami współczynnika Seebecka wraz z temperaturą. Zakłada on stały współczynnik Seebecka $41,276 \mu\text{V}/^\circ\text{C}$ w tym zakresie temperatur.

Korekta temperatury

Na **rysunku 2** pokazano wymaganą korektę temperatury. Ponownie, temperatura otoczenia zimnego złącza przesuwają krzywą od 0°C . Jako przykład pokazujemy krzywą zimnego złącza 25°C . Z wykresu można się zorientować, jaką wartość należy dodać lub odjąć od odczytu, aby uwzględnić zmiany Seebecka wraz z temperaturą.

Na przykład, jeżeli sonda mierzy rzeczywistą temperaturę 0°C przy temperaturze złącza zimnego wynoszącej 25°C , to aby uzyskać prawidłowy wynik 0°C do odczytu należy dodać $-1,55^\circ\text{C}$ (tj. odjąć $1,55^\circ\text{C}$).

Poprawki linearyzacji zostały wprowadzone do oprogramowania termometru. Obejmują one zakres od -161°C do $+1311^\circ\text{C}$, gdy zimne złącze ma temperaturę 0°C . Zazwyczaj temperatura zimnego złącza będzie nieco wyższa niż 0°C . Gdy temperatura zimnego złącza wynosi 25°C , zakres wynosi od -136°C do $+1336^\circ\text{C}$.

Linearyzacja opiera się na standardowych tabelach napięcia termoelektrycznego termopary typu K w zależności od temperatury; patrz siliconchip.au/link/abmo

Do wprowadzania poprawek można stosować różne metody. Jedną z nich jest opisanie napięcia termoelektrycznego w funkcji temperatury jako wielomianu matematycznego, a następnie obliczenie wymaganej korekty dla odczytu. Potrzebne może być wykonanie wielu obliczeń.

Opis tej i innych technik można znaleźć w referencyjnym dokumencie projektowym Texas Instruments „TIDA-00468 – Optimized Sensor Linearization for Thermocouple”, który można pobrać ze strony siliconchip.au/link/abmp i wybrać projekt referencyjny TIDA-00468.

Inną metodą jest korzystanie z tabeli zawierającej korekty względem napięcia uzyskiwanego z wyjścia termopary. Jest to podejście, na które się zdecydowaliśmy. Układ MAX31855 udostępnia napięcie wyjściowe termopary z uwzględnieniem kompensacji zimnego złącza, dlatego wartość zimnego złącza musi zostać usunięta z wartości przed zastosowaniem tabeli kompensacji dla termopary.

Cechy

- Szeroki zakres pomiaru temperatury (typowo od -50°C do $+1200^\circ\text{C}$)
- Dokładna rozdzielczość $0,25^\circ\text{C}$ dla wszystkich pomiarów i ustawień
- Dokładność do $\pm 2^\circ\text{C}$ w zakresie od -200°C do $+700^\circ\text{C}$; $\pm 4^\circ\text{C}$ do $+1350^\circ\text{C}$
- Kompaktowe urządzenie zasilane napięciem 12 V DC
- Niski pobór prądu – 75 mA przy pełnej jasności wyświetlacza i włączonym przełączniku
- Linearyzacja odczytów termopary
- Przełącznik termostatu
- Regulowana temperatura przełączania termostatu i histereza
- Działanie termostatu ogrzewania lub chłodzenia
- Regulowana jasność podświetlenia wyświetlacza
- Opcje uśredniania odczytów termometru
- Sygnalizacja błędów połączenia termopary
- Przełącznik przełącza do 30 V przy 10 A
- Zewnętrzny przełącznik może być używany do przełączania sieci lub wyższych prądów (patrz tekst).

Specyfikacja

- Zakres pomiarowy: zależny od termopary; od -270°C do $+1800^\circ\text{C}$
- Zakres pomiaru otoczenia (zimne złącze): -40°C do $+125^\circ\text{C}$
- Dokładność zimnego złącza: $\pm 2^\circ\text{C}$ od -20°C do $+85^\circ\text{C}$; $\pm 3^\circ\text{C}$ od -40°C do $+125^\circ\text{C}$
- Próg termostatu: od poniżej -270°C do powyżej 1800°C
- Histereza termostatu: od 0°C do 60°C
- Regulacja offsetu: -7°C do $+7^\circ\text{C}$ (kompensacja błędów offsetu i zimnego złącza)
- Linearyzacja: korygowana w krokach co $0,5^\circ\text{C}$ z rozdzielczością $0,25^\circ\text{C}$ od -161°C do 1311°C (zimne złącze przy 0°C), od -136°C do 1336°C (zimne złącze przy 25°C).
- Uśrednianie odczytów: 1, 2, 4, 8, 16, 32, 64 lub 128 odczytów
- Wskazanie termostatu: animowany wykres słupkowy w górę lub w dół podczas ogrzewania lub chłodzenia
- Regulacja jasności wyświetlacza: 10 stopni jasności plus wyłączenie
- Automatyczny powrót menu do opcji odczytu termometru
- Wskazanie błędów termopary: przerwanie obwodu, zwarcie do masy lub zwarcie do zasilania.

Po dokonaniu korekty poprzez dodanie lub odjęcie odpowiedniej wartości, wartość charakteryzująca zimne złącze jest dodawana z powrotem, dzięki czemu uzyskiwany jest pełny odczyt temperatury.

Linearyzacja odbywa się w krokach co 0,5°C. Po linearyzacji dokładność temperatury będzie ograniczona głównie przez błędy i offset układu scalonego MAX31855 i samej termopary.

Szczegóły schematu

Schemat termometru pokazano na **rysunku 3**. Konstrukcja została oparta na układzie MAX31855KASA+T z kompensacją zimnego złącza, pełniącym funkcję cyfrowego przetwornika termoparowego dla termopary typu K (IC1), oraz na 8-bitowym mikrokontrolerze PIC16F1459 (IC2). Mikrokontroler steruje również dwuwierszowym 16-znakowym wyświetlaczem LCD prezentującym wyniki pomiarów.

Gniazdo termopary (CON1) zostało zaprojektowane specjalnie dla termopary typu K, aby nie wytwarzać dodatkowego napięcia z powodu odmiennych połączeń metalowych. Napięcie przechodzi przez koraliki ferrytowe (FB1 i FB2) z kondensatorami 100 nF, które odprowadzają zakłócenia do masy.

Koraliki ferrytowe w połączeniu z kondensatorami działają jako wysokoczęstotliwościowe filtry tłumiące zakłócenia napięcia termopary wchodzącego do układu IC1. Diody tłumiące stany przejściowe TVS1 i TVS2 również zwiększają nadmierne napięcia wejściowe podawane do IC1.

Układ IC1 jest zasilany napięciem 3,3 V, a układ IC2 napięciem 5 V. Są one wyprowadzane z wejścia zasilania 12 V na CON2 z zabezpieczeniem przed odwrotną polaryzacją przez diodę D1. W rezultacie przez rezystor 100 Ω na wejście REG1 jest podawane napięcie 11,4 V, a wszelkie przepięcia z wejścia 12 V są ograniczane do 12 V przez diodę Zenera ZD1.

Elementy te zapewniają podstawową ochronę w przypadku przyłożenia znacznie wyższego napięcia do złącza CON2. Rezystor 100 Ω dzieli również rozpraszanie ciepła z układu REG1, aby rozprzaskaczać ciepło bardziej równomiernie wewnątrz obudowy

termometru. Pomaga to lepiej utrzymać stałą temperaturę zimnego złącza.

Układ REG2 zapewnia zasilanie 3,3 V dla IC1. IC1 pobiera maksymalnie 1,5 mA, więc moc wydzielana w REG2 jest bardzo mała, około 2,6 mW. Można ją obliczyć z zależności: $(5\text{ V} - 3,3\text{ V}) \times 1,5\text{ mA}$. Moc rozpraszana przez IC1 wynosi 5 mW ($3,3\text{ V} \times 1,5\text{ mA}$).

Biorąc pod uwagę współczynnik temperatury złącza do otoczenia wynoszący 170°C/W, oznacza to wzrost temperatury o 0,84°C, więc możemy oczekiwać, że pomiar zimnego złącza będzie wyższy od rzeczywistej temperatury otoczenia o tę wartość, plus wszelkie ciepło dostarczane przez rezystor 100 Ω, REG1, REG2 i IC2.

Układ MAX31855 zapewnia cyfrowy odczyt termopary z zastosowaną kompensacją zimnego złącza. Dane są wysyłane za pośrednictwem interfejsu szeregowego z pinem 5 dla zegara, pinem 6 dla wyboru układu i pinem 7 dla wyjścia danych szeregowych.

Dane szeregowo są monitorowane na wejściu RA5 układu IC2 (pin 2), natomiast układ IC2 steruje przebiegiem zegarowym i liniami wyboru układu z wyjść RC4 i RC5 (piny 6 i 8). Dzięki dzielnikom rezystancyjnym 1,1 kΩ/2,2 kΩ napięcie wyjściowe 5 V z układu IC2 jest zredukowane do poziomów 3,3 V odpowiednich dla IC1.

Układ IC2 odczytuje dane temperatury dostarczone przez IC1, bit po bicie. Uzyskiwane są w ten sposób informacje o temperaturze termopary z kompensacją zimnego złącza w postaci 14-bitowej wartości binarnej

ze znakiem, odczytywana jest również temperatura zimnego złącza w postaci 12-bitowej wartości binarnej ze znakiem oraz wszelkie warunki błędów termopary.

Wykrywane błędy to rozwarcie obwodu, zwarcie do masy i zwarcie do napięcia dodatniego.

Oprócz odczytywania danych z IC1, układ IC2 steruje modulem LCD i podświetleniem, monitoruje przełączniki Menu, Up i Down (S1...S3) oraz steruje przekaźnikiem termostatu, RLY1.

Moduł LCD jest sterowany za pomocą 4-bitowego interfejsu równoległego do jego wejść danych D4...D7. Są one podłączone do wyjść cyfrowych RB4...RB7 układu IC2. Wejścia Enable (EN) i Register Select (RS) wyświetlacza LCD są sterowane odpowiednio z wyjść RC2 i RC1 układu IC2.

Dane są wysyłane jako dwa zestawy po cztery bity, tworząc pełne 8-bitowe dane do generowania znaków na wyświetlaczu LCD. Nieużywane wejścia D0...D3 wyświetlacza LCD są dołączone do masy. Wyświetlacz LCD mógłby być sterowany za pomocą 8-bitowego interfejsu równoległego. Wymagałoby to dołączenia wszystkich wejść D0...D7 do IC2, a tyłu nie ma dostępnych.

Podświetlenie LCD

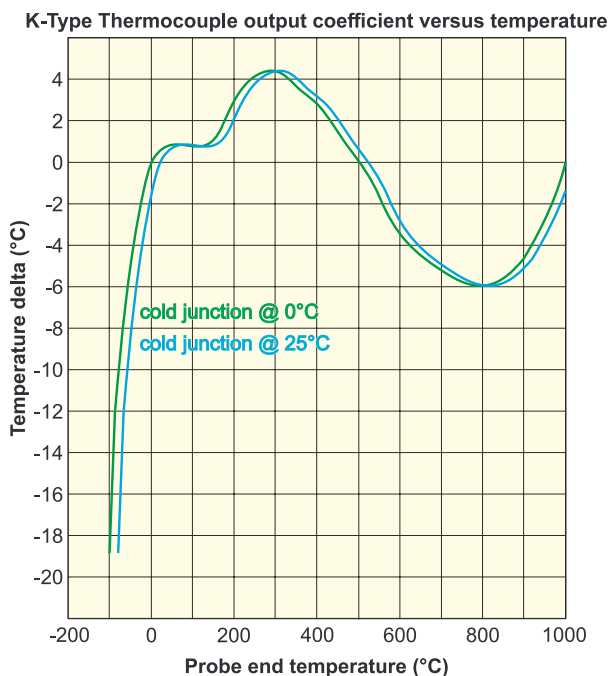
Podświetlenie modułu LCD jest zapewniane przez diody LED umieszczone za ekranem LCD. Anoda diody LED połączona z zaciskiem BLA (pin 15). Katodę podświetlenia (BLK – pin 16) podłączamy do drenu MOSFET-a Q2 poprzez rezystor ograniczający prąd 68 Ω. Diody LED są włączone, gdy Q2 jest aktywowany przez podanie wysokiego napięcia na jego bramkę z wyjścia RC5 układu IC2.

Gdy bramka jest wysterowana wysokim poziomem, napięcie drenu spada.

Wyjście RC5 jest szybko włączane i wyłączane w celu przyciemnienia wyświetlacza. Jasność świecenia wynika z wartości współczynnika wypełnienia przebiegu sterującego bramką. Gdy jest on równy 50%, diody LED są zasilane średnio połową maksymalnego prądu. Wyższe wartości współczynnika wypełnienia zapewniają większą jasność.

Wyjście RC5 (pin 5) dostarcza sygnał z modulacją szerokości impulsu (PWM) o częstotliwości 976 Hz. Jest to wystarczająco szybko, aby włączanie i wyłączanie diod LED nie było zauważalne.

Przełączniki S1 do S3 to przyciski chwilowe. Łączą się one z wejściami RC0, RA1 i RA0 układu IC2, które są podciągane do napięcia 5 V za pomocą rezystorów 10 kΩ. Naciśnięcie przełącznika jest wykrywane jako niski poziom na tym pinie (blisko 0 V), na co oprogramowanie



Rysunek 2. Błąd w odczytach temperatury, jeśli są one dokonywane przy założeniu stałej czułości. W celu uzyskania dokładniejszych wyników można odjąć te błędy od regularnych odczytów

IC2 reaguje, wybierając menu lub zmieniając wartość menu.

Przełącznik RLY1 jest sterowany przez tranzystor Q1, który z kolei jest sterowany z wyjścia cyfrowego RA4 układu IC2 (pin 3). Gdy wyjście to jest w stanie wysokim (5 V), tranzystor jest włączany poprzez prąd bazy przez rezystor 1 kΩ. Następnie kolektor przechodzi w stan niski, co powoduje zasilenie cewki przełącznika i połączenie styku wspólnego (C) i normalnie otwartego (NO).

Gdy RA4 osiągnie stan niski, Q1 wyłącza się. Przełącznik nie jest zasilany, a styki C i NC są połączone. Dioda D2 gasi wysokie napięcie wsteczne EEMF generowane przez cewkę przełącznika podczas wyłączania, chroniąc Q1 przed uszkodzeniem.

Regulacja

Termometr ma kilka ustawień wyświetlania i regulacji, które są przechowywane w pamięci nieulotnej. Wartości te pozostają zachowane po wyłączeniu zasilania. Za pomocą przycisku Menu przełącza się między poszczególnymi pozycjami menu, natomiast przyciski Up i Down dostosowują ustawienia.

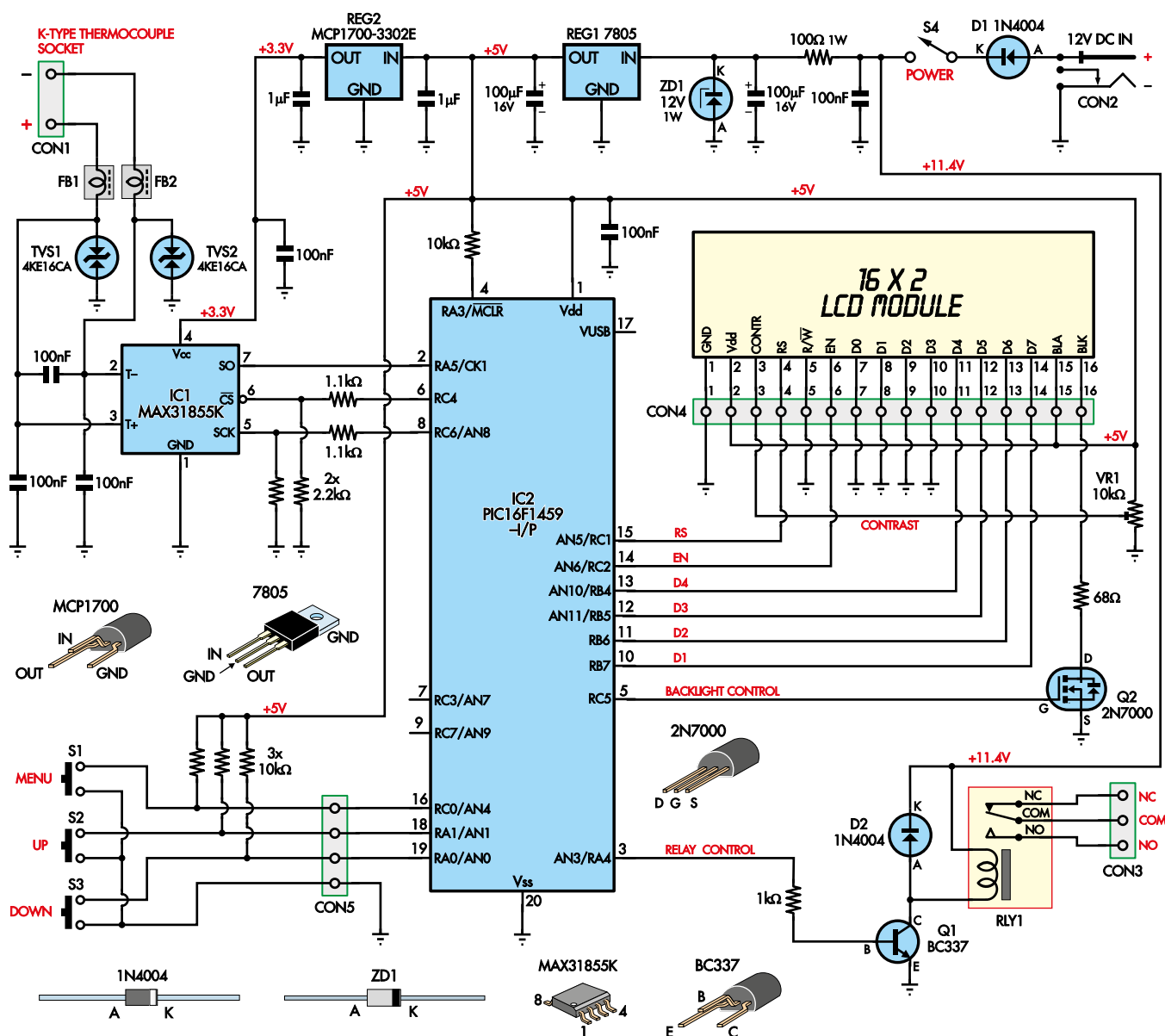
W przypadku ustawień temperatury, które można zmienić, przycisk Up zwiększa wartość, a przycisk Down zmniejsza ją w krokach co 0,25°C po krótkim naciśnięciu. Przytrzymanie przycisku powoduje zmianę wartości w coraz szybszym tempie. Pozwala to na osiągnięcie dużych wartości w rozsądnym czasie, jednocześnie umożliwiając mniejsze kroki o 0,25°C.

Jeśli w danym menu dostępne są dwie opcje, można użyć przycisku Up lub Down, aby wybrać drugą opcję. Szczegóły każdego menu znajdują się w ramce „Skrócony opis menu”.

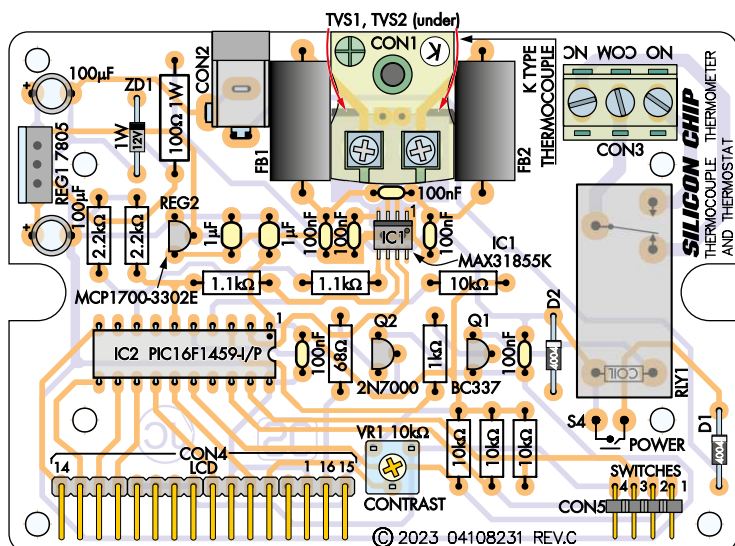
Animacje

Działanie termostatu podczas chłodzenia lub ogrzewania jest wskazywane za pomocą animowanego prostokątnego paska, który przesuwają się w prawym dolnym rogu wyświetlacza w dół w przypadku chłodzenia lub w górę w przypadku ogrzewania. Animacja jest wyświetlana dla menu Termometr, Ustawienie termostatu i Histereza (Thermometer, Thermostat Set and Hysteresis).

Gdy termostat jest wyłączony, animacja paska jest wyłączona.



Rysunek 3. Schemat jest prosty, ponieważ układ MAX31855 (IC1) mierzy temperaturę i przekazuje ją cyfrowo do mikrokontrolera IC2, który następnie aktualizuje ekran LCD za pośrednictwem magistrali czterobitowej. Pozostała część układu obejmuje trzy przyciski sterujące, przełącznik termostatu (RLY1) i liniowy zasilacz prądu stałego



Rysunek 4. Należy zwrócić uwagę, że dwa złącza kątowe (CON4 i CON5) są zamontowane w różny sposób. Jedynymi elementami na spodzie są dwie diody TVS, które nie są spolaryzowane. Ich pozycje są pokazane na warstwie opisowej PCB. Dwa duże koraliki ferrytowe mają wiele zwojów emaliowanego drutu miedzianego przechodzącego przez nie (patrz instrukcje w tekście)

Budowa

Termometr składa się z dwóch dwustronnych, metalizowanych płytek drukowanych, z główną płytką o wymiarach 98 mm × 70 mm o kodzie 04108231 i płytki drukowanej panelu przedniego o wymiarach 19 mm × 22 mm o kodzie 04108232. Są one umieszczone w półprzezroczystej czarnej obudowie Ritec ABS o wymiarach 105 mm × 80 mm × 40 mm.

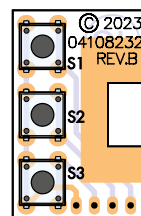
Przełącznik RLY1 zapewnia przełączanie wyjścia gniazda CON3. Może on obsługiwać prąd do 10 A przy napięciu do 30 V. Jeśli konieczne będzie przełączanie napięć sieciowych, wymagany będzie dodatkowy przełącznik zewnętrzny. Szczegóły dotyczące okablowania zewnętrznego przełącznika podamy później.

Budowę głównej płytki drukowanej należy rozpocząć od przyłutowania układu IC1. Jest to 8-pinowy układ scalony w obudowie SOIC, jednej z najprostszych do montażu

powierzchniowego. Zaczynij od prawidłowego zorientowania układu scalonego nad padami PCB (rysunek 4) i przylutuj pin 1.

Sprawdź ułożenie układu scalonego względem pozostałych padów, ponownie rozpuść lut, i jeśli okaże się konieczne, skoryguj ewentualnie wyrównanie pinów do pół lutowniczych. Jeśli natomiast ułożenie będzie prawidłowe, przylutuj pozostałe piny. Za pomocą odrobiny topnika w postaci pasty i plecionki lutowniczej możesz usunąć wszelkie powstałe mostki lutownicze.

Kolejnymi elementami do montażu są rezystory, diody i tłumiki przepięć TVS1 i TVS2. Upewnij się, że diody D1, D2 i ZD1 są zamontowane zgodnie z ułożeniem pokazanym na schemacie montażowym i opisie PCB. Nie pomył ich. TVS1 i TVS2 można zamontować w dowolny sposób. Przylutuj podstawkę pod IC2, upewniając się, że jest prawidłowo zorientowana.



Rysunek 5. Trzy dotykowe przyciski są jedynymi elementami na małej płytce drukowanej panelu przedniego. Łączy się ona z główną płytką drukowaną za pomocą prostokątnego złącza CON5

Koraliki ferrytowe FB1 i FB2 mają pięć zwojów emaliowanego drutu miedzianego o średnicy 0,8 mm każdy. Przed montażem na płytce drukowanej należy zdjąć izolację z końcówek za pomocą ostrego noża lub podobnego narzędzia.

Teraz mogą być zamontowane złącza kątowe CON4 i CON5. Są to złącza 4- i 16-pinowe. Jeśli masz dłuższą listwę, możesz podzielić ją na listwy 4- i 16-pinowe.

Należy pamiętać, że złącza CON4 i CON5 są instalowane w różny sposób. Złącze CON4 (dla LCD) jest instalowane prostą stroną pinów do PCB, natomiast złącze CON5 (dla PCB panelu przedniego) jest instalowane wygiętymi pinami do PCB. Pozwala to na wymagane pozycjonowanie modułu LCD i przełączników na panelu przednim.

Zamontuj dwa piny w punktach podłączenia zasilania S4, gotowe do późniejszego podłączenia do przełącznika.

Teraz zamontuj VR1, kondensatory, tranzystory Q1 i Q2 oraz stabilizatory REG1 i REG2. Kondensatory elektrolityczne muszą być zamontowane z zachowaniem prawidłowej polaryzacji – dłuższe wyprowadzenie to plus, natomiast pasek na obudowie kondensatora wskazuje jego wyprowadzenie ujemne. Upewnij się, że Q1, Q2 i REG2 nie są pomieszczone, ponieważ są to różne typy, które występują w podobnych obudowach TO-92.

Następnie można zamontować gniazdo zasilania, CON2 i gniazdo dla termopary typu K (CON1). Na koniec zamontuj przełącznik, jeśli zamierzasz go wykorzystać. Zapoznaj



Tył obudowy z podłączoną termoparą typu K

Przełącznik jest zamontowany w wycięciu na pionowej płytce drukowanej przycisku (rysunek 5)

się z fragmentem dotyczącym wykorzystania układu do sterowania zasilaniem urządzeń zewnętrznych, jeśli planujesz takie zastosowanie.

Montaż płytki PCB panelu przedniego

Montaż płytki drukowanej przednich przełączników (rysunek 5) jest prosty, obejmuje głównie instalację trzech elementów: S1, S2 i S3. Przełącznik S4 jest instalowany później, po przymocowaniu go do panelu przedniego.

Moduł LCD i panel przycisków można teraz podłączyć i przylutować do złączy kątowych na głównej płytce drukowanej (rysunek 7).

Wycięcia w panelu

Wywierć i wytnij przedni i tylny panel, jak pokazano na rysunku 6. Schemat można pobrać również ze strony siliconchip.com.au/Shop/11/294, a następnie wydrukować go w rozmiarze rzeczywistym i użyć jako szablonu. Prostokątne wycięcia można wykonać, wierząc serię drobnych otworów wzdłuż ich obwodu, a następnie usuwając środkową część i ostrożnie dopasowując kształt do wymiarów docelowych.

Gdy panele są gotowe, przymocuj przełącznik S4 do panelu przedniego za pomocą jednej nakrętki za panelem, a drugiej z przodu. Następnie umieść panel przedni nad wyświetlaczem LCD i z wystającymi przełącznikami S1...S3 i zamocuj w obudowie zespół

Na zdjęciu z tyłu płytki drukowanej widać wiele uzwojeń FB1 i FB2 oraz sposób zamontowania wyświetlacza LCD

składający się z panelu, płytki drukowanej przełącznika i głównej płytki drukowanej.

Przymocuj główną płytkę drukowaną do podstawy obudowy za pomocą śrub dostarczonych z obudową. Przełącznik S4 można teraz przymocować za pomocą żywicy epoksydowej do płytki drukowanej przełącznika. Przed dalszym montażem należy poczekać na utwardzenie kleju.

Jako alternatywę dla klejenia, przełącznik można przytwierdzić za pomocą podwójnego złącza widelkowego 6,3 mm do montażu w obudowie (Jaycar PT4916 lub Altronics H2261) lub pojedynczego złącza przylutowanego do przedniej części płytki drukowanej panelu przedniego.

Do zamontowania przełącznika konieczne będzie wywiercenie otworu, natomiast końcówki złącza widelkowego należy przyciąć do odpowiedniego rozmiaru i wygiąć. Po prawidłowym zamontowaniu prostokątna część obudowy przełącznika będzie znajdować się 2 mm od powierzchni płytki drukowanej panelu przedniego.

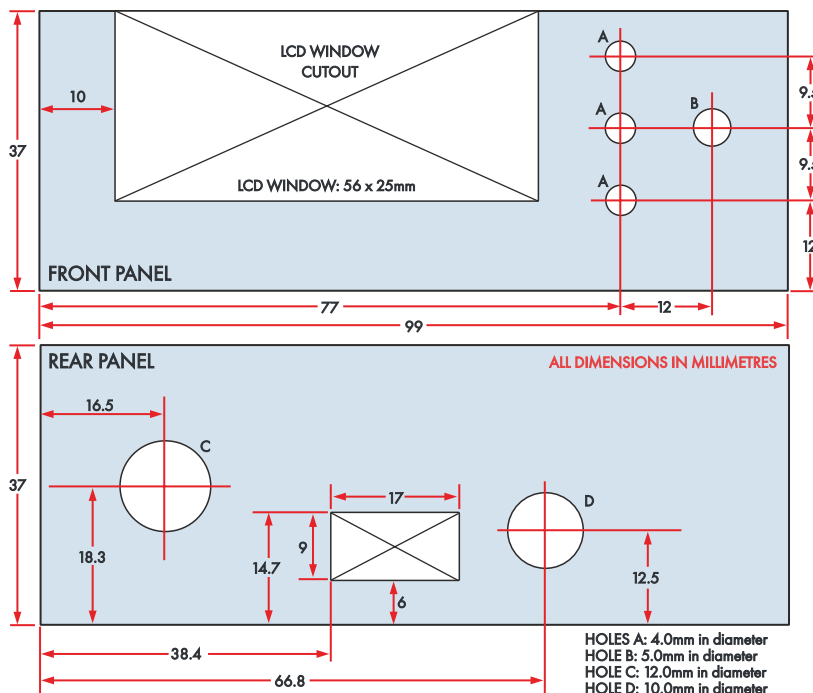
Przewody z dwóch górnych zacisków przełącznika powinny być teraz podłączone do pinów przełącznika na głównej płytce drukowanej.

Tworzenie etykiet paneli

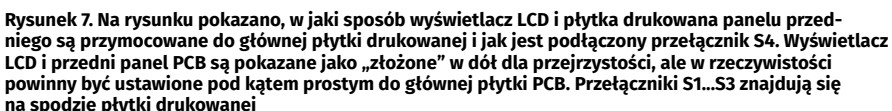
Na rysunku 8 przedstawiono etykiety na panel przedni, które można pobrać, wydrukować i przymocować do panelu przedniego oraz tylnego. Grafikę można wydrukować na etykiecie samoprzylepnej Avery „Heavy Duty White Polyester – Inkjet” w formacie A4, odpowiedniej do drukarek atramentowych lub etykiecie samoprzylepnej „Datapol” do drukarek laserowych. Wyciąć otwory, w tym otwór na wyświetlacz za pomocą ostrego noża modelarskiego.

Etykiety są dostępne na stronach:

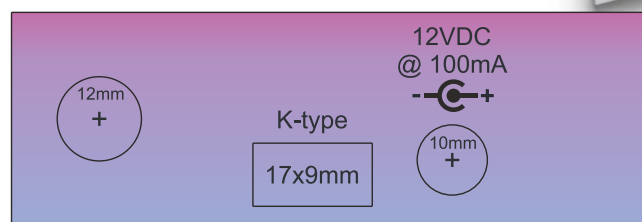
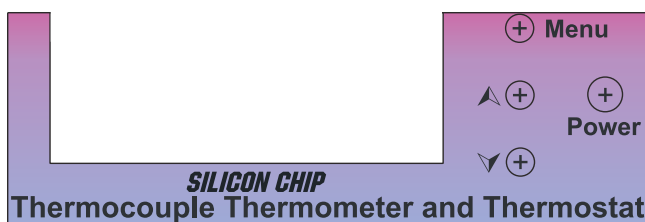
- www.blanklabels.com.au
- www.averyproducts.com.au



Rysunek 6. Wykonaj otwory i wycięcia w przednim i tylnym panelu, jak pokazano na tym rysunku. Możesz go również pobrać w formacie PDF ze strony internetowej Silicon Chip, a następnie wydrukować w rzeczywistym rozmiarze, wyciąć szablony i przykleić je do paneli



Jeśli wymagana jest wtyczka i gniazdo sieciowe, a obudowa jest metalowa, należy podłączyć uziemienie sieciowe do obudowy.



Rysunek 8. Grafika etykiet na panel przedni i tylny. Można je wydrukować na papierze samoprzylepnym lub fotograficznym, zgodnie z opisem w sekcji „Tworzenie etykiet na panel”

W przeciwnym razie należy podłączyć uzziemienie wejścia sieciowego bezpośrednio do uzziemienia gniazda wyjściowego ogólnego przeznaczenia (GPO).

W przypadku obudowy z tworzywa sztucznego nie jest wymagane uzziemienie obudowy, ale na zewnątrz obudowy nie

mogą znajdować się żadne odsłonięte metalowe śruby. Aby zapewnić bezpieczeństwo, należy używać śrub nylonowych.

Można zastosować przekaźniki SPST 12 V DC 30 A 240 V AC sprzedawane przez Jaycar (SY4040) i Altronics (S4211). Można również użyć

przełączników półprzewodnikowych przystosowanych do przełączania napięcia sieciowego AC. Potrzebne będą również dodatkowe elementy, takie jak opaski kablowe, zaciski P, dławiki kablowe, śruby, nakrętki, złącza widelkowe, przewód sieciowy 10 A itp.

Skrócony opis menu

Ustawienia początkowe pokazane w nawiasach na końcu każdego opisu menu poniżej, to ustawienia domyślne przed ich zmianą za pomocą menu. Wszelkie zmiany wartości lub ustawień spowodują ich zastąpienie.

Thermometer (Termometr)

Pokazuje odczyt temperatury sondy po kompensacji zimnego złącza. Chociaż może wyświetlać wartości od -270°C do $+1800^{\circ}\text{C}$, sonda może mieć węższy zakres roboczy. Ten ekran jest wyświetlany po włączeniu zasilania.

Offset Adjust (0,00°C) (Regulacja offsetu)

Powoduje zastosowanie korekty offsetu temperatury do odczytów termometru. Może kompensować wszelkie początkowe przesunięcia w odczycie termopary, błąd odczytu zimnego złącza i efekty samonagrzewania się układu scalonego. Offset można regulować w krokach co $0,25^{\circ}\text{C}$ powyżej i poniżej zera, od -7°C do $+7^{\circ}\text{C}$. Nie ma to wpływu na wartość nastawy termostatu ani odczyt temperatury złącza zimnego.

Thermostat Set (0,00°C) (Ustawienie termostatu)

Jest to ustawienie progu temperatury wyłączenia termostatu. Można go regulować w zakresie od -270°C do $+1800^{\circ}\text{C}$ w krokach co $0,25^{\circ}\text{C}$. Podczas pracy przekaźnik termostatu włączy się lub wyłączy tylko wtedy, gdy trzy odczyty temperatury osiągną lub przekroczą próg. Zapobiega to fałszywym odczytom powodującym przełączenie przekaźnika z powodu szumów. Należy pamiętać, że przełączanie termostatu będzie bardziej opóźnione w przypadku wybrania wyższych wartości uśredniania (patrz dalszy opis).

Hysteresis (4,00°C) (Histereza)

Dodanie histerezy zapobiega gwałtownemu przełączaniu termostatu, gdy temperatura jest bliska wartości progowej. W przypadku ogrzewania, po wyłączeniu termostatu temperatura musi spaść o wartość histerezy, zanim termostaat włączy się ponownie, aby wznowić ogrzewanie. W przypadku chłodzenia, po wyłączeniu termostatu temperatura musi wzrosnąć o wartość histerezy, zanim termostaat włączy się ponownie, aby rozpocząć chłodzenie. Histerezę można ustawić w zakresie od 0°C do 60°C w krokach co $0,25^{\circ}\text{C}$.

Brightness (50%) (Jasność)

Jasność podświetlenia wyświetlacza można ustawić na jeden z dziesięciu stopni, od niskiej do pełnej. Ustawienie jest wyświetlane na wykresie słupkowym, a jasność zmienia się wraz z modyfikacją ustawienia.

Averaging (1) (Uśrednianie)

Wyższe wartości uśredniania spowalniają aktualizację odczytu termometru, ale umożliwiają stabilniejszy odczyt temperatury, gdy

sonda temperatury jest narażona na szumy i zakłócenia sieciowe. Dostępne są opcje uśredniania dla 1, 2, 4, 8, 16, 32, 64 lub 128 pomiarów.

Gdy uśrednianie jest ustawione na osiem pomiarów lub więcej, odwrotny ukośnik przed słowem „Thermometer” w menu głównym zmienia się z jednej pozycji na drugą (górną lub dolną), aby wskazać, kiedy wartość temperatury jest aktualizowana. Jeśli ustawiona jest wartość 128 próbek, aktualizacja nowej uśrednionej wartości może potrwać do 10 sekund. Czasy te skracają się dla niższych wartości uśredniania (około pięciu sekund dla 64, 2,5 s dla 32 itd.).

Thermostat (chłodzenie) (Termostat)

Termostat można skonfigurować do ogrzewania lub chłodzenia. W przypadku ogrzewania termostat jest włączany, gdy temperatura jest niższa od ustawionej i wyłączany, gdy osiągnie ustawioną wartość. Alternatywnie, w przypadku chłodzenia termostat jest włączany, gdy temperatura jest wyższa od ustawionej, a wyłączany, gdy spadnie poniżej progu.

Auto Return (wyłączone) (Automatyczny powrót)

Włączenie tej funkcji powoduje powrót do głównego ekranu termometru, jeśli przez cztery sekundy nie zostanie naciśnięty żaden przycisk. Oszczędza to konieczności przechodzenia przez wszystkie poziomy menu, aby dotrzeć do głównego menu termometru.

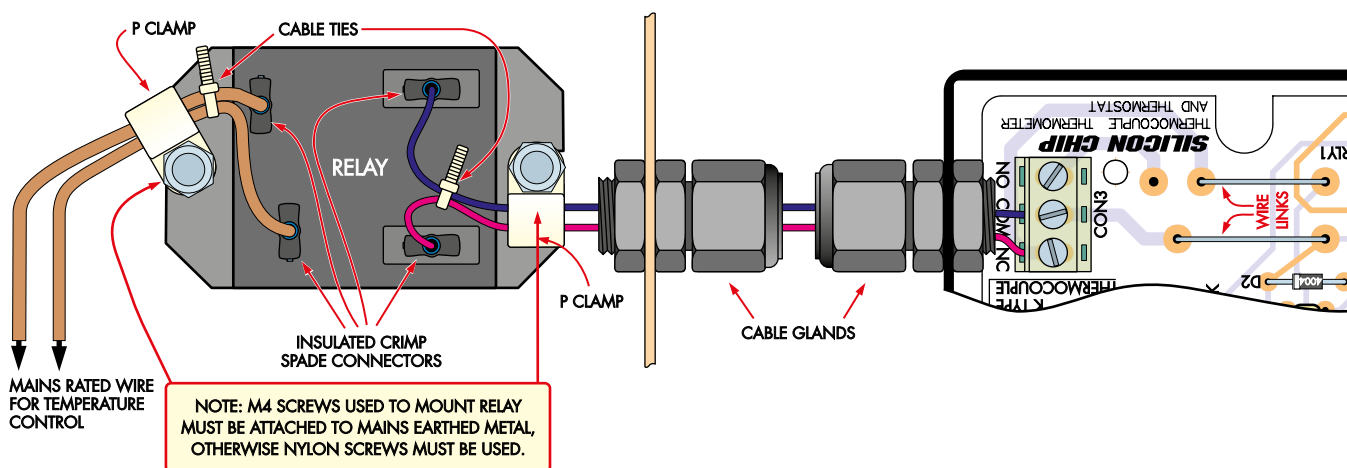
Linearisation (włączona) (Linearyzacja)

Określa, czy odczyty termopary są linearyzowane (korygowane) ze względu na zmianę współczynnika Seebecka w zależności od temperatury. Można ją włączyć lub wyłączyć.

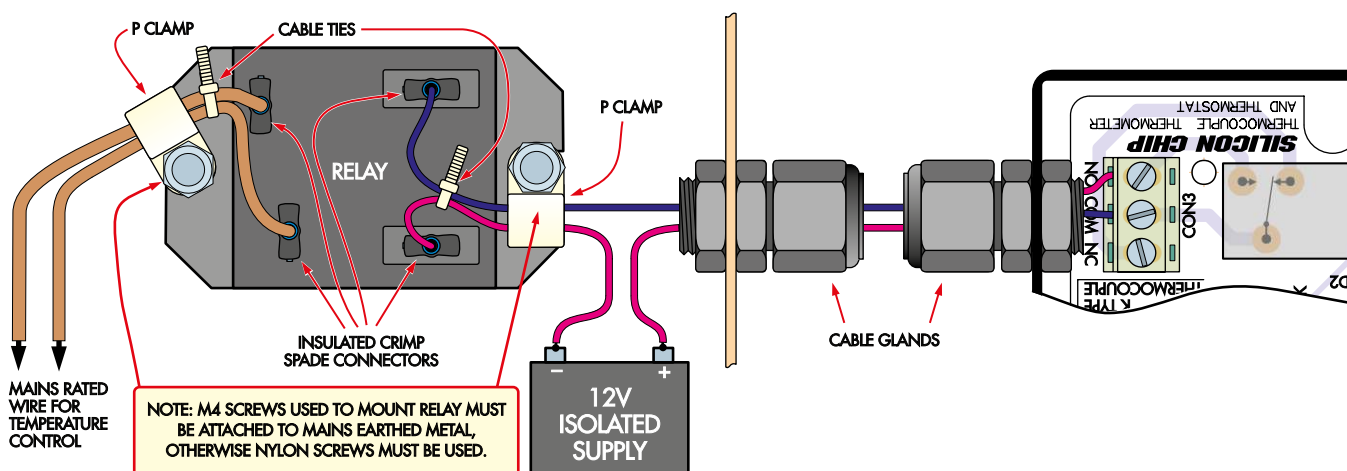
Po włączeniu, jeśli odczyt wykracza poza zakres temperatury, w którym przeprowadzana jest linearyzacja, na wyświetlaczu pojawi się komunikat „Linearisation Range Error” (Błąd zakresu linearyzacji). Ponadto, po włączeniu, nieliniowy odczyt można wyświetlić na głównym wyświetlaczu temperatury, naciskając przycisk Down.

Cold Junction (Zimne złącze)

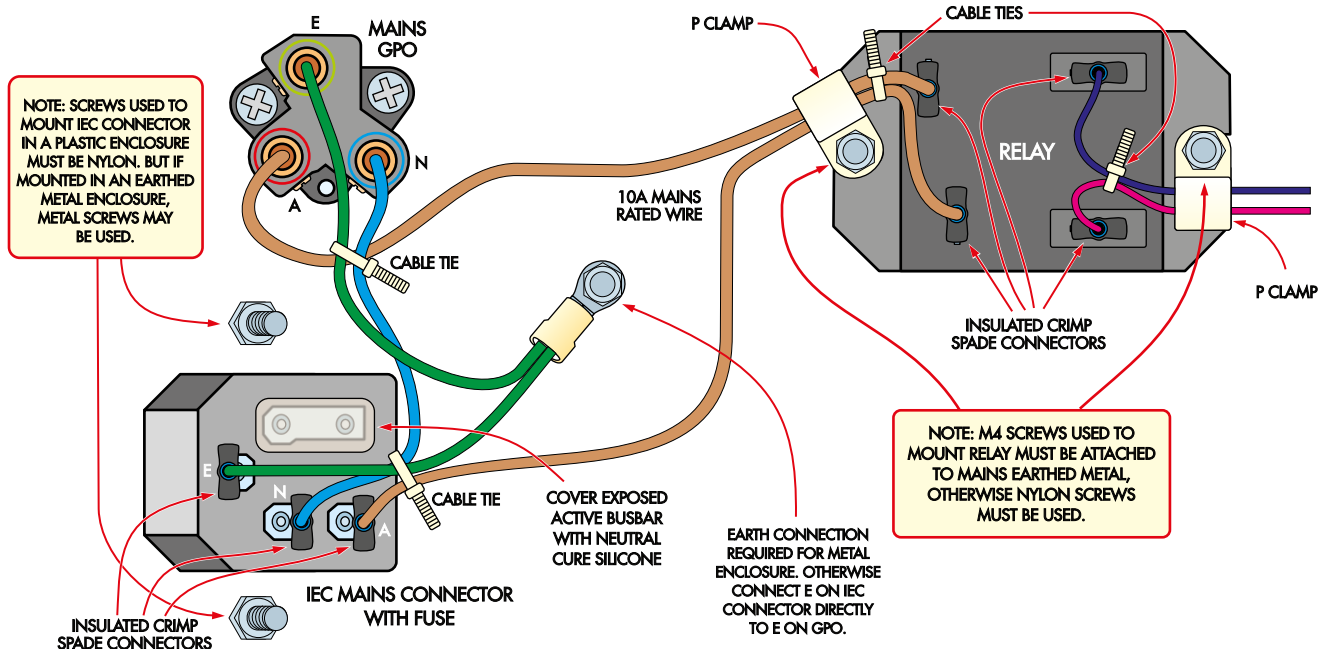
Pokazuje temperaturę zimnego złącza zmierzoną przez układ scalony MAX31855. Może ona wynosić od -40°C do $+125^{\circ}\text{C}$ w krokach co $0,25^{\circ}\text{C}$. Zazwyczaj pokazuje temperaturę otoczenia, ale uwzględnia błędy odczytu wynikające z samonagrzewania i dokładności pomiaru.



Rysunek 9. Należy zwrócić uwagę na połączenia przewodowe na płycie głównej w miejscu RLY1, tak aby napięcie 12 V było doprowadzane do zacisku wyjściowego przekaźnika w celu sterowania cewką przekaźnika zewnętrznego



Rysunek 10. Dodatkowe okablowanie do sterowania urządzeniem sieciowym za pomocą termostatu. Musi on znajdować się w odpowiedniej obudowie z prawidłowo zaizolowanym okablowaniem. Zakłada się, że masz zewnętrzne źródło 12 V DC; w przeciwnym razie użyj rysunku 9



Rysunek 11. W przypadku korzystania z zewnętrznego przekaźnika sieciowego można podłączyć go do gniazda wejściowego IEC i wyjścia GPO zamontowanego w obudowie zawierającej przekaźnik sieciowy w sposób przedstawiony na rysunku. Należy użyć prawidłowych kolorów przewodów i nie zdejmować opasek kablowych

Konfiguracja

W ramce „Skrócony opis menu” została zawarta lista dostępnych poziomów menu i ich funkcji. Należy je ustawić zgodnie z zastosowaniem. Aby wskazanie temperatury utrzymywało się stabilnie, wartość uśredniania musi być zazwyczaj większa niż jeden, zwłaszcza gdy podczas dotyknięcia sondy termopary pojawi się przydźwięk i szum.

Ustawienia termostatu wymagają wybrania trybu ogrzewania lub chłodzenia oraz dostosowania temperatury progowej i histerezy. Histereza ma na celu zapobieganie

gwałtownemu przełączaniu termostatu w pobliżu progu przełączania. Aby zapobiec wystąpieniu takiej sytuacji, należy ustawić ją wystarczająco wysoko.

Kalibracja

Ze względu na offset układu MAX31855, uzyskania prawidłowego odczytu temperatury wymaga kalibracji. Jest to spowodowane tym, że temperatura wewnątrz obudowy jest wyższa niż temperatura otoczenia. Menu „Offset” umożliwia regulację w celu skorygowania tych początkowych błędów. Najlepiej

zrobić to poprzez kalibrację termometru przy użyciu roztworu referencyjnego 0°C.

Można to zrobić przy użyciu słoika z czystą, świeżą wodą, do której wysypano wystarczającą ilość pokruszonego lodu, aby temperatura osiągnęła 0°C. Powinno być możliwe dostosowanie odczytu termometru za pomocą regulacji offsetu, tak aby wyświetlacz pokazywał 0°C. Konieczne będzie sprawdzenie, czy linearyzacja jest włączona (zobacz, jak to sprawdzić w ramce „Skrócony opis menu”).

Przed wykonaniem kalibracji termometr powinien być włączony na jakiś czas (np. pół godziny lub dłużej), co pozwoli mieć pewność, że ustabilizował się termicznie.

W celu sprawdzenia kalibracji w wyższej temperaturze można przyjąć wartość odniesienia 100°C, uzyskiwaną podczas ciągłego wrzenia wody w warunkach normalnego ciśnienia atmosferycznego (1013 hPa). Należy jednak pamiętać, że temperatura wrzenia wody maleje wraz ze spadkiem ciśnienia atmosferycznego – średnio o około 0,325°C na każde 100 m wzrostu wysokości nad poziomem morza. W praktyce oznacza to, że na wysokości 1000 m woda wrze w temperaturze około 96,7°C, a na 2000 m – około 93,5°C. ■

John Clarke

Wykaz elementów:

- 1 dwustronna płytka drukowana o kodzie 04108231, 98 mm × 70 mm
- 1 dwustronna płytka drukowana o kodzie 04108232, 19 mm × 22 mm
- 1 obudowa Ritec 105 mm × 80 mm × 40 mm czarna, półprzezroczysta ABS [Altronics H0192].
- 1 alfanumeryczny wyświetlacz LCD 2×16 znaków [Altronics Z7013].
- 1 sonda termopary typu K [Jaycar QM1283 (-50°C do +250°C), QM1282 (-50°C do +900°C), element14 2947102 (0°C do +800°C)].
- 1 dławik kablowy dla kabla o średnicy 3 mm...6 mm
- 3 mikroprzyciski SPST do montażu na płytce drukowanej 6 mm (S1...S3) [Jaycar SP0603, Altronics S1124]
- 1 subminiaturowy przełącznik SPDT (S4) [Jaycar ST0300]
- 1 zestaw wtyczek 12 V DC 100 mA+ z wtyczką cylindryczną 2,1 mm lub 2,5 mm ID
- 1 przełącznik 12 V SPDT 10 A (RLY1) [Jaycar SY4050, Altronics S4197]
- 1 gniazdo termopary typu K (CON1) [element14 3810628]
- 1 gniazdo zasilania do montażu na płytce drukowanej, 2,1 mm lub 2,5 mm ID (w zależności od zasilacza CON2) [Jaycar PS0520, Altronics P0621A]
- 13-drożne złącze śrubowe, raster 5,08 mm (CON3)
- 1 16-drożne złącze kątowe, raster 2,54 mm (CON4)
- 1 4-drożne złącze kątowe, raster 2,54 mm (CON5)
- 1 20-pinowa podstawka DIL IC (dla IC2)
- 2 duże ferrytowe koraliki tłumiące (FB1, FB2) [Jaycar LF1256 (opakowanie 6 sztuk), Altronics L4710A]
- 1 gniazdo zasilania przewód miedziany o długości 250 mm i średnicy 0,8 mm (dla FB1 i FB2)
- 2 odcinki lekkiego przewodu połączeniowego o długości 50 mm (dla S4)
- 2 piny/kolki lutownicze
- 1 jednoobrotowy potencjometr montażowy 10 kΩ (VR1) [Jaycar RT4600, Altronics R2597]
- 1 niewielka ilość żywicy epoksydowej lub złącze widelkowe 6,3 mm do montażu w obudowie (do montażu S4) [Jaycar PT4916, Altronics H2261]

Półprzewodniki

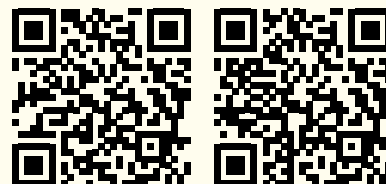
- 1 MAX31855KASA-T układ scalony termopary z kompensacją zimnego złącza do przetwornika cyfrowo-analogowego dla termopary typu K (IC1) [element14 2515622]
- 18-bitowy mikrokontroler PIC16F1459-I/P zaprogramowany kodem 0410823A.hex, DIP-20 (IC2)
- 1 7805 1 A stabilizator 5 V, TO-220 (REG1)
- 1 MCP1700-3302-E/TO lub AMS1117-3.3 stabilizator liniowy 3,3 V o niskim spadku napięcia, TO-92 (REG2) [Silicon Chip SC2782, element14 1296588]
- 1 BC337 tranzystor NPN 45 V 500 mA, TO-92 (Q1)
- 1 2N7000 MOSFET N-kanalowy 60 V 200 mA, TO-92 (Q2)
- 2 (P)4KE15CA lub (P)4KE16CA diody tłumienia stanów nieustalonych 400 W 12,8 V...13,6 V (TVS1, TVS2) [Jaycar ZR1162]
- 1 dioda Zenera 12 V 1 W (ZD1) [1N4742]
- 2 diody 1N4004 400 V 1 A (D1, D2)

Kondensatory

- 2 elektrolityczne okrągłe 100 µF 16 V PC
- 2 1 µF 50 V X5R lub X7R ceramiczne okrągłe
- 6 100 nF 50 V X5R lub X7R ceramiczne okrągłe

Rezystory (wszystkie 1/4 W, 1%, o ile nie zaznaczono inaczej)

- 4 10 kΩ 1 1 kΩ 2 2,2 kΩ 1 100 Ω 1 W 2 1,1 kΩ 1 68 W 1/2 W lub 0,6 W



Materiały dodatkowe dostępne są na stronie:
<https://www.siliconchip.com.au/Shop/8/6806>
<https://www.siliconchip.com.au/Shop/8/6807>

Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022.
www.siliconchip.com.au

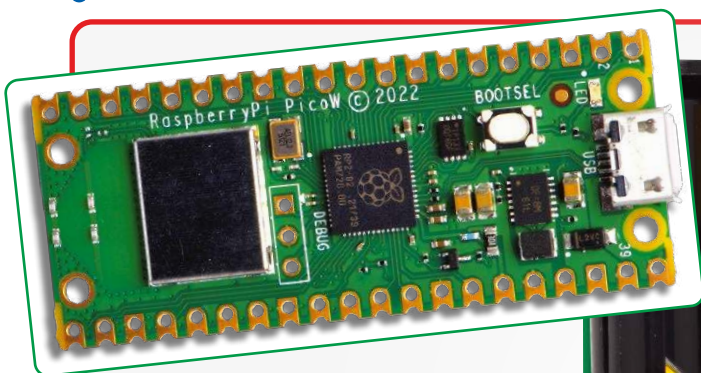
Mikołaj, Bałwanek, Elf i... wiele innych ozdób świątecznych kupisz i zmontujesz



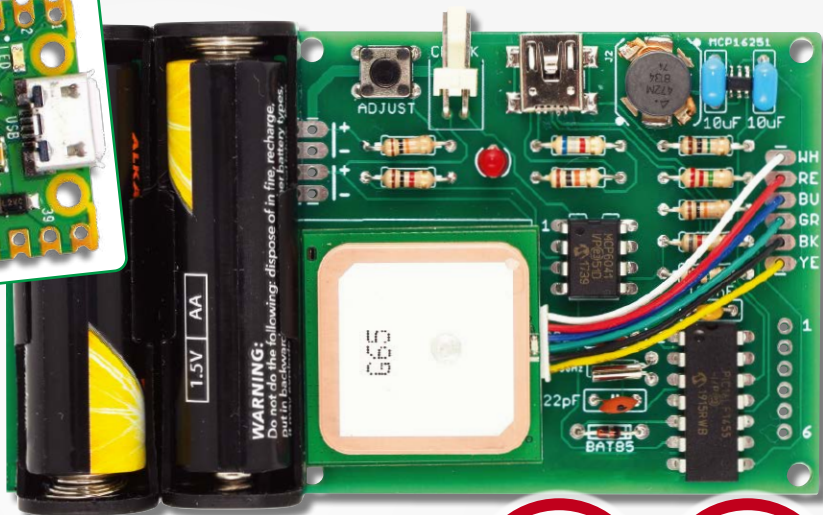
<https://sklep.avt.pl/pl/menu/ozdoby-485.html>



REKLAMA



Raspberry Pi Pico W może być używany jako zamiennik modułów GPS w istniejących projektach utrzymywania czasu, gdy nie można uzyskać niezawodnego sygnału GPS. Pobiera on czas z internetowego serwera NTP przez Wi-Fi a następnie udostępnia czas z dokładnością do ułamka sekundy.



Źródło Czasu Wi-Fi dla zegarów GPS

Odkąd moduły GPS stały się dostępne dla hobbystów, używamy ich jako dokładnych źródeł czasu. Chociaż GPS i inne podobne systemy satelitarne zrewolucjonizowały nawigację i mapowanie, zapewniają jednocześnie łatwy globalny dostęp do bardzo dokładnych źródeł czasu.

Każdy satelita GPS jest wyposażony w dwa zegary atomowe, które co sekundę przesyłają bardzo dokładny sygnał czasu. Do tej pory używaliśmy tego sygnału w wielu projektach, w tym w niedawnym, bardzo popularnym sterowniku zegara analogowego GPS z września 2022 r. (siliconchip.au/Series/391) – EdW 4/2025.

GPS był pierwszym systemem GNSS (Globalnego Systemu Nawigacji Satelitarnej), ale obecnie istnieje kilka innych systemów, w tym rosyjski GLONASS, europejski Galileo i chiński Beidou.

Indyjski Regionalny System Nawigacji Satelitarnej (IRNSS) i japoński System

Satelitarny Quasi-Zenith (QZSS) zostały zaprojektowane w celu poprawy pozycjonowania na skalę krajową, przy czym QZSS jest również dostępny w Australii, ponieważ orbity satelitów pokrywają również ten kraj.

Chociaż w wymienionych systemach wykorzystywane są subtelnie różniące się technologie (nawet GPS ewoluował w ciągu 50 lat istnienia), ustanowiono wspólny interfejs zewnętrzny. W rzeczywistości moduł GPS VK2828U7G5LF, którego używamy w wielu projektach, może odbierać sygnały z satelitów GPS, GLONASS i Galileo.

Na potrzeby tego artykułu będziemy używać GPS jako terminu obejmującego wszystkie różne systemy nawigacji satelitarnej. Należy jednak pamiętać, że niektóre z tych systemów nie są prawdziwie globalne, ponieważ satelity zwykle nie zapewniają zasięgu na dużych szerokościach geograficznych (w pobliżu biegunów).

Poprzednie wersje Źródeł Czasu GPS

W numerze z kwietnia 2018 roku opublikowaliśmy Źródło Czasu GPS Claytona (siliconchip.au/Article/11039). Jak sama nazwa wskazuje, urządzenie to nie używa żadnej technologii GPS, ale może być stosowane jako źródło sygnałów czasu podobnych do GPS, gdy rzeczywisty sygnał tego systemu jest niedostępny. Często zaleca się go jako zamiennik modułu GPS w projektach zegarów.

Motywacją dla tej koncepcji był fakt, że wiele zegarów jest używanych w pomieszczeniach, gdzie bardzo słabe sygnały GPS są trudne do odebrania. Z drugiej strony, w pomieszczeniach zwykle dostępna jest bezprzewodowa sieć Wi-Fi.

Rzeczywista warstwa sprzętowa projektu źródła czasu z roku 2018 to po prostu moduł mikrokontrolera D1 Mini Wi-Fi ESP8266. Moduł ma wgrane oprogramowanie układowe, pozwalające łączyć się z siecią Wi-Fi i aktualizować wewnętrzny zegar z Internetu za pomocą NTP (sieciowy protokół czasu).

Tak uzyskane dane są następnie wykorzystywane do generowania „zdań” przekazania czasu. Generowany jest również sygnał 1PPS, chociaż nie będzie on miał precyzji rzeczywistego modułu GPS.

Z jakimi projektami współpracuje Źródło Czasu Wi-Fi?

- Nowy Zegar Analogowy z Synchronizacją GPS, wrzesień 2022; siliconchip.au/Article/15466 (EdW 4/2025)
- Zegar Analogowy z synchronizacją GPS, luty 2017; siliconchip.au/Article/10527 (EdW 9/2022)
- 6-cyfrowy zegar GPS LED o wysokiej widoczności, grudzień 2015 – styczeń 2016; siliconchip.au/Series/294
- 6-cyfrowy Retro Nixie Clock Mk2, luty – marzec 2005; siliconchip.au/Series/282
- 6-cyfrowy zegar z blokadą GPS, maj – czerwiec 2009; siliconchip.com.au/Series/37

Funkcje Źródła Czasu Wi-Fi

- Dostarcza dane NMEA 0183 symulujące Źródło Czasu GPS
- Regulowana szybkość transmisji
- Poziomy logiczne 3,3 V współpracują z systemami 3,3 V i 5 V
- Syntetyzowany sygnał 1PPS
- Pobiera czas z serwerów NTP przez Wi-Fi
- Generuje przybliżoną szerokość i długość geograficzną na podstawie adresu IP
- Może również wyświetlać fikcyjne współrzędne geograficzne
- Może skanować do ośmiu sieci Wi-Fi (SSID)
- Możliwość konfiguracji poprzez wirtualny port szeregowy USB, niezależnie od strumienia danych
- Wykorzystuje kompaktowy i niedrogi moduł Raspberry Pi Pico W.
- Zintegrowana przetwornica buck/boost działa wydajnie w zakresie od 1,8 do 5,5 V.
- Generator kwarcowy zapewnia dokładność lepszą niż 30 ppm pomiędzy aktualizacjami
- Pobiera prąd o natężeniu 50 mA lub do 100 mA podczas transmisji Wi-Fi (zasilanie 3,0 V)

Aktualizacja Pico W

Opisany w artykule projekt jest aktualizacją oryginalnego GPS Claytona, ale używa modułu Raspberry Pi Pico W zamiast D1 Mini. Chociaż mogliśmy przerobić kod, aby działał z Pico W, istnieje kilka powodów, dla których tego nie zrobiliśmy.

W ciągu ostatnich pięciu lat otrzymaliśmy wiele sugestii dotyczących ulepszeń, więc sensowne było uwzględnienie ich tam, gdzie było to możliwe.

Zdecydowaliśmy się na użycie C SDK, ponieważ stwierdziliśmy, że daje nam to lepszy dostęp do funkcji niskiego poziomu, dzięki czemu programy działają szybciej. Niektóre z nowych funkcji były możliwe (i znacznie łatwiejsze do wdrożenia) dzięki rozwiązaniom pakietu C SDK i jego bibliotek programowych.

Nie ma wątpliwości, że Pico W ma bardzo dobrą cenę, co powoduje, że jest to atrakcyjna opcja, gdy moduł stanowi całość lub większość wymaganego sprzętu. Rzeczywiście, jest tańszy niż moduł GPS, który może zastąpić. Ale szczególne cechy jego mikrokontrolera RP2040 pomogły nam stworzyć moduł źródła czasu Wi-Fi.

Na przykład można zaimplementować wirtualny port szeregowy USB, co oznacza, że menu konfiguracji jest oddzielone od strumienia danych NMEA (National Marine Electronics Association). Ze względu

na charakter portu szeregowego w D1 Mini, były one współdzielone przez GPS Claytona, więc korzystanie z menu konfiguracji przewidywało strumień danych.

Pico W implementuje również wirtualny napęd USB do programowania pamięci Flash. Niektórzy użytkownicy mieli trudności z załadowaniem pamięci Flash do D1 Mini z różnych powodów. Na przykład wymagana jest do tego specjalna aplikacja do programowania lub Arduino IDE.

Z drugiej strony, Pico W może być flashowany przez prawie każdy komputer z portem USB. Proces jest bardzo prosty – wystarczy skopiować plik do wirtualnego napędu USB.

Procesor RP2040 w Pico W ma dwa rdzenie, więc jeden może być przeznaczony do wysyłania danych NMEA i nie może być blokowany przez aktywność drugiego rdzenia, który obsługuje konfigurację i połączenia Wi-Fi.

Pico W ma również wbudowaną przetwornicę impulsową, która jest bardziej wydajna niż stabilizator liniowy znajdujący się w D1 Mini. Niektórzy Czytelnicy zgłaszali problemy z zasilaniem D1 Mini, więc jest to oczekiwana aktualizacja. Nie tylko zmniejsza zapotrzebowanie na prąd przy wyższych napięciach zasilania, ale także umożliwia pracę przy zasilaniu wynoszącym zaledwie 1,8 V.

Podobnie jak wcześniejsze źródła czasu, moduł źródła czasu Wi-Fi emituje trzy komunikaty NMEA: „RMC” (zalecane minimalne dane dla GPS), „GGA” (informacje o pozycji) i „GSA” (dane satelitarne).

Większość naszych projektów zegarów GPS używa tylko komunikatów RMC, a niektóre również GGA. Dane te są więc całkowicie wystarczające do sterowania tymi zegarami.

Komunikaty NMEA

Praktycznie wszystkie moduły GPS przesyłają dane zgodnie ze standardem NMEA 0183. Standard ten w rzeczywistości określa dane szeregowo 4800 bodów przy użyciu sygnału symetrycznego zgodnego ze standardem elektrycznym RS-422.

Nowszy standard NMEA 2000 wykorzystuje sieć magistrali CAN z prędkością 250 kb/s. Pełna treść tego standardu nie jest publicznie dostępna, więc prostszy NMEA 0183 jest nadal szeroko stosowany, ponieważ jest dobrze znany.

Większość odbiorników używa obecnie sygnałów o pojedynczym poziomie logicznym (zazwyczaj 3,3 V) z szybkością transmisji 9600 lub nawet wyższą. Wiele modułów zapewnia również sygnał 1PPS (impuls na sekundę), który jest zsynchronizowany z zegarami atomowymi satelitów.

Dane szeregowo składają się z linii znaków ASCII zwanych zdaniami (sentence). Dla naszych celów każde zdanie jest oznaczone na początku znakiem „\$”, po którym następują dwa znaki identyfikujące „rozmówcę”. Zazwyczaj jest to „GP” dla systemów GPS, chociaż widzieliśmy kilka modułów, które używają „GN”, gdy dane z wielu systemów satelitarnych są łączone.

Kolejne trzy znaki identyfikują typ wiadomości, a następnie dane specyficzne dla zdania i kod sumy kontrolnej zapewniający pewien stopień ochrony przed uszkodzonymi danymi.

Najpopularniejsze zdania kodujące czas zawierają również dane o lokalizacji, więc moduł źródła czasu Wi-Fi może generować fikcyjne dane o lokalizacji lub nawet korzystać z usługi danych geolokalizacyjnych adresu IP w celu wygenerowania przybliżonej lokalizacji. W każdym razie, dobrym pomysłem jest wygenerowanie takich danych na wypadek, gdyby urządzenie odbierające oczekiwało, że w tej lokalizacji znajdują się prawidłowe dane, nawet jeśli nie są one używane.

To przybliżenie nigdy nie będzie wystarczająco dobre do celów nawigacyjnych. Zwykle jednak wystarcza do określenia strefy czasowej, co jest idealne dla zegarów wykorzystujących w tym celu dane o lokalizacji GPS.

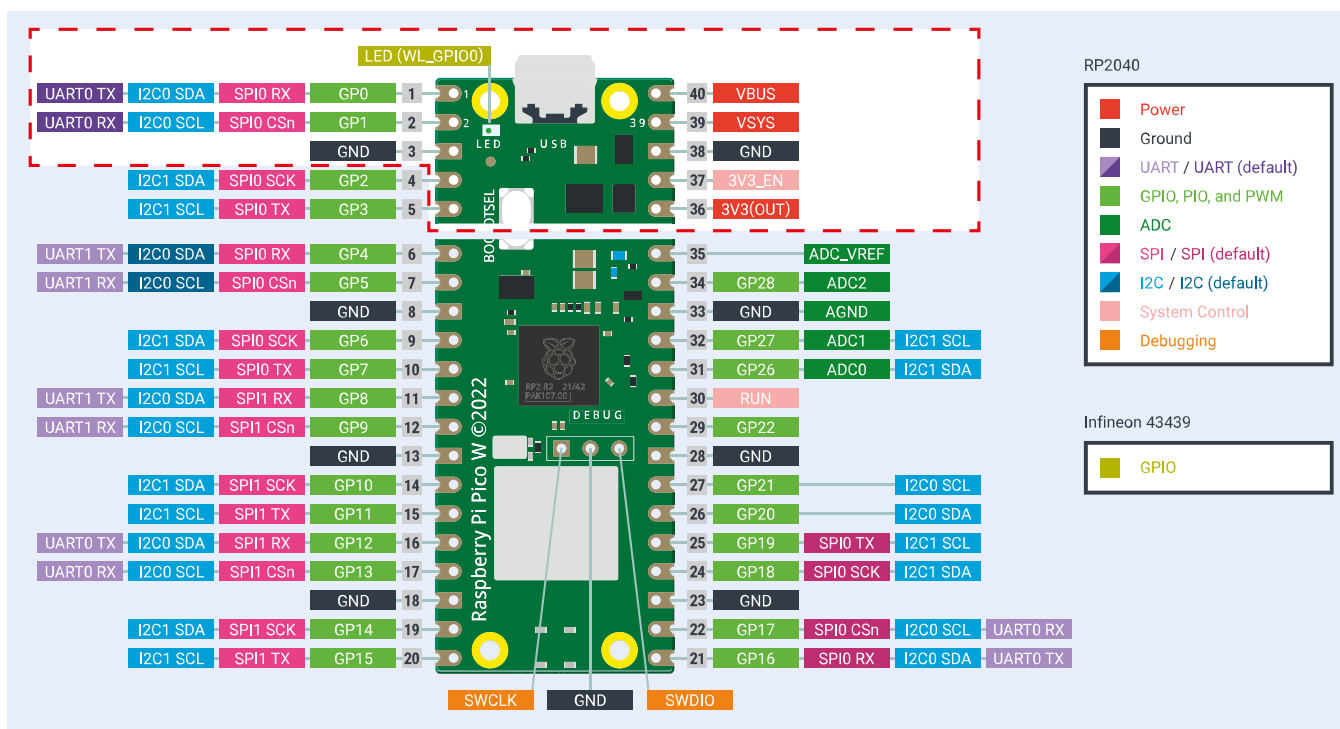
Przykładowo, 6-cyfrowy zegar GPS LED High Visibility z grudnia 2015 i stycznia 2016 (siliconchip.au/Series/294) używa danych o lokalizacji do automatycznego ustawiania strefy czasowej i zasad oszczędzania światła dziennego.

Ponieważ większość naszych projektów GPS używa sygnału GPS do pomiaru czasu zegara, moduł źródła czasu Wi-Fi doskonale nadaje się do użytku z zegarami wewnętrznymi, które mogą nie mieć widoku na niebo, a tym samym na sygnał GPS, można za to łatwo podłączyć je do sieci Wi-Fi.

Sprzęt

Warstwa sprzętowa modułu źródła czasu Wi-Fi jest minimalna. Na rysunku 1 w przewidywanej ramce pokazano rozkład wyprowadzeń Pico W po jego zaprogramowaniu. Pozostała część tego rysunku obejmuje pełną mapę wszystkich wyprowadzeń wraz z ich funkcjami.

Jak widać, zachowaliśmy wszystkie przydatne piny na jednym końcu. Dobrze by było skrócić płytkę poprzez odcięcie niepotrzebnej sekcji. Niestety, potrzebna jest cała płytka i nie można jej znacznie zmniejszyć, zwłaszcza że antena Wi-Fi znajduje się na końcu naprzeciwko złącza USB.



Rysunek 1. Wyprowadzenia modułu Pico W, które można wykorzystać w Źródle Czasu Wi-Fi, zaznaczono przerywaną czerwoną ramką. Pin 1 (GP0) pełni funkcję wyjścia UART TX – znajduje się najbliższej krawędzi płytki z gniazdem USB i leży w pobliżu pinów zasilania. Prawdopodobnie dla większości projektów zegarowych nie będziesz potrzebował wszystkich pokazanych tu połączeń (rysunki 3..6). Trzy lub cztery linie sygnałowe są często wystarczające. Pin 1 – szeregowo dane NMEA, pin 2 – sygnał 1PPS, pin 3 – masa, pin 36 – 3,3 V, pin 37 – włączenie 3,3 V (aktywny stan wysoki), pin 38 – masa, pin 39 – wejście 1,8 V do 5,5 V, pin 40 – zasilanie USB

Wyprowadzenia zasilania są zlokalizowane po prawej stronie, w pobliżu złącza USB. Są to: pin 40 (VBUS), pin 39 (VSYS) i pin 38 (GND). W rzeczywistości istnieje kilka pinów GND (rysunek 1), ale piny 3 i 38 znajdują się najbliżej innych ważnych wyprowadzeń.

Nóżka 37 (3V3_EN) jest wejściem sterującym stabilizatora w Pico W. Jest ono utrzymywane w stanie wysokim przez rezystor 100 kΩ, ale przez ściągnięcie tego wejścia do stanu niskiego można wyłączyć stabilizator, a tym samym wyłączyć Pico W.

Nóżka 1 (GP0) jest źródłem generowanych danych szeregowych NMEA, które na bieżąco mają poziom logiczny 3,3 V.

Sprzętowe peryferia UART Pico W są dostępne tylko na określonych wyprowadzeniach. Pin 1 został wybrany, ponieważ jest to wyjście UART TX znajdujące się najbliżej złącza USB i pinów zasilania. Wybraliśmy sąsiedni pin 2 (GP1) dla wyjścia 1PPS. Mógł to być dowolny pin z pozostałych wyprowadzeń GPIO.

Pokazaliśmy wyjście 3,3 V tylko dlatego, że może być przydatne, jeśli potrzebujesz stabilizowanego zasilania 3,3 V do swojego projektu. Stabilizator w Pico W może dostarczyć prąd do 2 A, chociaż część tego prądu jest pobierana przez Pico W.

Na **rysunku 2** został przedstawiony układ zasilania Pico W. Pomoże Ci zdecydować,

jak podłączyć moduł źródła czasu Wi-Fi w Twoim układzie.

Większość użytkowników będzie musiała po prostu podłączyć zasilanie między pinami VSYS i GND. Należy jednak pamiętać, że między VUSB i VSYS znajduje się dioda, więc jeśli jest podłączony kabel USB, może on zasilать VSYS, zwłaszcza wtedy, gdy napięcie VSYS jest niższe niż 5 V z USB.

O ile nie masz pewności, że nie podłączysz niczego do VSYS, gdy zasilanie jest podłączone do VUSB (na przykład przez gniazdo USB), najbezpieczniejszą opcją będzie podłączenie przychodzącego zasilania do VSYS za pomocą diody Schottky'ego, która zapobiegnie przepływowi prądu z VBUS do zasilania.

Biorąc pod uwagę, że większość osób będzie używać portu USB do programowania, konfigurowania i testowania Pico W, najprostszym rozwiązaniem jest odłączenie kabla USB przed podłączeniem do układu docelowego. W takim przypadku bezpośrednie połączenia z wyprowadzeniami Pico W będą prawidłowe.

W dalszej kolejności pokażemy również, jak podłączyć moduł źródła czasu Wi-Fi do niektórych z naszych najnowszych zegarów.

Prace nad oprogramowaniem

Raspberry Pi C SDK wciąż ewoluuje, zwłaszcza jego fragmenty dotyczące funkcji

Wi-Fi Pico W. Jest jednak dobrze udokumentowany, a zainteresowanie jest wystarczające, aby społeczność internetowa była również bardzo pomocna.

Napotkaliśmy pewne drobne trudności podczas prac nad oprogramowaniem, ale udało nam się je obejść. Używaliśmy wersji 1.5.0 SDK. Wersje wcześniejsze niż 1.4.0 nie obsługiwały Pico W, a późniejsze mogą się różnić.

Jak już wspomnieliśmy, Pico W ma dwa rdzenie procesora. Jeden z nich (drugi rdzeń) jest zaprogramowany do generowania danych NMEA i impulsów 1PPS. Ma to kluczowe znaczenie, ponieważ odkryliśmy, że D1 Mini (używany w projekcie źródła czasu z roku 2018) czasami blokował się (był zajęty i nie był w stanie uruchomić innych fragmentów swojego programu) podczas operacji Wi-Fi.

Dzięki skonfigurowaniu jednego rdzenia do wykonywania krytycznych czynności, moduł źródła czasu Wi-Fi może nadal działać, nawet w skrajnym przypadku, gdy jeden rdzeń procesora całkowicie się zawiesi. Rdzeń ten może nawet zresetować Źródło Czasu w pewnych warunkach.

Podczas resetowania niektóre dane są przechowywane w pamięci RAM, aby zachować bieżący czas podczas resetowania. Jest to możliwe, ponieważ pamięć RAM pozostaje zasilana podczas procesu miękkiego resetu.

Tabela 1. Polecenia konfiguracji źródła czasu Wi-Fi

Komenda	Funkcja	Uwagi
1	Skanuje sieci i wyświetla listę w kolejności malejącego RSSI	Na liście jest również wyświetlany kanał i typ uwierzytelniania. Numer wyświetlany w kolumnie n jest używany dla polecenia 4
2	Wyświetlanie bieżącej listy sieci	Lista jest aktywna, ale może nie odzwierciedlać zawartości pamięci Flash, chyba że zapis został zakończony
3n	Usunięcie elementu n z listy wyświetlanej przez polecenie 2	
4n	Dodaje sieć n z polecenia 1	Wyświetla również monit o hasło. Jeśli wszystkie gniazda są zajęte, zostanie wysłany komunikat o błędzie i trzeba będzie użyć polecenia 3, aby zwolnić gniazdo
5	Wprowadź nazwę sieci ręcznie	
6n	Wprowadź hasło dla sieci, używając n z listy wyświetlanej przez polecenie 2	Nie powinno być używane, chyba że konieczna jest zmiana hasła
7	Testowanie sieci na liście	Skanuje listę i próbuje połączyć się z każdą siecią po kolei. Może to trochę potrwać, a sukces jest zgłaszany tylko w przypadku uzyskania adresu IP
8	Zapisz wszystkie ustawienia w pamięci Flash	Dobrym pomysłem jest ponowne uruchomienie sterownika. Pozwoli to mieć pewność, że wszystkie ustawienia zostały poprawnie załadowane
9	Ustaw dwuliterowy kod kraju	Wczytywane tylko podczas uruchamiania, więc po ustawieniu i użyciu polecenia 8 w celu zapisania należy ponownie uruchomić sterownik
A	Ustaw URL API do pobierania informacji o szerokości i długości geograficznej na podstawie adresu IP	Powinny być zwrócone dwa wiersze tekstu z dziesiętną szerokością geograficzną w jednym wierszu i długością geograficzną w następnym. Ustaw URL na pusty, aby wyłączyć
B	Ustaw domyślną szerokość geograficzną	Wprowadź w formacie dziesiętnym
C	Ustaw domyślną długość geograficzną	Wprowadź w formacie dziesiętnym
D	Ustaw szybkość transmisji NMEA	Domyślnie jest to 9600, ale można użyć dowolnej szybkości z zakresu od 300 do 921600
E	Ustawianie kodu rozmówcy	Domyślnie jest to „GP”, ale mogą to być dowolne dwa znaki. „GP” działa dla wszystkich naszych zegarów
F	Ustawianie limitu czasu ważności NTP w minutach	Najdłuższy okres, przez który czas może być uznany za ważny bez (zazwyczaj cogodzinnej) aktualizacji NTP, od 60 min do 50000 min (około miesiąca)
G	Ustawianie adresu URL serwera NTP	Domyślnym adresem jest „pool.ntp.org”, który automatycznie przekierowuje do serwera znajdującego się w pobliżu geograficznym. Można użyć innych, takich jak „time.nist.gov”. Adres IP może nie być prawidłowy do momentu podłączenia do sieci
H	Ustaw domyślny rok	Rok używany podczas uruchamiania, gdy nie są dostępne żadne inne dane czasowe, od 1970 do 4095. Zobacz osobny artykuł na temat błędu Y2K38, aby dowiedzieć się, dlaczego jest to ważne
I	Przełączanie debugowania danych wyjściowych NMEA do portu szeregowego USB	Służy do sprawdzania i debugowania danych NMEA. To ustawienie jest zapisywane na wypadek, gdyby dane te miały być zawsze dostępne przez port szeregowy USB
J	Ponowne uruchomienie Pico W	Zaleca się ponowne uruchomienie po zapisaniu ustawień, aby mieć pewność, że wszystkie ustawienia zostaną ponownie załadowane podczas uruchamiania. Jeśli przytrzymasz przycisk BOOTSEL podczas ponownego uruchamiania, możesz użyć tej metody, aby przejść do trybu bootloadera

utruty połączenia zostaje podjęta próba ponownego nawiązania połączenia.

Ponieważ wiele docelowych zastosowań Źródła Czasu Wi-Fi wymaga jego krótkotrwałego uruchomienia (dla oszczędzania energii), urządzenie przy starcie skanuje sieć, by upewnić się, że pierwsza próba połączenia dotyczy dostępnej sieci.

Wirtualny port szeregowy generuje również informacje o stanie, głównie dotyczące statusu Wi-Fi i czasu od ostatniej aktualizacji NTP. Jedną z opcji menu umożliwia zrzucanie danych NMEA do wirtualnego portu szeregowego w celu łatwego debugowania.

Pierwszy rdzeń jest również odpowiedzialny za sterowanie wbudowaną diodą LED Pico W, która służy do migania wskaźnika stanu.

Dioda LED włącza się na stałe po podłączeniu zasilania, wskazując, że Źródło Czasu jest włączone prawidłowo. Może również migać raz, dwa lub trzy razy na sekundę. Jedno mignięcie oznacza, że urządzenie jest podłączone do sieci Wi-Fi, a dwa mignięcia oznaczają, że czas jest uznawany za prawidłowy. Trzy błyski występują, gdy oba te warunki są spełnione.

W ogólnym przypadku, czas jest prawidłowy, jeśli w ciągu ostatnich kilku godzin otrzymano aktualizację NTP, chociaż limit ten można dostosować.

Generator kwarcowy używany w Pico W ma tolerancję 30 ppm, co oznacza, że może dryfować nawet o jedną sekundę na osiem godzin. W praktyce widzieliśmy korekty NTP do 200 ms, więc jesteśmy pewni, że przy ustawieniach domyślnych czas będzie dokładny do pół sekundy.

Programowanie Pico W

Warto najpierw zaprogramować Pico W i sprawdzić, czy działa zgodnie z oczekiwaniami. Przytrzymaj przycisk BOOTSEL w Pico W i podłącz moduł do komputera. Powinien pojawić się napęd USB o nazwie „RPI-RP2”. Skopiuj do niego plik NEW_CLAYTONS_1.UF2. Po około sekundzie dioda LED powinna się zaświecić.

Następnie, aby uzyskać dostęp do menu, można użyć programu monitora szeregowego. W systemie Windows używamy programu TeraTerm, natomiast w systemie Linux można użyć programu minicom. Dostęp do interaktywnego menu uzyskuje się po otwarciu wirtualnego portu szeregowego Pico W.

Upewnij się, że program terminala używa CR lub CR+LF jako zakończenia linii. Ponieważ jest to wirtualny port szeregowy, szybkość transmisji jest nieistotna i każde ustawienie szybkości transmisji powinno działać.

Podstawowa konfiguracja

Po wszystkich poleceniach należy wpisać Enter.

Pico W implementuje kody krajów pozwalające zapewnić, że do komunikacji używane są prawidłowe (legalne) kanały Wi-Fi. Domyślne ustawienie „XX” to podzbiór, który jest bezpieczny na całym świecie, ale nie pozwala na korzystanie z niektórych kanałów Wi-Fi. Ustawienie to powinno więc działać, ale może nie być optymalne.

Dobłą praktyką jest wybór kodu odpowiedniego dla kraju użytkownika. Wybierz polecenie 9 i wpisz dwuliterowy kod kraju (np. AU, NZ, US, UK, PL). Następnie zapisz

Command 1

```
1
Scanning
Scan complete
Scanned network list:
n  SSID          RSSI  Chan  Auth
0  AndroidAP4AA0  -44   1     PASS
1  APV Admin Only  -65   3     PASS
2  APHV Conference -66   3     PASS
3  TPW4G_ZeB426   -82   11    PASS
4  Wi-Fi-5E5EE1   -84   8     PASS
5  NTGR_4E0C      -93   11    PASS
```

Command 43

```
43
2 TPW4G_ZeB426
Added OK
Enter password.
PASSWORD
password saved.
```

Command 2

```
2
Saved network list:
0 AndroidAP4AA0
1 Tim
2 TPW4G_ZeB426
```

Command 32

```
32
SSID deleted.
Saved network list:
0 AndroidAP4AA0
1 Tim
```

Command 7

```
7
Testing networks.
0 AndroidAP4AA0
>connected with IP:192.168.208.140
1 Tim
>SSID not found
2 Networks tested, 1 OK
```

Ekran 2. Na tym edytowanym zrzucie ekranowym pokazano dane wyjściowe niektórych bardziej złożonych poleceń. Należy zauważyć, że polecenia te zostały wydane w pokazanej kolejności, aby dodać, a następnie usunąć identyfikator SSID. Polecenia 3 i 4 wymagają drugiego parametru, którym jest liczba zwrócona przez polecenie 2 i 1 (odpowiednio) wydane wcześniej

ustawienia poleceniem 8 i uruchom ponownie Pico W (polecenie JJ).

Uwaga edytora: kody powinny być w formacie alpha-2, patrz: <https://w.wiki/4kP>

W razie potrzeby połącz się ponownie z Pico W. TeraTerm zwykle robi to automatycznie.

Teraz, w celu uruchomienia skanowania Wi-Fi użyj opcji menu 1. Powinno się to zakończyć w ciągu około sekundy. Sieci są wymienione w kolejności RSSI (siła sygnału), więc powinieneś znaleźć swój identyfikator SSID na samej górze.

Należy pamiętać, że polecenia wymienione z przyrostkiem n wymagają drugiego argumentu numerycznego. Na przykład, jeśli Twoja sieć pojawia się jako pierwsza (obok numeru 0), wprowadź polecenie 40. Zostanie wyświetlony monit o hasło dla tej sieci. Wpisz je i naciśnij Enter.

Można wprowadzić wiele sieci bez ponownego skanowania. Jeśli sieć nie zostanie wyświetlona, użyj polecenia 5 pozwalającego ręcznie wprowadzić nazwę – zostaniesz również poproszony o podanie hasła.

Samo polecenie 6 służy do zmiany lub ustawienia hasła, jeśli na przykład zostało ono wprowadzone nieprawidłowo.

Do przetestowania zapisanej sieci wypróbuj następnie polecenie 7. Powinien pojawić się komunikat „Connected with IP”, a następnie adres IP dla każdego SSID. Jeśli nie pojawi się, spróbuj ponownie. Jeśli pojawi się komunikat „Auth failed (password?)”, wprowadzone hasło może być nieprawidłowe. Do ponownego wprowadzenia możesz użyć polecenia 6.

Jeśli do portu szeregowego nie zostało nic nadane będzie on co około 15 sekund wysyłał aktualizacje. Ma to na celu zapobieganie zakłócaniu procesu konfiguracji przez aktualizacje.

Jeśli wszystko przebiega prawidłowo, użyj polecenia 8 zapisującego ustawienia w pamięci Flash i ponownie uruchom urządzenie, żeby upewnić się, czy ustawienia zostały załadowane. Jest to konieczne, ponieważ niektóre parametry można ustawić tylko raz, a najprostszym sposobem obejścia tego problemu jest ponowne uruchomienie urządzenia.

Powinno to stanowić minimum potrzebne do skonfigurowania Źródła Czasu Wi-Fi do pracy z większością naszych zegarów. Szczegółowa lista poleceń, wraz z ich zastosowaniem i przeznaczeniem, została przedstawiona w tabeli 1.

Na ekranie 2 pokazano typowe odpowiedzi na bardziej powszechne i złożone polecenia. Większość innych poleceń wymaga prostej odpowiedzi i zgłosi komunikat, jeśli wystąpi problem.

Z kolei na ekranie 3 widoczny jest typowy przebieg procesu uruchamiania, chociaż zdarzenia mogą nie występować w tej kolejności. Możesz także zauważyć znacznie większe korekty NTP, ale jest to normalne.

Można przełączyć przesyłanie zdań GPS przez port szeregowy USB za pomocą polecenia I. Jest to widoczne na ekranie 4. Oczywiście dane mogą być inne. Jeśli posiadasz program komputerowy, który może przetwarzać dane GPS, możesz go użyć do weryfikacji danych.

```
Time is 04:01:30 on 13/02/2023. NO NTP.
Connect failed
Connecting to 0 AndroidAP4AA0
Skip IPAPI fetch, no Wi-Fi.
```

NTP adjustment: 11953

```
Connected with IP: 192.168.130.138
Time is 04:01:45 on 13/02/2023. NTP OK. Last updated 0 minutes ago.
IPAPI start.
Headers of 170 bytes report 18 bytes of content.
Received 18 bytes.
HTTP finished:200
OK
Lat/Lon=-27.467899,153.032501
Time is 04:02:00 on 13/02/2023. NTP OK. Last updated 0 minutes ago.
Time is 04:02:15 on 13/02/2023. NTP OK. Last updated 0 minutes ago.
Time is 04:02:30 on 13/02/2023. NTP OK. Last updated 0 minutes ago.
Time is 04:02:45 on 13/02/2023. NTP OK. Last updated 1 minutes ago.
Time is 04:03:00 on 13/02/2023. NTP OK. Last updated 1 minutes ago.
Time is 04:03:15 on 13/02/2023. NTP OK. Last updated 1 minutes ago.
Time is 04:03:30 on 13/02/2023. NTP OK. Last updated 1 minutes ago.
Time is 04:03:45 on 13/02/2023. NTP OK. Last updated 2 minutes ago.
Time is 04:04:01 on 13/02/2023. NTP OK. Last updated 2 minutes ago.
```

Ekran 3. Kilka ostatnich wierszy na tym ekranie (przy użyciu programu terminala szeregowego TeraTerm) pokazuje, że moduł źródła czasu Wi-Fi potoczył się z Wi-Fi i zaktualizował swój czas z serwerów NTP. Poprzednie linie są wyświetlane typowo podczas normalnego uruchamiania

```
$GPRMC,050215.000,A,2728.0004,S,15301.0057,E,0.00,000.00,130223,,,3F
$GPGGA,050215.000,2728.0004,S,15301.0057,E,1.04,1.0,0.0,M,0.0,M,,*78
$GPGSA,A,3,,,,,,,,,,,,,1.00,1.00,1.00,*2F
$GPRMC,050216.000,A,2728.0004,S,15301.0057,E,0.00,000.00,130223,,,3C
$GPGGA,050216.000,2728.0004,S,15301.0057,E,1.04,1.0,0.0,M,0.0,M,,*7B
$GPGSA,A,3,,,,,,,,,,,,,1.00,1.00,1.00,*2F
$GPRMC,050217.000,A,2728.0004,S,15301.0057,E,0.00,000.00,130223,,,3D
$GPGGA,050217.000,2728.0004,S,15301.0057,E,1.04,1.0,0.0,M,0.0,M,,*7A
$GPGSA,A,3,,,,,,,,,,,,,1.00,1.00,1.00,*2F
$GPRMC,050218.000,A,2728.0004,S,15301.0057,E,0.00,000.00,130223,,,32
$GPGGA,050218.000,2728.0004,S,15301.0057,E,1.04,1.0,0.0,M,0.0,M,,*75
$GPGSA,A,3,,,,,,,,,,,,,1.00,1.00,1.00,*2F
$GPRMC,050219.000,A,2728.0004,S,15301.0057,E,0.00,000.00,130223,,,33
$GPGGA,050219.000,2728.0004,S,15301.0057,E,1.04,1.0,0.0,M,0.0,M,,*74
$GPGSA,A,3,,,,,,,,,,,,,1.00,1.00,1.00,*2F
$GPRMC,050220.000,A,2728.0004,S,15301.0057,E,0.00,000.00,130223,,,39
$GPGGA,050220.000,2728.0004,S,15301.0057,E,1.04,1.0,0.0,M,0.0,M,,*7E
$GPGSA,A,3,,,,,,,,,,,,,1.00,1.00,1.00,*2F
$GPRMC,050221.000,A,2728.0004,S,15301.0057,E,0.00,000.00,130223,,,38
$GPGGA,050221.000,2728.0004,S,15301.0057,E,1.04,1.0,0.0,M,0.0,M,,*7F
$GPGSA,A,3,,,,,,,,,,,,,1.00,1.00,1.00,*2F
$GPRMC,050222.000,A,2728.0004,S,15301.0057,E,0.00,000.00,130223,,,3B
$GPGGA,050222.000,2728.0004,S,15301.0057,E,1.04,1.0,0.0,M,0.0,M,,*7C
$GPGSA,A,3,,,,,,,,,,,,,1.00,1.00,1.00,*2F
```

Ekran 4. Polecenie I wysyła zdania GPS do wirtualnego terminala szeregowego umożliwiając potwierdzenie generowanych danych. Ustawienie to można zapisać w pamięci Flash, dzięki czemu dane GPS będą nadal wysyłane do wirtualnego portu szeregowego USB nawet po ponownym uruchomieniu

Podłączanie do zegara

Moduł źródła czasu Wi-Fi może połączyć się na poziomie logicznym z prawie każdym systemem, który oczekuje danych NMEA 0183, jednak brak dokładnych danych o prędkości i lokalizacji oznacza, że w niektórych przypadkach nie jest to najlepszy wybór.

Nie zalecamy używania go jako źródła dla żadnego z naszych wzorców częstotliwości opartych na GPS. Sygnał 1PPS dostarczany przez opisywany moduł źródła czasu Wi-Fi nie ma na celu uzyskania niezbędnej precyzji, a ponieważ zawsze podaje tylko prędkość 0 węzłów, nie będzie zbyt przydatny w GPS Finesaver (siliconchip.au/Article/11673).

Nie będzie też zbyt przydatny jako pomoc nawigacyjna, wykluczając zaawansowany komputer GPS z 2021 roku (siliconchip.au/

Series/366), więc zakładamy, że używasz Źródła Czasu Wi-Fi z jednym z naszych zegarów GPS.

Poniżej znajdują się instrukcje dotyczące korzystania ze źródła czasu w niektórych projektach zegarów GPS, które opublikowaliśmy w ciągu ostatnich dziesięciu lat. W tabeli 2 podsumowano również, w jaki sposób opisywane tu Źródło Czasu Wi-Fi może zastąpić niektóre popularne moduły GPS.

Należy pamiętać, że te połączenia mogą nie być optymalne, zwłaszcza w przypadku zegarów zasilanych bateryjnie. Warto poeksperymentować z alternatywnymi konfiguracjami. Z tego powodu sugerowane okablowanie dla najnowszych zegarów zasilanych bateryjnie różni się od informacji podanych w tabeli 2.

Problem polega na tym, że moduł źródła czasu Wi-Fi ma wyższe zapotrzebowanie na prąd niż większość modułów GPS, a układy czasami nie są w stanie zapewnić wystarczającego prądu do jego zasilania.

Nowy zegar analogowy z synchronizacją GPS – EdW 4/2025.

Najnowszy zegar z synchronizacją GPS został opublikowany we wrześniu 2022 r. (siliconchip.au/Series/391) oraz EdW 4/2025, a następnie w wydaniu z listopada 2022 r. pojawiła się aktualizacja opisująca sposób podłączenia oryginalnego Źródła Czasu GPS Claytona.

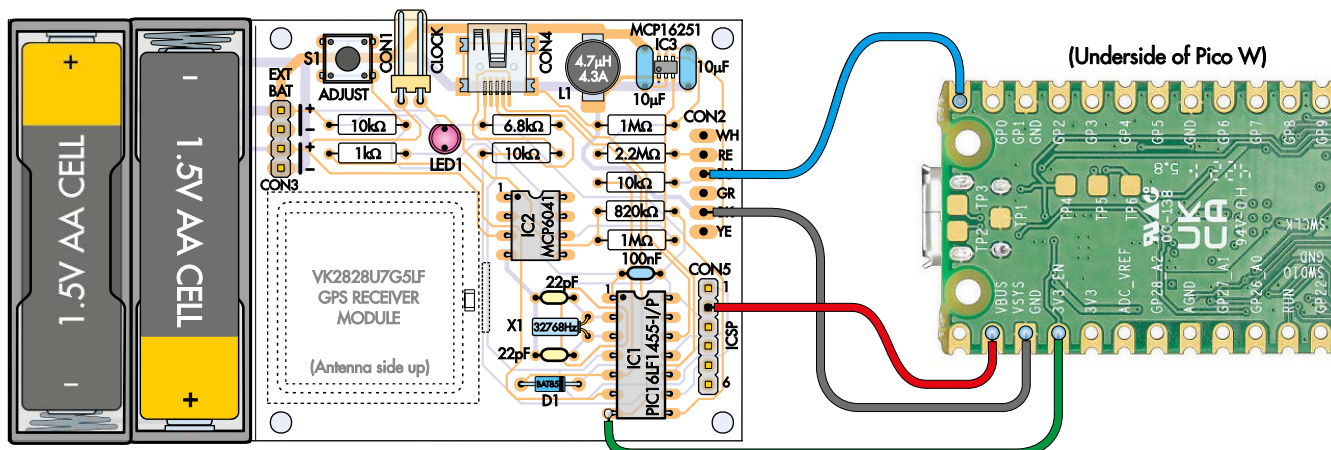
Podobnie jak wiele naszych ostatnich projektów GPS, wykorzystuje on moduł GPS VK2828U7G5LF. W rzeczywistości zalecamy ten moduł jako zamiennik wszystkich poprzednich modułów GPS, których używaliśmy w projektach zegarów.

Moduł VK2828U7G5LF ma sześć złączy, ale dla Źródła Czasu potrzebne są tylko cztery. Wszystkie połączenia są dość proste, chociaż nie wszystkie sygnały w złączu modułu GPS są wykorzystane (rysunek 3).

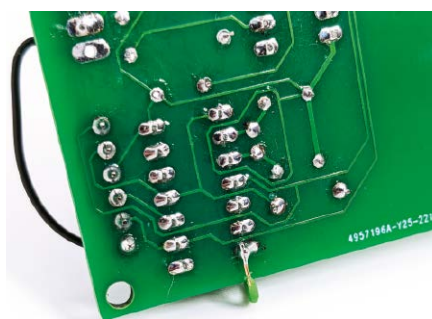
Czarne i niebieskie przewody są podłączone do oczywistych pinów złącza modułu GPS. Czerwony przewód doprowadza zasilanie z baterii bezpośrednio do pinu VSYS płytki Pico W – można w tym celu wykorzystać jeden z pinów złącza ICSP.

Zwróć uwagę, że zasilanie jest doprowadzone przez złącze, dzięki czemu można je łatwo odłączyć podczas programowania przez USB. Chroni to baterię przed przypadkowym podaniem napięcia 5 V z kabla USB.

Przy użyciu tylko tych trzech linii Pico W działałby nieprzerwanie. Zielony przewód łączy pin 3V3_EN z pinem 7 układu IC1 na płycie drukowanej zegara. Pin ten jest zwykle używany do sterowania przetwornicą



Rysunek 3. Podłączenie do nowego zegara analogowego zsynchronizowanego z GPS za pomocą pinu 3V3_EN na Pico W zapewnia najbardziej efektywne wykorzystanie wbudowanej przetwornicy podwyższającej Pico W omijając własny stabilizator boost zegara (Pico W jest dla przejrzystości pokazany na rysunkach 3...6 w rozmiarze większym niż rzeczywisty). W tym przypadku można pominąć IC3, L1 i dwa kondensatory 10 µF



Moduł źródła czasu Wi-Fi podłączony do nowego zegara analogowego z synchronizacją GPS z 2022 roku (EdW 4/2025). Aby oszczędzać baterię, przetwornica podwyższająca na płytce drukowanej zegara jest pominięta. W rzeczywistości te wbudowane komponenty można całkowicie pominąć. Zdjęcie w lewym górnym rogu pokazuje zielony przewód podłączony bezpośrednio do pinu 7 układu IC1 na odwrocie płytki drukowanej



podwyższającą zegara. Dla wygody podłącza się go pod płytkę drukowaną, jak pokazano na zdjęciu.

Na schemacie pominięto przetwornicę podwyższającą (boost) w nowym zegarze analogowym z synchronizacją GPS, co jest możliwe, ponieważ Pico W ma własną przetwornicę buck/boost. Oznacza to również, że jeśli budujesz płytkę zegara od podstaw, możesz pominąć układ scalony przetwornicy podwyższającej wraz z powiązanymi elementami.

Dzięki takiemu rozwiązaniu Pico W będzie zasilany nawet wtedy, gdy napięcie baterii spadnie do 2 V, czyli dolnego limitu zegara. Na tym etapie nie było wystarczającego napięcia, aby zasilić niebieską diodę LED w zegarze, ale pozostałe elementy działały zgodnie z oczekiwaniami.

Zdjęcia pokazują Źródło Czasu podłączone za pomocą krótkich przewodów, a następnie zamontowane na układach scalonych za pomocą dwustronnej taśmy klejącej. Zwróć uwagę, że antena Wi-Fi Pico W jest wyraźnie widoczna na płytce drukowanej poniżej.

Moduł źródła czasu Wi-Fi zazwyczaj potrzebuje około 25 sekund, aby „nawiązać

połączenie”, często trwa to szybciej, a czasami dłużej, jeśli nie połączy się natychmiast z siecią Wi-Fi. Powinno to dotyczyć większości zegarów korzystających ze Źródła Czasu.

Po włączeniu zegara z podłączonym Źródłem Czasu, zegar zamiga swoją diodą LED raz lub dwa razy, po czym dioda LED Źródła Czasu włączy się i zacznie migać z taką samą częstotliwością jak dioda LED zegara. Po kilku kolejnych sekundach dioda LED źródła czasu zgaśnie, wskazując, że zegar uzyskał prawidłowy czas i wyłączył Źródło Czasu.

Zazwyczaj dioda LED zegara powinna również zgasnąć po maksymalnie pół godzinie (a zegar powinien zacząć tykać), więc jeśli miga dłużej, należy to sprawdzić.

Ogólnie stwierdził, że jeśli dane wyświetlane za pośrednictwem terminala szeregowego USB były prawidłowe, Źródło Czasu działało poprawnie po podłączeniu do zegara.

Sterownik zegara analogowego z synchronizacją GPS – EdW 9/2022.

W artykule dotyczącym Sterownika Zegara Analogowego z Synchronizacją GPS z lutego 2017 r. (siliconchip.au/Article/10527)

– przedruk w EdW 9/2022 – również zalecano użycie modułu GPS VK2828U7G5LF. Należy pamiętać, że nie testowaliśmy tego układu ani żadnego z poniższych z zegarami sprzed 2022 r.

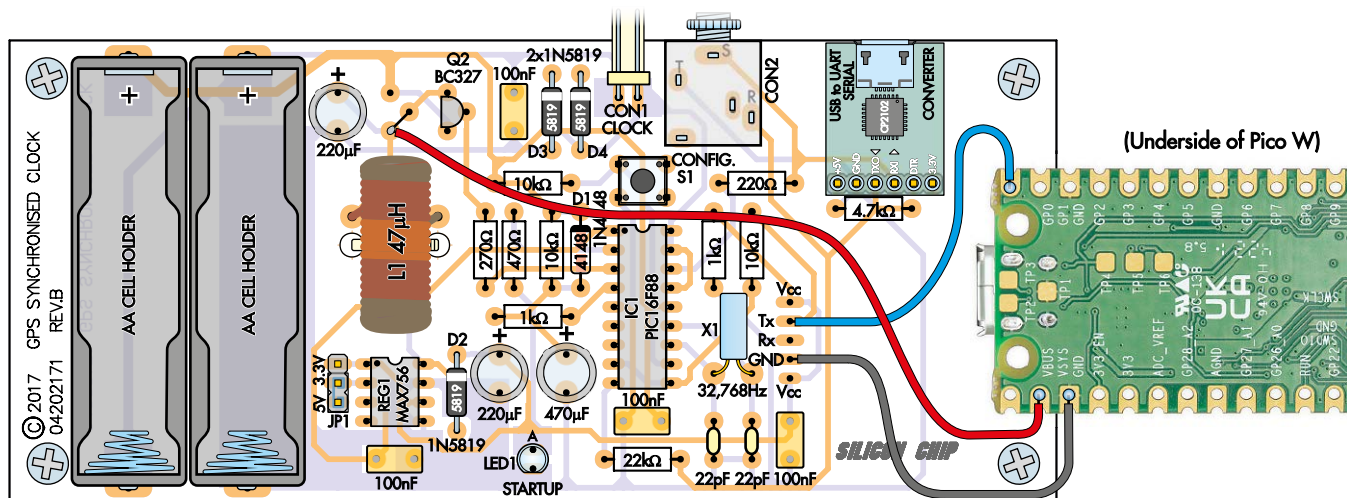
Tutaj proponujemy odmianę, która pozwoli uniknąć niewielkiej nieefektywności poprzez ominięcie przetwornicy podwyższającej sterownika zegara. Ponieważ Pico W może pracować z napięciem do 1,8 V w VSYS, pobieramy 3 V bezpośrednio z wejścia przetwornicy podwyższającej, jak pokazano na **rysunku 4**.

Zegary GPS od 2015 r.

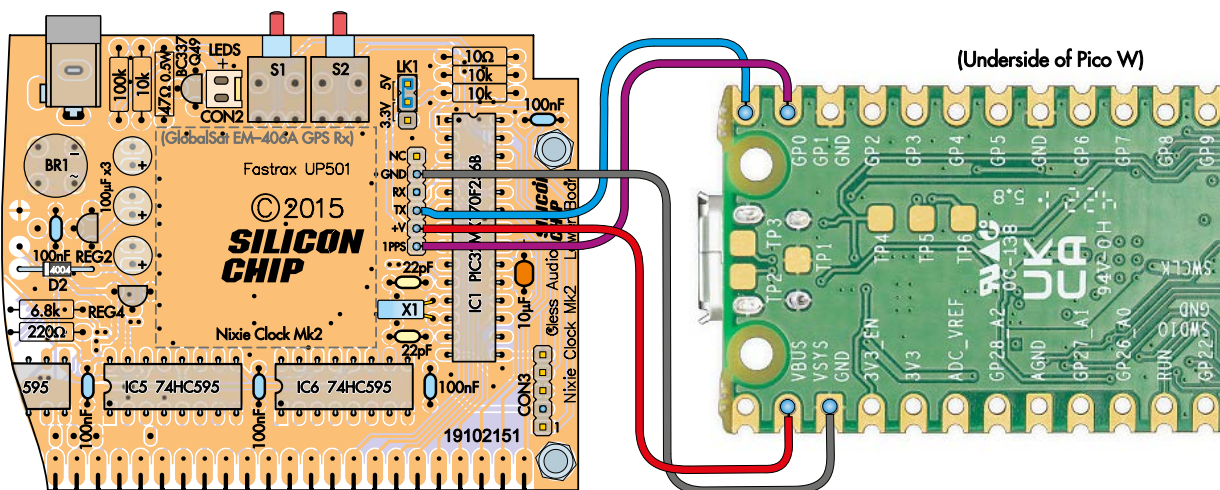
Wszystkie wcześniej opublikowane przez nas zegary GPS były zasilane z zewnętrznego źródła, dlatego nie występują w nich problemy typowe dla układów zasilanych z baterii.

Na **rysunkach 5 i 6** pokazano, jak podłączyć moduł źródła czasu Wi-Fi odpowiednio do 6-cyfrowego zegara Retro Nixie Clock Mk.2 i 6-cyfrowego zegara LED GPS o wysokiej widoczności.

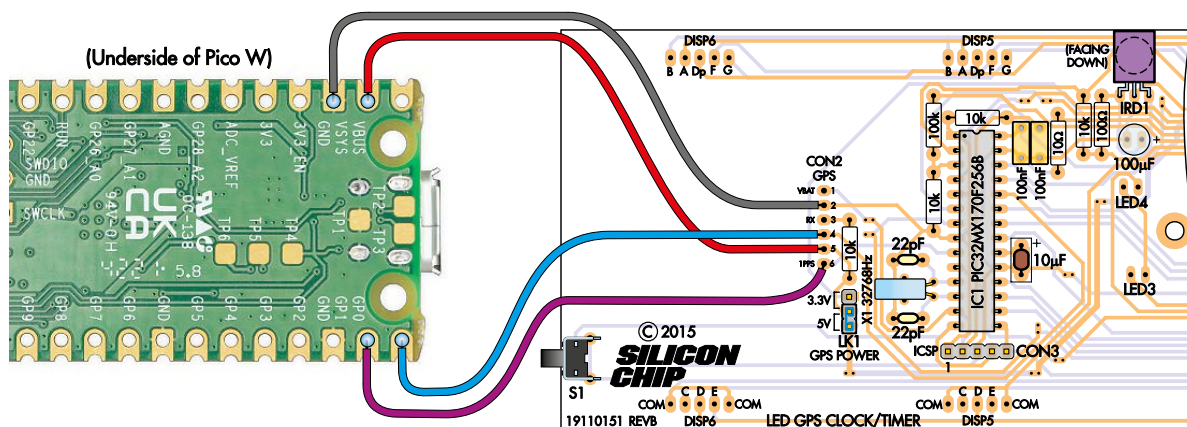
Warto zauważyć, że w obu urządzeniach zastosowano identyczny układ pinów złączy do współpracy z modułami GPS – zgodny z połączeniami pokazanymi w tabeli 2.



Rysunek 4. Podłączenie Źródła Czasu do sterownika Zegara Analogowego z Synchronizacją GPS z 2017 r. (EdW 9/2022). Pozwala również ominąć wbudowany stabilizator zegara w celu zasilania Pico W. Należy pamiętać, że nie testowaliśmy tej konfiguracji



Rysunek 5. Połączenia z zegarem Nixie 2015. LK1 (który pozwala ustawić zasilanie na 3,3 V lub 5 V dla podłączonego modułu) powinien być ustawiony na napięcie 5 V. Mimo to, niniejszy projekt nie jest zasilany z baterii, więc wydajność jest mniej krytyczna



Rysunek 6. W 6-cyfrowym zegarze GPS LED o wysokiej widoczności zastosowano ten sam układ pinów złącza, jaki obowiązywał w zegarze Nixie, więc okablowanie jest takie samo, podobnie jak wybór ustawienia LK1 w pozycji 5 V

Ze względu na wydajność, napięcia zasilania GPS w tych projektach (LK1 dla obu projektów) powinno być ustawione w pozycji 5 V, ponieważ Pico W bez problemu działa z napięciem 5 V na wejściu VSYS.

Jeśli po podłączeniu Źródła Czasu do jednego z pozostałych zegarów wystąpią jakiegokolwiek problemy, najprawdopodobniej będą one

związane z zasilaniem. Sprawdź, czy napięcie na pinie 3V3_OUT jest bliskie 3,3 V. Jeśli nie, układ może nie być w stanie dostarczyć wystarczającego prądu dla Pico W.

Wnioski

Płytki Pico W zapewnia przydatne funkcje w zastosowaniach takich, jak zintegrowany

zasilacz buck-boost, specjalne urządzenie peryferyjne USB umożliwiające oddzielną konsolę konfiguracyjną i dobre wsparcie oprogramowania.

Moduł źródła czasu Wi-Fi jest naturalnym rozwinięciem oryginalnego Źródła Czasu GPS Claytona z 2018 roku i jest to urządzenie o równie prostej budowie i niedrogie. Wersja Pico W zapewnia dodatkowe funkcje, w szczególności możliwość automatycznego łączenia się z jedną z kilku sieci Wi-Fi.

W chwili pisania tego tekstu obsługa Bluetooth jest na wczesnym etapie (beta), więc sprawdzimy, czy możliwe jest dodanie interfejsu Bluetooth do konfiguracji. Byłoby to bardzo przydatne do aktualizacji ustawień, ponieważ wyeliminowałoby potrzebę podłączania kabla USB. ■

Tim Blythman

Tabela 2. Mapowanie pinów Źródła Czasu Wi-Fi w porównaniu do modułów GPS			
Pico W	VK2828	EM408	Uwagi
Pin 1 GP0 (dane NMEA)	TxD (4, niebieski)	Tx (4)	
Pin 2 GP1 (1PPS)	1PPS (6, biały)	Niepodłączony	Niepotrzebne w większości zastosowań
Pin 3/38 GND	GND (2, czarny)	GND(2)	
Pin 39 VSYS	VCC (5, czerwony)	V+(5)	Lub inne źródło od 1,8 V do 5,5 V
Pin 40 VBUS	Niepodłączony	Niepodłączony	
Niepotrzebne	EN (1, żółty)	EN (1)	
Niepotrzebne	RxD (3, zielony)	RX(3)	

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022.
www.siliconchip.com.au

AT-AD269S
Mikroskop cyfrowy
z ekranem 10 cali,
powiększenie do 5000×,
5 obiektywów i endoskop
ANDONSTAR AD269S-M



AT-AD409PRO
Mikroskop do lutowania
z profesjonalnym
metalowym stojakiem,
ekran 10,1 cala,
powiększenie do 300×, HDMI
ANDONSTAR AD409Pro



BESTSELLERY sklepu AVT – sklep.avt.pl

Mikroskopy cyfrowe dla elektroników

Rabat dla Czytelników EdW
przy zakupie podaj kod **EdW2505MC**

Kod ważny do 30.09.2025

-3%

Rabat dla Prenumeratorów EdW
przy zakupie podaj numer prenumeraty

-6%

AT-AD246S-M
Mikroskop cyfrowy 7 cali
z powiększeniem:
60...240×, 18...720×,
1560...2040×
ANDONSTAR AD246S-M



AT-AD407
Mikroskop cyfrowy 7 cali,
powiększenie do 270×
ANDONSTAR AD407

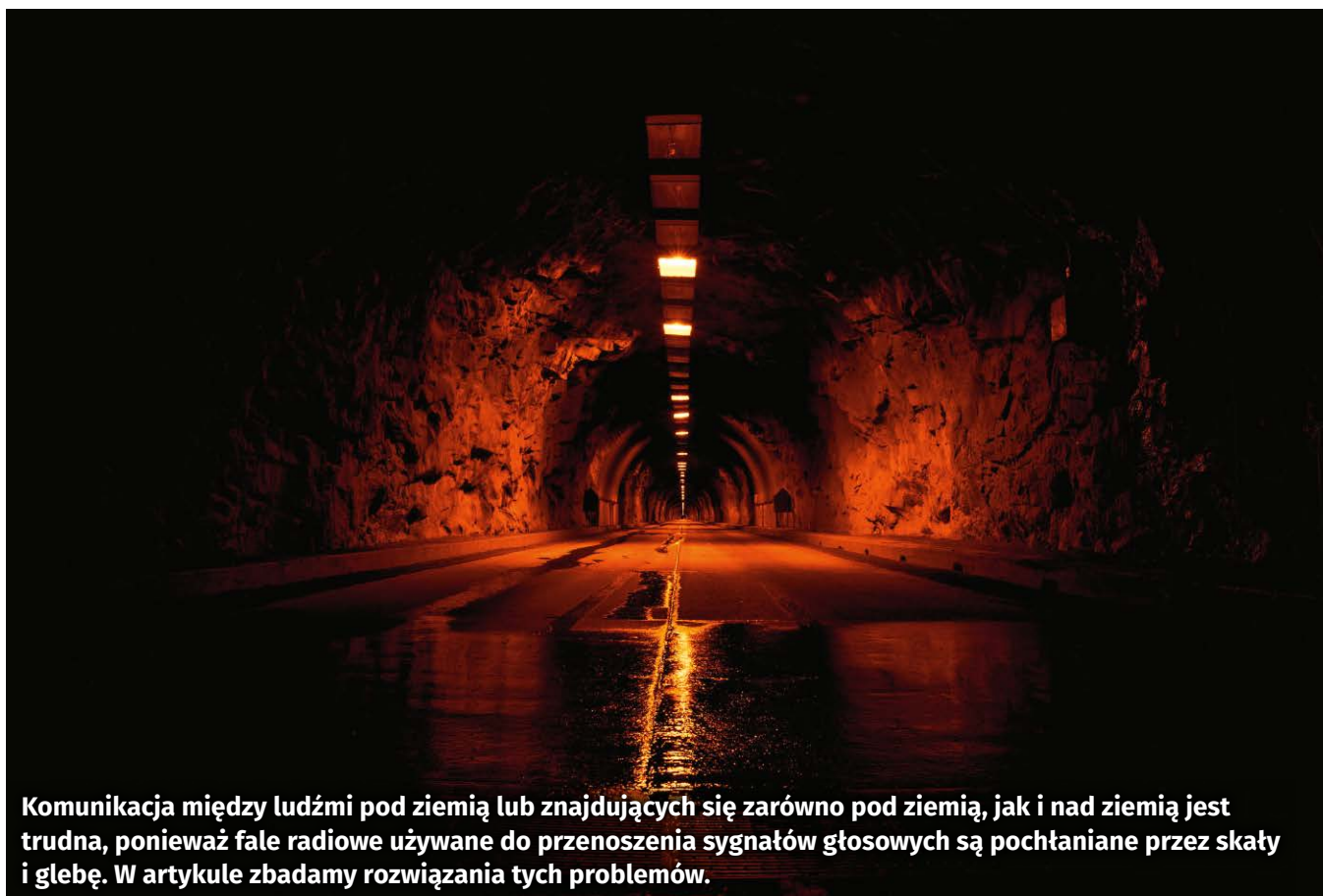


AT-AD249S-M
Mikroskop cyfrowy 10 cali
z powiększeniem:
60...240×, 18...720×, 1560...2040×
ANDONSTAR AD249S-M



AT-AD210
Mikroskop cyfrowy 5...260×
z wyświetlaczem 10,1 cala
ANDONSTAR AD210





Komunikacja między ludźmi pod ziemią lub znajdujących się zarówno pod ziemią, jak i nad ziemią jest trudna, ponieważ fale radiowe używane do przenoszenia sygnałów głosowych są pochłaniane przez skały i glebę. W artykule zbadamy rozwiązania tych problemów.

Źródło obrazu: <https://unsplash.com/photos/5p-37k8hKc>

Łączność w różnych mediach, część 2 – ziemia

Komunikacja podwodna omawiana we wcześniejszym artykule koncentrowała się głównie na okrętach podwodnych i innych podobnych jednostkach. Pod ziemią jest również miejsc, takich jak kopalnie, tunele i systemy jaskiń, w których ludzie mogą potrzebować komunikować się ze sobą lub z osobami pozostającymi na powierzchni.

Są też przypadki, takie jak lawiny, w których ludzie mogą zostać zasypani śniegiem, co stwarza podobne wyzwania. Nawet pozornie niezwiązana z tym łączność radiowa w kabinach samolotów oparta jest na tej samej technologii i rozwiązaniach. Zaczniemy od opisanie niektórych koncepcji stosowanych we wszystkich tych przypadkach.

Promieniujące linie zasilające (feedery)

Promieniujące linie zasilające, znane również jako „nieszczelne linie koncentryczne”

lub „nieszczelne linie zasilające”, są ważne dla komunikacji pod ziemią lub w dowolnym zamkniętym obszarze osłoniętym przed nadajnikami radiowymi. Mogą być stosowane w jaskiniach, tunelach, kopalniach, parkingach, a nawet wewnątrz samolotów lub statków.

Promieniująca linia zasilająca (feeder) działa jak niedoskonały kabel koncentryczny, w którego oplocie ekranującym (zewnętrznym) wykonano szczeliny lub przerwy umożliwiające emisję promieniowania elektromagnetycznego (**rysunek 30**). Jest to przeciwieństwo zwykłego kabla koncentrycznego, który został zaprojektowany tak, aby zatrzymywać lub blokować jak najwięcej promieniowania elektromagnetycznego.

Ponieważ siła sygnału jest tracona podczas przesyłania go linią promieniującą, konieczne jest stosowanie wzmacniania rozmieszczonych w regularnych odstępach, co około 350...500 m.

Promieniujące feedery stosuje się między innymi w tunelach drogowych, takich jak Sydney Harbour Tunnel, Lane Cove Tunnel, Burnley Tunnel, AirportlinkM7 czy Northbridge Tunnel. Dzięki nim można tam odbierać stacje radiowe i sygnał sieci komórkowych nawet w miejscach odległych od wlotów tunelu.

Komunikacja w tunelach drogowych i kolejowych

Sygnały radiowe nie docierają zbyt daleko w głąb tuneli. Sygnały AM mają długość fali od 175 m do 555 m, więc nie dotrą daleko w głąb tunelu, biorąc pod uwagę, że jego średnica będzie znacznie mniejsza niż długości fali.

Sygnały FM o długości fali od 2,8 m do 3,4 m mogą przemieszczać się przez wystarczająco szeroki tunel, ale aby sygnał mógł czyściej wejść do tunelu, musiałby znajdować się

na linii wzroku z wnętrza tunelu. Odbity sygnał z zewnątrz byłby znacznie słabszy. Bez promieniującego feedera sygnał radiowy zostałby w dużej mierze pochłonięty wskutek wielokrotnych odbić od ścian tunelu, chyba że tunel byłby idealnie prosty i miał bezpośredni „widok” na nadajnik.

Długości fal dla częstotliwości DAB mieszczą się w zakresie od 1,3 m do 1,6 m i zachowują się podobnie do sygnałów FM. Sygnały telefonii komórkowej mają jeszcze mniejsze długości fal, od 43 cm do milimetrów w przypadku 5G. Sygnał mógłby docierać w głąb tunelu, jeśli byłby on idealnie prosty i zachowana byłaby widoczność nadajnika.

Warunki te są rzadko spełnione, więc kontakt radiowy jest zwykle utrzymywany wewnątrz tunelu dzięki „retransmisji”. Retransmisja komercyjnych kanałów AM, FM i DAB zwiększa zadowolenie kierowców i zmniejsza rozproszenie uwagi, ponieważ nie przerywa ich ulubionego programu radiowego. Biorąc pod uwagę drogie opłaty, które płaci się za korzystanie z tych tuneli, jest to minimum, którego mogliby oczekiwać!

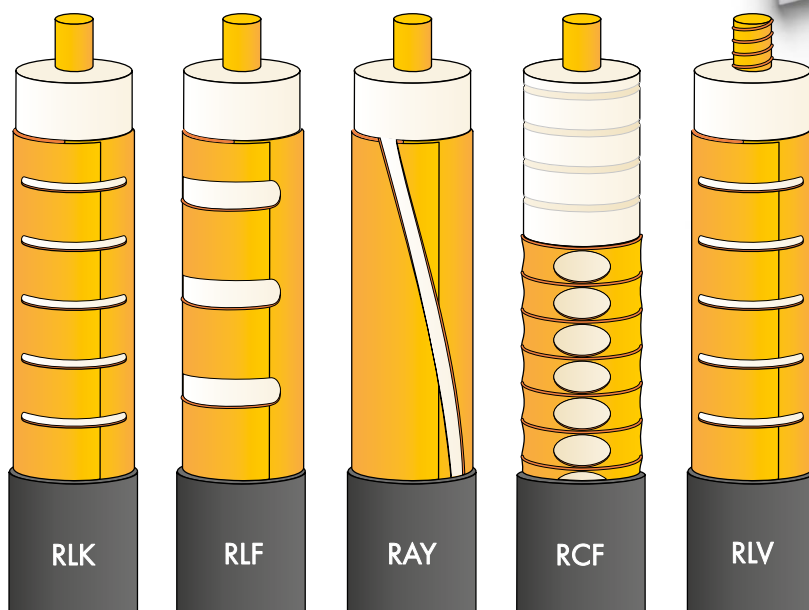
Takie systemy retransmisji zazwyczaj mają również funkcję zwaną „włamaniem audio”, dzięki której komunikaty alarmowe lub serwisowe mogą być nadawane we wszystkich retransmitowanych programach radiowych, niezależnie od kanału, na który dostrojone jest radio pojazdu. W nagłych wypadkach zwykle pojawiają się znaki z napisem „włącz radio” (lub podobnym), aby kierowcy mogli zostać poinformowani o najlepszym sposobie postępowania.

Retransmisja pasywna kontra aktywna

Sygnały radiowe mogą być retransmitowane pasywnie lub aktywnie. Retransmisja pasywna (**rysunek 31**) polega na podłączeniu anteny zewnętrznej do jednej lub wielu anten wewnętrznych w celu retransmisji sygnału w innym kierunku; w tym przypadku wzduż tunelu.

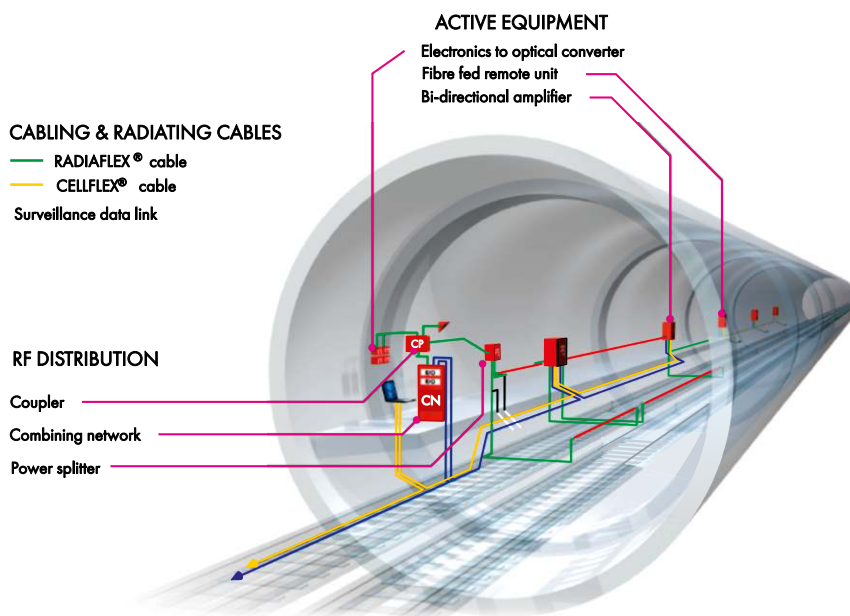
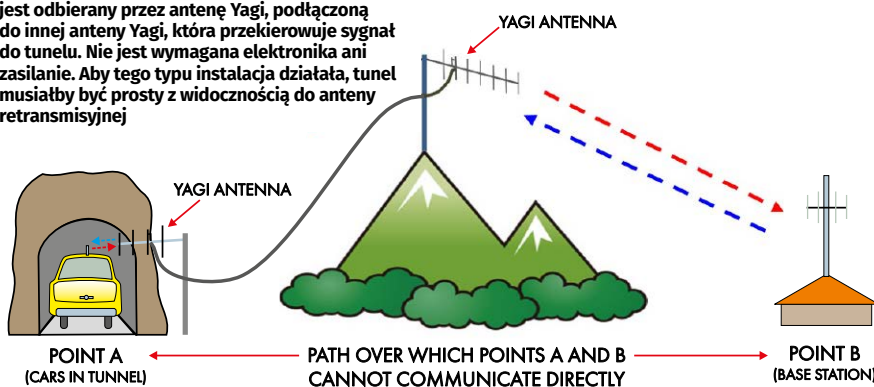
W przypadku fal o krótszej długości, takich jak FM lub DAB, stosuje się anteny Yagi rozmieszczone wzduż tunelu w określonych odstępach. W przypadku dłuższych fal może to być nieuszczelny feeder.

Retransmisja pasywna jest odpowiednia tylko dla prostych tuneli w zasięgu wzroku do anteny (anten) retransmisyjnej. Jeżeli jest więcej niż jedna antena, wymagane są rozgałęźniki sygnału. Bez nich sygnał byłby nadmiernie osłabiony. Za pomocą rozgałęźników można go wzmocnić tak, jak robi się to w bloku mieszkalnym z jedną anteną i wieloma gniazdkami.



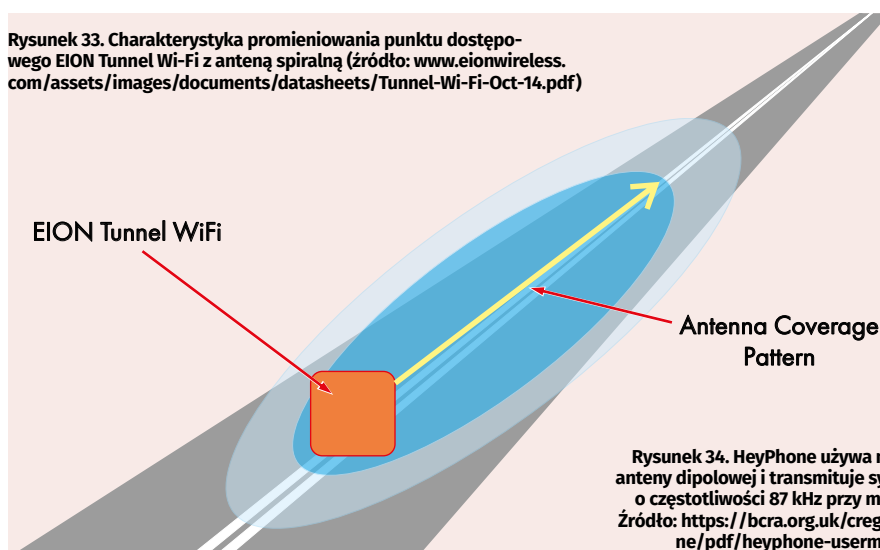
Rysunek 30. Widoki przekroju różnych kabli promieniujących oferowanych przez Exlanta (<http://exlanta.com>)

Rysunek 31. Pasywny repeater używany w niektórych instalacjach tunelowych. Sygnał zewnętrzny jest odbierany przez antenę Yagi, podłączoną do innej anteny Yagi, która przekierowuje sygnał do tunelu. Nie jest wymagana elektronika ani zasilanie. Aby tego typu instalacja działała, tunel musiałby być prosty z widocznością do anteny retransmisyjnej



Rysunek 32. Przykład tunelu z promieniującym feederem i urządzeniami pomocniczymi (źródło: <https://alliancecorporation.ca/manufacture/rfs-radio-frequency-systems/>)

Rysunek 33. Charakterystyka promieniowania punktu dostępowego EION Tunnel Wi-Fi z anteną spiralną (źródło: www.eionwireless.com/assets/images/documents/datasheets/Tunnel-Wi-Fi-Oct-14.pdf)



Rysunek 34. HeyPhone używa naziemnej anteny dipolowej i transmituje sygnał USB o częstotliwości 87 kHz przy mocy ~10 W (źródło: <https://bcra.org.uk/creg/heyphone/pdf/heyphone-usermanual.pdf>)

Częściej stosowane jest retransmitowanie aktywne. Odbiorniki odbierają i dekodują sygnały za pomocą anten na zewnątrz tunelu. Następnie podają zdekodowane sygnały (np. audio) przez elektronikę audio do wzmacniaczy i nadajników, które ponownie emitują je na oryginalnych częstotliwościach za pomocą anten w całym tunelu. Typowa konfiguracja została przedstawiona na **rysunku 32**.

W zależności od jednostki (jednostek) retransmisji i konfiguracji, możliwe jest nadawanie AM i FM, DAB, przywoływanie VHF/UHF/800 MHz i dwukierunkowy dostęp radiowy w tunelu.

W przypadku telefonów komórkowych zwykle łatwiej jest zainstalować małe maszty komórkowe w tunelu połączone z siecią dosyłową, niż próbować zachować dwukierunkową komunikację między telefonami w tunelu a wieżami poza nim. Telefony są „przekazywane” między masztami wewnątrz i na zewnątrz tunelu, tak jak podczas przemieszczania się między standardowymi masztami.

Oprócz tuneli, takie systemy mogą być stosowane w innych podziemnych strukturach, takich jak parkingi, kopalnie i wewnątrz budynków, gdzie odbiór może być słaby z powodu metalowej folii na oknach lub z innych powodów.

Wi-Fi w tunelach

Wi-Fi można instalować w tunelach i innych przestrzeniach podziemnych. Najskuteczniejsze jest w tym przypadku użycie punktów dostępowych Wi-Fi ze specjalnie zaprojektowanymi antenami spiralnymi, które mają rozszerzoną charakterystykę promieniowania w kierunku tunelu, zamiast tradycyjnej charakterystyki dookoła.

Specjalnie skonstruowane punkty dostępowe do tego zastosowania są dostępne w firmie EION Inc (**rysunek 33**).

Komunikacja w jaskini

W przypadku radia dostępnego w jaskiniach, sygnał radiowy jest transmitowany przez ziemię (Through-the-Earth, TtE) lub w bezpośrednim zasięgu wzroku (Line of Sight, LoS) z przekaźnikami lub wykorzystaniem wielu coraz słabszych odbić. Zwykle radiotelefony mogą być używane w jaskiniach na krótkich dystansach w zasięgu wzroku, ale rzadko są odpowiednie, ponieważ w jaskiniach niezbyt często spotykamy długie i proste korytarze.

Radio może być również transmitowane i odbierane za pośrednictwem promieniujących feederów, ale oczywiście wiąże się to z prowadzeniem kabli w jaskiniach, podobnie jak w przypadku konwencjonalnej telefonii 1- lub dwuprzewodowej (z uziemionym obwodem powrotnym).

Molefone (TtE)

Molefone (**rysunek 35**) był radiotelefonem opracowanym do ratownictwa jaskiniowego i ogólnego użytku przez Boba Mackina z Lancaster University w latach 70. XX w i był szeroko stosowany w latach 80. i późniejszych. Wykorzystywał wieloobrotową antenę pętlową o średnicy około 41 cm



Rysunek 35. Obsługa Molefone w jaskiniach Matienzo, Hiszpania. Zwróć uwagę na antenę pętlową wykonaną z komputerowego przewodu taśmowego. (źródło: <http://matienzocaves.org.uk/ugpics/2366-2007e-molepe.htm>)

Korzystanie z łączności naziemnej 87 kHz w Australii

Pomimo tego, że łączność 87 kHz została ustanowiona jako międzynarodowy standard łączności ratunkowej w jaskiniach, najwyraźniej nie została zatwierdzona przez ACMA (Australian Communications and Media Authority) i jej używanie w Australii w tym celu byłoby nielegalne.

Dlatego Speleונים produkuje tylko przewodowe urządzenia Michiephone, a nie urządzenia bezprzewodowe. Ponieważ jest to międzynarodowy standard, a ryzyko zakłóceń jest niskie lub nieistniejące, ACMA powinna ponownie rozważyć swój sprzeciw wobec takiego użycia i zrobić wyjątek, przynajmniej w celach ratowniczych lub eksploracyjnych.

osiągając zasięg około 150...200 m przez skałę przy mocy 10 W.

Urządzenie pracowało z częstotliwością 87 kHz, wykorzystując modulację Upper Side Band (USB, górna wstęga boczna). Schematy układów nie są dostępne. Częstotliwość 87 kHz stała się standardem dla innych systemów komunikacji jaskiniowej, takich jak HeyPhone i System Nicola (oba będą wspomniane dalej), w celu zachowania kompatybilności. Nie są one obecnie produkowane ze względu na niedostępność niektórych elementów, i wynikający z tego brak możliwości napraw uszkodzonych jednostek.

HeyPhone (TtE)

HeyPhone (<https://bcra.org.uk/creg/heyphone/> i **rysunek 34**) został zaprojektowany przez Johna Hey'a. Stanowi swego rodzaju zamiennik Molefone'a. W 1999 roku Brytyjska Rada Ratownictwa Jaskiniowego (BCRC) zainicjowała projekt we współpracy z Johnem Heyem.

W przeciwieństwie do systemu Molefone, w HeyPhone jako główną antenę stosuje się

dipol ziemny (ground dipole), a nie antenę pętlową, choć urządzenie może również współpracować z pętlą.

Dipol ziemny składa się z dwóch uziemionych elektrod oddalonych od siebie o 25...100 m. Anteny tego typu mają większą zdolność penetracji gruntu niż anteny pętlowe używane w systemie Molefone.

Podobnie jak Molefone, HeyPhone korzystał z łączności 87 kHz z modulacją USB przy mocy około 10 W, i jak się okazało oba radia były kompatybilne.

Projekt ten nie jest już aktywny ani wspierany, ale jeśli jesteś eksperymentatorem, możesz zdobyć schematy układów i inne dokumenty pozwalające na zbudowanie tych urządzeń we własnym zakresie:

<https://bcra.org.uk/creg/heyphone/documentation.html>

Instrukcję obsługi urządzenia można również znaleźć pod adresem:

<https://bcra.org.uk/creg/heyphone/pdf/heyphone-usermanual.pdf>

Mówi się, że telefony HeyPhones były używane podczas akcji ratunkowej w jaskini Tham Luang (Tajlandia, czerwiec-lipiec 2018 r.), wraz z urządzeniami radiowymi Maxtech mesh.

System Nicola (TtE)

Po śmierci Nicolii Dollimore w wypadku w jaskini w 1996 roku, zebrano fundusze na stworzenie „najlepszego radia jaskiniowego”. Był to wspólny wysiłek Francuzów, Szwajcarów i Brytyjczyków, oparty na telefonie HeyPhone. Radio Mk2 zostało wprowadzone do użytku w 1998 roku i jest systemem używanym w całej Francji. Cyfrowa wersja Mk3 została opracowana na początku XXI wieku, podczas gdy Mk4 jest obecnie w fazie rozwoju (**rysunek 36**).

Radio Mk2 działa na częstotliwości około 87 kHz z mocą 3 W i modulacją USB. Ziemna antena dipolowa ma dwie

elektrody w ziemi oddalone od siebie o około 40...80 m. Zasięg transmisji przez skały wynosi około 500...1200 m, w zależności od warunków.

Niestety nie ma zbyt wielu informacji na temat tego radia. System Nicola nie ma strony domowej w Internecie, ma jednak profil na Facebooku, www.facebook.com/AssociationNicola/.

Cave-Link (TtE)

Cave-Link (www.cavelink.com/cl3x_neu/index.php/en/) to system łączności w jaskiniach, który do przesyłania danych tekstowych wykorzystuje fale bardzo niskiej częstotliwości (VLF), zamiast transmisji głosowej. Łączność jest utrzymywana do głębokości 1300 m, a nawet większej. Naziemna część systemu, którą producent nazywa „modemem prądu ziemnego” (**rysunek 37**), może być również podłączona do systemu telefonii komórkowej w celu przesyłania wiadomości SMS.

Niektóre europejskie organizacje zajmujące się ratownictwem jaskiniowym używają Cave-Link i jest on również używany do rejestrowania danych z czujników umieszczonych wewnątrz jaskiń (np. rejestrujących przepływ wody, głębokość wody, temperaturę, poziom CO₂, pH, ciśnienie itp.) Działa w zakresie od 20 kHz do 140 kHz z modulacją 4PSK (QPSK – kwadraturowe kluczowanie przesunięcia fazowego) i protokołu korekcji błędów ARQ (automatyczne żądanie powtórzenia).

Anteny na powierzchni i w jaskini składają się z dwóch metalowych płytek, z których



Rysunek 36. Dwa radiotelefony Nicola Mk4 (ułożone jeden na drugim), będące obecnie w fazie opracowywania. Źródło: Strona System Nicola na Facebooku



Rysunek 37. Terminal Cave-link do wysyłania danych tekstowych przez VLF przez ziemię. Źródło: https://expo.survex.com/expofiles/documents/hardware/Cavelink2.13_en_2014-3.pdf



Rysunek 38. Jill Rowling ze Speleונים używająca jednoprzewodowego telefonu Michiephone. Źródło: www.speleונים.com.au/business/ (powielone za zgodą)



Rysunek 39. Radiotelefon Entel/Maxtech MaxMesh SDR używany podczas akcji ratunkowej w tajlandzkiej jaskini. Źródło: www.entelkorea.com/assets/resources/brochures/HT786-MaxMesh.pdf

każda jest podłączona do jednego przewodu linii zasilającej z nadajnika lub odbiornika, zakopanych w ziemi i połączonych kablem. Tworzy to antenę znaną jako dipol ziemny (rysunek 16 – EdW 11/2025). Odległość między płytami odpowiada pionowej głębokości transmisji, która jest około dziesięć razy większa od poziomej odległości między nimi.

HF Radio (TtE)

Radio HF ma pewne możliwości penetracji ziemi, głównie przez suche skały w suchych regionach. Przeprowadzono pewne eksperymenty na częstotliwości 1,8750 MHz przy użyciu transceivera Elecraft KX3 (<https://youtu.be/WTnrDwIPKrl>).

Inne zgłoszone eksperymenty to:

1. Paul Jorgensen, KE7HR, z transceiverem FT817ND pracującym na częstotliwości 3,9 MHz i wykorzystującym jednowstęgową modulację AM (Single Side Band, SSB) przy mocy 5 W, zademonstrował komunikację głosową do głębokości 238 m w Carlsbad Cavern, Nowy Meksyk, USA.
2. W 2015 r. brytyjska grupa Cave Radio and Electronics Group nawiązała łączność na głębokości 100 m przy skośnej odległości 692 m, wykorzystując moc 20 W na częstotliwości 7,135 MHz SSB za pomocą transceivera IC-706.
3. W numerze 97 czasopisma BCRA Cave Radio and Electronics Group Journal (marzec 2017) opisano odbiór sygnałów WSPR (Weak Signal Propagation Reporter)

na częstotliwości 7 MHz, zarejestrowanych 100 m pod ziemią w Wielkiej Brytanii – pochodzących z dziewięciu krajów.

Telefon dwuprzewodowy

W przeszłości grotołazi korzystali z wojskowych telefonów polowych typu dwuprzewodowego, pochodzących z nadwyżek sprzętu. Z czasem zostały one w większości zastąpione telefonami jednoprzewodowymi, znanymi jako Michiephone'y.

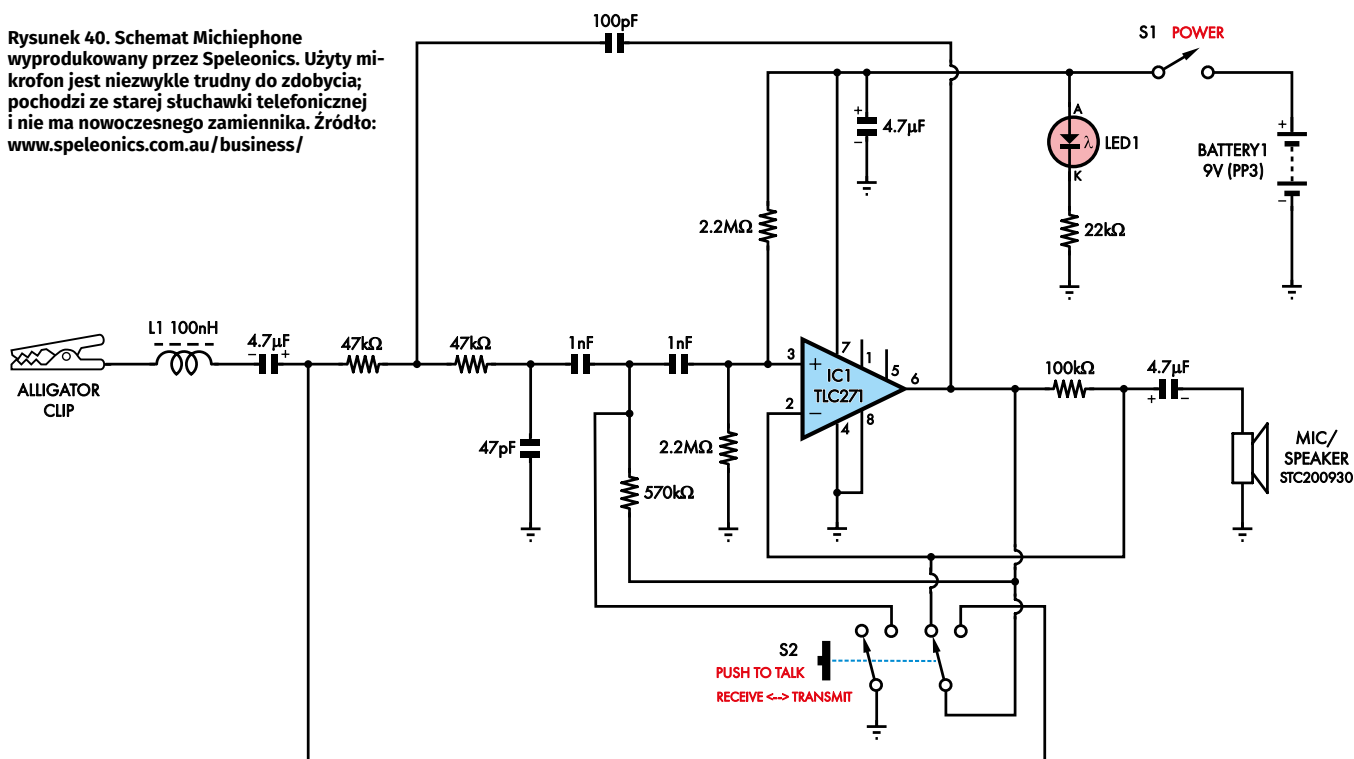
Telefon jednoprzewodowy

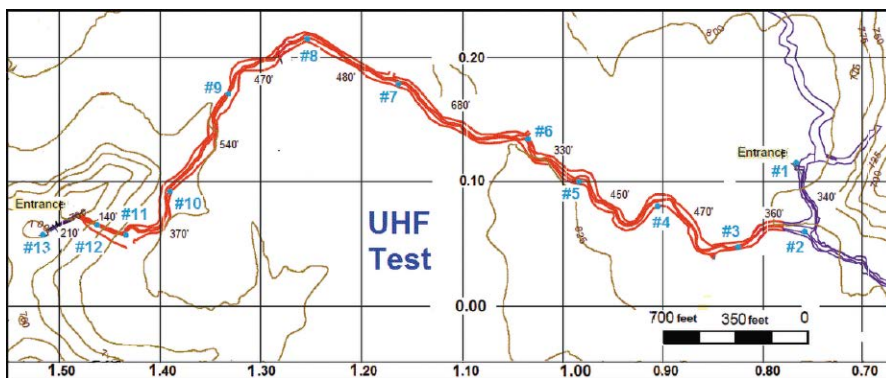
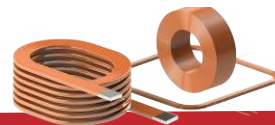
Jednoprzewodowy telefon jaskiniowy, znany również jako Michiephone, używa tylko jednego przewodu zamiast dwóch używanych w klasycznych telefonach analogowych. Obwód powrotny przebiega przez ziemię (rysunek 38). Dzięki jednemu przewodowi szpula waży mniej i łatwiej jest ją rozłożyć. Działają przez wiele dni na bateriach (www.spleonics.com.au/business/michiephones/).

Na rysunku 40 przedstawiono klasyczny schemat dla takiego typowego urządzenia, zaprojektowanego przez Australijczyka Neville'a Michie w latach 70. XX w. Jest to bardzo prosta konstrukcja, której głównym elementem jest wzmacniacz operacyjny.

Speleonics to australijski producent tych urządzeń, choć wydaje się, że obecnie nie produkuje już jakiegokolwiek sprzętu tego typu. Główną różnicą między urządzeniami tej firmy a oryginalnym projektem jest zastosowanie

Rysunek 40. Schemat Michiephone wyprodukowany przez Speleonics. Użyty mikrofon jest niezwykle trudny do zdobycia; pochodzi ze starej słuchawki telefonicznej i nie ma nowoczesnego zamiennika. Źródło: www.spleonics.com.au/business/





Rysunek 41. Wyniki testu radiowego APRS UHF w jaskini Mammoth, USA, pokazujące lokalizację radiotelefonów (ponumerowanych na niebiesko) w systemie jaskini, ścieżkę komunikacji na czerwono i odległości w stopach (700 stóp = 213 m). Źródło: www.aprs.org/cave-link.html

filtru usuwającego przydźwięk sieciowy 50 Hz, którego nie uwzględniono w oryginale.

VHF i UHF Mesh Radio (LoS)

Podczas akcji ratunkowej tajskiej młodzieżowej drużyny piłkarskiej uwięzionej w jaskini w 2018 roku, ratownicy nawiązali łączność za pomocą sprzętu przywiezionego z Izraela, wyprodukowanego przez Maxtech Networks (<https://max-mesh.com/>). Sprzęt mieścił się w jednej walizce i składał się z podobnych do walkie-talkie radiotelefonów definiowanych programowo (SDR – software defined radio).

Przywieziono 17 lub 19 radiotelefonów (w zależności od publikowanego raportu), ale tylko 11 zostało ostatecznie wykorzystanych do ustanowienia połączenia komunikacyjnego 4 km w głąb jaskini poprzez utworzenie sieci mesh. Oprogramowanie mesh opracowała firma Maxtech, a platforma radiowa jest autorstwa brytyjskiej firmy Entel (rysunek 39). Radiotelefony działają w zakresach VHF i UHF (225 MHz...470 MHz).

Bez sieci mesh komunikacja w jaskini między dwoma radiotelefonami na tych częstotliwościach odbywałaby się w zasięgu widoczności lub poprzez ograniczoną liczbę odbić. Tymczasem w sieci mesh każde radio może działać jako stacja przekaźnikowa dla następnego.

Poszczególne radiotelefony nadal komunikują się ze sobą w zasięgu widoczności lub wykorzystując odbicia. Pomimo tego, radiotelefon na początku sieci radiotelefonów (np. przy wejściu do tunelu) może płynnie komunikować się z radiotelefonem na drugim końcu. Każdy kolejny radiotelefon w sieci mesh przekazuje wiadomość do następnego, nawet jeśli nie ma bezpośredniego połączenia między komunikującymi się radiotelefonami (pierwszym i ostatnim).

Podczas akcji ratunkowej w jaskini nawiązano łączność audio i wideo przy użyciu

11 radiotelefonów (siliconchip.au/link/abir), z których każdy miał baterię wystarczającą na 10 godzin pracy. W niektórych miejscach jedyna droga prowadziła przez wodę, więc położono podwodne kable do transmisji danych, aby połączyć pary radiotelefonów definiowanych programowo.

Sieć mesh utworzona przez radiotelefony była samoformująca się, samoroutingowa, samonaprawiająca się i nie wymagała żadnej innej infrastruktury. Była to „mobilna sieć ad hoc” (MANET) używająca schematu kontroli dostępu do mediów (MAC) z podziałem czasu (TDMA) z innowacyjnym algorytmem routingu.

Należy pamiętać, że radiotelefony te nie zostały zaprojektowane wyłącznie do ratownictwa jaskiniowego, byłyby pomocne w każdym nieprzyjawnym środowisku, na przykład takim jak zawalone budynki po trzęsieniu ziemi.

System może również tworzyć bramy do telefonów 3G i 4G, radiotelefonów analogowych i innych sieci.

Więcej informacji można znaleźć w materiale wideo na stronie siliconchip.au/link/abis oraz w materiale zatytułowanym „Maxtech networks video over radio” w serwisie

YouTube pod adresem <https://youtu.be/C2q9L8iAOyA>.

UHF Mesh Radio (LoS)

Wideo „Underground & Through-The-Earth Communications” na stronie <https://youtu.be/WTnrDwIPKrI> opisuje eksperymentalną sieć mesh, zbudowaną z mostów bezprzewodowych Ubiquity M2 (2,4 GHz) i M900 (900 MHz) w technologii MIMO (Multiple Input Multiple Output – wiele wejść i wiele wyjść), z wykorzystaniem niestandardowego oprogramowania układowego dostępnego na stronie <http://hsmm-mesh.org/>.

W rezultacie powstała sieć mesh IEEE 802.148 do komunikacji w jaskini. Transmisja głosu była realizowana przez łącze HF, a następnie cyfrowo przesyłana przez sieć mesh w głąb jaskini.

APRS (LoS)

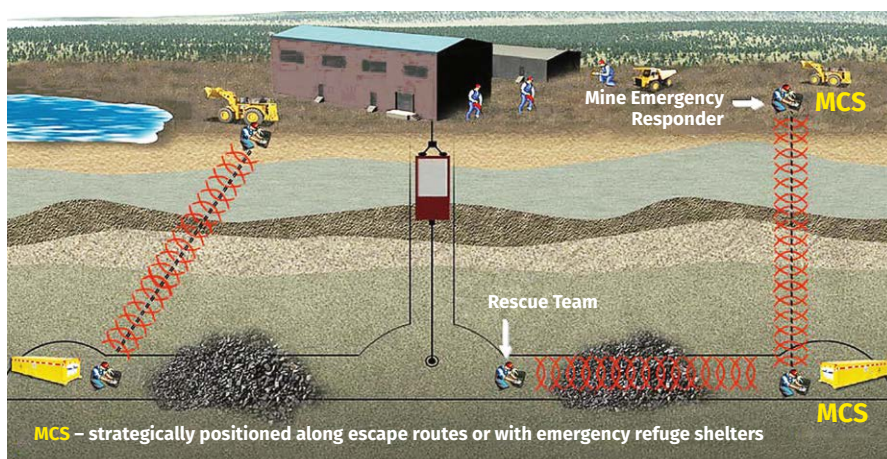
W jaskiniach możliwe jest również korzystanie z radiotelefonów APRS (Automatic Packet Reporting System). APRS jest amatorskim protokołem radiowym, więc nie jest obecnie dostępny do ogólnego użytku w jaskiniach, ale operatorzy radiowi, którzy są również grotołazami, badają jego zastosowanie.

Podobnie jak w przypadku omówionych już sieci radiowych typu mesh, częstotliwości VHF i UHF są użyteczne tylko wtedy, gdy zachowana jest widoczność anten lub poprzez ograniczone odbicia. W przeciwieństwie do radiotelefonów Maxtech, za pomocą APRS można przysyłać tylko dane. Podobnie jak w Maxtech, poszczególne radiotelefony mogą działać jako stacje przekaźnikowe („digipeaters”) dla kilku radiotelefonów w łańcuchu.

Eksperyment z radiami APRS został przeprowadzony w dniach 2–3 marca 2013 r. w Mammoth Cave, Kentucky, USA, najdłuższym znanym systemie jaskiń na świecie (rysunek 41). Stwierdzono, że w przypadku



Rysunek 42. Urządzenie MagneLink obok górnika. Źródło: www.teslasociety.ch/info/magnetlink/2.pdf



Rysunek 43. Magnetyczny System Łączności MagneLink (MCS) w zastosowaniach ratowniczych
 Źródło: www.teslasociety.ch/info/magnetlink/2.pdf

radiotelefonów VHF średnia długość przekroku między sąsiednimi węzłami sieci wynosiła 119 m, a maksymalna 162 m. W przypadku UHF średnia długość przekroku wynosiła 134 m, a maksymalna 207 m.

Stwierdzono ponadto, że sygnały mogły pokonywać zakręt 90° w korytarzu jaskini bez znaczącej różnicy w zasięgu w porównaniu z odcinkiem prostym. Zwiększenie mocy do 50 W nie robiło większej różnicy w porównaniu do 5 W lub mniejszej. Nawet moc 0,5 W była satysfakcjonująca. Korytarze jaskiń były dość duże, o szerokości od 9 m do 15 m i wysokości od 3 m do 6 m.

Promieniujący feeder w jaskiniach (przewodowa)

Podobnie jak w kopalniach i tunelach, komunikacja radiowa może być prowadzona również w oparciu o promieniujący feeder, pod warunkiem pozostawania w zasięgu jego widoczności. Takie rozwiązanie jest zwykle stosowane w jaskiniach turystycznych, jednak feedery były zastosowane eksperymentalnie w innych jaskiniach.

Ze względu na wysoki koszt specjalnie wykonanego kabla promieniującego, celem

eksperymentu opisanego na stronie siliconchip.au/link/abit było znalezienie taniego zamiennika drogiego kabla. Zauważono, że tani domowy kabel satelitalny był wystarczająco nieszczelny (w sposób niezamierzony), że nadawał się do tego typu zastosowań.

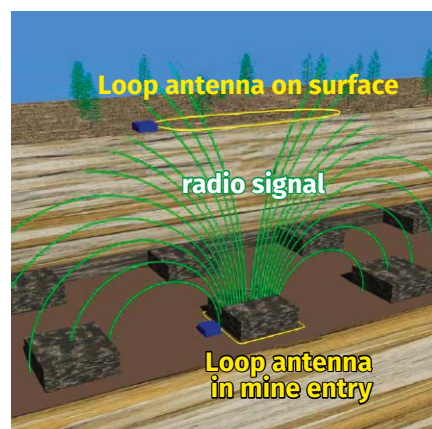
Komunikacja w kopalniach

Bezprzewodowa łączność radiowa w kopalniach może odbywać się za pośrednictwem ziemi, promieniujących feederów lub przewodowych systemów telefonicznych.

MagneLink (bezprzewodowy, przez ziemię)

Magnetic Communication System (MCS) firmy Lockheed Martin (**rysunek 42**) to system łączności awaryjnej używany w kopalniach do komunikacji z uwięzionymi górnikami i zespołami ratowniczymi zapewniający dwukierunkową komunikację głosową i tekstową.

Uwięzieni górnicy z dostępem do MagneLink mogą go aktywować, aby wysłać sygnał nawigacyjny pomagający zespołom ratunkowym znaleźć uwięzionych górników. Może być używany pionowo między ziemią a kopalnią lub poziomo wzdłuż chodnika kopalni (**rysunek 43**).



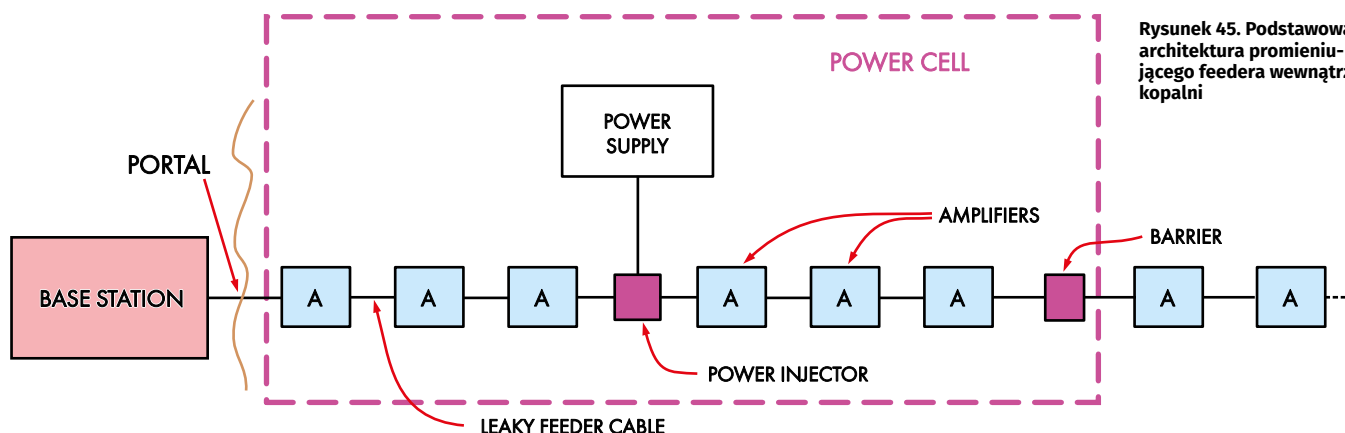
Rysunek 44. Jak działa MagneLink i inne systemy łączności naziemnej używające anteny pętlowe.
 Źródło: www.cdc.gov/niosh/mining/UserFiles/Works/pdfs/2013-105.pdf

W systemie są zastosowane anteny pętlowe, więc komunikacja odbywa się za pośrednictwem składowej pola magnetycznego sygnału radiowego, a nie składowej pola elektrycznego (**rysunek 44**). Pozwala to na stosowanie znacznie mniejszych anten niż alternatywny typ, dipol naziemny, który może mieć dziesiątki lub setki metrów długości.

Część systemu zainstalowana w kopalni jest przeznaczona do umieszczenia w wyznaczonych miejscach, takich jak „strefy schronienia”. Antena pętlowa jest owinięta poziomo wokół filarów lub innych elementów konstrukcyjnych kopalni, na przykład wokół filara nośnego (niewydrążonego obszaru podtrzymującego strop). W testach system MagneLink osiągnął głębokość komunikacji 457 m dla głosu i 610 m dla tekstu.

Promieniujące feedery w kopalniach (przewodowe + bezprzewodowe)

Promieniujące feedery (nieszczelne linie zasilania anten) działają w kopalniach podobnie jak w innych miejscach, takich jak tunele. Są one przeznaczone do komunikacji



Rysunek 45. Podstawowa architektura promieniującego feedera wewnątrz kopalni

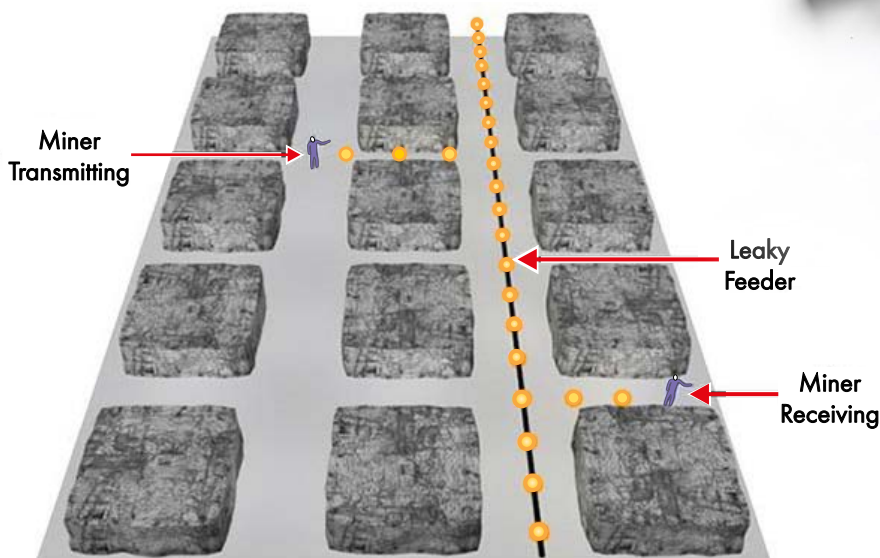
dwukierunkowej przy użyciu urządzeń przenośnych. Na powierzchni lub w innym centrum dowodzenia stacja bazowa jest odpowiedzialna za wysyłanie i odbieranie transmisji (rysunki 45 i 46).

Co około 350...500 m znajdują się również wzmacniacze, których zasilanie może być przenoszone przez sam feeder, zwykle przy napięciu 12 V. Używane częstotliwości to najczęściej pasma VHF i UHF. Podstawowym elementem składowym promieniującego feedera w kopalni jest ogniwo zasilające, z jednym ogniwem na sekcję kopalni, przy czym wiele z nich może być połączonych ze sobą. Zwielokrotniona liczba ogniw zapewnia redundancję w przypadku uszkodzenia jednej sekcji (rysunek 48).

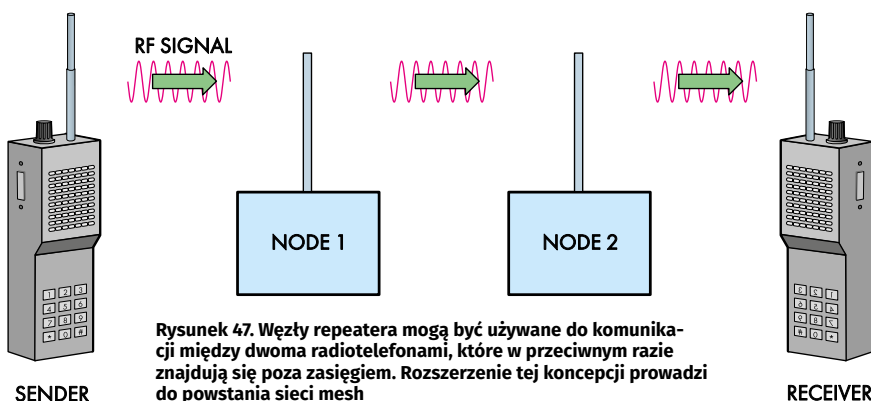
Węzły/sieć mesh (bezprowadowa)

Innym sposobem wykorzystania ręcznych radiotelefonów w kopalni jest użycie ich jako części systemu opartego na węzłach. Zasięg radiotelefonów pod ziemią jest wprawdzie ogólnie ograniczony, ale małe stacje przekaznikowe lub węzły mogą znacznie go zwiększyć (rysunek 47). Węzły te i używane z nimi radiotelefony są sterowane za pomocą systemu mikroprocesorowego.

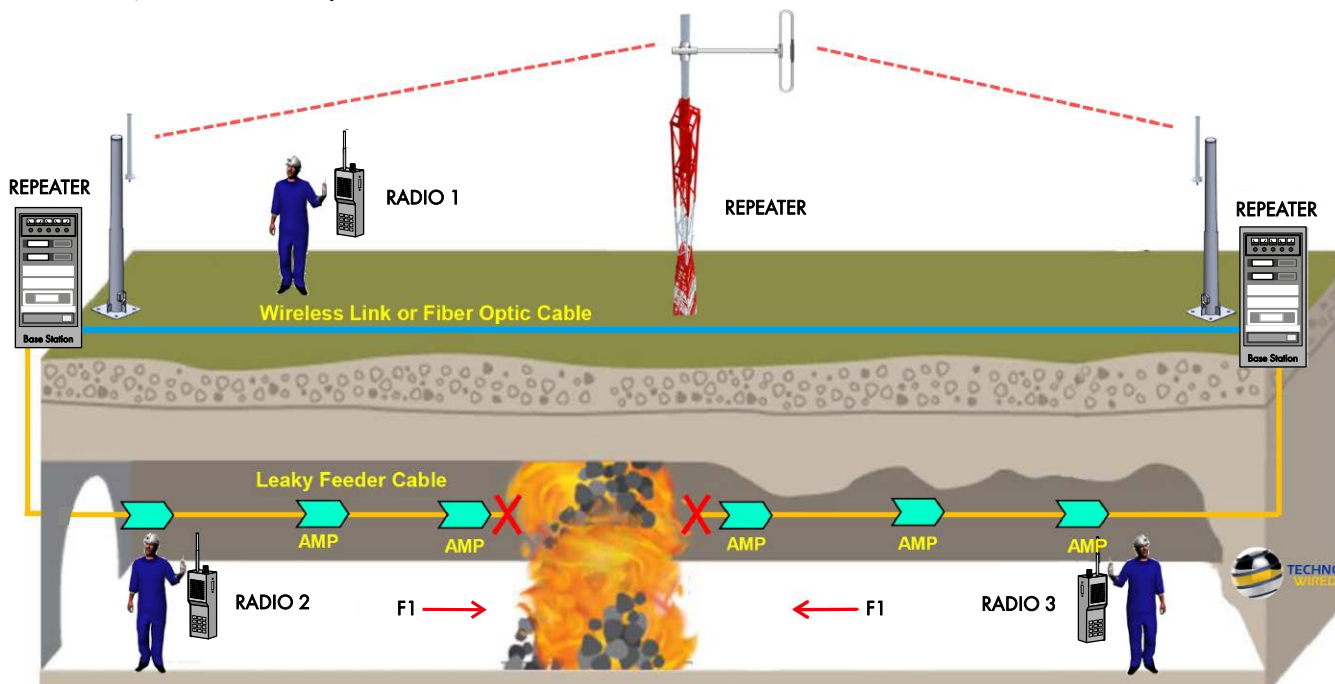
Jak wspomniano wcześniej, system tworzy sieć typu mesh, gdy używanych jest wiele węzłów. Sieć taka kieruje sygnały między węzłami, gdy uzna to za stosowne (rysunek 49). Jeśli jeden z węzłów nie działa, ustanawiana jest ścieżka alternatywna.



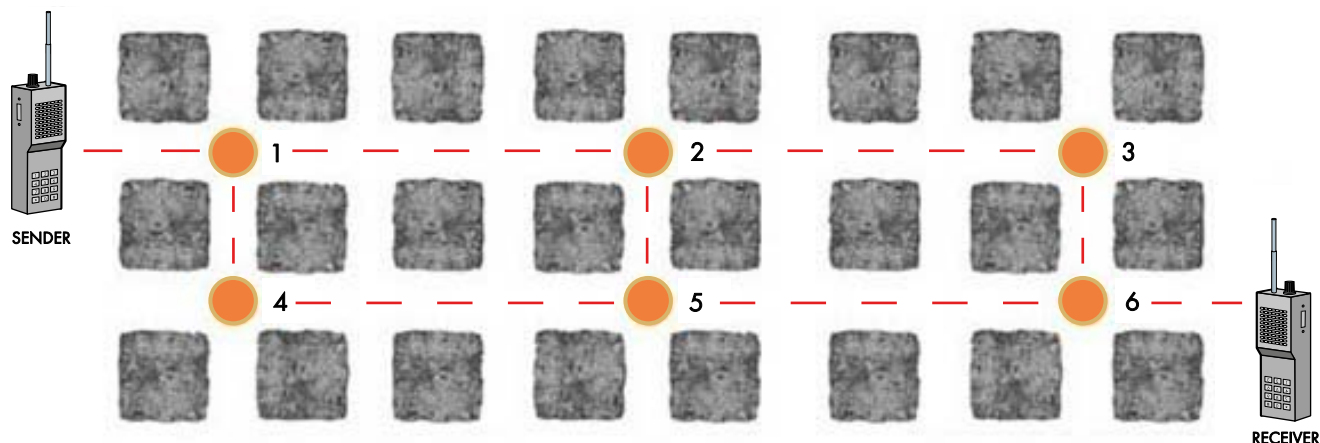
Rysunek 46. Ułożenie odcinka promieniującego feedera w kopalni. Kropki pokazują sygnał między dwoma górnikami



Rysunek 47. Węzły repeatera mogą być używane do komunikacji między dwoma radiotelefonami, które w przeciwnym razie znajdują się poza zasięgiem. Rozszerzenie tej koncepcji prowadzi do powstania sieci mesh



Rysunek 48. Przykład tego, jak promieniujący feeder z nadmiarowością nad i pod ziemią, może nadal działać po katastrofie
Źródło: www.technowired.net/wp-content/uploads/2017/02/4.-Sistema-MCA1000-Digital-en.pdf



Rysunek 49. Do komunikacji między dwoma radiotelefonami w różnych lokalizacjach może być używanych wiele węzłów repeaterów, które w przeciwnym razie byłyby poza zasięgiem. Razem węzły te tworzą sieć mesh

System średniej częstotliwości (przewodowy + bezprzewodowy)

Fale radiowe średniej częstotliwości w zamkniętych przestrzeniach podziemnych będą spręgać się z dowolnymi istniejącymi przewodami, takimi jak linie energetyczne, kable danych lub promieniujące feedery. W przeciwieństwie do radia VHF i UHF, średnie częstotliwości mogą wykorzystywać dowolny istniejący przewód. Jeśli więc odpowiedni przewód znajduje się w pobliżu, radiotelefony MF mogą być używane na dole w kopalni na większą odległość i do komunikacji nie jest wymagany żaden przemiennik.

Wadą jest to, że ręczne radiotelefony MF są znacznie większe niż radiotelefony VHF i UHF. Rozwiązaniem jest użycie

konwerterów VHF/UHF na MF. Umożliwia to użycie małego ręcznego radia w zasięgu konwertera, który następnie retransmituje sygnał do pasma MF, sprzęgając go z pobliskimi przewodami. Na drugim końcu łączy sygnał MF jest konwertowany do VHF/UHF, umożliwiając innemu górnikowi odbiór transmisji.

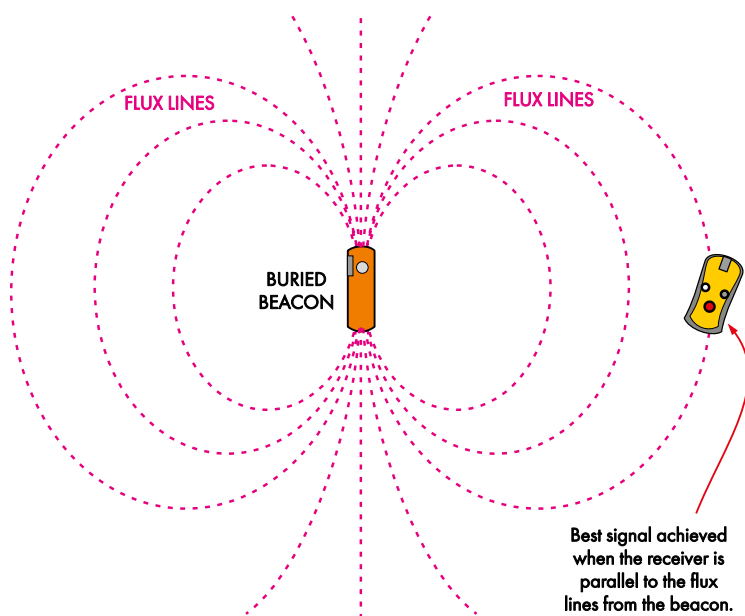
Sygnalizatory lawinowe

Lawiny powstają, gdy niestabilna warstwa śniegu odrywa się i stacza z góry, grzebiąc na swojej drodze nieszczęsnych narciarzy lub wędrowców. Występują powszechnie w Europie i Ameryce Północnej. Osoby przebywające w strefach zagrożonych lawinami często noszą ze sobą transceiver nazwany lokalizatorem lawinowym.

Lawinowe lokalizatory awaryjne zostały wynalezione w 1968 roku, a pierwsze egzemplarze komercyjne zostały po raz pierwszy sprzedane w 1971 roku. Działały na częstotliwości 2,275 kHz (ULF). W 1986 roku jako standardową częstotliwość przyjęto 457 kHz (MF).

Częstotliwość 457 kHz (656 m) została przyjęta, ponieważ nie jest podatna na silne tłumienie przez śnieg, skały, drzewa, gruz lub ludzi i jest mniej podatna na problemy wynikające z odbić wielościeżkowych w porównaniu ze znacznie niższą częstotliwością 2,275 kHz.

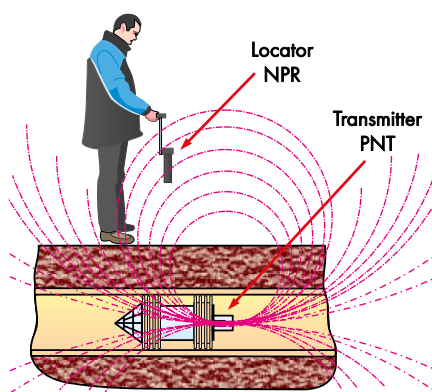
Z konieczności długość anteny stanowi jedynie niewielką część długości fali. Sprawia to, że transmisja charakteryzuje się małą wydajnością. Efektywną długość elektryczną można jednak zwiększyć, stosując antenę o wielu zwojach z rdzeniem ferrytowym.



Rysunek 50. Kształt emitowanego sygnału wpływa na schemat poszukiwań podczas akcji ratunkowych pod lawiną. Aby szybko zlokalizować osoby zasypane pod śniegiem za pomocą nadajników, wymagana jest praktyka. Oryginalne źródło: <https://youtu.be/tXpEUBDzbu0>



Rysunek 51. Sygnalizator lawinowy Mammut Barryvox S do wyszukiwania zasypanych ofiar. Posiada funkcję wspomagającą wyszukiwanie, W-Link i trzy anteny. Źródło: <https://varuste.net/p77030/mammut-barryvox-s>



Rysunek 52. Tzw. „pig” może być zlokalizowany przez stal, glebę i beton, odbierając nadawany przez niego sygnał 22 Hz

Podczas użytkowania, gdy każdy członek grupy wychodzi na obszar zagrożony lawiną, włącza swój nadajnik, który emituje sygnał dźwiękowy przez radio raz na sekundę.

Jeśli któryś z członków grupy zostanie zasypany przez lawinę, pozostali członkowie przełączą swoje urządzenia z nadawania na odbiór, co umożliwi odbieranie sygnałów od zaspanych członków wyprawy.

Zasięg nadajników lokalizatorów wynosi 40...80 m. Ze względu na kształt emitowanego sygnału, istnieje specyficzna technika znajdowania osoby zakopanej w śniegu. Do jej opanowania wymagana jest pewna praktyka. W takich przypadkach zawsze najważniejszy jest czas. Na **rysunku 50** pokazana jest charakterystyka promieniowania, a w serwisie YouTube można znaleźć różne filmy objaśniające wymaganą technikę poszukiwań.

W nowoczesnych lokalizatorach lawinowych stosuje się cyfrowe tryby transmisji. Niektóre z nich, oprócz standardowego sygnału 457 kHz, wykorzystują dodatkowy kanał W-Link, pracujący na częstotliwości 869,8 MHz lub 916...926 MHz, zależnie od regionu. W-Link przesyła dodatkowe informacje, między innymi identyfikator urządzenia i pozwala na ignorowanie sygnałów osób już uratowanych.

Nowoczesne lokalizatory (**rysunek 51**) używają również dwie lub trzy anteny 457 kHz w trybie odbioru, aby odbiornik był bardziej czuły w niektórych kierunkach, w zależności od wzajemnego ustawienia nadajnika i odbiornika.

Jeśli nosisz taki lokalizator, trzymaj go pod odzieżą wierzchnią, aby zapobiec zamarznięciu baterii i oderwaniu urządzenia w przypadku zejścia lawiny.

Wiele osób nie zdaje sobie sprawy, że lawiny mogą występować w Australii. Choć są one rzadkie i nie tak duże jak za oceanem, zdarzają się w niektórych regionach alpejskich, choć zazwyczaj nie na obszarach uczęszczanych przez narciarzy i nie są tak niebezpieczne jak te, które występują w obu Amerykach, Azji czy Europie. Australijska organizacja Mountain Safety Collective (<https://mountainsafetycollective.org/>) prowadzi szkolenia i dysponuje zespołami ratowniczymi na wypadek wystąpienia lawin.

Komunikacja pigging (rurociągi)

Pigging polega na wprowadzeniu „piga” (swego rodzaju tłoka) do rurociągu, w celu jego wyczyszczenia lub przeprowadzenia inspekcji (**rysunki 52 i 53**). Tłok jest urządzeniem, które szczelnie wypełnia wewnętrzną średnicę rury i jest popychany przez ciśnienie płynu lub gazu za nim. Niektóre z nich są wyposażone w elektronikę, która przekazuje ich pozycję lub inne dane do otoczenia zewnętrznego.

Standardowa częstotliwość komunikacji między pigami wynosi 22 Hz. Takie sygnały przenikają przez metal rurociągu i glebę lub żelbeton nad nim.

Nieszczelne feedery w samolotach

Samoloty wprawdzie nie znajdują się pod ziemią (jestem pewien, że pasażerowie odetchnęli z ulgą czytając to!), ale niektóre

problemy spotykane w łączności w tunelach dotyczą również odbioru radiowego na ich pokładzie. Systemy nieszczelnych feederów zostały opracowane przez firmy takie jak W. L. Gore & Associates do użytku w kabinach samolotów szerokokadłubowych i jednokadłubowych – zapoznaj się z dokumentem PDF na stronie siliconchip.au/link/abiq.

Takie systemy powietrzne zapewniają „pikokomórki” dla zasięgu telefonii komórkowej, punkty dostępu dla Wi-Fi i obsługują protokoły Bluetooth, DECT, DECT2, Globalstar, GSM, Sat, MMS, PDC i TETRA. Zmniejszają martwe strefy i zmniejszają wagę wymaganego sprzętu.

Anteny są odpowiednie dla częstotliwości od 400 MHz do 6 GHz. Zobacz film „GORE Leaky Feeder Antennas” zamieszczony na stronie <https://youtu.be/ZK7wBCfJJa0>.

Wnioski

Podsumowując, istnieją dwie główne techniki komunikacji podziemnej bez konieczności prowadzenia przewodów w całej zamkniętej przestrzeni: wykorzystanie niskich częstotliwości (zazwyczaj VLF lub LF, od 3 kHz do 300 kHz) w celu lepszej penetracji skał i gleby lub użycie repeaterów (skonfigurowanych ewentualnie w sieć kratową – mesh) w celu przezwyciężenia trudności związanych z zasięgiem w granicach widoczności w zakrzywionych tunelach lub szeregu jaskiń/sztolni kopalnianych.

Główną zaletą podejścia VLF/LF jest to, że wymagane są tylko dwa radia, jednak niskie częstotliwości wymagają zazwyczaj używania stosunkowo dużych anten (nieco złagodzone przez zastosowanie pętli). W przypadkach takich jak tunele lub kopalnie, gdzie występuje aktywność częstotliwości i istnieje już znaczna infrastruktura, sieci mesh lub nieszczelne feedery zapewniają większą elastyczność.

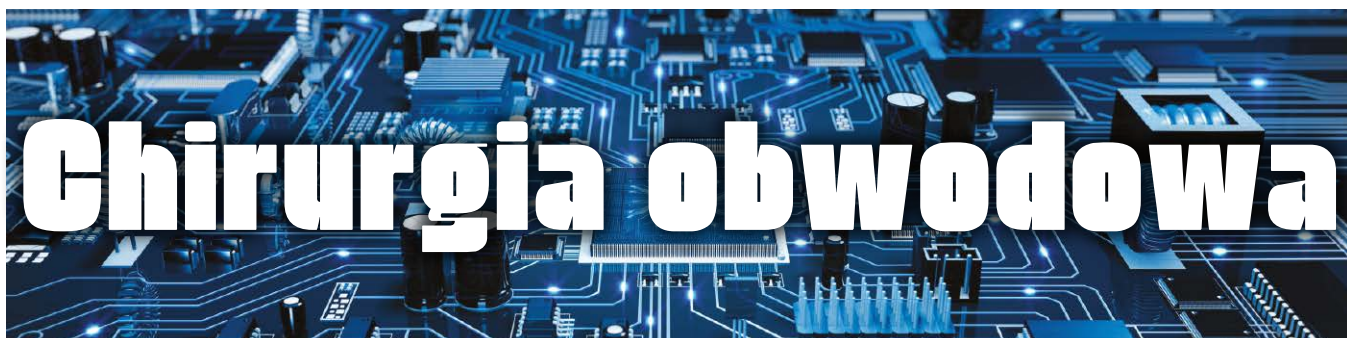
W zastosowaniach ratowniczych prawdopodobnie będzie wymagane połączenie tych dwóch podejść. Radiotelefony VLF/LF mogą być początkowo używane do czasu zbudowania sieci mesh, umożliwiając ratownikom komunikację za pomocą małych ręcznych radiotelefonów. Biorąc pod uwagę niski koszt potężnych chipów RF w dzisiejszych czasach, prawdopodobnie nie minie dużo czasu, zanim tanie radia mesh będą powszechnie dostępne. Być może będą to nawet projekty open-source. ■

dr David Maddison

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022.
www.siliconchip.com.au



Rysunek 53. Pig czyszczący rurociąg w wyciętym odcinku rury. Niektóre mają elektronikę i komunikują się na częstotliwości 22 Hz. Źródło: <https://w.wiki/6Exp> (CC BY-SA 2.0)



Chirurgia obwodowa

Zniekształcenia i obwody zniekształcające, część 1

Temat zniekształceń należy do podstawowych zagadnień w elektroakustyce i wielu innych dziedzinach elektroniki. Jak już sama nazwa wskazuje, zniekształcenia to coś, co zmienia kształt przebiegu. To dość szeroka definicja, obejmująca również efekt działania idealnych układów filtrujących; na przykład filtr dolnoprzepustowy „zaokrągla” ostre zbocza fali prostokątnej. Jednak termin „zniekształcenia” jest używany najczęściej w odniesieniu do wpływu nieliniowości obwodów na sygnały. Takie właśnie zjawiska wyjaśnimy szczegółowo w niniejszym artykule. Zniekształcenia we wzmacniaczach, filtrach itp. stanowią zwykle wadę. Istnieją jednak przypadki, gdy są zjawiskiem pożądanym.

Zniekształcenia wynikają oczywiście z niedoskonałości układów wzmacniających, które powinny tworzyć dokładną kopię sygnału wejściowego przy dużym poziomie mocy. Sygnał wyjściowy ma jednak w praktyce nieco inny kształt niż sygnał wejściowy.

W przypadku wzmacniaczy audio ilość wytwarzanych zniekształceń jest zwykle charakteryzowana przez wynik pomiaru „całkowitych zniekształceń harmonicznych” (THD). Zniekształcenia spowodowane nieliniowością dodają do sygnału wyjściowego nowe składowe o częstotliwościach, których na wejściu nie było. Używane jest tutaj określenie „harmoniczne”, ponieważ w najprostszym przypadku fali sinusoidalnej (pojedyncza częstotliwość), zniekształcenia wprowadzają do sygnału składowe o częstotliwościach będących całkowitymi wielokrotnościami częstotliwości oryginalnej, czyli jej harmoniczne. „THD” jest zatem miarą ilości harmonicznych wprowadzanych przez zniekształcenia i służy jako wskaźnik jakości wzmacniaczy i innych układów.

Idealny, liniowy filtr, o którym wspominałem wyżej, nie dodaje do sygnału wyjściowego żadnych nowych częstotliwości, których by nie było w sygnale wejściowym, więc nie wprowadza zniekształceń harmonicznych. Zmienia on jedynie względne amplitudy częstotliwości składowych obecnych

w sygnale wejściowym, i przez to powoduje zmianę kształtu fali.

Zniekształcenia muzyczne

Zniekształcenia nieliniowe są zwykle niepożądane. Aczkolwiek mają one także użyteczne zastosowania, a szczególnie znanym tego przypadkiem są efekty stosowane w elektronice muzycznej. Znaczny poziom zniekształceń nieliniowych w układach elektroakustycznych brzmi zwykle szorstko i nieprzyjemnie. Odpowiednio użyte, mogą być jednak przydatnym narzędziem twórczym. Rozpatrując to w kategoriach muzycznych – zmiana kształtu fali zmienia jej barwę. Barwa to cecha dźwięku muzycznego inna niż wysokość (częstotliwość podstawowa) czy głośność. To właśnie barwa odróżnia dźwięki różnych instrumentów grających tę samą nutę z tą samą głośnością, a możliwość zmiany barwy instrumentu zwiększa zakres ekspresji muzycznej. Jeśli wyjdziemy od dźwięku „czystego” (zbliżonego do fali sinusoidalnej), to niewielka ilość zniekształceń doda mu ciepła i ziarnistości. Większa ilość zniekształceń spowoduje, że dźwięk stanie się bardziej rozmyty, szorstki, zgrzytliwy. Jeśli zamiast czystego tonu poddamy zniekształceniu dźwięk złożony, powstały efekt zabrzmi zwykle o wiele ostrzej. Użycie zniekształceń w muzyce kojarzy się powszechnie

z gitarą elektryczną i muzyką rockową. Są tam szeroko stosowane zniekształcające efekty gitarowe („fuzz” i „distortion”; przypis redaktora). Zniekształcenia można jednak stosować do dowolnego instrumentu, również do głosu ludzkiego, i w różnych gatunkach muzycznych. **Przypis redaktora: silnie zniekształcony dźwięk, odpowiednio filtrowany i w kontrolowany sposób dodawany do sygnału wyjściowego, wykorzystuje urządzenie „exciter”, chętnie używane do obróbki np. głosu wokalisty czy instrumentów perkusyjnych. Użycie tego urządzenia „rozjaśnia”, „nasyca” i „ożywia” brzmienie instrumentu.**

Zastosowania w efektach muzycznych

W kilku numerach czasopisma „Practical Electronics” (kwiecień/maj 2022) opublikowano projekt Digital FX autorstwa Johna Clarke’a. Projekt zawiera różnorodne efekty dźwiękowe dla muzyków. Był ukierunkowany głównie na efekty stosunkowo złożone, takie jak phasing, chorus i pitchshifting, ale znalazło się też miejsce na efekty oparte na zniekształceniach. Mniej więcej rok wcześniej John opublikował projekt podłogowego efektu przesterowania („Nutube Guitar Overdrive and Distortion Pedal”; „Practical Electronics”, marzec 2021). Zainspirowani tymi projektami

przyjrzymy się podstawom powstawania zniekształceń (w tym podstawom pomiaru THD), zbadamy parę symulacji w SPICE związanych ze zniekształceniami oraz – w części 2. – rozpatrzmy kilka podstawowych układów używanych do tworzenia efektów opartych na zniekształceniach.

Układy liniowe i nieliniowe

Układ liniowy to taki, w którym sygnał wyjściowy jest powiązany z sygnałem wejściowym poprzez pomnożenie go przez prosty współczynnik skalowania G_1 . Na przykład dla sygnału wejściowego V_{in} możemy zapisać sygnał wyjściowy V_{out} jako:

$$V_{out} = G_1 \cdot V_{in}$$

Jest to funkcja przenoszenia układu. Jeśli G_1 jest większe lub równe 1, układ jest

wzmacniaczem o wzmacnieniu G_1 , w przeciwnym razie jest tłumikiem. Określenie „liniowy” nabiera sensu, jeśli sporządzimy wykres zależności między V_{in} i V_{out} . Jest to idealna linia prosta przechodząca przez początek układu współrzędnych, jak widać na wykresie górnym na **rysunku 1**. Nachylenie tej linii jest równe wzmacnieniu. Każdy układ, dla którego wykres charakterystyki wejście-wyjście nie jest linią prostą, jest nieliniowy i wprowadza zniekształcenia.

W każdym układzie rzeczywistym wzmacnienie G_1 zmienia się wraz z częstotliwością. Jeśli zmiana ta została specjalnie ukształtowana tak, by układ przepuszczał lub tłumił sygnały o określonych częstotliwościach, to mamy filtr. Może być on liniowy – wtedy zależność wartości sygnału

wejściowego od wyjściowego jest linią prostą przy dowolnej częstotliwości, nawet jeśli wzmacnienie (nachylenie linii) jest dla różnych częstotliwości różne.

Może się okazać, że układ rzeczywisty ma odpowiedź doskonale liniową, ale na wyjściu występuje stałe przesunięcie („offset”). Zapisalibyśmy wtedy funkcję przenoszenia jako:

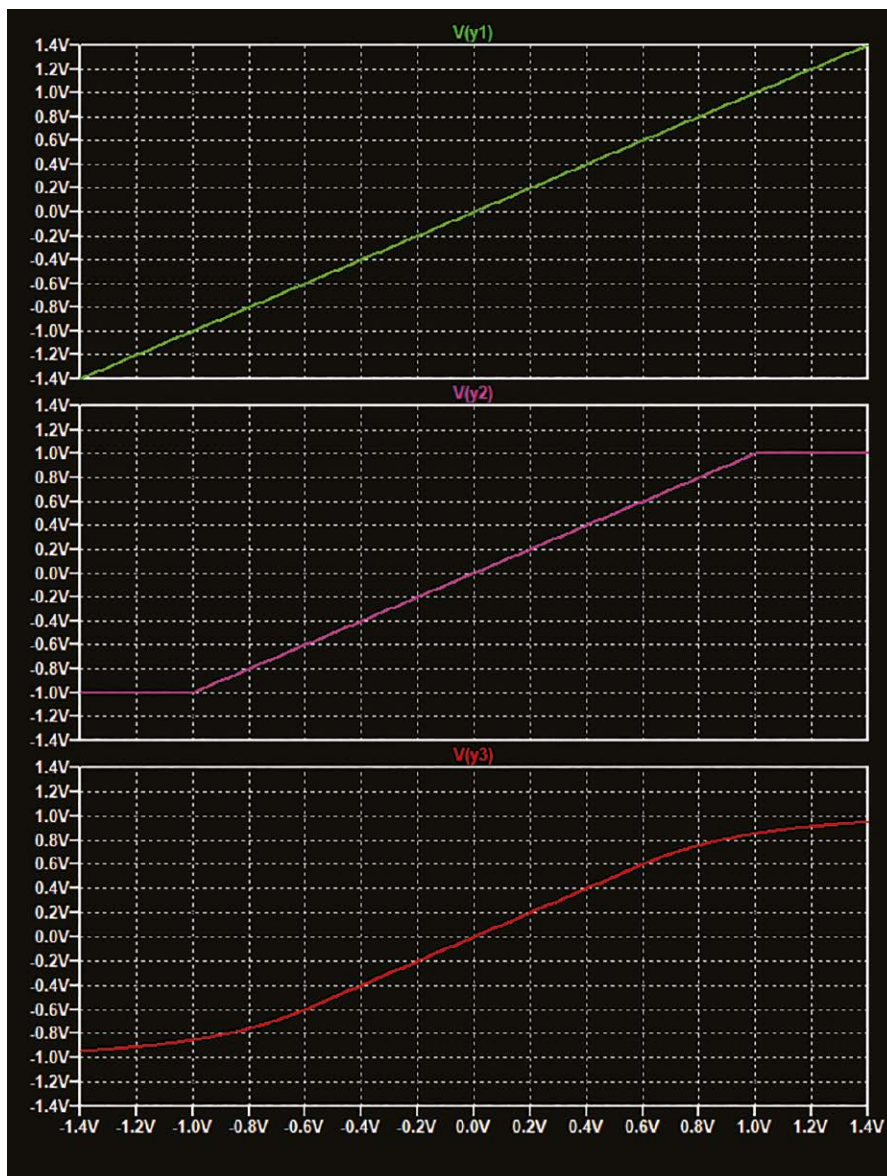
$$V_{out} = G_0 + G_1 \cdot V_{in}$$

G_0 jest przesunięciem. Już nie można powiedzieć, że układ nie zmienia kształtu przebiegu, aczkolwiek przy rozważaniu zniekształceń nie to jest najważniejsze. Ważne, że wykres wciąż jest linią prostą – choć nie przechodzi już przez punkt początkowy układu współrzędnych.

Ograniczenia od napięcia zasilania

Istnieje wiele powodów, dla których odpowiedź wzmacniacza może nie być w pełni liniowa. Chyba najbardziej oczywistym powodem jest, że poziom sygnału wyjściowego jest zawsze ograniczony. Ograniczenie to jest przede wszystkim spowodowane skończoną wartością napięcia zasilania. Przykłady zależności wejście-wyjście dla wzmacniaczy o skończonych maksymalnych poziomach wyjściowych pokazano na środkowym i dolnym wykresie na **rysunku 1**. W obu przypadkach amplituda napięcia wyjściowego jest ograniczona do 1 V. Środkowy wykres ilustruje skokowe przejście od obszaru pracy liniowej do stanu obciążenia, a wykres dolny – łagodne przejście w stan ograniczania.

Wykresy charakterystyk wejścia-wyjścia na **rysunku 1** zostały uzyskane przy użyciu behawioralnych źródeł napięcia i symulacji przemiatania stałoprądowego w symulatorze LTspice, przy użyciu schematu z **rysunku 2**. Źródło B1 stanowi prostą kopię napięcia wejściowego (napięcie w węzle x od źródła V1). Jest to odpowiedź liniowa (wykres górny). Źródło B2 reprezentuje „twarde” ograniczenie przy 1 V. Jego napięcie to wartość mniejsza (minimum, min) z dwóch wartości: napięcia wejściowego i poziomu 1 V. Aby uzyskać charakterystykę symetryczną względem 0 V, argumentem funkcji minimum jest wartość bezwzględna napięcia wejściowego (abs), a wynik funkcji minimum jest dalej mnożony przez „funkcję znaku” (sgn), która zwraca +1 lub -1 w zależności od znaku swojego argumentu. Źródło B3 tworzy łagodniejsze ograniczenie (wykres dolny). Zamiast funkcji min(x,y) używana jest funkcja uplim(x,y,z). Funkcja uplim jest podobna do funkcji min, ale jej pierwsza pochodna jest ciągła, a szerokość przejścia



Rysunek 1. Zależności wejście-wyjście wzmacniacza. Górny wykres to zależność liniowa, środkowy wykazuje ograniczenie „twarde” (obcinanie) przy ± 1 V, a wykres dolny ilustruje obcinanie łagodniejsze – „ograniczanie” – przy ± 1 V

wynosi „z”. Wygląda to przejście w zakresie wyznaczonym przez „z”. Termin „ciągła pierwsza pochodna” oznacza, że wykres nie ma skokowych zmian nachylenia; nie ma ostrych rogów, takich jak te na wykresie środkowym.

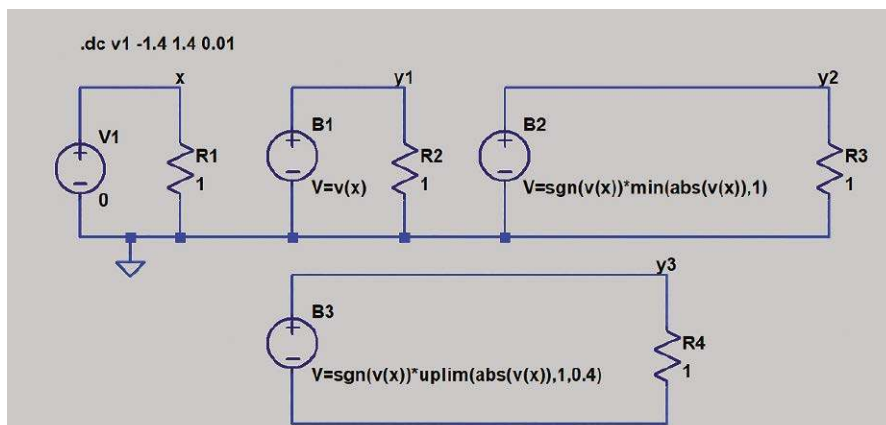
Wpływ wszystkich trzech funkcji wejścia-wyjścia z rysunku 1 na przebieg wejściowy pokazano na **rysunku 3**. Górny przebieg to sinusoida idealna (niezniekształcona). Przebieg środkowy wykazuje „ściśnięcie” – części sinusoidy o najwyższej amplitudzie ulegają spłaszczeniu. Charakterystyka o łagodniejszych przejściach (wykres dolny) powoduje „ograniczenie”, określane również jako „kompresja”.

Przebiegi uzyskano poprzez zmianę źródła V1 na schemacie z rysunku 2 na źródło sygnału sinusoidalnego 1 kHz, 1,3 V i uruchomienie symulacji przejściowej (rysunek 3).

Matematyczny opis zniekształceń

Jeśli zależność między wejściem a wyjściem wzmacniacza nie jest prostoliniowa, wówczas pojawia się problem matematycznej reprezentacji sygnału wyjściowego w celu określenia poziomu zniekształceń. Moglibyśmy podjąć się bardzo szczegółowej analizy układu, biorąc pod uwagę charakterystyki wszystkich jego elementów, ale byłoby to w ogólnym przypadku bardzo trudne i mogłoby dać równanie zbyt nieporęczne do dalszej analizy. Poza tym stosowałoby się to do tylko jednego analizowanego układu. Z każdym nowym musielibyśmy zaczynać od zera.

Potrzebujemy metody bardziej ogólnej, która mogłaby znaleźć zastosowanie do dowolnego układu wytwarzającego zniekształcenia. Receptą na nasz problem jest „twierdzenie Taylora”, opublikowane przez angielskiego matematyka Brooka Taylora w 1715 roku. Mówiąc prosto – twierdzenie to mówi, że każda gładka funkcja matematyczna może być przybliżona wielomianem (tzw. „szeregiem Taylora”). Wielomian jest równaniem utworzonym przez sumę potęg interesującej nas zmiennej (w tym przypadku V_{in}). W tej sumie znajduje się sam V_{in} (podniesiony do potęgi 1), V_{in}^2 (podniesiony do kwadratu), V_{in}^3 (podniesiony do sześcianu), V_{in}^4 (podniesiony do potęgi 4.) itd. Każda z tych potęg jest mnożona przez inny stały współczynnik (G_1 , G_2 , G_3 itd.). Indeks liczbowy współczynnika odnosi się do odpowiadającej mu potęgi V_{in} . Możemy również uwzględnić stałe przesunięcie, dodając współczynnik G_0 . Otrzymujemy w pełni ogólne równanie funkcji przejściowej wzmacniacza ze zniekształceniami:



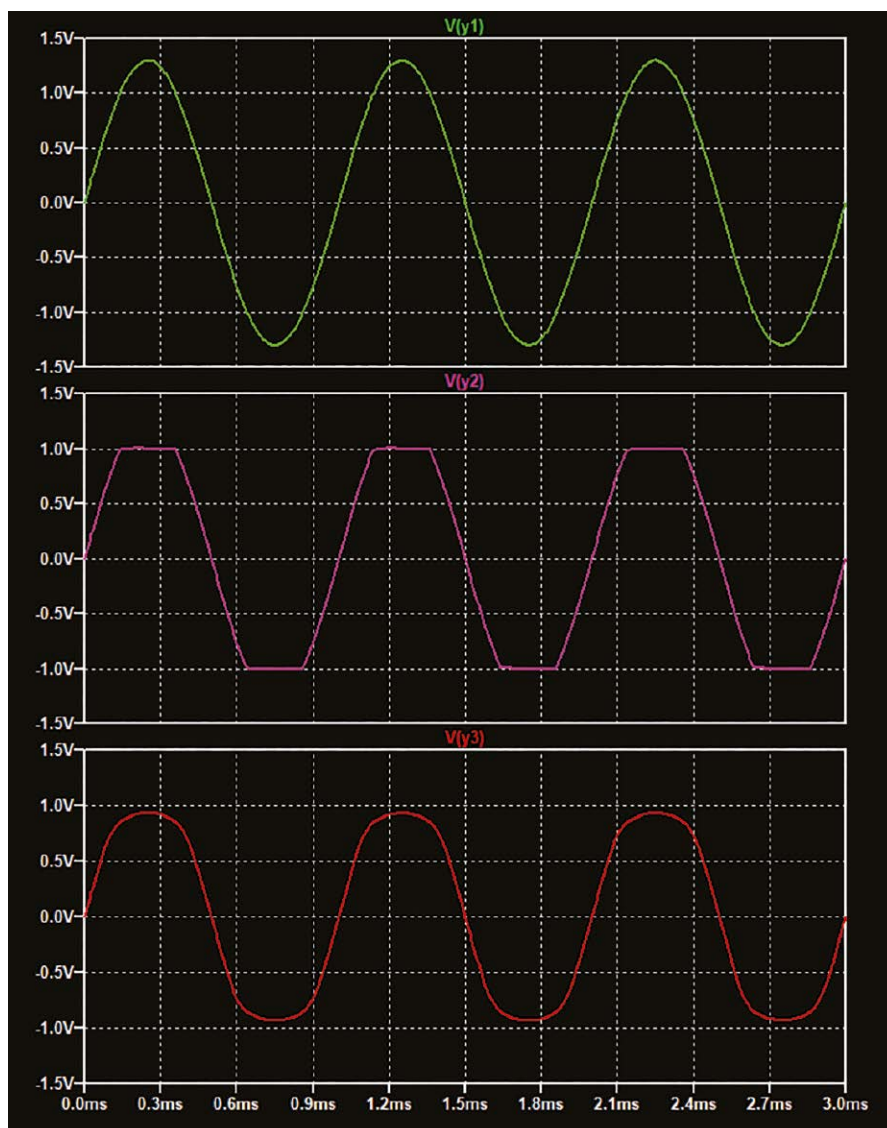
Rysunek 2. Schemat w LTspice użyty do sporządzenia wykresów z rysunku 1

$$V_{out} = G_0 + G_1 \cdot V_{in} + G_2 \cdot V_{in}^2 + G_3 \cdot V_{in}^3 + G_4 \cdot V_{in}^4 + \dots$$

Równanie to jest naturalnym rozszerzeniem równania idealnego wzmacniacza liniowego, przytoczonego poprzednio:

$$V_{out} = G_0 + G_1 \cdot V_{in}$$

Składniki, które sumujemy, jak np. $G_1 \cdot V_{in}$ czy $G_2 \cdot V_{in}^2$, są określane jako „wyrazy” sumy. Zasadniczo suma może się składać z bardzo dużej ilości wyrazów. W praktyce jednak wartości współczynników przy wyższych



Rysunek 3. Wpływ charakterystyk wejście-wyjście z rysunku 1 na przebieg sinusoidalny o amplitudzie 1,3 V

potęgach V_{in} stają się bardzo małe. Jest zatem mało prawdopodobne, abyśmy do obliczenia parametru THD potrzebowali ilości wyrażen większej niż np. 10.

Wyrażenie na sygnał wyjściowy V_{out} możemy rozdzielić na trzy składniki: przesunięcie (offset), sygnał idealny oraz zniekształcenia:

przesunięcie: G_0

sygnał idealny: $G_1 \cdot V_{in}$

zniekształcenia: $G_2 \cdot V_{in}^2 + G_3 \cdot V_{in}^3 + G_4 \cdot V_{in}^4 + \dots$

Gdybyśmy znali V_{in} i wartości wszystkich współczynników, moglibyśmy obliczyć względne poziomy sygnał idealnego i zniekształceń, a tym samym znaleźć zawartość procentową zniekształceń w sumarycznym sygnale wyjściowym. Prowadzi to do kilku problemów – nie zdefiniowaliśmy V_{in} i nie znamy wartości współczynników.

Przydatna byłaby możliwość porównywania różnych układów pod względem wytwarzanych przez nie zniekształceń. Porównywanie ma sens tylko wtedy, jeśli można to zrobić w spójny sposób. Dlatego w takim porównaniu konieczne jest użycie tego samego rodzaju sygnału V_{in} dla wszystkich typów zniekształceń. Jako podstawa do pomiaru zniekształceń idealnym wyborem będzie fala sinusoidalna dzięki swoim specjalnym własnościom.

Widmo sygnału

Falę okresową o dowolnym kształcie można utworzyć poprzez zsumowanie szeregu fal sinusoidalnych o różnych częstotliwościach i amplitudach. Ta suma fal sinusoidalnych znana jest jako „szereg Fouriera”. Teoria ta została opracowana przez Jeana Baptiste’a Josepha Fouriera (1768–1830), francuskiego matematyka i fizyka. Na przykład przebieg prostokątny może mieć częstotliwość

1 kHz, ale jest to tylko częstotliwość podstawowa. W przebiegu tym występują również inne częstotliwości. Jeśli fala ta ma amplitudę ± 1 V, to tworzy ją suma fal sinusoidalnych o amplitudach i częstotliwościach: 900 mV/1 kHz, 300 mV/3 kHz, 180 mV/5 kHz, 129 mV/7 kHz – i tak dalej do nieskończoności. Każda sinusoida ma częstotliwość będącą nieparzystą wielokrotnością n częstotliwości podstawowej, a jej amplituda wynosi $1/n$ amplitudy sinusoidy o podstawowej częstotliwości.

Falę okresową o dowolnym kształcie możemy rozłożyć na sinusoidy składowe, wyznaczając amplitudę każdej z nich. Jeśli narysujemy wykres amplitud składowych fal sinusoidalnych w funkcji częstotliwości, to otrzymamy widmo kształtu fali. Fragment widma fali prostokątnej pokazano na **rysunku 4**. Wykreślając widmo używamy często skal logarytmicznych dla częstotliwości i dla amplitudy (wyrażanej w decybelach; dB). Na rysunku 4 zastosowano jednak skale liniowe, co ułatwia dostrzeżenie zależności typu $1/n$ obowiązującej dla amplitud. Tworzenie wykresów widma przebiegów w programie LTspice omówimy w kolejnym odcinku.

Jeśli na wejście idealnego, liniowego wzmacniacza podamy sinusoidę, to na wyjściu otrzymamy sinusoidę o tej samej częstotliwości. Widmo sygnału zarówno wejściowego jak i wyjściowego będzie zawierać tylko pojedynczą składową. Jeśli jednak wzmacniacz wprowadza jakiegokolwiek zniekształcenia, to wówczas kształt fali wyjściowej nie będzie już idealną sinusoidą – jak na rysunku 3. Przebieg na wyjściu będzie zawierał dodatkowe składowe sinusoidalne o częstotliwościach innych niż częstotliwość wejściowa. Analiza Fouriera ujawni

te częstotliwości w widmie sygnału wyjściowego. **Rysunek 5** przedstawia widma wszystkich trzech przebiegów z rysunku 3. Wykres górny – dla czystej fali sinusoidalnej – pokazuje pojedynczy prążek przy częstotliwości tej fali. Widmo jest takie samo jak dla sygnału wejściowego. Pozostałe dwa wykresy pokazują dodatkowe częstotliwości, będące skutkiem zniekształcenia fali sinusoidalnej. Użyta jest tutaj decybelowa (logarytmiczna) skala amplitudy, ponieważ niektóre harmoniczne są zbyt słabe, by można je było dostrzec przy skalowaniu liniowym.

Zniekształcenia harmoniczne

Przy założeniu, że V_{in} jest falą sinusoidalną, możemy przystąpić do wykorzystania wzoru na funkcję przejściową wzmacniacza do obliczenia zniekształceń. Na początek dla uproszczenia założymy, że jedynym składnikiem powodującym zniekształcenia jest składnik z kwadratem V_{in} , czyli $G_2 \cdot V_{in}^2$. Pominiemy również błąd stałoprądowy (przesunięcie). Jako sygnału wejściowego użyjemy kosinusoidy, ponieważ jest ona nieco dogodniejsza do obliczeń niż sinusoida, a podstawowe właściwości odnośnie liczenia widma są takie same. Mamy zatem:

$$V_{out} = G_1 \cdot V_{in} + G_2 \cdot V_{in}^2$$

oraz

$$V_{in} = V_m \cdot \cos(\omega t)$$

gdzie: V_m to amplituda wejściowej fali kosinusoidalnej, Ω to jej pulsacja w radianach na sekundę (2π razy częstotliwość w hercach), a t to czas. Podstawiając wzór na V_{in} do równania na V_{out} otrzymujemy:

$$V_{out} = G_1 \cdot V_m \cdot \cos(\omega t) + G_2 \cdot V_m^2 \cdot \cos^2(\omega t)$$

Aby poradzić sobie z wyrażeniem $\cos^2$, zastosujemy jedną z podstawowych tożsamości trygonometrycznych, znaną jako wzór na podwojony kąt, a wynikającą z twierdzenia Pitagorasa. W tym przypadku:

$$\cos(2x) = 2\cos^2(x) - 1;$$

stąd:

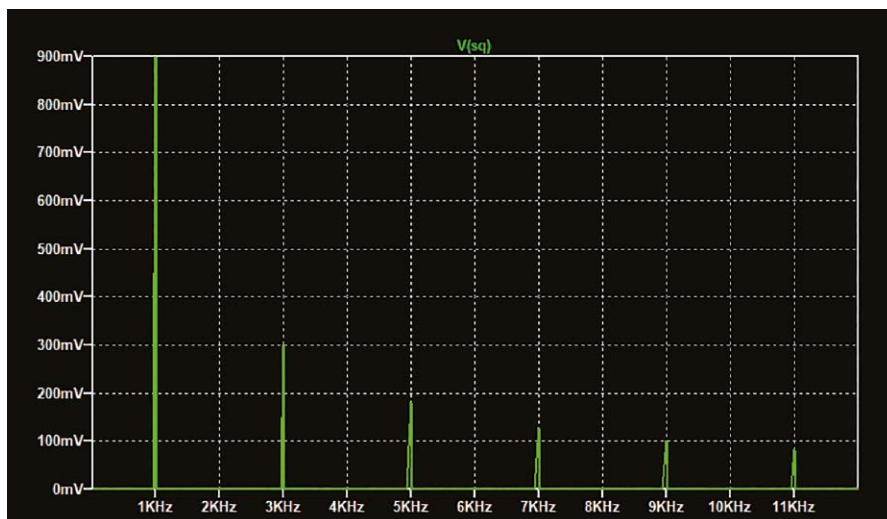
$$\cos^2(x) = \frac{1}{2} \cdot [1 + \cos(2x)]$$

Tak więc równanie na V_{out} przybiera postać:

$$V_{out} = G_1 \cdot V_m \cdot \cos(\omega t) + \frac{1}{2} \cdot G_2 \cdot V_m^2 + \frac{1}{2} \cdot G_2 \cdot V_m^2 \cdot \cos(2\omega t)$$

$\frac{1}{2} \cdot G_2 \cdot V_m^2$ jest tutaj składową stałą – w wyrażeniu tym nie ma kosinusa ani innej wartości zależnej od częstotliwości. Składnik ten stanowi stałe przesunięcie i jako taki zostanie przez nas pominięty. Bardziej interesujący jest składnik $\frac{1}{2} \cdot G_2 \cdot V_m^2 \cdot \cos(2\omega t)$, który reprezentuje składową o dwukrotnie większej częstotliwości (2ω) niż oryginalna (ω).

Ponieważ szukamy równań ogólnych, szczegółowe wartości współczynników takich jak $\frac{1}{2} \cdot G_2 \cdot V_m^2$ nie są w tym momencie bardzo istotne i możemy uprościć równanie,



Rysunek 4. Fragment widma fali prostokątnej. Liniowe skalowanie częstotliwości i amplitud

wprowadzając symbol $B_2 = \frac{1}{2} \cdot G_2 \cdot V_m^2$ i podobne symbole dla innych współczynników:

$$V_{out} = B_0 + B_1 \cdot \cos(\omega t) + B_2 \cdot \cos(2\omega t)$$

Równanie to otrzymaliśmy, zaczynając od uproszczonej wersji równania na zniekształcony sygnał wyjściowy, w której jako składnik wnoszący zniekształcenia uwzględniliśmy tylko $G_2 \cdot V_m^2$. Jeśli uwzględnimy więcej składników, obliczenia będą przebiegać w podobny sposób. Z różnymi potęgami kosinusa poradzimy sobie, używając odpowiednich tożsamości trygonometrycznych. W miarę jak równania będą się stawać coraz dłuższe, będzie o wiele więcej roboty, ale wynik będzie miał uporządkowaną formę:

$$V_{out} = B_0 + B_1 \cdot \cos(\omega t) + B_2 \cdot \cos(2\omega t) + B_3 \cdot \cos(3\omega t) + B_4 \cdot \cos(4\omega t) + \dots$$

Pamiętajmy, że jest to wzór na sygnał wyjściowy wzmacniacza zniekształcającego z podaną na wejście pojedynczą falą sinusoidalną (ściślej: w tym przypadku kosinusoidalną).

Podobnie jak w przypadku ogólnego wzoru na sygnał wyjściowy, możemy rozdzielić V_{out} na trzy składniki: przesunięcie, sygnał idealny oraz zniekształcenia:

przesunięcie: B_0

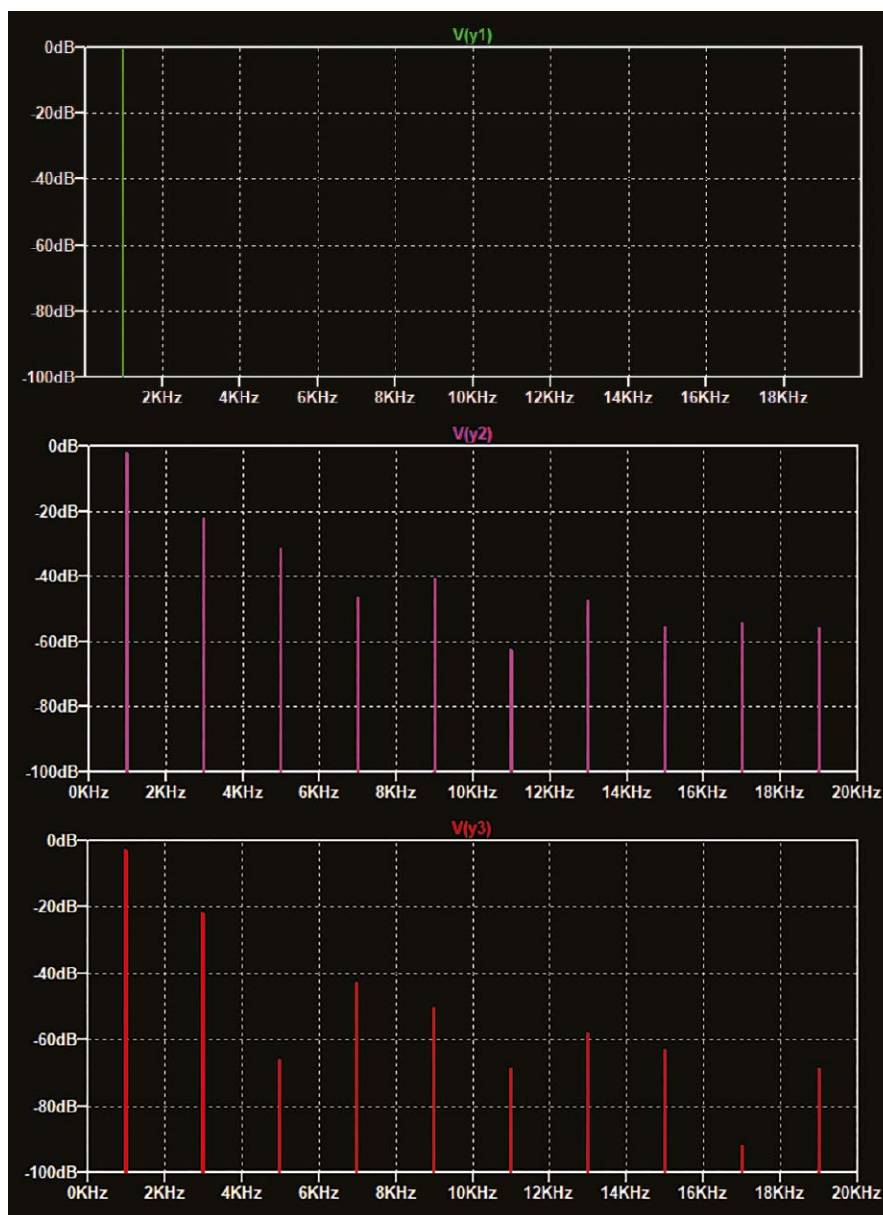
sygnał idealny: $B_1 \cdot \cos(\omega t)$

zniekształcenia:

$$B_2 \cdot \cos(2\omega t) + B_3 \cdot \cos(3\omega t) + B_4 \cdot \cos(4\omega t) + \dots$$

Widzimy, że zniekształcenia wejściowej fali sinusoidalnej skutkują powstaniem szeregu fal sinusoidalnych (kosinusów) o częstotliwościach dwukrotnie (2ω), trzykrotnie (3ω), czterokrotnie (4ω) itd. większych niż częstotliwość sygnału wejściowego. Częstotliwości będące całkowitymi wielokrotnościami danej częstotliwości nazywane są harmonicznymi tej częstotliwości, a ona sama jest określana jako częstotliwość podstawowa. Wynik pokazuje nam, że przy sinusoidalnym sygnale wejściowym zniekształcenia polegają na pojawieniu się wyłącznie harmonicznym sygnału wejściowego. Dlatego też używamy terminu „zniekształcenia harmoniczne”.

Jeśli sygnał wejściowy jest bardziej złożony niż pojedyncza fala sinusoidalna, wówczas zniekształcenia spowodują powstanie częstotliwości innych niż harmoniczne. Na przykład jeśli sygnał wejściowy składa się z dwóch przebiegów sinusoidalnych o różnych pulsacjach (ω_1 i ω_2), zniekształcenia będą obejmować, oprócz harmonicznym obu przebiegów, sumę ($\omega_1 + \omega_2$) i różnicę ($\omega_1 - \omega_2$) obu pulsacji oraz inne ich kombinacje. Jest to efekt znany jako zniekształcenia intermodulacyjne. Poziom zniekształceń intermodulacyjnych ma wielkie znaczenie w układach radiowych. Zniekształcenia intermodulacyjne stanowią też istotny parametr sprzętu do odtwarzania muzyki.



Rysunek 5. Widma przebiegów z rysunku 3. Częstotliwość w skali liniowej, amplituda w decybelach

Składnik $B_2 \cdot \cos(2\omega t)$ w równaniach mówi nam, jaki jest poziom zniekształceń związany z drugą harmoniczną. Sam w sobie nie jest jednak zbyt użyteczny, ponieważ tak naprawdę chcemy wiedzieć, jak duże są zniekształcenia w porównaniu z przypadkiem idealnym. Definiujemy więc D_2 – poziom zniekształceń dla drugiej harmonicznej – jako:

$$D_2 = \frac{|B_2|}{|B_1|}$$

W podobny sposób możemy zdefiniować wartości zniekształceń dla innych harmonicznym.

Wyrażenia typu $B_2 \cdot \cos(2\omega t)$ są zależne od czasu i wyrażają poziom składnika zniekształceń w dowolnej chwili. Natomiast przy porównywaniu układów bardziej przydatna byłaby znajomość

jakiegoś uśrednionego poziomu zniekształcenia. Wartość średnia sinusoidy bez przesunięcia wynosi zero. Używamy więc wartości średniokwadratowej (RMS). Wartość ta, opisująca napięcie zmienne, jest równa napięciu stałemu, które powodowałoby wydzielanie się w stałej rezystancji takiej samej mocy jak przy danym napięciu zmiennym. Dla sinusoidy o amplitudzie V_m wartość RMS jest dobrze znana i wynosi $V_{RMS} = V_m / \sqrt{2} = 0,707 \cdot V_m$. Natomiast idealny sygnał wyjściowy to $V_1 = B_1 \cdot \cos(\omega t)$. Jego wartość średniokwadratowa wynosi $V_{1,RMS} = B_1 / \sqrt{2}$. Moc P_1 wydzielana w dowolnej rezystancji R zasilanej z idealnego (niezniekształconego) sygnału wyjściowego wynosi:

$$P_1 = \frac{V_{1,RMS}^2}{R} = \frac{B_1^2}{2R}$$

Analogicznie możemy znaleźć moc drugiej harmonicznej, a następnie stosunek obu mocy, a tym samym poziom zniekształceń dla drugiej harmonicznej – D_2 :

$$P_2 = \frac{B_2^2}{2R}$$

$$\frac{P_2}{P_1} = \frac{B_2^2}{B_1^2} = D_2^2$$

Całkowite zniekształcenia harmoniczne

Przydatna byłaby znajomość łącznej mocy zniekształceń i możliwość przyrównywania jej do mocy sygnału idealnego. Moglibyśmy ulec pokusie i w celu uzyskania mocy całkowitej zsumować po prostu $P_2, P_3, P_4, \dots$. Matematyk mógłby w tym momencie zgłosić obiekcie i powiedzieć, że proste wyliczenie sumy mocy harmonicznych niekoniecznie musi tu być sposobem właściwym. Na szczęście, korzystając z czegoś, co się nazywa Twierdzeniem Parsewala, można wykazać, że wartości mocy można jednak zwyczajnie sumować. Twierdzenie to wykazuje, że możemy obliczyć moc lub energię sygnału poprzez jego widmo

(szereg Fouriera) i uzyskać tę samą odpowiedź jak w przypadku bezpośredniego obliczania mocy sygnału. Całkowita moc zniekształceń, $P_D = P_2 + P_3 + P_4 + \dots$, wynosi zatem:

$$P_D = \frac{B_2^2}{2R} + \frac{B_3^2}{2R} + \frac{B_4^2}{2R} + \dots$$

Po podzieleniu tej sumy przez moc sygnału niezniekształconego mamy:

$$\frac{P_D}{P_1} = \frac{B_2^2}{B_1^2} + \frac{B_3^2}{B_1^2} + \frac{B_4^2}{B_1^2} + \dots$$

$$\frac{P_D}{P_1} = D_2^2 + D_3^2 + D_4^2 + \dots$$

Wyrażając całkowite zniekształcenia pojedynczym symbolem (D), możemy napisać:

$$D^2 = \frac{P_D}{P_1} = \frac{P_2 + P_3 + P_4 + \dots}{P_1}$$

D jest współczynnikiem całkowitych zniekształceń harmonicznych (THD) i jest definiowany przez współczynniki zniekształceń poszczególnych harmonicznych:

$$D = \sqrt{D_2^2 + D_3^2 + D_4^2 + \dots}$$

Moc jest proporcjonalna do kwadratu napięcia, więc THD możemy również zapisać jako:

$$D = \frac{\sqrt{P_D}}{\sqrt{P_1}} = \frac{\sqrt{V_2^2 + V_3^2 + V_4^2 + \dots}}{V_1}$$

Tutaj V_1 jest amplitudą (lub wartością średniokwadratową) sygnału podstawowego, a $V_2, V_3, \dots$ są amplitudami (lub wartościami średniokwadratowymi) harmonicznych będących produktami zniekształcenia. THD można wyrazić w procentach lub w decybelach: $20 \cdot \log_{10}(D)$.

Niestety, podana definicja THD nie jest jedyną stosowaną. Alternatywnie współczynnik D definiuje się bezpośrednio jako stosunek mocy P_D/P_1 . Wersja „mocowa”, wyrażona w decybelach jako $10 \cdot \log_{10}(D)$, daje tę samą wartość, co wersja „napięciowa” w decybelach. Wartości procentowe są jednak inne. Dlatego do podawanych wartości THD należy podchodzić ostrożnie i sprawdzać, w jaki sposób zostały one zdefiniowane. ■

Ian Bell

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, czerwiec 2022 (www.epemag3.com)

REKLAMA

m.technik

Ciekawi świata są zawsze młodzi

przejrzyj i kupisz w prezencie
na każdą okazję na stronie
www.ulubionykiosk.pl





Migające diody LED i śliniacy się inżynierowie (27)

Nie wiem jak Wam, ale mnie dni zdają się uciekać coraz szybciej. Oprócz codziennej pracy jako niezależny konsultant techniczny i publicysta („piszę za wikt”), zajmuję się zazwyczaj różnymi projektami związanymi z moim hobby. Nieustannie jestem też olśniewany i zniewalany przez mnogość zachwycającego materiału, na który przypadkowo natrafiam, gdy krążę – niektórzy powiedzą „błądzą” – po Internecie. Niestety, umiem się koncentrować tylko przez krótki czas i rzadko udaje mi się skończyć jakąś rzecz, zanim pojawi się coś nowego i olśniewającego, co odwraca moją uwagę... Zupełnie jak wiewiórka!

Wydrukuj sobie

Jeszcze nie tak dawno temu amatorzy elektroniki musieli rysować schematy ręcznie i wytrawiać płytki drukowane w szkodliwych związkach chemicznych, pracując gdzieś w piwnicach jak konspiratorzy. Obecnie mamy wiele niedrogich narzędzi projektowych, które pozwalają hobbystom rysować schematy i projektować płytki PCB. Nawet narzędzia bezpłatne, takie jak DesignSpark PCB (<https://bit.ly/3pBQuWE>), albo KiCad (<https://www.kicad.org/>) – przypis redaktora, są niezwykle wydajne i zaawansowane. W wielu przypadkach narzędzia te mają możliwości, o których

jeszcze kilka lat temu można było tylko pomarzyć nawet w przypadku profesjonalnych produktów z najwyższej półki. Po sporządzeniu projektu PCB można skorzystać z oferty jednego z wielu zakładów produkcyjnych, które chętnie i niedrogo wykonują prototypy i dostarczają je na czas do rąk zainteresowanych klientów.

Z drugiej strony, podobnie jak można w domu używać drukarki 3D do tworzenia ciekawych gadżetów i bibelotów, istnieje możliwość wytwarzania płytek drukowanych w zaciszu własnego warsztatu (w moim przypadku zwanego również jadalnią, o czym rzadko omieszkam przypomnieć moją żonę – Gina Wspaniałą). Akurat jakimś dziwnym zbiegiem okoliczności rozmawiałem ostatnio z założycielami firmy BotFactory (<https://bit.ly/35yQsYH>). Ich najnowsza propozycja to system SV2 (**rysunek 1**), który z łatwością mieści się na stole warsztatowym (choć niekoniecznie mieści się w budżecie...). System łączy w sobie drukarkę z atramentem przewodzącym, wytłaczarkę pasty lutowniczej, maszynę do montażu powierzchniowego i piec do lutowania rozpliwowego. Pozwala na utworzenie płytki drukowanej w ciągu kilku minut.

Inne podejście reprezentuje interesujący gadżet o nazwie Print-a-Sketch, na który natknąłem się niedawno na stronie Hackaday

(<https://bit.ly/3vHdQOc>). To niewielkie urządzenie pozwala rysować połączenia elektryczne tuszem przewodzącym na niemal każdej powierzchni. Dzięki optycznemu czujnikowi ruchu (podobnemu do tych stosowanych w myszkach komputerowych) Print-a-Sketch wie, gdzie się znajduje i z jaką prędkością się porusza. Pozwala to jego układowi sterującemu kompensować niestabilne ruchy i na przykład prostować linie rysowane odręcznie drżącą ręką, zapewniając w ten sposób precyzję drukowania poniżej 0,5 mm.

Życie niejedno ma imię

Jak już zapewne wiecie, mój kumpel Steve Manley odgrywa ogromną rolę w wielu projektach hobbystycznych, o których piszę w artykułach Migające diody LED i śliniacy się inżynierowie w Practical Electronics. Steve mieszka w Wielkiej Brytanii, a ja w Stanach Zjednoczonych, ale dzięki cudowi Internetu możemy współpracować tak, jakbyśmy mieszkali tuż obok siebie. Szczercie mówiąc, Steve często wykonuje większość roboty, ja natomiast zbieram owoce. Każdy z nas ma swoją rolę do odegrania... Jako przykład można podać animatroniczną głowę robota, którą opracowujemy.

Pomysł na ten projekt zrodził się pod koniec 2021 roku, kiedy zbudowałem kilka



Rysunek 1. Przedstawiamy SV2 (zdjęcie: BotFactory)



Rysunek 2. Sprytne dzieło Jamesa Brutona



Rysunek 3. Oto nasza animatroniczna głowa SMAD (zdjęcie: Steve Manley)

głów pseudorobotycznych. Każda z nich jako oczy wykorzystywała dwa nasze SMAD-y („Steve and Max’s Awesome Display”). Patrz Practical Electronics, wrzesień, październik i listopad 2021; EdW 4/2025 oraz EdW 5/2025. Nieco później (Practical Electronics, grudzień 2021 r.; EdW 7/2025) nasz znakomity wydawca Matt Pulzer zaproponował, abym dodał do moich robotów jakiś rodzaj ruchu. Nie jestem już pewien, kto co powiedział, kiedy i dlaczego, ale w ramach tego projektu zacząłem przeglądać i analizować interesujący artykuł na stronie www.Hackaday.com, napisany przez Danie Conrادية (<https://bit.ly/3KjHzRv>). Artykuł Danie zapoznał mnie z genialnym, zautomatyzowanym artefaktem stworzonym przez Jamesa Brutona (**rysunek 2**).

Ta bestyjka (urządzenie, nie James) ma służyć jako podstawa realizacji realistycznych animacji. Problem, którym zajął się James, wynika z faktu, że amatorskie serwomechanizmy (o których będziemy mówić w przyszłym odcinku) poruszają się w stronę celu ze stałą prędkością. W wielu przypadkach to żaden problem, ale w zastosowaniach animatronicznych może umniejszać realizm. Jak się nad tym zastanowić, to bardzo niewiele ruchów naszego ciała odbywa się ze stałą prędkością. Kiedy poruszamy takimi częściami ciała jak głowa i oczy, to one zawracają albo przyspieszają, albo zwalniają. Jeśli na przykład z jednej strony mojej głowy rozlegnie się głośny dźwięk, to mięśnie szyi szybko przyspieszą ruch głowy w kierunku celu, a następnie, gdy głowa zacznie osiągać pożądaną pozycję, stopniowo go zwolnią. Nie mam wątpliwości, że techniki opisane przez Jamesa z czasem znajdą zastosowanie w naszych projektach animatronicznych. A skoro o tym mowa...

Z głową w dłoniach

Kiedy rozpoczęliśmy prace nad animatroniczną głową, zacząłem eksperymentować z prostymi, gotowymi mechanizmami obrotu i pochylenia. W międzyczasie Steve uruchomił swoje zintegrowane oprogramowanie CAD, CAM, CAE i PCB Fusion 360 (<https://autode.sk/3hVxGNv>). Nie minęło dużo czasu, a Steve na swojej niezawodnej drukarce 3D stworzył coś pięknego i przyjemnego dla oka.

W poprzednim artykule (Practical Electronics, marzec 2022; EdW 10/2025) pokazałem Wam zdjęcia mojej prymitywnej platformy prototypowej i porównałem ją z fascynującą konsolą sterującą Steve’a. Przedstawiłem również zdjęcie ilustrujące niesamowite podejście Steve’a do wykorzystania dwóch serwomechanizmów w celu realizacji obrotu i pochylenia jednego z oczu.

Animację tego ruchu można obejrzeć na moim kanale na YouTube (<https://bit.ly/3GcVTJK>).

Tak na marginesie – po przeczytaniu tego artykułu i obejrzeniu filmu Ken Wood, członek społeczności Practical Electronics (który jeszcze chętniej niż Wasz skromny narrator stosuje nawiasy), wysłał mi e-maila z następującą treścią: „Cześć Max, jeśli chodzi o mechanizm obrotu i pochylenia Twojego kolegi: innym rozwiązaniem może być zdalne sterowanie przegubowymi »mocowaniami« na końcu ramienia poprzez dwie linki – cięgła (takie jak w hamulcach rowerowych, a nawet w ręcznych skrzyniach biegów – zanim pojawiły się skomplikowane systemy stosowane dzisiaj). Mechanizm opisany w Twoim artykule mógłby mieć obrotownicę z dala od serwomechanizmów, a ich ramiona łączyć cięgłami (o ile serwomechanizmy są wystarczająco mocne). Cięgła stanowią alternatywę dla systemów hydraulicznych (tyle, że nie da się nimi wzmacniać sił) i przy starannym zaprojektowaniu mogą zakreślać wokół przegubów pośredniczących (można również użyć więcej linek). Zasadniczą rzeczą jest utrzymanie wszystkich cięgier w tej samej płaszczyźnie, równoległej do osi obrotu, tak aby podczas zginania przegubu nie dochodziło do różnicowej zmiany długości. Z pomocą podobnych mechanizmów można realizować jeszcze bardziej interesujące systemy przegubowe. Bierzemy rurkę falistą (taką, która się zgina) i prowadzimy przez nią cztery linki (bez osłony), rozmieszczone na obwodzie. Pociągnięcie dowolnej linki spowoduje wygięcie rury w odpowiednim kierunku (podobnie jak wygina się trąba słonia). Skręcenie linek w przeciwnych kierunkach spowoduje skręcenie rury. W tym przykładzie linki działają jak ścięgna, a sprężystość rury zapewnia siłę przywracającą (podobnie jak mięśnie przeciwstawne)”.

No cóż, to wszystko na pewno daje do myślenia, ale odbiegamy od tematu...

Od czasu, gdy pisałem poprzedni artykuł, Steve ukończył pracę nad

swoją animatroniczną głową (**rysunek 3**). Abyście mogli się nią zachwycić, Steve zamieścił na YouTube film, który tę wspaniałą piękność prezentuje w całej okazałości (<https://bit.ly/3sMi6Kk>).

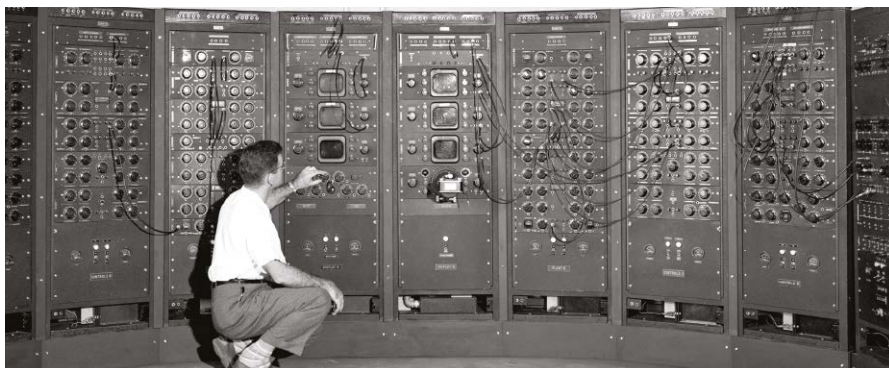
Jeśli jesteście zainteresowani stworzeniem jednego z tych niesamowitych artefaktów dla siebie, w następnym artykule omówię więcej szczegółów dotyczących serwomotorów i płytek rozgałęźnych (BOB), których użyliśmy. Ponadto Steve bardzo uprzejmie udostępnił pliki swoich projektów 3D. Plik skompresowany CB-May22-01.zip, zawierający te materiały, możecie pobrać ze strony internetowej Practical Electronics z maja 2022 r. (<https://bit.ly/3oouhbl>).

Następnie zajmiemy się oprogramowaniem, w tym realizacją realistycznych technik animacji przedstawionych przez Jamesa Brutona, o których wspominałem wcześniej. Wspaniałą wiadomością jest jednak to, że Steve bardzo uprzejmie stworzył drugą głowę animatroniczną SMAD i konsolę sterującą dla mnie. Kiedy piszę te słowa, od zaledwie kilku dni mam w domu to чудо. Dzięki temu mogę z dumą ogłosić, że siedzę trzymając w dłoniach głowę – a nie co dzień można coś takiego powiedzieć.

Nie mów językiem specjalistów!

Na pewno pamiętacie, że czasami wspominałem o moim kumplu Paulu Parrym. Paul jest właścicielem firmy Bad Dog Designs, gdzie tworzy wspaniałe arcydzieła z lampami Nixie (<https://bit.ly/3sMFWps>). Kilka lat temu Paul przedstawił mnie dr Lucy Rogers, która jest autorką, wynalazczynią i profesorem wizytującym na wydziale inżynierii, kreatywności i komunikacji uczelni Brunel University London.

Jak napisałem w ostatnim wpisie na blogu: „Nawet jeśli jesteś największym specjalistą, nie radzę Ci mówić językiem specjalistów. Nie potrafię powiedzieć, ile razy widziałem, jak oczy ludzi szklą się, kiedy z nimi rozmawiam... a sytuacja jest jeszcze gorsza, kiedy



Rysunek 4. Komputer analogowy, Laboratorium Napędów Lotniczych Lewis, około 1949 r. (zdjęcie: NASA)

używam języka specjalistów”. Chodzi o to, że właśnie usłyszałem, że dr Lucy oferuje zestaw „kursów komunikacji online prowadzonych przez osobę z branży technologicznej dla osób z branży technologicznej” (<https://bit.ly/35xEG0w>).

Jak to było kiedyś

W dzisiejszych czasach jesteśmy wprost rozpieszczani mocą obliczeniową i sprawnością komputerów. Nadal trudno mi uwierzyć, że można kupić płytkę rozwojową mikrokontrolera Seeeduino XIAO – która ma wymiary zaledwie 24×18 mm, jest wyposażona w mikrokontroler z 32-bitowym rdzeniem Arm Cortex-M0+, taktowany zegarem 48 MHz,

z 256 KB pamięci Flash i 32 KB pamięci SRAM – za cenę zbliżoną do sumy drobnych monet, jakie mam w portfelu (<https://bit.ly/3IPpSc2>).

To ogromna różnica w porównaniu z sytuacją z 1975 roku, kiedy to stawiałem pierwsze kroki na drodze do uzyskania tytułu licencjata w dziedzinie inżynierii sterowania na Uniwersytecie Sheffield Hallam. W tamtym czasie jedynym dostępnym dla nas komputerem cyfrowym był potężny komputer mainframe, który zajmował cały budynek. Pamiętam, jak w języku programowania Fortran (FORMula TRANslator) pisałem swój pierwszy program dla tego monstera. Pierwszym krokiem było napisanie programu ołówkiem na kartce papieru. Następnie wprowadzałem ten program do Teletype'u, który każdy wiersz tekstu przekształcał w układ otworów na karcie perforowanej. Potem osobiście zanosilem mój plik kart do budynku komputerowego i przekazywałem go jakiejś gburowatej osobie czającej się za ladą recepcyjną. **Przypis redaktora:** określenie „plik”, oznaczające program lub dane w pamięci masowej, wzięło się właśnie od plików kart perforowanych, używanych w epoce wielkich komputerów. Kiedy tydzień później wracałem, otrzymywałem z powrotem moją oryginalną „talię” oraz kartkę papieru z odręczną notatką „w wierszu 2 brakuje przecinka”. Wrrr! – że powiem dosadnie – czy nie mogli dodać tego przecinka za mnie? Efektem takiego sposobu działania było to, że nawet najprostsze programy pochłaniały większość semestru, zanim dały sensowne wyniki. **Przypis redaktora:** sam pamiętam, jak w połowie lat 80. na uczelni męczyłem się przez wiele tygodni na dużym komputerze, próbując uruchomić pewien program. W końcu poszedłem do Koła Naukowego, gdzie na mikrokomputerze 6502 z interpreterem BASICa program ten napisałem, uruchomiłem i uzyskałem wyniki w ciągu kilku godzin...

Jedynym komputerem, który fizycznie znajdował się na Wydziale Inżynierii, była analogowa piękność, taka jak ta pokazana na **rysunku 4**. Komputer analogowy polega na tym, że mamy sporo małych modułów analogowych, z których każdy pełni jakąś funkcję. Na przykład porównuje dwa sygnały, sprawdzając, który jest większy (ma wyższe napięcie), sumuje kilka sygnałów, mnoży dwa sygnały czy też całkuje sygnał w czasie. Było też kilka potencjometrów, które służyły do zadawania wartości współczynników. Wejścia i wyjścia modułów były łączone między sobą kablami z wtyczkami jack. Kto nigdy sam nie wykonywał takich połączeń, nie ma pojęcia, jak trudne było

potem znajdowanie błędów konfiguracji, która nie działała zgodnie z planem.

Jedną z wad takiej analogowej formy obliczeń jest brak powtarzalności, ponieważ wyniki kolejnych uruchomień, choć zbliżone do siebie, zawsze wykazują niewielkie różnice. Ponadto na kolejnych etapach działania drobne błędy początkowe mogą się kumulować i wzmacniać. Mimo to osoby, które nie miały okazji korzystać z tego typu maszyn, mogą być zaskoczone faktem, że obliczenia analogowe i „analogowe przetwarzanie sygnałów” (ASP) mogą być niezwykle skuteczne w modelowaniu różnych systemów dynamicznych. O ile dobrze pamiętam, moim pierwszym zadaniem było utworzenie modelu chłodzenia wnętrza lodówki po jej pierwszym załączeniu, a także symulacja zmian temperatury wewnętrznej po otwarciu drzwi lodówki.

Analog Thing (THAT)

Mówię o tym tutaj, ponieważ dowiedziałem się właśnie o nowo powstającym komputerze analogowym o nazwie The Analog Thing (THAT) (dosłownie „Rzecz Analogowa”; **przypis redaktora**) (<https://bit.ly/370WEZH>), który widać na **rysunku 5**. Chyba się starzeję, ponieważ najpierw myślałem, że „THAT” jest częścią nazwy tego urządzenia. Zajęło mi dobrych parę sekund, zanim zdałem sobie sprawę, że „THAT” jest w rzeczywistości akronimem całej nazwy („The Analog Thing”).

Mamy tu osiem potencjometrów do ustawiania współczynników, pięć integratorów, cztery sumatory, dwa komparatory, dwa bloki mnożące oraz miernik panelowy. Są również porty wyjściowe X, Y, Z i U, które mogą sterować urządzeniami zewnętrznymi, np. oscyloskopem, a ponadto porty Master i Minion, które umożliwiają połączenie wielu urządzeń THAT w łańcuch w celu realizacji dowolnie dużego systemu. Muszę przyznać, że bardzo kusi mnie, aby „wydać kasę” na takie małe cudenka. Odkryłem, że mam dziwną ochotę stworzyć jakąś skomplikowaną funkcję, a następnie spędzać sporo czasu na rozpracowywaniu płataniny kabli. Jeśli nie dałoby mi to nic innego, to przynajmniej zabrałoby mnie w sentymentalną podróż do czasów młodości.

Ale to nie wszystko, ponieważ jeden z powodów, który przyciągnął mnie do THAT, jest taki, że zastanawiam się nad pewną analogową zagadką związaną z potencjometrami, których Steve i ja używamy do sterowania naszymi animatronicznymi głowami. Za chwilę wspólnie pomyślimy nad tym problemem, ale najpierw...

Wędrowka w maliny

Być może zauważyliście, że czasami jakbym zapędzał się w maliny, kiedy zaglądam



Rysunek 5. The Analog Thing (zdjęcie: Anabrid)

się w pozornie przypadkowy temat, który mnie interesuje. Nie dajcie się zwieść, ponieważ za moim szaleństwem kryje się metoda (naprawdę!). Projekty opisywane w Practical Electronics, oprócz tego, że mają walor edukacyjny, są zabawne i interesujące same w sobie, dają naszym Czytelnikom solidne podstawy do opracowywania w przyszłości własnych układów i systemów. A kiedy przedstawiamy porady, sztuczki i sposoby dotyczące jakiejś sytuacji, zwykle stosują się one również do innych przypadków.

Przykład – we wcześniejszym artykule („Practical Electronics”, październik 2021 r.; EdW 05/2025) opisałem sposób wykorzystania dwuwymiarowych tablic 8-bitowych liczb całkowitych do definiowania grup LED-ów używanych do wzorów na naszych SMADach. Dwa miesiące później („Practical Electronics”, grudzień 2021 r.; EdW 07/2025) wspominałem o tym, że członek społeczności „Practical Electronics”, którego nazwiemy Simon (no bo tak właśnie ma na imię), wysłał mi e-maila, w którym napisał, że przeczytanie mojego oryginalnego artykułu i zapoznanie się z tą techniką bardzo mu pomogło przy tworzeniu trenera kodu Morse’a opartego na Arduino. Wspominał o tym tylko dlatego, że zamierzam zrobić krótką dygresję na temat, który w tej chwili może nie wydawać się zbyt ważny. Jednak metody, które tutaj omawiamy, z pewnością okażą się przydatne w przyszłości.

Lewą marsz!

Kilka miesięcy temu („Practical Electronics”, styczeń 2022 r.; EdW 08/2025) przedstawiałem 4-osiowe joysticki JH-D400X-R4 10K, których używamy do sterowania naszymi sztucznymi głowami. W przyszłości, aby głowy mogły same sobą sterować (drżycie

ze strachu!), planujemy wykorzystać czujniki oraz sztuczną inteligencję – ale to już temat na inną okazję.

JH-D400X-R4 to element tak naprawdę 3-osiowy, ale z przyciskiem na górze, a twórcy tego joysticka pomysłowo uznali ten przycisk za czwartą oś sterowania. Joystick zawiera trzy potencjometry 10 kΩ, z których każdy jest wycentrowany w pozycji środkowej (5 kΩ). Jeśli przesuniemy potencjometr wzdłuż dowolnej osi i puścimy go, powróci on do pozycji środkowej. Pchnięcie joysticka do przodu lub pociągnięcie go do tyłu obraca jednym z potencjometrów. Przesunięcie go w lewo lub w prawo – innym potencjometrem. Obracanie joystickiem w prawo lub w lewo wpływa na trzeci potencjometr.

W obecnej konfiguracji głową steruje domyślnie lewy joystick. Jeśli stoicie twarzą do głowy, to przesunięcie drążka w lewo/prawo spowoduje obrót głowy w lewo/prawo (z punktu widzenia głowy – w prawo/lewo). Natomiast odsunięcie drążka od siebie spowoduje odchylenie głowy do góry, a do siebie – pochYLENIE głowy w dół. Dla nas ten sposób działania jest sensowny. Innym użytkownikom może jednak wydać się mało intuicyjny. Musimy również brać pod uwagę sytuację, w której operator stoi nie przed głową, a za nią. Dlatego na wierzchu konsoli sterującej umieściliśmy trzy przełączniki. Jeden odwraca kierunek obrotów lewo-prawo, drugi odwraca kierunek wychyleń góra-dół... A trzeci? Nie powiem!

Joystick prawy domyślnie steruje w podobny sposób jednocześnie obydwoma oczami SMADa, sprawiając, że patrzą one w lewo, w prawo, w górę i w dół. Jednak po naciśnięciu przycisku na lewym joysticku system przechodzi w inny tryb. Głowa pozostaje w aktualnej pozycji, a lewy i prawy joystick sterują odpowiednio lewym i prawym okiem SMADa. Ponowne naciśnięcie przycisku lewego joysticka spowoduje powrót do trybu domyślnego.

Dlaczego, ach dlaczego?

Kiedy po raz pierwszy zacząłem eksperymentować z wczesną wersją mojej animatronicznej głowy, zdecydowałem się użyć modułu Teensy (https://bit.ly/3J7pfeh). Wybrałem Teensy 3.6, ponieważ podobały mi się jego możliwości w postaci 32-bitowego procesora Arm Cortex-M4F z zegarem 180 MHz, 256 KB pamięci RAM, 1 MB pamięci Flash, 58 linii GPIO (wejścia/wyjścia cyfrowe) i 25 wejść analogowych oraz b) akurat miałem taki pod ręką.

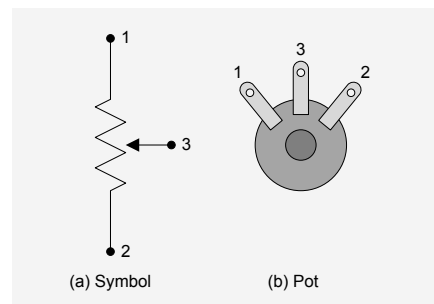
Twórcy Teensy udostępniają świetne diagramy funkcji pinów (https://bit.ly/3J9lq8a).

Uproszczony diagram dla górnego rzędu pinów Teensy 3.6 widzimy na **rysunku 6**. Piny cyfrowe od 40 do 57 oraz wejścia analogowe A10, A11, A23 i A24 są dostępne od spodu płytki.

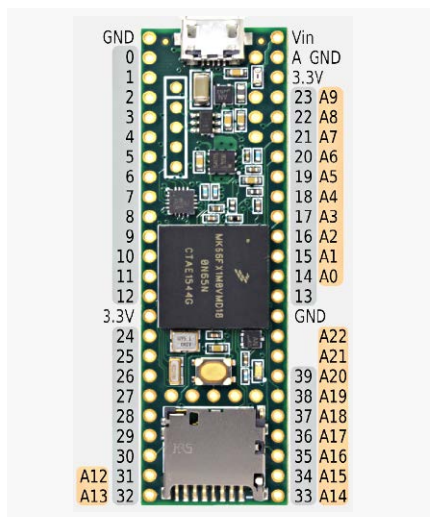
Zacząłem od dołączenia suwaków sześciu potencjometrów w joystickach do pinów analogowych A12 i A13 po lewej stronie Teensy oraz A14, A15, A16 i A17 po prawej stronie. Następnie utworzyłem mały szkic (program), który cyklicznie odczytywał wartości z potencjometrów i wyświetlał je na monitorze szeregowym wchodzącym w skład zintegrowanego środowiska programistycznego (IDE) Arduino. Wciąż pamiętam dreszczyk emocji, jaki towarzyszył mi, gdy chwyciłem lewy joystick i poruszałem nim w lewo/prawo i w przód/tył – a także... uczucie klęski, przygnębienia i rozczarowania, gdy nic się nie wydarzyło.

Na szczęście, zamiast rzucać się na oślep w szaleńczą walkę z błędami, najpierw spróbowałem poruszać w lewo/prawo i w przód/tył prawym joystickiem. O radości! Liczby na ekranie odpowiednio się zmieniały. Następnie spróbowałem oba joysticki obracać w prawo/lewo i znowu wartości wyświetlane w monitorze szeregowym zmieniały się zgodnie z oczekiwaniami.

Mój niezawodny multimetr pokazał, że każda oś obu joysticków reagowała tak, jak powinna, a do wejść Teensy docierały prawidłowe sygnały. Nie zajęło też dużo czasu ustalenie, że dwie nieodczytywane osie były dołączone do wejść analogowych A12 i A13. Dlaczego wejścia te nie działały prawidłowo? Zastanawiałem się nad możliwymi przyczynami, ale nic nie przyszło mi do głowy. W końcu, zamiast walić głową w mur, po prostu odlutowałem przewody doprowadzone do A12 i A13, dołączyłem je do A18 i A19 i wszystko działało tak, jak powinno. Dlaczego, ach dlaczego tak się działo? Nie mam pojęcia. Oczywiście nie miałbym nic przeciwko, żeby się tego dowiedzieć (jeśli macie jakieś pomysły, chętnie je poznam), ale nie spędza mi to snu z powiek. Życie jest na to zbyt krótkie.



Rysunek 7. Symbol i rysunek potencjometru



Rysunek 6. Wyprowadzenia na górnej stronie Teensy 3.6

Małe kroczki

Aby mieć pewność, że wszyscy tańczymy w tym samym rytmie, przypomnijmy sobie, że potencjometr ma trzy zaciski (**rysunek 7**). Dwa z nich są dołączone do końców ścieżki rezystancyjnej (w naszym przypadku 10 kΩ), a trzeci do ruchomego suwaka. Używamy naszych potencjometrów jako dzielników napięcia, więc do zacisku (1) dołączamy zasilanie, a do zacisku (2) masę (0 V) – lub odwrotnie (kolejność nie ma znaczenia, ponieważ w oprogramowaniu możemy łatwo manipulować odczytami).

Jeśli używacie modułu z mikrokontrolerem, takiego jak Arduino Uno, który jest zasilany napięciem 5 V, a jego porty GPIO działają w zakresie napięć 0...5 V, to do zasilania potencjometrów użyjecie napięcia 5 V. Dla porównania: mój Teensy 3.6 jest zasilany napięciem 5 V, ale jego porty tolerują tylko napięcie od 0 do 3,3 V, więc do zasilania potencjometrów używam zasilania 3,3 V, wytwarzanego przez stabilizator wbudowany w Teensy.

Teensy 3.6 ma 10-bitowy przetwornik analogowo-cyfrowy (ADC), wspólny dla wszystkich wejść analogowych. Dziesięć bitów może obsługiwać $2^{10} = 1024$ różne wartości dwójkowe, numerowane od 0 do 1023. Oznacza to, że napięcia od 0 V do 3,3 V z suwaków moich potencjometrów będą postrzegane przez Teensy jako liczby całkowite w zakresie od 0 do 1023.

Wróćmy teraz do mojego programu, który w pętli odczytuje wartości z sześciu potencjometrów i wyświetla je na monitorze szeregowym Arduino (**rysunek 8**). W odniesieniu do adnotacji LLR, LFR, LRT, RLR, RFB i RRT pierwsza litera oznacza potencjometr lewy (L) lub prawy (R), natomiast pozostałe litery odzwierciedlają oś: LR = lewo/prawo, FR = przód/tył, RT = obrót.

Szczerze mówiąc, podczas łączenia przewodów nie poświęciłem czasu na podłączanie danych potencjometrów do określonych wejść analogowych mikrokontrolera, ponieważ, jak zwykle, tego typu rzeczy łatwo jest rozwiązać w oprogramowaniu. To wyjaśnia, dlaczego zaczynamy od zdefiniowania stałej NUM_POTS równej 6, a następnie deklarujemy tablicę liczb całkowitych o nazwie PinsPots[], której używamy do uporządkowania wejść w żądanej kolejności:

```
PinsPots[NUM_POTS] = {A18, A14, A19, A17, A15, A16};
```

Deklarujemy również tablicę liczb całkowitych o nazwie PotVals[], której użyjemy do przechowywania wartości odczytanych z potencjometrów. Cały program opiera się na pętli wywołującej dwie funkcje, z opóźnieniem 1 sekundy na końcu pętli. Pierwsza funkcja, **ReadPots()**, odczytuje wejścia analogowe i wpisuje odczytane wartości do PotVals[], a następnie druga funkcja, **DisplayPotValues()**, odczytuje wartości z PotVals[] i wyświetla je na monitorze szeregowym. Serce funkcji **ReadPots()** jest prosta instrukcja **for()**:

```
for (int iPot = 0; iPot < NUM_POTS; iPot++)
{
    PotVals[iPot] = analogRead(PinsPots[iPot]);
}
```

Jedynym powodem, dla którego może to być interesujące, jest sposób, w jaki zamierzamy to rozwijać. Jeśli chcecie, możecie jak zwykle pobrać ten program ze strony internetowej „Practical Electronics” (plik CB-May22-02.txt).

Chwiejność

Pamiętamy, że potencjometry w naszych joystickach domyślnie ustawiają się w pozycji środkowej, a zatem spodziewalibyśmy się, że w tych położeniach wszystkie odczyty będą wynosić 511. W rzeczywistości – jak można zaobserwować na rysunku 8 – każdy potencjometr ma jakąś własną wartość środkową (LLR = ~506, LFB = ~520, LRT = ~532 itd.), co będziemy musieli uwzględnić na pewnym etapie w przyszłości. Kolejnym aspektem, którego nie brałem pod uwagę, jest to, że po popchnięciu i puszczeniu jednego z joysticków może on nie powracać do dokładnie tej samej pozycji środkowej. Jest to kolejna kwestia, nad którą będziemy musieli się w najbliższych dniach zastanowić.

Chciałbym jednak skupić się na tym, że nawet gdy nie poruszamy joystickami, ich wartości nieco się wahają z powodu szumu w układzie i jego otoczeniu. np. sygnał RRT na rysunku 8 waha się między 509, 510 a 511. Szum ten może pochodzić z różnych źródeł, np. z zasilacza albo od promieniowania elektromagnetycznego z innych urządzeń znajdujących się w pobliżu.

Z jednej strony, z powodów, które staną się jasne później, niewielkie wahania, które tutaj obserwujemy, w przypadku naszej aplikacji animatronicznej nie mają znaczenia. Z drugiej strony, jestem z zawodu specjalistą od techniki cyfrowej (stoję twarde na gruncie zer i jedynek) i uważam, że chwiejne sygnały są nieco niepokojące, więc lubię się ich pozbywać, gdy tylko mam okazję.

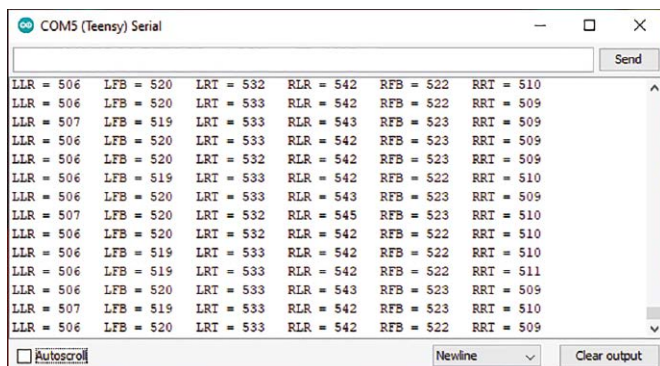
Pamiętam, jak kiedyś żartowałem ze znajomym inżynierem-analogowcem, że to dziwne, że wahania zawsze pojawiają się na najmniej znaczących bitach (LSB), a on szybko odpowiedział: „Te bity od tego właśnie są!”.

Gdzie jest źródło?

Zanim przystąpimy do działania, warto sprawdzić, czy możemy określić źródło wahań, zwłaszcza jeśli poziom szumu, jaki obserwujemy, jest znaczny. Jednym ze sposobów jest przesunięcie joysticka tak, aby odczytywana wartość wynosiła 0, co oznacza, że jego suwak znajduje się na końcu ścieżki dołączonym do masy. Jeśli szum nadal obserwujemy, prawdopodobnie ma on pochodzenie elektromagnetyczne. Następnie poruszamy joystickiem tak, aby wartość wynosiła 1023, co oznacza, że jego suwak znajduje się na końcu ścieżki dołączonym do zasilania. Jeśli nadal obserwujemy zakłócenia, mogą one pochodzić z zasilacza. Innym szybkim sprawdzeniem zakłóceń z zasilania byłoby zastąpienie zasilacza bateriami; jeśli zakłócenia znikną, to przyczyna jest oczywista.

Czy receptę stanowią kondensatory?

Co do odczytów z naszych potencjometrów – możemy uznać wahania na najmłodszych bitach za szum o wysokiej (stosunkowo) częstotliwości. Jednym ze sposobów usunięcia tego rodzaju szumu jest dodanie kondensatorów do układu. Kondensatory, poza tłumieniem szumu, powinny również wygładzać szybkie zmiany napięcia, filtrować wysokie częstotliwości i odprowadzać do masy ładunki



Rysunek 8. Jednokrotne odczytywanie potencjometrów

elektrostatyczne (ESD). Te oczekiwania sprawiają, że użyte przez nas elementy powinny być fizycznie małe i zamontowane jak najbliżej pinów wejściowych mikrokontrolera.

Szczerze mówiąc, na te tematy można by napisać niejedną książkę. Niektórzy takie książki rzeczywiście napisali. W naszym przypadku, gdybyśmy zdecydowali, że najlepszym rozwiązaniem są kondensatory, dla każdego wejścia analogowego dodalibyśmy parę połączonych równolegle kondensatorów (0,01 μ F i 1 μ F) między linią sygnału a masą, oraz taką samą parę między sygnałem a zasilaniem (rysunek 9).

Ja jednak zdecydowałem, że zupełnie nie chcę spędzać niedzielnego popołudnia na przylutowywaniu kondensatorów do mojej płytki prototypowej (nie mówiąc o tym, że nie przewidziałem dla nich miejsca), więc postanowiłem wybrać rozwiązanie programowe.

Fantastyczne filtry

Obawiam się, że musimy zrobić małą dygresję. W fizyce termin „częstotliwość” odnosi się do ilości powtarzających się zdarzeń w jakiejś jednostce czasu. Częstotliwość wyrażamy w hercach (Hz). 1 Hz odpowiada jednemu zdarzeniu na sekundę. Oznacza to, że sygnał o niskiej częstotliwości (np. głos nucącego mężczyzny) ma mniejszą liczbę herców, a sygnał o wyższej częstotliwości (np. mój pisk, gdy widzę mysz) ma więcej Hz.

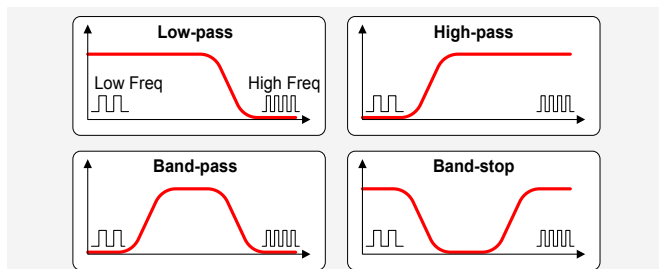
Jeśli mamy sygnał składający się z całego szeregu częstotliwości i chcemy usunąć niektóre z nich, możemy to zrobić jakimś filtrem. Cztery podstawowe typy filtrów to: filtr dolnoprzepustowy (który przepuszcza niskie częstotliwości i blokuje wysokie), filtr górno-przepustowy (przepuszcza wysokie częstotliwości i blokuje niskie), filtr pasmowy oraz filtr pasmowozaporowy (rysunek 10). Filtry takie możemy realizować sprzętowo lub programowo.

Nie bądź przeciętny

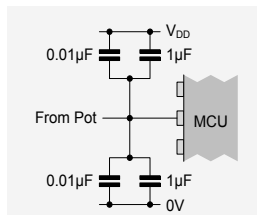
Jedno z rozwiązań naszego problemu polega na pobieraniu po kilka odczytów z każdego potencjometru, a następnie uśrednianie ich (czyli sumowanie „n” wartości, a następnie dzielenie wyniku przez „n”). Na przykład możemy zmodyfikować rdzeń naszej funkcji `ReadPots()` w następujący sposób (plik CB-May22-03.txt):

```
for (int iPot = 0; iPot < NUM_POTS; iPot++)
{
    PotVals[iPot] = 0;

    for (int iSmp = 0; iSmp < NUM_POT_SAMPLES; iSmp++)
    {
        PotVals[iPot] += analogRead(PinsPots[iPot]);
    }
}
```



Rysunek 10. Cztery podstawowe rodzaje filtrów



Rysunek 9. Użycie kondensatorów do tłumienia szumów

```
}
```

```
PotVals[iPot] /= NUM_POT_SAMPLES;
}
```

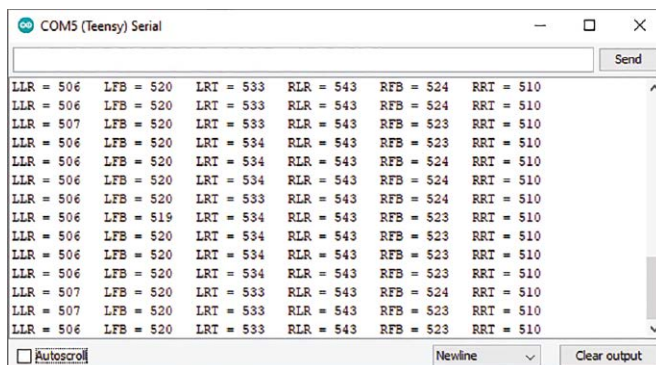
Wykonujemy tutaj komplet odczytów pierwszego potencjometru, następnie drugiego, potem trzeciego i tak dalej. Podejściem alternatywnym byłoby rozpoczęcie od pojedynczego odczytu z każdego potencjometru, następnie wykonanie drugiego odczytu z każdego potencjometru i tak dalej (plik CB-May22-04.txt). Nie jestem pewien, czy oficjalnie uznano, które podejście jest najsukuteczniejsze. Mam jednak „przezczucie”, że ta druga technika może mieć niewielką przewagę.

Jeśli ktoś się zastanawia, co my tutaj robimy – to wprowadzamy bardzo prymitywną formę filtra dolnoprzepustowego o skończonej odpowiedzi impulsowej (FIR). Określenie „skończona” bierze się stąd, że operujemy na skończonej ilości próbek. Na potrzeby tej dyskusji zacząłem od ustawienia `NUM_POT_SAMPLES` równego 5 (rysunek 11).

Jak widać, wyniki są znacznie lepsze niż w przypadku brania próbek pojedynczych, ale nadal w niektórych sygnałach występują niewielkie wahania. Moglibyśmy zwiększyć rozmiar danych – powiedzmy do dziesięciu próbek – ale spowodowałoby to pewien problem.

Do sterowania animatroniczną głową Steve używa płytki Teensy LC, wyposażonej w 32-bitowy mikrokontroler z rdzeniem Arm Cortex-M0+, taktowany zegarem 48 MHz, z 8 KB pamięci RAM i 62 KB pamięci Flash. Istotne jest to, że Teensy LC ma do dyspozycji mniej pinów GPIO, więc Steve postanowił dołączyć napięcia z potencjometrów do dwóch 4-kanałowych płytek przetworników 12-bitowych firmy Adafruit (<https://bit.ly/3w8xYt3>). Teensy LC komunikuje się z tymi płytkami poprzez 2-przewodowy interfejs I²C, co oznacza, że odczyt napięć ze wszystkich sześciu potencjometrów odbywa się po zaledwie dwóch pinach. Jest to okupione długim czasem obsługi, co jest nieodłączną cechą interfejsu I²C.

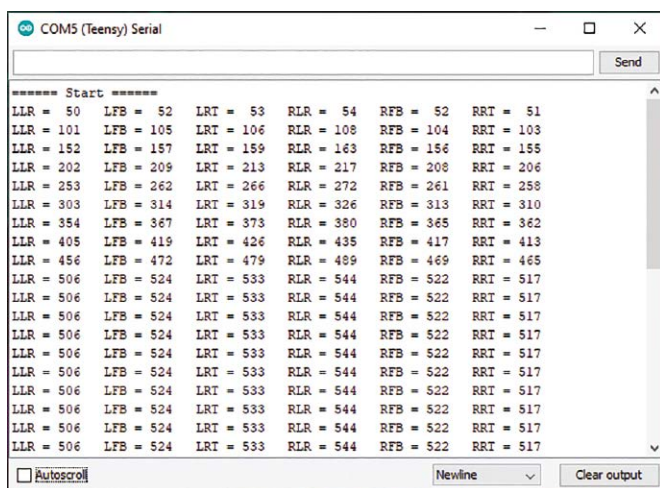
Aby zrozumieć, o jakich czasach tutaj mówimy, Steve stworzył prosty program testowy (CB-May22-05.txt) porównujący czasy (w mikrosekundach) potrzebne do odczytania napięć z sześciu potencjometrów 1- lub dziesięciokrotnie, bezpośrednio z wejścia analogowego bądź z pośrednictwem płytki ADC z łączem I²C.



Rysunek 11. Pięciokrotne odczytywanie potencjometrów

	Teensy LC	Teensy 3.6
6 Standard Analog Reads	69	43
60 Standard Analog Reads	644	298
6 I ² C Analog Reads	12,055	12,297
60 I ² C Analog Reads	120,538	123,035

Rysunek 12. Czasy odczytów wartości analogowych (w μ s)



Rysunek 13. Dziesięciokrotne odczytywanie potencjometrów z uśrednianiem kroczącym

Trochę dla zabawy użyłem programu Steve'a do przetestowania mojego Teensy 3.6. Wyniki dla obu procesorów pokazano na **rysunku 12**. Przypis redaktora: przecinki na rysunku nie są znakami dziesiętnymi, lecz, zgodnie z konwencją panującą w krajach anglosaskich, rozdzielają grupy trzycyfrowe. Wyniki z płytką ADC przekraczają zatem dziesięć tysięcy.

Nie jest zaskoczeniem – biorąc pod uwagę szybkość zegara – że Teensy 3.6 wykonuje odczyty bezpośrednie szybciej niż Teensy LC. Natomiast odczyty przez I²C zajmują mniej więcej tyle samo czasu. Ściśle rzecz biorąc, Teensy 3.6 jest w tym przypadku nieco wolniejszy, ale nie mam pojęcia, dlaczego tak się dzieje.

Kontynuacja

Nasz dylemat polega więc na tym, że chcemy zwiększać liczbę próbek, ale nie chcemy się pogodzić z wydłużaniem czasu. Istnieje na szczęście rozwiązanie tego dylematu. Polega ono na wyznaczeniu tzw. średniej kroczącej (znanej również jako średnia bieżąca lub średnia ruchoma). Średnia krocząca jest powszechnie stosowana w przypadku serii danych w celu wygładzania wahań krótkoterminowych i uwypuklania trendów długoterminowych i cykli. Jest to również rodzaj prostego filtra FIR.

W tym przypadku zaczynamy od ustawienia wszystkich wartości na 0. Za każdym razem, gdy przechodzimy przez pętlę, pobieramy tylko jedną próbkę napięcia każdego potencjometru. Trik polega na tym, że „n” ostatnich próbek przechowujemy w pamięci. Kiedy przychodzi czas na pobranie wyniku, sumujemy te „n” próbek i dzielimy wynik przez „n” (CB-May22-06.txt).

Wyniki średniej kroczącej przy użyciu rozmiaru pamięci próbek równego 10 pokazano na **rysunku 13**. Zwróćcie uwagę,

REKLAMA



Mikołaj, Bałwanek, Elf i...
wiele innych ozdób świątecznych
kupisz i zmontujesz

<https://sklep.avt.pl/pl/menu/ozdoby-485.html>



że ponieważ zaczynamy od zainicjowania wszystkich wartości próbek na 0, w ciągu pierwszych 10 powtórzeń pętli głównej wyniki stopniowo rosną. Dopiero potem utrzymują się na stabilnym poziomie.

To wszystko jest bez znaczenia

Wspomniałem wcześniej, że niewielkie wahania, które obserwowaliśmy przed zastosowaniem średniej kroczącej, w naszej aplikacji animatronicznej nie mają istotnego znaczenia. Oznaczałoby to w zasadzie, że nasze powyższe dyskusje na temat wahań są w tym kontekście nieistotne. Zaznaczałem jednak, że omawiane tutaj metody mogą w przyszłości okazać się interesujące w innych zastosowaniach – więc nie wszystko jest stracone.

Skąd powyższe stwierdzenie? Załóżmy, że nie stosujemy żadnego filtrowania i wracamy do wykorzystywania wartości nieobrobionych (rysunek 8). W jednym z poprzednich artykułów („Practical Electronics”, marzec 2020 r.; EdW 10/2025) wprowadziliśmy pojęcie modulacji szerokości impulsu (PWM). Załóżmy, że chcemy wykorzystać odczyty z potencjometrów w zakresie 0...1023 do generowania sygnałów PWM w celu sterowania diodami LED, przy czym sygnały PWM mają na przykład zakres 0...255. Oznaczałoby to, że musimy odwzorować wartości źródłowe 0...1023 na odpowiadające im wartości docelowe 0...255. W tym przypadku najłatwiejszym sposobem odwzorowania jest dzielenie wartości źródłowych przez 4 – co jest równoważne przesunięciu ich o dwa bity w prawo. Powoduje to obcięcie dwóch najmłodszych bitów. Jest to jakby forma modulacji amplitudy, odrzucająca chwilową część naszych sygnałów.

Jeśli chodzi o hobbystyczne serwomechanizmy, których używamy do sterowania naszymi animatronicznymi głowami, to nie chcemy, aby jakiegokolwiek drgania sygnałów z potencjometrów przekładały się na nerwowe drgania głowy i oczu. Zastanówmy się więc nad tym. Serwomechanizmy obecnie obracamy w zakresie tylko od około -45° do +45°, chociaż ich zakres ruchu jest większy. Nawet zakładając dokładność 0,5°, co moim zdaniem jest założeniem dość optymistycznym, oznacza to, że będę przekładał wartości 0...1023 z potencjometrów na wartości 0...179 używane do sterowania serwomechanizmami. A to oznacza, że nie zobaczę drgań nawet o rozmiarze dwóch najmłodszych bitów. W analogiczny sposób Steve będzie przekładał wartości 0...4095 ze swoich 12-bitowych przetworników ADC na te same wartości 0...179 używane do sterowania serwomechanizmami, co oznacza, że odrzuci aż cztery chwilowe bity.

Jak zwykle tylko poruszyliśmy pewien rozległy temat, ale przynajmniej mamy coś do przemyślenia, zanim spotkamy się ponownie. ■

Clive „Max” Maxfield

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, maj 2022 (www.epemag3.com)

Patronat EdW nad szkołami i uczelnianymi Kołami Naukowymi rozkwita i daje redakcji EdW impulsy zachęcające do wspierania edukacji szkolnej i uczelnianej. Działa sprzężenie zwrotne. Dostajemy mnóstwo wiadomości od uczniów, nauczycieli i studentów. Dla nich jest ta rubryka.

Wykład 36

Szumy w układach elektronicznych

Szum to na ogół niepożądany sygnał, z którym każdy elektronik ma do czynienia – i z którym musi nauczyć się żyć. Istnieje wiele rodzajów, a nawet „kolorów” szumów.

Wprowadzenie do pojęcia szumu Czym jest szum?

Słowo szum jest pojęciem ogólnym, odnoszącym się do szeregu zjawisk fizycznych i elektrycznych, które mają jednak ten sam skutek.

Szum elektroniczny to w większości przypadków niepożądany sygnał, obecny we wszystkich punktach układu, przez które przepływa prąd elektryczny lub w których występuje napięcie.

Sygnał szumu miesza się z użytecznym sygnałem, który ma zostać przetworzony przez układ, przez co pojawia się na jego wyjściu w sposób zakłócający.

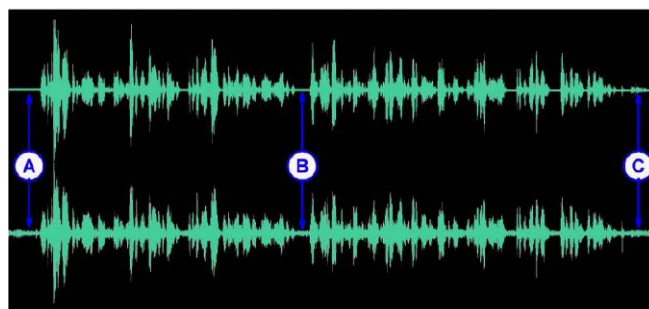
Na oscylogramach przedstawiono przykład sygnału audio bez szumu (górze) oraz z szumem (dół). Gdy sygnał użyteczny ma wartość 0 V (punkty A i B), najlepiej widać obecność szumu. Szum może nawet doprowadzić do całkowitego „zalanía” małych sygnałów użytecznych (punkt C) – wówczas oryginalny sygnał znika całkowicie, a słyszalny pozostaje jedynie szum.

Jeśli osiągnąłeś już pewien wiek i pamiętasz czasy analogowej telewizji, z pewnością znasz doskonale charakterystyczny obraz „śnieżącego” ekranu, pojawiający się po wybraniu kanału, na którym nie było nic nadawane. Szare, migoczące punkty wypełniające ekran były skutkiem szumu odbieranego przez wejście antenowe i wzmacnianego przez układy wzmacniaczy telewizora.

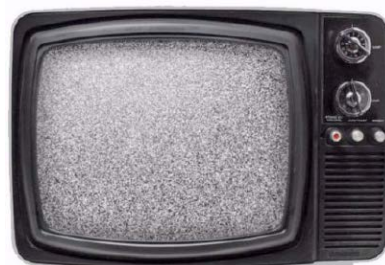
Podstawowe źródła szumów

Na podstawie dwóch wcześniej omówionych przykładów można wskazać dwie zasadnicze przyczyny powstawania szumów. Szum może pojawić się w postaci napięcia szumowego w układzie elektronicznym – przykład zakłóconego sygnału audio doskonale to ilustruje.

Szum może jednak również powstawać pierwotnie jako promieniowanie elektromagnetyczne obecne w przestrzeni wokół nas. Takie elektromagnetyczne zakłócenia zostają przekształcone w napięcie szumowe przez przewód elektryczny, który zachowuje się jak antena. Klasyczny „śnieg” na ekranie starego telewizora analogowego jest właśnie tego typowym przykładem.



Sygnał audio z szumem i bez szumu (© Mediaprocessing)



Znany obraz szumów na starej telewizji analogowej (© Ziggo)

Nieprzewidywalna amplituda

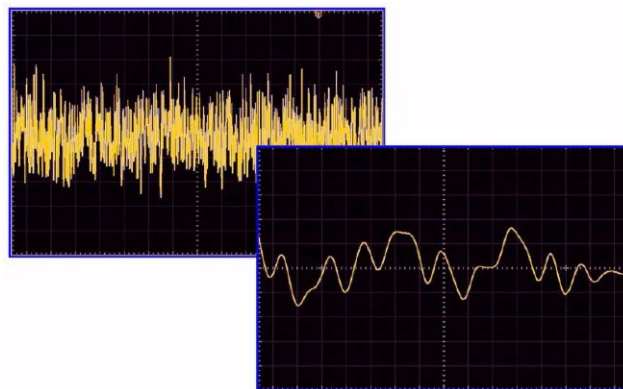
Pojęcie amplitudy odnosi się do chwilowej wartości wielkości elektrycznej – napięcia lub prądu – w danym momencie.

Szum ma zawsze postać sygnału zmiennego o całkowicie przypadkowo zmieniającej się amplitudzie. Oznacza to, że absolutnie nie da się przewidzieć ani obliczyć, jaką wartość napięcia będzie miał sygnał szumu w danym momencie. W jednej chwili może wynosić 0,1 mV, a milisekundę później już 2,3 mV.

Szum nie jest więc sygnałem okresowym – nie zawiera żadnych powtarzających się cykli czy wzorców.

Jest to doskonale widoczne podczas obserwacji na ekranie oscyloskopu: na poniższym oscylogramie nie sposób dostrzec żadnego powtarzającego się fragmentu przebiegu. Zmiany sygnału są całkowicie nieprzewidywalne.

Jeśli rozciągniemy obraz, ustawiając krótszą podstawę czasu, widać wyraźnie, jak przypadkowo zmieniają się amplitudy w funkcji czasu.



Sygnał szumowy na ekranie oscyloskopu (© 2022 Jos Verstraten)

Bardzo szerokie pasmo częstotliwości

Z przedstawionego wyżej oscylogramu można wywnioskować jeszcze jedną właściwość szumu. Widać wyraźnie, że szum charakteryzuje się nie tylko nieprzewidywalną amplitudą, ale również nieprzewidywalną częstotliwością. Krótki okres (czyli wysoka częstotliwość) może być natychmiast zastąpiony przez znacznie dłuższy (niską częstotliwość), a ten z kolei – znowu przez bardzo krótki.

Mówi się zatem, że sygnał szumowy ma bardzo szerokie pasmo.

Z czysto teoretycznego punktu widzenia w skład szumu wchodziłyby sygnały o częstotliwościach od 0 Hz aż po niemal nieskończenie wysokie wartości. To, w jaki sposób moc sygnału szumowego jest rozłożona w praktycznym zakresie częstotliwości, decyduje o tak zwanym kolorze szumu.

Do tego zagadnienia powrócimy w dalszej części artykułu.

Stosunek sygnału do szumu (SNR)

W języku angielskim używa się określenia Signal to Noise Ratio – w skrócie SNR.

Parametr ten określa stosunek mocy sygnału użytecznego do mocy szumu obecnego w tym samym sygnale:

$$\text{SNR} = \frac{P_{\text{sygnału}}}{P_{\text{szumu}}}$$

Ponieważ jest to iloraz dwóch wielkości o tych samych jednostkach, SNR jest wielkością bezwymiarową – po prostu liczbą.

Skala logarytmiczna

W praktyce moc sygnału jest zwykle o kilka rzędów wielkości większa od mocy szumu, dlatego stosunek ten przedstawia się najczęściej w skali logarytmicznej, w tzw. decybelach (dB):

$$\text{SNR}_{\text{dB}} = 10 \cdot \log_{10} \left(\frac{P_{\text{sygnału}}}{P_{\text{szumu}}} \right) \text{dB}$$

Ujęcie napięciowe

Ponieważ w technice audio częściej operuje się napięciem niż mocą sygnału, wzór ten przekształca się, korzystając z zależności:

$$P = \frac{U^2}{R}$$

Wynika stąd, że stosunek sygnału do szumu wyrażony w decybelach dla wartości skutecznych napięć ma postać:

$$\text{SNR}_{\text{dB}} = 20 \cdot \log_{10} \left(\frac{U_{\text{sygnału}}}{U_{\text{szumu}}} \right) \text{dB}$$

Aby zapewnić dobrą zrozumiałość rozmowy telefonicznej, przyjęto międzynarodowo, że na końcu łącza telefonicznego stosunek sygnału do szumu (SNR) powinien wynosić około 30, co odpowiada wartości około 15 dB.

W przypadku wzmacniacza audio przyjmuje się, że przy SNR=60 dB szum jest jeszcze bardzo cichy i ledwie słyszalny, natomiast SNR=90 dB oznacza bardzo dobrą jakość dźwięku.

Współczynnik szumów, czyli Noise Figure (NF)

Szum może być obecny w sygnale już na wejściu układu, ale może też zostać wprowadzony przez sam układ elektroniczny. Aby określić tę zależność, wprowadzono pojęcie współczynnika szumów (Noise Figure, NF).

Współczynnik szumów jest wielkością bezwymiarową, która określa stosunek SNR na wejściu układu do SNR na jego wyjściu:

$$\text{SNR}_{\text{we}} = \text{NF} \cdot \text{SNR}_{\text{wy}}$$

Każdy układ elektroniczny dodaje pewną ilość szumu do sygnału wejściowego, dlatego szum na wyjściu zawsze będzie większy niż na wejściu. Oznacza to, że stosunek sygnału do szumu na wyjściu jest zawsze mniejszy niż na wejściu, a więc NF z definicji musi być większe od 1.

Wartość współczynnika szumów często wyraża się również w decybelach (dB):

$$NF_{dB} = 10 \cdot \log_{10}(NF)$$

Rodzaje szumów

Szum może mieć różne źródła. Wyróżnia się między innymi:

- szum termiczny (Johnson-Nyquist noise),
- szum atmosferyczny,
- szum galaktyczny,
- szum wyładowań (tzw. „hagelruis”),
- szum wytwarzany przez człowieka,
- szum kwantyzacji,
- szum tranzystorowy,
- szum podziału prądu,
- szum impulsowy (burst noise),
- szum migotania (flicker noise, 1/f),
- szum czasu przelotu nośników (transit-time noise).

Wszystkie te rodzaje szumów zostaną omówione w dalszej części artykułu.

Kolory szumów

Podobnie jak światło może mieć różne kolory zależnie od długości fali promieniowania elektromagnetycznego, tak również szумы mogą mieć różne „kolory”, w zależności od ich widma częstotliwościowego.

Wyróżnia się:

- szum biały,
- szum różowy,
- szum brązowy,
- szum niebieski,
- szum fioletowy,
- szum szary.

O tych odmianach szumów również opowiemy szerzej w dalszej części artykułu.

Szum termiczny (thermal noise, Johnson–Nyquist noise)

Historia

Ten rodzaj szumu został odkryty w 1926 roku przez Johna B. Johnsona. Jednak to Harry Nyquist był tym, który potrafił zjawisko to opisać matematycznie.

Wszystko się porusza i drga...

Jak wiadomo, materia zbudowana jest z atomów, których jądra otaczają krążące wokół elektrony. Pod wpływem temperatury elektrony zaczynają drgać, co sprawia, że mogą łatwo przeskakiwać z jednego atomu na inny. Te całkowicie przypadkowe ruchy powodują, że w przewodniku chwilowo i lokalnie powstaje nadmiar elektronów. Ułamek sekundy później nadmiar ten znika. Takie nieustanne, losowe przemieszczanie się elektronów powoduje powstawanie bardzo małych różnic napięcia pomiędzy końcami przewodnika. Zjawisko to nazywa się szumem termicznym.

Pasmo częstotliwości

Szum termiczny ma bardzo szerokie pasmo – praktycznie obejmuje cały zakres częstotliwości. W rzeczywistości jego pasmo jest ograniczone jedynie przez pasmo zastosowanego układu elektronicznego, w którym występuje.

Energia, moc i napięcie szumu

Energia zawarta w szumie jest wprost proporcjonalna do temperatury i opisana wzorem:

$$E_{szumu} = k \cdot T [J]$$

gdzie:

k – stała Boltzmanna,

$$k = 1,38 \cdot 10^{-23} \frac{J}{K}$$

T – temperatura w kelwinach [K].

Moc szumu wyraża się zależnością:

$$P_{szumu} = k \cdot T \cdot B [W]$$

gdzie

B – oznacza szerokość pasma w hercach [Hz].

Ponieważ między mocą a napięciem istnieje zależność kwadratowa:

$$P_{\text{szumu}} = \frac{(U_{\text{szumu}})^2}{R}$$

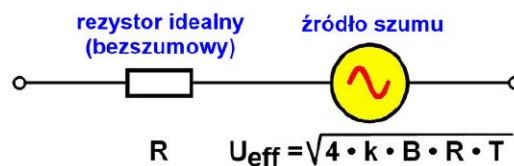
można obliczyć wartość skuteczną napięcia szumowego ze wzoru:

$$U_{\text{szumu}} = \sqrt{4 \cdot k \cdot B \cdot R \cdot T}$$

gdzie:

R – rezystancja elementu, w którym powstaje szum.

To napięcie szumowe można sobie wyobrazić jako źródło napięcia przemiennego połączone szeregowo z idealnie pozbawionym szumów rezystorem.



Schematyczne przedstawienie szumu termicznego na rezystorze (© 2022 Jos Verstraten)

Teoretyczna minimalna wartość napięcia szumu termicznego

Szum termiczny jest zjawiskiem nieuniknionym. Tylko w temperaturze 0 kelwinów napięcie szumowe również wynosiłoby zero – sytuacja taka jednak nigdy nie występuje w praktyce, więc z szumem termicznym trzeba po prostu nauczyć się żyć.

Dla orientacji: napięcie szumowe wynosi około 1,8 μV przy temperaturze pokojowej 20°C (293 K), dla rezystora 10 kΩ, przy szerokości pasma 20 kHz.

Równoważna rezystancja szumowa

Powstawanie napięcia szumowego na rezystorze pod wpływem temperatury było jednym z pierwszych zjawisk szumowych, które udało się w pełni wyjaśnić teoretycznie. Z tego powodu wprowadzono pojęcie równoważnej rezystancji szumowej (equivalent noise resistance), pozwalające sprowadzić wszelkie rodzaje szumów do postaci jednej „szumiącej” rezystancji.

Na poniższym rysunku przedstawiono to w sposób uproszczony: po lewej stronie znajduje się układ, który generuje pewną ilość szumu o nieznanym pochodzeniu, a po prawej – ten sam układ, ale pozbawiony własnych szumów, zakończony rezystorem R_zast, którego wartość tak dobrano, aby wytwarzał taką samą ilość szumu, jak układ po lewej.

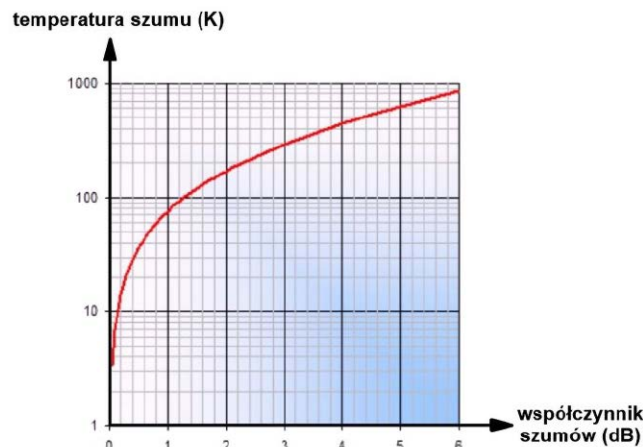
Wartość R_zast określa więc równoważną rezystancję szumową badanego układu.



Pojęcie „równoważnej rezystancji szumowej” (© 2022 Jos Verstraten)

Temperatura szumowa

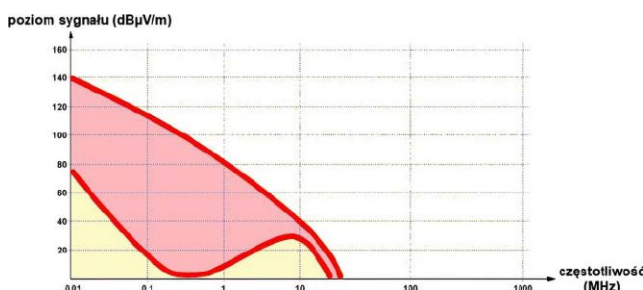
Ponieważ szum termiczny generowany w rezystorze zależy od jego temperatury, wprowadzono inny sposób określania jego wielkości – pojęcie temperatury szumowej (noise temperature). Matematyczne wyprowadzenie tego pojęcia wymaga dość zaawansowanej wiedzy, dlatego w tym artykule ograniczymy się do wyjaśnienia istoty zjawiska słowami. Temperatura szumowa układu lub elementu, który wytwarza szum, jest taką temperaturą, jaką musiałby mieć rezystor, aby wytwarzał dokładnie tyle samo szumu co dany układ. Poniższy wykres pokazuje zależność pomiędzy temperaturą szumową a współczynnikiem szumów (w dB).



Zależność między temperaturą szumu termicznego a współczynnikiem szumów (© 2022 Jos Verstraten)

Szum atmosferyczny (atmospheric noise, static noise) Powstawanie

Szum atmosferyczny powstaje w atmosferze otaczającej Ziemię, a jego głównym źródłem są wyładowania atmosferyczne – błyskawice. Każde wyładowanie generuje pole elektromagnetyczne, które indukuje napięcie w każdym przewodniku elektrycznym znajdującym się w jego zasięgu. Ponieważ pola elektromagnetyczne mogą rozprzestrzeniać się na duże odległości, szum ten jest wynikiem sumowania się wpływu wielu błyskawic – nie tylko w pobliżu odbiornika, lecz również w bardzo odległych rejonach. Źródłem są zarówno wyładowania między chmurami, jak i pomiędzy chmurami a ziemią – a tych jest znacznie więcej, niż mogłoby się wydawać! Na całym świecie dochodzi każdego dnia do milionów wyładowań atmosferycznych, co odpowiada mniej więcej stu uderzeniom pioruna na sekundę. Energia elektromagnetyczna z tych wyładowań dociera do nas różnymi drogami propagacji, co powoduje zjawisko dyspersji – typowe dla wszystkich fal, także świetlnych. Dlatego odległe burze odbieramy nie



Rozkład energii szumu atmosferycznego w funkcji częstotliwości (© 2022 Jos Verstraten)

jako pojedyncze impulsy, lecz jako ciągle tło szumowe. Do szumu atmosferycznego przyczyniają się również tzw. wyładowania koronowe, zachodzące wysoko w atmosferze.

Pasmo częstotliwości

Szum atmosferyczny nie ma bardzo szerokiego pasma – jego energia skupia się głównie w zakresie niskich częstotliwości, do około 20 MHz. Na wykresie pokazano, jak energia szumu jest rozłożona w funkcji częstotliwości.

Szum galaktyczny (cosmic noise, solar noise) Powstawanie szumu galaktycznego

Szum galaktyczny pochodzi z przestrzeni kosmicznej i ma wiele różnych źródeł. Wszystkie gwiazdy – w tym również Słońce – emitują różne formy promieniowania elektromagnetycznego oraz naładowane cząstki, które docierają do Ziemi. Fale te indukują napięcia szumowe w antenach oraz w długich przewodach elektrycznych. Znaczną część tego szumu powoduje tzw. promieniowanie Czerenkowa (Cherenkov radiation). W górnych warstwach atmosfery (na wysokości około 10 km i więcej) wysokoenergetyczne cząstki są pobudzane przez ekstremalnie energetyczne promieniowanie gamma, pochodzące z galaktyk i supernowych. Podczas wędrówki w kierunku powierzchni Ziemi cząstki te wywołują deszcz wtórnych, wysokoenergetycznych cząstek naładowanych, które z kolei powodują krótkie błyski światła w górnych warstwach atmosfery. Promieniowanie Czerenkowa można obserwować jedynie z przestrzeni kosmicznej lub przy pomocy specjalnych teleskopów Czerenkowa umieszczonych na powierzchni Ziemi.

Pasmo częstotliwości szumu galaktycznego

Szum galaktyczny ma bardzo szerokie pasmo – sięga powyżej 100 MHz – jednak jego natężenie jest znacznie mniejsze niż w przypadku szumu atmosferycznego.



Rozkład energii szumu galaktycznego w funkcji częstotliwości (© 2022 Jos Verstraten)

Szum wyładowań (shot noise, szum Schottky'ego)

Odkrycie szumu wyładowań

Szum wyładowań został po raz pierwszy opisany w 1918 roku przez Waltera Schottky'ego, który badał niewyjaśnione wcześniej wahania prądu w lampach elektronowych.

Szum wyładowań jako zjawisko kwantowe

Zjawiska tego nie da się wyjaśnić klasyczną mechaniką. W ujęciu klasycznym przepływ prądu jest procesem ciągłym: gdy między dwoma punktami występuje różnica potencjałów, elektrony płyną przez przewodnik w sposób ciągły – niczym woda płynąca ze stałą prędkością w rzece. Dzięki mechanice kwantowej wiemy dziś, że taki obraz jest nieprawidłowy. Uwolnienie elektronu z atomu jest procesem statystycznym – nie można dokładnie przewidzieć, kiedy dany atom „wypuści” elektron. W efekcie może się zdarzyć, że w chwili t_1 zostanie wyemitowana pewna liczba n elektronów, a w chwili t_2 – już tylko m.

W skali makroskopowej, obserwowanej przyrządami pomiarowymi, wydaje się, że prąd jest stały. Jednak w skali mikroskopowej liczba elektronów przepływających przez przewodnik ciągle się zmienia. To właśnie te losowe fluktuacje liczby elektronów powodują pozornie niewytłumaczalne wahania prądu w lampach elektronowych i w półprzewodnikach – i zjawisko to nazwano szumem wyładowań (shot noise).

Wielkość szumu wyładowań

W praktyce szum wyładowań stanowi jedynie niewielką część całkowitego szumu w typowym układzie elektronicznym. Dla orientacji: prąd 1 A odpowiada przepływowi około $6,24 \cdot 10^{18}$ elektronów na sekundę.

Nawet jeśli – na skutek zjawisk kwantowych – liczba ta zmieni się o kilka miliardów na sekundę, efekt ten będzie niemal niezauważalny i nie wpłynie istotnie na stałość prądu. Szum wyładowań jest jednak całkowicie niezależny od temperatury i szerokości pasma. W bardzo czułych układach, które są silnie chłodzone w celu zminimalizowania szumu termicznego, szum wyładowań może stać się czynnikiem dominującym.

Szum wytwarzany przez człowieka (man-made noise)

Szum generowany przez człowieka pochodzi dziś z niezliczonych źródeł, obecnych w każdym domu i na każdej ulicy. Należą do nich między innymi: zakłócenia powodowane przez urządzenia gospodarstwa domowego sterowane elektronicznie, ściemniacze oświetlenia, lampy LED, zasilacze impulsowe (w tym adaptory sieciowe), różnego rodzaju urządzenia elektroniczne z przetwornicami, oraz iskrzenie powstające przy pracy pociągów i tramwajów. Wszystkie te źródła zakłóceń mogą powodować szumy tylko wtedy, gdy ich energia jest promieniowana w przestrzeń przez anteny. A anten – w nowoczesnym domu – naprawdę nie brakuje! Większość urządzeń elektrycznych i elektronicznych ma nieekranowany przewód zasilający lub połączenie z innym urządzeniem, a takie przewody stanowią idealne anteny dla częstotliwości generowanych przez te źródła zakłóceń.

Zakres częstotliwości szumów pochodzenia ludzkiego

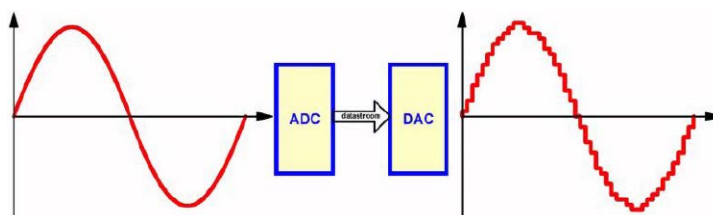
Szum wytwarzany przez człowieka występuje głównie w zakresie od 0,1 MHz do 4 MHz. Nawet w spokojnej, wiejskiej okolicy, z dala od dużych miast, jego poziom jest przeciętnie o około 50 dB wyższy od poziomu szumów naturalnych. W gęsto zabudowanych obszarach miejskich wartość ta może wzrosnąć o kolejne 20...30 dB.

Szum kwantyzacji (quantization noise) Powstawanie szumu kwantyzacji

Sygnał cyfrowy, z definicji, może przyjmować tylko dwie wartości: L (niski poziom logiczny) lub H (wysoki poziom logiczny). Oznacza to, że nawet kombinacja wielu sygnałów cyfrowych może reprezentować jedynie ograniczoną liczbę możliwych stanów. Kiedy sygnał analogowy jest przekształcany na kod cyfrowy, niezależnie od jego złożoności, otrzymany zapis cyfrowy może jedynie przybliżać chwilową wartość sygnału analogowego. Nieskończona liczba możliwych wartości napięcia w świecie analogowym zostaje odwzorowana w skończoną liczbę kombinacji kodów cyfrowych.

To właśnie różnica między przetwarzaniem analogowym a cyfrowym stanowi istotę działania przetworników ADC i DAC i określana jest mianem kwantyzacji (quantization). Wyjściowe kody przetwornika ADC są w każdej chwili jedynie przybliżeniem rzeczywistego napięcia wejściowego. Jeśli następnie te kody zostaną ponownie przekształcone w sygnał analogowy przez DAC, otrzymany sygnał będzie również przybliżeniem oryginału, a nie jego idealnym odwzorowaniem. Dlatego mówi się, że proces kwantyzacji zawsze wytwarza „schodkowe” przybliżenie przebiegu sygnału analogowego.

Na rysunku pokazano tę schodkową aproksymację (dla przejrzystości mocno przerysowaną). W praktyce oczywiście stosuje się więcej niż cztery bity oraz znacznie większą liczbę próbek na sekundę.



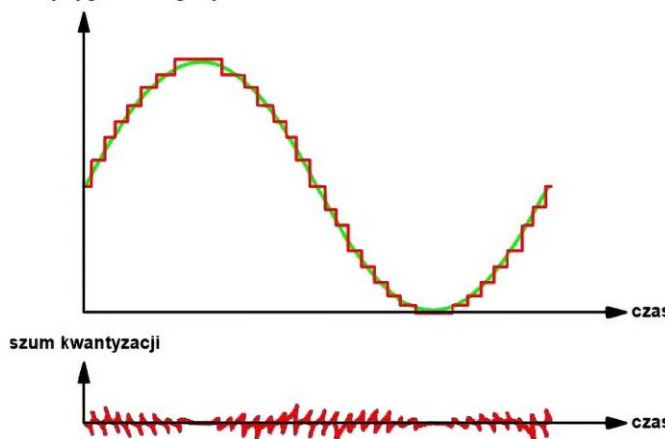
Powstawanie szumu kwantyzacji (© 2022 Jos Verstraten)

Szum kwantyzacji

Pomiędzy oryginalnym sygnałem analogowym a sygnałem analogowym odtworzonym po przetworzeniu zawsze występuje pewna różnica – czyli zniekształcenie. Można je wprowadzić zminimalizować, ale nigdy całkowicie wyeliminować. Oznacza to, że każdy sygnał analogowy, który zostaje przetworzony cyfrowo przez układ ADC → DAC, jest z samej zasady zniekształcony. To wrodzone ograniczenie nazywa się szumem kwantyzacji (quantization noise). Szum kwantyzacji stanowi jedną z najważniejszych charakterystyk systemów cyfrowego przetwarzania sygnałów.

Na rysunku widać: na górnym wykresie: oryginalny sygnał analogowy (zielony) oraz sygnał odtworzony (czerwony), na dolnym wykresie: samą składową szumową, czyli różnicę między tymi dwoma sygnałami.

odtworzony sygnał analogowy



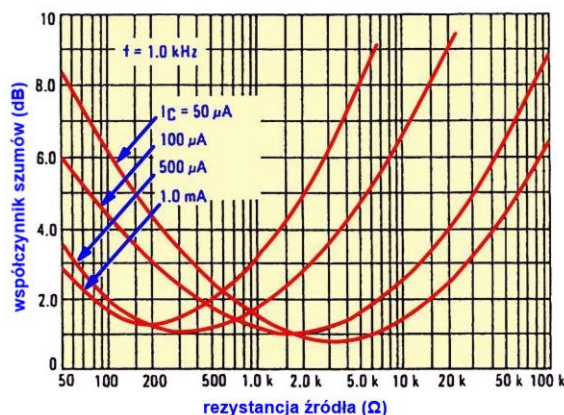
Szum tranzystorowy (semiconductor noise) Prąd kolektora a szum

Tranzystory wytwarzają wewnętrznie pewną ilość szumu, który zależy od ich warunków pracy. Istnieje określona zależność między prądem kolektora a poziomem generowanego szumu, a zależność tę przedstawia charakterystyka szumowa tranzystora, zwykle publikowana w jego dokumentacji technicznej. Z wykresu (patrz poniżej) wynika jednoznacznie, że poziom szumu może zmieniać się nawet czterokrotnie, w zależności od wartości prądu kolektora. Dla typowo niskoszumnego tranzystora, takiego jak BC109, przy impedancji źródła 2 kΩ, minimalny poziom własnego szumu półprzewodnika występuje, gdy prąd kolektora wynosi około 100 μA. Jeśli impedancja źródła spada do 200 Ω, należy zwiększyć prąd kolektora do około 1 mA, aby uzyskać najniższy możliwy poziom szumu tranzystorowego.

Szum a impedancja źródła

Z przedstawionego wykresu jasno wynika, że poziom szumu gwałtownie rośnie, gdy impedancja źródła sygnału jest bardzo niska. Niektóre

Przykład szumu kwantyzacji w odtworzonym sygnał analogowym (© 2022 Jos Verstraten)



Zależność między impedancją źródła, prądem kolektora a poziomem szumu tranzystora (© 2022 Jos Verstraten)

źródła, takie jak wkładki typu moving coil w gramofonach, mają impedancję rzędu kilku omów. Są to elementy magnetodynamiczne, w których cewka składa się z kilku zwojów cienkiego przewodnika wykonanego ze specjalnego stopu metalu i połączonego z igłą wkładki. Nietrudno zauważyć, że takie źródło sygnału ma bardzo małą impedancję wyjściową. Dodatkowo, napięcie generowane przez taką wkładkę jest niezwykle niskie – wartości rzędu 150 μV nie należą do rzadkości. Aby uzyskać użyteczny poziom sygnału, trzeba go silnie wzmacnić, co sprawia, że szum pierwszego stopnia wzmacniacza staje się wyjątkowo istotny – ponieważ zostaje on również wzmacniony przez kolejne stopnie. Z tego powodu bardzo niskoomowych i bardzo niskonapięciowych źródeł sygnału nie można wzmacniać za pomocą standardowych układów. W takich przypadkach stosuje się specjalne rozwiązania, na przykład wzmacniacze pracujące równolegle, które pozwalają ograniczyć wpływ szumu własnego.

Szum podziału prądu (partition noise) Jak powstaje szum podziału

Szum podziału powstaje wtedy, gdy prąd elektryczny może rozdzielać się pomiędzy dwa przewodniki. W skali makroskopowej wszystko wydaje się w porządku: jeśli prąd o wartości dokładnie 1,00 A rozdziela się między dwa identyczne rezystory, to w każdej gałęzi popłynie 0,50 A. Jednak w skali mikroskopowej sytuacja wygląda zupełnie inaczej – niektóre elektrony, które „powinny” przepłynąć jednym torem, pod wpływem przypadkowych czynników zmieniają kierunek i trafiają do drugiego toru, i odwrotnie. Na te losowe przejścia znaczny wpływ ma temperatura, a gdy rozmiary przewodników są bardzo małe, a prądy niskie, zaczynają dominować efekty kwantowe, które dodatkowo wzmacniają to zjawisko.

Zjawisko znane z czasów lamp elektronowych

Szum podziału był dobrze znanym problemem w układach lampowych, szczególnie w przypadku pentod, gdzie prąd emitowany przez katodę rozdziela się między anodę a siatkę ekranującą. Dodatkowo elektrony emitowane z katody mają dużą energię kinetyczną („są pobudzone”), co sprawia, że ich tor ruchu jest trudny do przewidzenia, a zjawisko szumu podziału – bardziej wyraźne.

Szum impulsowy (burst noise, popcorn noise) Odkrycie szumu impulsowego

Szum impulsowy po raz pierwszy zaobserwowano w pierwszych diodach stykowych, a dokładnie zbadano go, gdy okazało się, że pierwszy zintegrowany wzmacniacz operacyjny – $\mu\text{A}709$ – jest szczególnie podatny na to zjawisko. Z punktu widzenia matematyki zjawisko to można opisać za pomocą łańcucha Markowa, opracowanego przez rosyjskiego matematyka Andrieja Markowa. Opisuje on system, który porusza się pomiędzy skończoną liczbą stanów, wykonując losowe, krokowe przejścia z jednego stanu do drugiego.

Powstawanie szumu impulsowego

Szum impulsowy występuje w cienkich warstwach izolacyjnych tlenku bramki, powszechnie spotykanych w półprzewodnikach typu FET oraz w strukturach układów scalonych, gdzie warstwy te izolują od siebie przewodzące ścieżki. W takich warstwach mogą się pojawiać nieprzewidywalne przejścia między dwoma poziomami napięcia. Zmiany te mają charakter skokowy – amplituda skoku wynosi zaledwie kilkadziesiąt mikrowoltów (μV), jednak są łatwo wykrywalne, jeśli wiadomo, czego należy szukać. Na poniższym oscylogramie przedstawiono przykład szumu impulsowego. Widać wyraźnie, że zjawisko to występuje nieregularnie, ale zawsze przełącza się pomiędzy dwoma stałymi poziomami napięcia.



Szum impulsowy (burst noise) uwidoczniony na oscylogramie (© 2022 Jos Verstraten)

Przyczyny szumu impulsowego

Jedną z głównych przyczyn powstawania tego rodzaju szumu jest obecność zanieczyszczeń w materiale, z którego wykonane są warstwy tlenkowe (oxide barriers) stosowane w strukturach półprzewodnikowych.

„Popcorn noise”

Szum impulsowy bywa też nazywany „szumem popcornowym” (popcorn noise) ze względu na charakterystyczny dźwięk, przypominający odgłos pękającego ziarna kukurydzy, jaki można usłyszeć w wzmacniaczach audio dotkniętych tym zjawiskiem.

Szum migotania (flicker noise, 1/f noise) Czym jest szum migotania?

Szum migotania, nazywany również szumem 1/f, to wciąż nie do końca poznane zjawisko, występujące praktycznie we wszystkich elementach elektronicznych. Po raz pierwszy zaobserwował je Walter Schottky, który w 1918 roku opisał je w czasopiśmie Annalen der Physik. Szum ten pojawia się od chwili, gdy przez element zaczyna płynąć prąd, a amplituda szumu maleje proporcjonalnie do wzrostu częstotliwości pomiaru. Dlatego właśnie stosuje się oznaczenie 1/f, ponieważ amplituda szumu jest odwrotnie proporcjonalna do częstotliwości.

Szum migotania a szum termiczny

W układach niskoczęstotliwościowych szum migotania dominuje w najniższym zakresie częstotliwości. Wraz ze wzrostem częstotliwości jego udział stopniowo maleje, aż przy pewnej częstotliwości $f(c)$ przewagę zaczyna mieć szum termiczny. Dla tranzystorów bipolarnych oraz JFET-ów częstotliwość ta wynosi zwykle około 2 kHz.

Proporcjonalność do prądu

Amplituda szumu rośnie proporcjonalnie do natężenia prądu, który przepływa przez dany element elektroniczny.

Szum czasu przelotu (transit-time noise) Znaczenie przy wysokich częstotliwościach

Ten rodzaj szumu staje się istotny, gdy czas przelotu nośnika ładunku (np. elektronu lub dziury) przez materiał półprzewodnika staje się porównywalny z okresem sygnału, który jest do tego elementu doprowadzany. Czas przelotu to czas potrzebny nośnikowi ładunku, aby pokonać warstwę półprzewodnika między elektrodami. Jeśli ten czas jest zbliżony do okresu sygnału, może dojść do interferencji pomiędzy zmianami amplitudy sygnału w czasie a ruchem nośników w półprzewodniku w tym samym przedziale czasowym. Zjawisko to jest więc szczególnie istotne w układach wysokoczęstotliwościowych.

Proporcjonalność do częstotliwości

Amplituda tego rodzaju szumu rośnie w przybliżeniu liniowo wraz ze wzrostem częstotliwości sygnału. Zjawisko to staje się istotnym czynnikiem przy przetwarzaniu sygnałów w zakresie mikrofalowym, czyli przy częstotliwościach powyżej 300 MHz.

Biały szum (white noise) Definicja

Jeśli przez bardzo długi czas obserwować sygnał białego szumu za pomocą analizatora widma o działaniu całkującym, otrzymany wykres na ekranie będzie prostą poziomą linią. Mówi się wówczas, że biały szum ma stałą gęstość widmową mocy, czyli że wszystkie pasma częstotliwości są reprezentowane w jednakowym stopniu. Graficznie można to przedstawić jako poziomą linię na wykresie amplitudy w funkcji częstotliwości. Nazwa biały szum pochodzi z analogii do białego światła, które również jest mieszaniną wszystkich długości fal widzialnych dla ludzkiego oka.

Biały szum w praktyce

Z czysto matematycznego punktu widzenia biały szum definiuje się jako sygnał losowy o nieskończenie szerokim paśmie częstotliwości. Oczywiście jest to jedynie model teoretyczny. W praktyce pasmo białego szumu jest ograniczone: przez pasmo źródła, które go generuje, przez pasmo układów, które ten sygnał przetwarzają, oraz przez pasmo przyrządów pomiarowych, które go rejestrują. Sygnał uznaje się za biały szum, jeśli jego charakterystyka amplituda–częstotliwość jest płaska w zakresie częstotliwości istotnym dla danej aplikacji. W technice audio oznacza to, że charakterystyka powinna być płaska w zakresie 20 Hz...20 kHz.

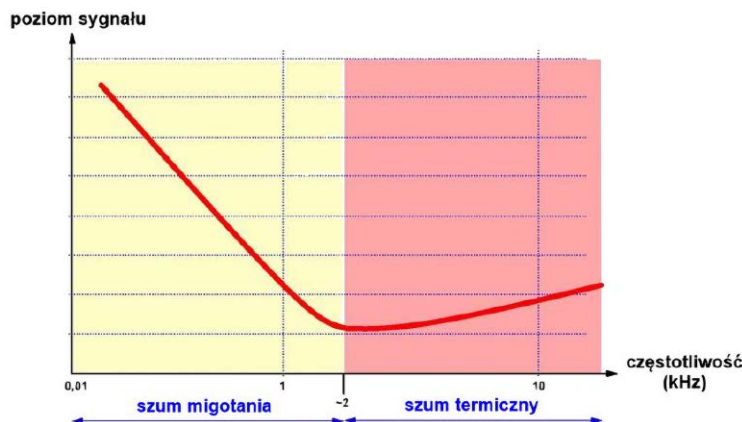
Zastosowania białego szumu

Takie sygnały stanowią doskonałe narzędzie do badania właściwości akustycznych różnych pomieszczeń i systemów. Jeśli biały szum zostanie odtworzony przez wysokiej jakości wzmacniacz i głośnik o płaskiej charakterystyce częstotliwościowej, a następnie zarejestrowany w różnych punktach sali mikrofonem pomiarowym i przeanalizowany przy użyciu analizatora widma, można w ten sposób określić wiele parametrów akustycznych pomieszczenia.

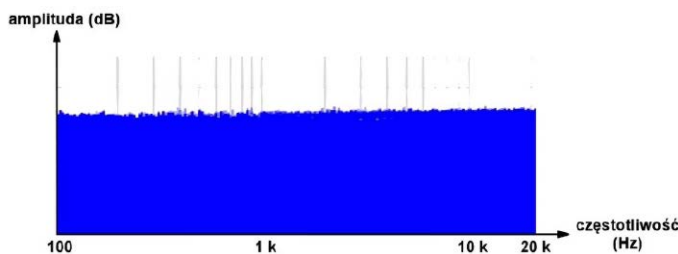
W informatyce biały szum bywa używany do generowania rzeczywiście losowych liczb, czego nie da się osiągnąć samymi algorytmami programowymi – te bowiem dają tylko liczby pseudolosowe.

Generatory białego szumu znajdują też zastosowanie w medycynie alternatywnej, na przykład w terapii tinnitus (szumów usznych) i bezsenności. Tinnitus to schorzenie, w którym pacjent słyszy nieistniejące dźwięki – szumy, piski lub gwizdy – pochodzące z wewnętrznych procesów w uchu lub w mózgu. Sama nazwa tinnitus pochodzi od łacińskiego tinnire, czyli „dzwonić”.

Niektóre badania porównawcze sugerują również, że odsłuch białego szumu w paśmie audio może poprawiać funkcje poznawcze – zwiększać koncentrację i zdolność przetwarzania informacji.



Przejście od szumu migotania do szumu termicznego (© 2022 Jos Verstraten)

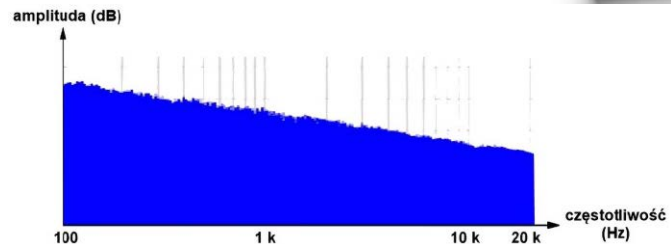


Charakterystyka amplitudowo-częstotliwościowa białego szumu (© 2012 Mick Ohrberg, opracowanie 2022: Jos Verstraten)

Różowy szum (pink noise, $1/f$ noise)

Definicja

Charakterystyka amplituda–częstotliwość różowego szumu, przedstawiona na logarytmicznej osi częstotliwości, wykazuje liniowy spadek o 10 dB na każdą dekadę częstotliwości. Gęstość widmowa mocy maleje w stosunku do białego szumu z nachyleniem 3 dB na oktawę, czyli jest odwrotnie proporcjonalna do częstotliwości ($1/f$). Z tego powodu różowy szum bywa często nazywany właśnie „szumem $1/f$ ”.



Charakterystyka amplitudowo–częstotliwościowa różowego szumu
(© 2012 Mick Ohrberg, opracowanie 2022: Jos Verstraten)

Czym jest oktawa?

Pojęcie „oktawa” pochodzi z teorii muzyki. Instrument muzyczny może wytwarzać dźwięki (nuty), z których każda ma określoną częstotliwość i nazwę: C, D, E, F, G, A, H – czyli do, re, mi, fa, sol, la, si. Ta sekwencja siedmiu nut powtarza się wielokrotnie, od najniższych do najwyższych dźwięków, które instrument jest w stanie zagrać. Poszczególne zakresy oznacza się cyframi. Częstotliwości podwajają się przy każdej kolejnej oktawie. Na przykład: nuta A_4 (la z czwartej oktawy) ma częstotliwość 440 Hz, nuta A_5 ma 880 Hz, a nuta A_3 – 220 Hz. Wszystkie te nuty tworzą skalę muzyczną. Odstęp częstotliwości między pierwszą a ósmą nutą tej skali nazywa się oktawą. Słowo oktawa pochodzi od greckiego okta, oznaczającego „osiem”. Dwie nuty oddalone o oktawę mają tę samą nazwę – na przykład od C do kolejnego C powyżej lub poniżej.

Jaki jest sens tego konkretnego tłumienia (czyli spadku o 3 dB na oktawę)?

Pisaliśmy już, że biały szum często wykorzystuje się do pomiaru charakterystyki przenoszenia pomieszczenia lub urządzenia audio. W tego typu pomiarach stosuje się dużą liczbę filtrów, które rozkładają badany sygnał szumu na wiele wąskich pasm częstotliwości. Następnie mierzy się, ile sygnału przypada na każde z tych pasm.

Problem polega na tym, że bardzo trudno jest zaprojektować zestaw filtrów o różnych częstotliwościach środkowych, ale o jednakowej szerokości pasma. W praktyce stosuje się filtry o stałym współczynniku dobroci Q, co oznacza, że im wyższa częstotliwość środkowa filtru, tym szersze jest jego pasmo przenoszenia. W efekcie, jeśli analizujemy biały szum za pomocą takiego zestawu filtrów o stałym Q, okaże się, że sygnał na wyjściu filtrów rośnie wraz ze wzrostem częstotliwości środkowej. To logiczne: im wyższa częstotliwość, tym szersze pasmo przepuszczane przez filtr, a więc tym więcej składowych szumu zostaje przepuszczonych. Można matematycznie wykazać, że podwojenie częstotliwości środkowej powoduje wzrost napięcia wyjściowego o 3 dB. Aby temu przeciwdziałać, należy przepuścić biały szum przez filtr dolnoprzepustowy o charakterystyce spadku 3 dB na oktawę. Oktawa to odstęp między daną częstotliwością a częstotliwością dwukrotnie wyższą, dlatego filtr powinien tłumić sygnał o 3 dB przy każdym podwojeniu częstotliwości. W wyniku takiego przetwarzania otrzymujemy różowy szum, który zawiera proporcjonalnie więcej składowych niskoczęstotliwościowych niż wysokoczęstotliwościowych. Zakres akustyczny od 20 Hz do 20 480 Hz obejmuje około dziesięć oktaw, więc najwyższe częstotliwości w różowym szumie będą tłumione o około 30 dB w stosunku do najniższych.

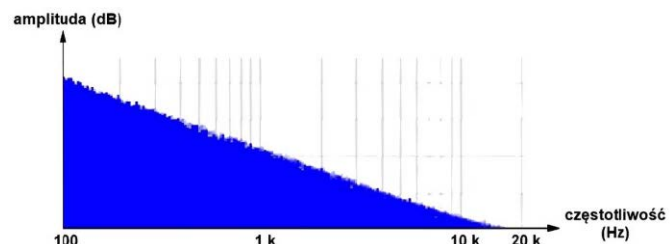
Wniosek

Różowy szum otrzymuje się z białego szumu po to, aby można było w prosty sposób wykonywać pomiary widma dźwięku przy użyciu nieskomplikowanych analogowych filtrów pasmowych.

Szum brązowy (Brownian noise, brown noise, $1/f^2$ noise)

Definicja

Szum brązowy powstaje z białego szumu po przepuszczeniu go przez filtr dolnoprzepustowy, który tłumia sygnał o 6 dB na każdą oktawę. W całym zakresie częstotliwości akustycznych od 20 Hz do 20 480 Hz biały szum powinien więc zostać stopniowo stłumiony o maksymalnie 60 dB. Gęstość mocy (czyli ilość energii przypadającej na jednostkę pasma) jest odwrotnie proporcjonalna do kwadratu częstotliwości ($1/f^2$). Oznacza to, że gęstość mocy maleje czterokrotnie, czyli o 6 dB, przy każdym podwojeniu częstotliwości (czyli o jedną oktawę).



Charakterystyka amplitudowo–częstotliwościowa szumu brązowego
(© 2012 Mick Ohrberg, edycja 2022: Jos Verstraten)

Nieporozumienia dotyczące nazwy

W języku niderlandzkim mówi się o „szumie brązowym” (bruine ruis), natomiast w języku angielskim można spotkać zarówno określenie „brown noise”, jak i „Brownian noise”. To drugie odnosi się do szkockiego botanika Roberta Browna, odkrywcy ruchów Browna (Brownian motion). Z matematycznego punktu widzenia ruchy Browna odpowiadają szumowi o widmie $1/f^2$, czyli właśnie szumowi brązowemu. Ponieważ pozostałe rodzaje szumów nazwano od kolorów, a słowo brown w języku angielskim oznacza również „brązowy”, było naturalne, że także ten rodzaj szumu otrzyma nazwę związaną z kolorem.

Zastosowania szumu brązowego

Oprócz zastosowań w badaniach naukowych, szum brązowy jest również wykorzystywany w terapiach relaksacyjnych i medytacyjnych, na przykład jako dźwięk tła podczas sesji mindfulness lub w różnego rodzaju praktykach redukujących stres. Zdarza się też, że szum brązowy stosuje się jako rodzaj „antydzwięku” w walce z uciążliwymi dźwiękami o bardzo niskiej częstotliwości (LFN – Low-Frequency Noise). Osoby narażone na ciągłe działanie takich dźwięków często skarżą się na szereg dolegliwości fizycznych i zdrowotnych, takich jak: problemy ze snem, zmęczenie, bóle głowy, duszności, uczucie ucisku w klatce piersiowej lub w uszach, kołatanie serca, wrażenie drgań w ciele.

Szum niebieski (blue noise)

Definicja

Gęstość mocy szumu niebieskiego rośnie o 3 dB na oktawę wraz ze wzrostem częstotliwości w określonym zakresie pasma. W całym paśmie akustycznym oznacza to, że amplituda zwiększa się łącznie o około 30 dB.

Zadziwiająca podobieństwa w naturze

Komórki siatkówki oka są rozmieszczone w układzie o statystycznym rozkładzie odpowiadającym szumowi niebieskiemu, co – jak się okazuje – zapewnia najlepszą możliwą rozdzielczość wzrokową. Z kolei promieniowanie Czerenkowa, występujące wysoko w atmosferze, jest naturalnym przykładem niemal idealnego szumu niebieskiego.

Zastosowania

Szum niebieski znajduje zastosowanie w muzyce elektronicznej i filmowej, między innymi do odtwarzania dźwięków przypominających szum wody.

Szum fioletowy (violet noise)

Definicja

Gęstość mocy szumu fioletowego rośnie o 6 dB na każdą oktawę wraz ze wzrostem częstotliwości w danym zakresie częstotliwości.

Zastosowanie

Szum fioletowy jest wykorzystywany w technice dźwiękowej do maskowania błędów kwantyzacji podczas przygotowywania dźwięku cyfrowego do zapisu na płytach CD audio. Próbkę dźwięku na płycie CD mają rozdzielczość 16 bitów, natomiast oryginalne nagrania studyjne są często rejestrowane z znacznie wyższą rozdzielczością. Podczas konwersji takich nagrań do formatu 16-bitowego powstają błędy kwantyzacji, które można częściowo zamaskować przez dodanie szumu fioletowego. Technikę tę nazywa się „ditheringiem”, jednak jej dokładne omówienie wykracza poza ramy tego – i tak już obszernego – opracowania poświęconego zagadnieniu szumów w elektronice.

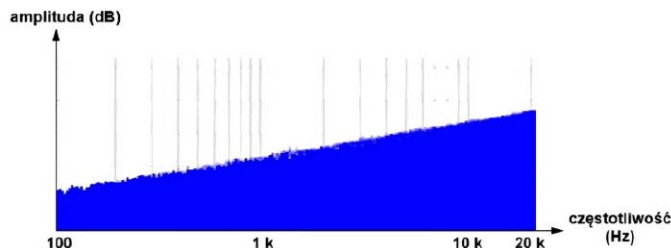
Szum szary (grey noise)

Definicja

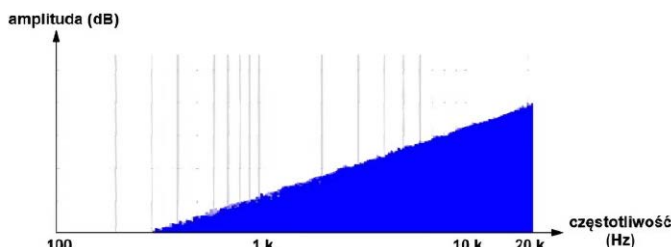
W przypadku szumu szarego biały szum jest przepuszczany przez specjalny filtr, którego charakterystyka odpowiada krzywej czułości ludzkiego słuchu. Dzięki temu słuchacz odnosi wrażenie, że wszystkie częstotliwości w sygnale mają jednakową głośność.

Zastosowania

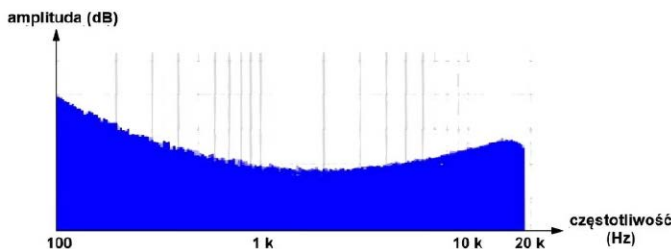
Szum szary jest wykorzystywany w różnych badaniach psychoakustycznych, dotyczących percepcji dźwięku przez różne grupy ludzi. W medycynie stosuje się go do leczenia szumów usznych (tinnitus) oraz nadwrażliwości słuchowej (hyperacusis). Słowo hyperacusis pochodzi z greki i oznacza dosłownie „słyszę zbyt wiele”. Osoby cierpiące na to schorzenie (szacuje się, że dotyczy ono około 3% populacji Holandii) odczuwają zwyczajne dźwięki jako uciążliwe, drażniące, a czasem nawet bolesne. ■



Charakterystyka amplitudowo-częstotliwościowa szumu niebieskiego
(© 2012 Mick Ohrberg, opracowanie z 2022 roku: Jos Verstraten)



Charakterystyka amplitudowo-częstotliwościowa szumu fioletowego
(© 2012 Mick Ohrberg, opracowanie z 2022 roku: Jos Verstraten)



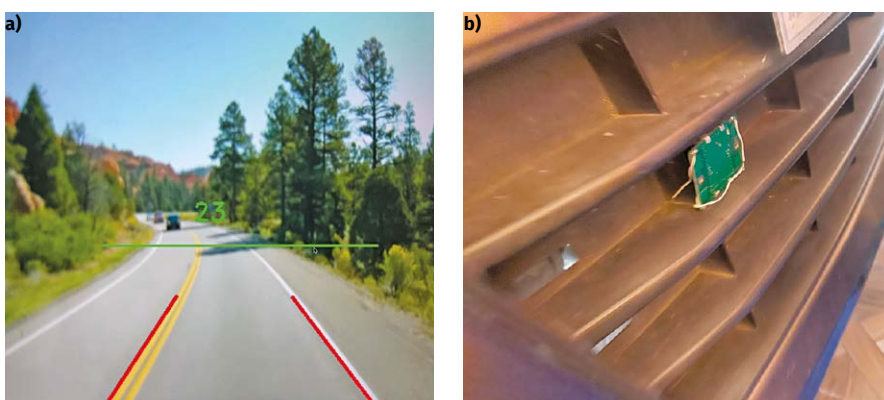
Charakterystyka amplitudowo-częstotliwościowa szumu szarego
(© 2012 Mick Ohrberg, opracowanie z 2022 roku: Jos Verstraten)

Jos Verstraten

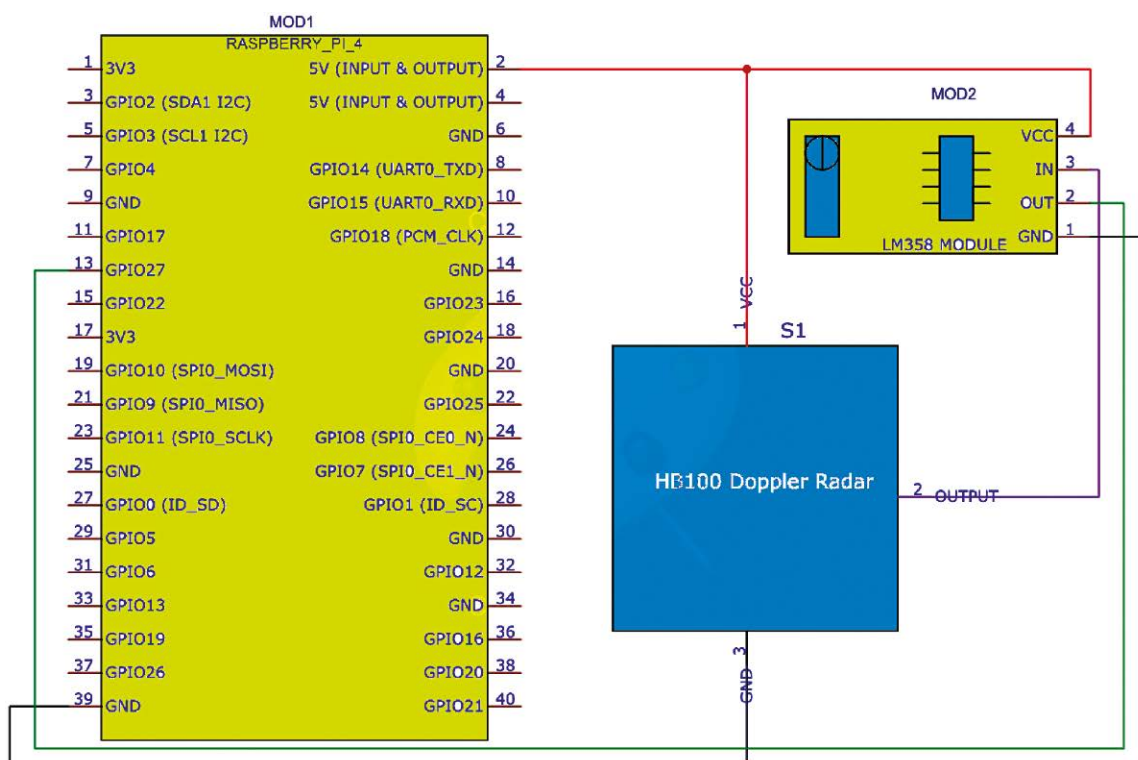
Zastosowanie radaru dopplerowskiego w systemie ADAS, część 4

W poprzednich odcinkach omawialiśmy system ADAS (zaawansowany system wspomagania kierowcy), w którym wykorzystano lidar, kamerę i system operacyjny ROS do tworzenia mapy. Teraz dodajemy ważny element bezpieczeństwa systemu ADAS, dzięki któremu można będzie sprawdzać względną prędkość poruszających się pojazdów. Element ten będzie ponadto wykrywać pas ruchu i z pomocą radaru na bieżąco wyznaczać bezpieczną odległość od pojazdu poprzedzającego (rysunek 1).

Na rynku jest wiele podsystemów lidarów i radarowych do systemu ADAS i systemów jazdy autonomicznej. Mogą one nie tylko wykrywać i mapować otoczenie, ale także informować o prędkości poruszających się pojazdów i pieszych oraz dokładnie zmierzyć odległość. Większość z nich opiera się jednak na laserach, co oznacza, że systemy te mogą zawodzić w obecności mgły lub przy bardzo słabym oświetleniu. Z tego powodu wykorzystujemy tutaj radar dopplerowski, który może pracować we wszystkich warunkach oświetleniowych i w mgłę. Jest on również w stanie mierzyć bardzo duże prędkości, czego większość innych systemów nie potrafi.



Rysunek 1. a) System z radarem i kamerą pokazujący prędkość, kierunek i pas ruchu; b) radar dopplerowski umocowany na samochodzie



Rysunek 2. Schemat układu z radarem do systemu ADAS

Podzespoły wymagane do zbudowania tego systemu są wymienione w Liście elementów.

Rysunek 2 przedstawia schemat układu radaru dopplerowskiego do systemu ADAS. Układ składa się z Raspberry Pi (MOD1), układu LM358, radaru dopplerowskiego HB100, kamery Raspberry Pi i kilku innych elementów.

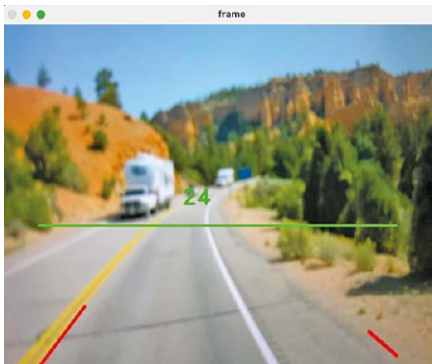
Przyjrzyjmy się najpierw jak działa radar dopplerowski, czym różni się od innych radarów i jak go wykorzystamy.

Wyznaczanie prędkości z użyciem radaru

Radar dopplerowski wykorzystuje efekt Dopplera, znany również jako „przesunięcie dopplerowskie”. Efekt ten opisuje zmianę częstotliwości i długości fali (na przykład dźwiękowej czy świetlnej) u obserwatora, względem którego źródło fali się porusza. Gdy obiekt wysyłający falę znajduje się w ruchu względem obserwatora, fale emitowane przez ten obiekt ulegają ściskaniu przed obiektem i rozciąganiu za nim. Powoduje to zmianę długości fali i częstotliwości mierzonych przez obserwatora.

Istnieją dwa podstawowe rodzaje efektu Dopplera, które opiszemy dla przypadku fali dźwiękowej:

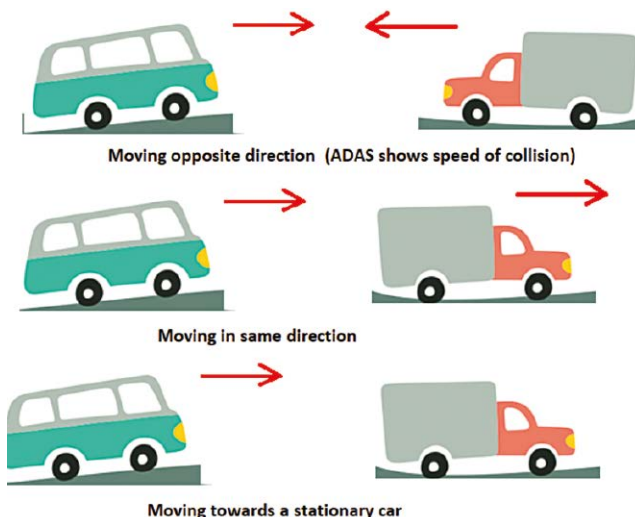
- Ruchome źródło i nieruchomy obserwator. Kiedy źródło dźwięku porusza się w kierunku obserwatora, fale dźwiękowe są ściskane, a obserwowana częstotliwość jest wyższa (wyższy dźwięk). Kiedy źródło dźwięku oddala się od obserwatora, fale dźwiękowe są rozciągane, a częstotliwość niższa (dźwięk niższy).
- Stojące źródło i poruszający się obserwator. Występuje ten sam efekt. Jeśli obserwator porusza się w kierunku źródła dźwięku, obserwowana częstotliwość jest wyższa. Jeśli obserwator oddala się od źródła, częstotliwość jest niższa.



Rysunek 5. Wyświetlacz w systemie ADAS pokazujący bezpieczny odstęp zieloną linią oraz prędkość, z jaką oba pojazdy zbliżają się do siebie



Rysunek 3. a) Niebieski samochód poruszający się w kierunku nieruchomego białego samochodu i zderzający go, b) samochody poruszające się w przeciwnych kierunkach, zmierzające do zderzenia, c) samochody poruszające się w tym samym kierunku



Rysunek 4. Trzy różne przypadki wzajemnego poruszania się samochodów

Beat Frequency and Speed

For a transmitted frequency of GHz = Hz
 the reflected frequency will be Hz
 and the beat frequency will be Hz

if the speed of the vehicle is $v_{\text{target}} =$ m/s = mi/hr = km/hr = ft/s

Note: If not entered as data, the transmitted frequency will default to 10.6 GHz, a typical X-band radar frequency. K-band units with frequencies on the order of 30 GHz are also used. Any parameters can be changed.

[Police RADAR](#)
[Expressions used in calculation](#)
[Doppler effect for electromagnetic waves](#)

[HyperPhysics**** Sound](#)
R Nave
[Go Back](#)

Rysunek 6. Kalkulator prędkości na podstawie częstotliwości dla radaru dopplerowskiego (źródło: HyperPhysics-radar)

```

ADAS_Radar2.py
1 import cv2
2 import numpy as np
3 import RPi.GPIO as GPIO
4 import time
5
6 # X-band radar sensing output connected to GPIO 27 of Rpi
7 HB100_INPUT_PIN = 27
8
9 # configure GPIO pins for input into GPIO 27
10 GPIO.setmode(GPIO.BCM)
11 GPIO.setup(HB100_INPUT_PIN, GPIO.IN)
12
13 # sample global vars for adjustment of sensitivity/measurements
14 # - A greater max-pulse-count will count more pulse inputs, leading
15 #   to more averaged doppler frequency - it will also take longer
16 # - A higher motion-sensitivity will reduce false motion readings
17 MAX_PULSE_COUNT = 10
18 MOTION_SENSITIVITY = 10
19
20
21 cap = cv2.VideoCapture(0) # initialize the camera
22
23 safe_distance_color = (0, 255, 0) # Green color
24 safe_distance_thickness = 2
25 safe_distance_position = 280 # Adjust the vertical position as needed
26
27
28 def count_frequency(GPIO_pin, max_pulse_count=10, ms_timeout=50):
29     """ Monitors the desired GPIO input pin and measures the frequency
30     of an incoming signal. For this example it measures the output of
31     a HB100 X-Band Radar Doppler sensor, where the frequency represents
32     the measured Doppler frequency due to detected motion.
33     """
34     start_time = time.time()
35     pulse_count = 0
36
37     # count time it takes for 10 pulses
38     for count in range(max_pulse_count):
39
40         # wait for falling pulse edge - or timeout after 50 ms
41         edge_detected = GPIO.wait_for_edge(GPIO_pin, GPIO.FALLING, timeout=ms_timeout)
42
43         # if pulse detected - iterate count
44         if edge_detected is not None:

```

Rysunek 7. Fragment programu w systemie ADAS do obsługi radaru

Opisywane w tym artykule systemy radarowe działają zasadniczo na tej samej zasadzie, co inne radary dopplerowskie, ale są specjalnie przystosowane do użytku w pojazdach. Oto z jakich elementów składa się działanie dopplerowskiego radaru w pojeździe:

- Wysyłanie sygnałów radarowych. Radar pojazdu emituje sygnały o częstotliwości radiowej (zazwyczaj w zakresie mikrofal) w określonym kierunku. Sygnały te są wysyłane z anten radarowych lub czujników umieszczonych

na przedniej maskownicy (grillu) lub zderzakach pojazdu.

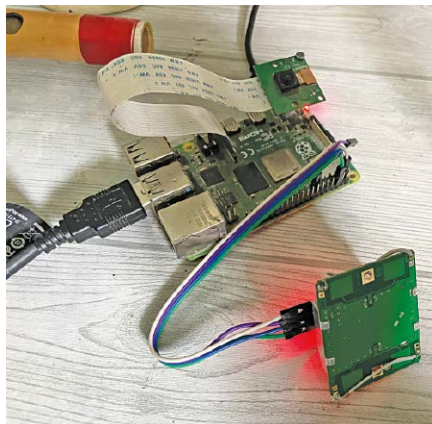
- Odbicie i przesunięcie dopplerowskie. Gdy sygnały radarowe napotykały obiekty na drodze pojazdu, takie jak inne pojazdy, piesi lub przeszkody, część energii fali jest odbijana z powrotem w kierunku anteny radaru. Jeśli obiekty te są nieruchome lub poruszają się z niewielką prędkością względem pojazdu, częstotliwość powracającego sygnału radarowego nie zmienia się lub zmienia się nieznacznie.

- Pomiar przesunięcia dopplerowskiego. Jeśli obiekt porusza się w kierunku pojazdu lub oddala się od niego, powoduje dopplerowskie przesunięcie w odbitym sygnale radarowym. Jeśli obiekt zbliża się, sygnał radarowy podlega dodatniemu przesunięciu dopplerowskiemu (wzrost częstotliwości). Jeśli obiekt oddala się, wywołuje to ujemne przesunięcie dopplerowskie (spadek częstotliwości).

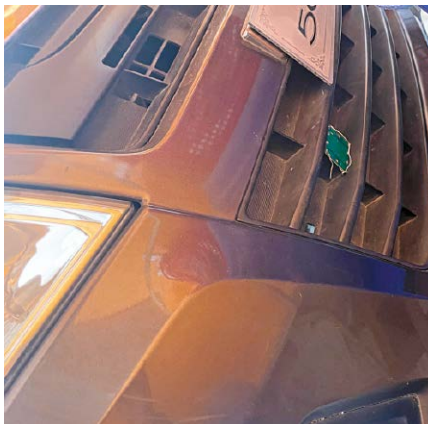
W radarze dopplerowskim wykorzystujemy te elementy do pomiaru prędkości, która w systemie ADAS podlega trzem różnym przypadkom (rysunek 3):

- Przypadek 1. Jeśli pojazd z radarem stoi w miejscu, a do niego zbliża się lub od niego oddala pieszy lub samochód, wówczas prędkość pokazywana przez radar jest prędkością samochodu czy pieszego zbliżającego/oddalającego się.
- Przypadek 2. Jeśli samochód z radarem porusza się, a przed nim znajduje się nieporuszający się samochód, ludzie lub ściana, wówczas pokazywana jest prędkość, z jaką samochód porusza się w kierunku do/od tego nieruchomego obiektu.
- Przypadek 3. Jeśli samochód porusza się w kierunku innego samochodu, pokazywana jest prędkość, z jaką pojazdy mogą zderzyć się ze sobą. Jeśli oba samochody poruszają się w tym samym kierunku, wyświetlana jest różnica ich prędkości. Jeśli oba samochody poruszają się w tym samym kierunku z taką samą prędkością, wówczas wyświetlana prędkość wynosi 0, co oznacza, że samochody się nie zderzą. Jeśli jeden z nich porusza się szybciej, wówczas wyświetlana jest prędkość, z jaką samochód ten oddala się od drugiego samochodu.

Rysunek 4 przedstawia wszystkie trzy przypadki. Jeśli samochód jadący z przodu naciśnie hamulec lub zwolni, podczas gdy zielony samochód jadący za nim nadal się

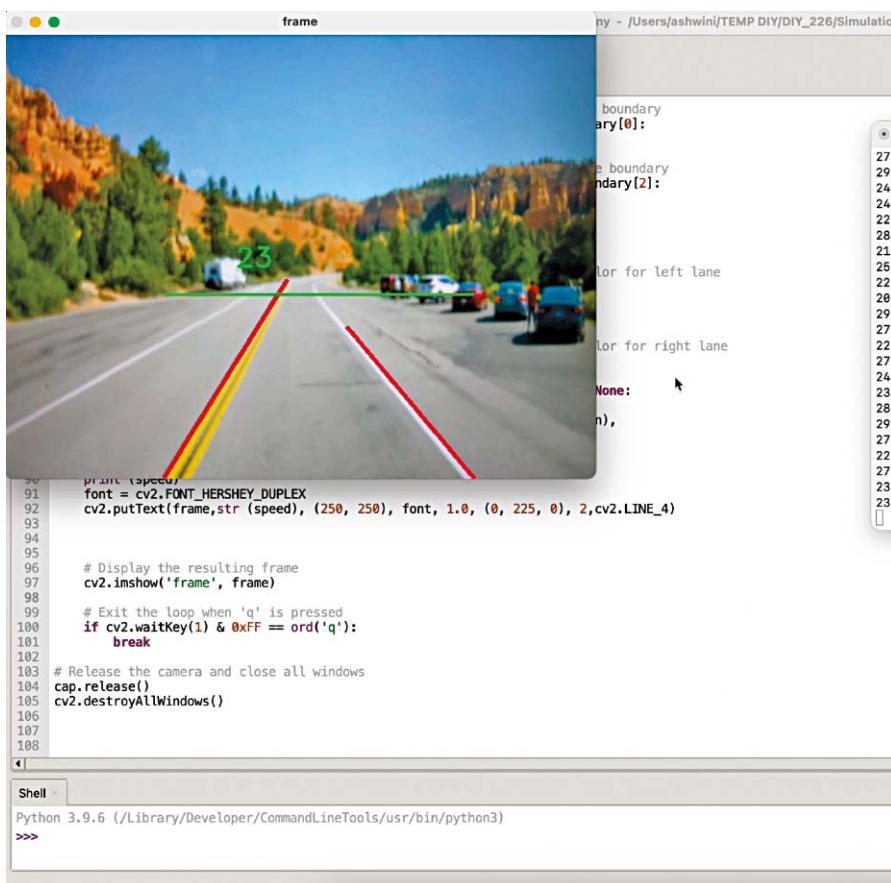


Rysunek 8. Konstrukcja Autora używana do prób



Rysunek 9. Radar i kamera zamontowane na zderzaku samochodu





Rysunek 10. Program wykrywający trasę i pas ruchu, wyświetlający prędkość zbliżających się pojazdów i zieloną linię bezpiecznego odstępu

porusza ze stałą lub rosnącą prędkością, wówczas wyświetlacz ADAS (rysunek 5) pokazuje wypadkową prędkość, przy której samochody mogą się zderzyć. Aby uniknąć wypadku, należy zachować odstęp zaznaczony na wyświetlaczu zieloną linią. Natomiast jeśli kierowca dopasowuje prędkość do samochodu jadącego z przodu, wyświetlana prędkość względna wyniesie 0, co oznacza, że oba samochody poruszają się z taką samą prędkością i nie dojdzie do kolizji.

Aby zobaczyć, jak działa przesunięcie dopplerowskie między poruszającymi się obiektami, można sprawdzić link <http://hyperphysics.phy-astr.gsu.edu/hbase/Sound/radar.html>. Prędkość jest obliczana na podstawie przesunięcia

dopplerowskiego z zależności między częstotliwościami fal. Rysunek 6 przedstawia przykład kalkulatora prędkości przeznaczanego dla radaru dopplerowskiego.

Oprogramowanie

Programowaniem zajmowaliśmy się już w poprzednich odcinkach tego cyklu. Zakładamy więc, że Czytelnik ma już działający Raspberry Pi/Nvidia Jetson z systemem operacyjnym i środowiskiem Python IDLE. Utworzyliśmy również program dla ADAS-a i do wykrywania pasa ruchu. Program ten będziemy kontynuować. Dodamy do niego funkcje obsługi radaru i przystosujemy go tak, aby można było skutecznie w czasie rzeczywistym wyznaczać odległości. Zaleca się przejrzeć

poprzednie odcinki cyklu i przypomnieć sobie cały proces.

W programie najpierw importujemy moduł Pythona RPI.GPIO, ponieważ będziemy używać bloku wejścia/wyjścia (GPIO) i odczytywać dane radaru. Następnie ustawiamy liczbę impulsów i czułość czujnika radarowego. Dalej musimy zdefiniować numer pinu wejścia/wyjścia, do którego podłączymy radar. Można tu użyć dowolnego pinu niewykorzystanego w Raspberry Pi. My użyliśmy GPIO 27. Następnie tworzymy funkcję do obliczania częstotliwości i wpisujemy wzór do przeliczania zmiany częstotliwości na prędkość, podany w przytoczonym wcześniej linku. Przypis redaktora: wzór ten nie jest tam wymieniony wprost, lecz można go otrzymać po prostych przekształceniach:

$$v_{\text{TARGET}} = \frac{\Delta f}{f} \cdot \frac{c}{2}$$

gdzie

v_{TARGET} = prędkość obiektu [m/s],
 Δf – zmiana częstotliwości fali [Hz],
 f – częstotliwość nadawanej fali [Hz],
 c – prędkość światła w próżni (299 792 458 [m/s]).

Teraz tworzymy pętlę „while”, w której w ramach wideo będziemy w czasie rzeczywistym lokalizować drogę i wykrywać pasy ruchu, a także rysować linię oznaczającą bezpieczny odstęp. Będziemy też wyświetlać prędkość obliczoną dzięki radarowi. Rysunek 7 przedstawia fragment programu obsługi radaru w systemie ADAS.

Konstrukcja i testowanie

W oparciu o schemat z rysunku 2 dołączamy radar HB100 do modułu LM358, a jego wyjście łączymy z pinem GPIO w Raspberry Pi/Nvidia Jetson. Podłączamy kamerę do portu kamery CSI w Raspberry Pi. Połączenia w prototypie pokazano na rysunku 8. Następnie mocujemy radar i kamerę z przodu pojazdu (rysunek 9). Zasilamy układ z zasilacza DC 5 V/2 A i uruchamiamy program. Będzie on wyświetlał prędkość i pokazywał bezpieczną odległość, tak jak to opisywaliśmy w artykule. ■

Ashwini Kumar Sinha

Lista elementów		
Element	Ilość	Opis
Raspberry Pi/Nvidia Jetson (MOD1)	1	SBC
Kamera Raspberry Pi	1	Kamera CSI
Radar HB100 (S1)	1	Czujnik radarowy 10,525 GHz
Wzmacniacz LM358 (MOD2)	1	Moduł wzmacniacza sygnału
Zasilacz 5 V	1	przetwornica DC/DC 5 V 2 A
Przewody	10	Cienkie przewody izolowane

Materiał filmowy do artykułu:
<https://youtu.be/Lm-HwF9gW1s>

• Strona projektu:
<https://www.electronicshobby.com/electronics-projects/doppler-radar-adas>

Materiały dodatkowe są dostępne na stronie elportal.pl/do-pobrania

Ten artykuł był wcześniej opublikowany na łamach „EFY”, październik 2023 (efymag.com)

Urządzenie do cheating'u z użyciem ChatGPT, o małych rozmiarach

ChatGPT jest obecnie jednym z najbardziej znanych i często używanych, zaawansowanych narzędzi AI. Może generować kod, tworzyć strony internetowe, obrazy i wiele więcej. Ale czy to potężne narzędzie może zmieścić się w niewielkim urządzeniu? Czy może zintegrować się z czujnikami IoT? Czy może współpracować z urządzeniami elektronicznymi? Tak, opisywane urządzenie właśnie do tego służy.

Można stworzyć urządzenie bazujące na ChatGPT, które po zadaniu pytania za pomocą szeregowego połączenia internetowego lub szeregowego połączenia USB wyświetla odpowiedzi na ekranie OLED. Zapewnia szybką pomoc i może być wykorzystane do cheating'u. Ma średnicę zaledwie 3 cm.

Na rysunku 1 przedstawiony jest autorski prototyp współpracujący z terminalem WebSerial. Komponenty potrzebne do budowy urządzenia są wymienione w zestawieniu materiałów.

Do łączenia się z siecią Wi-Fi i z ChatGPT API oraz do hostowania serwera szeregowego wykorzystana została płytki IndusBoard. Pytania zadawane za pośrednictwem terminala WebSerial są przesyłane do ChatGPT API. Odpowiedź jest wyświetlana na ekranie OLED podłączonym do płytki.

Konfiguracja

Najpierw za pomocą linku <https://openai.com/index/openai-api> prowadzącego do strony OpenAI należy utworzyć nowy klucz do API ChatGPT, zapisać go i nie udostępniać nikomu. Aby móc korzystać z API należy podać informacje rozliczeniowe dotyczące płatności.

Create new secret key

Owned by

☒ You ☐ Service account

This API key is tied to your user and can make requests against the selected project. If you are removed from the organization or project, this key will be disabled.

Name Optional

My Test Key

Project

Default project

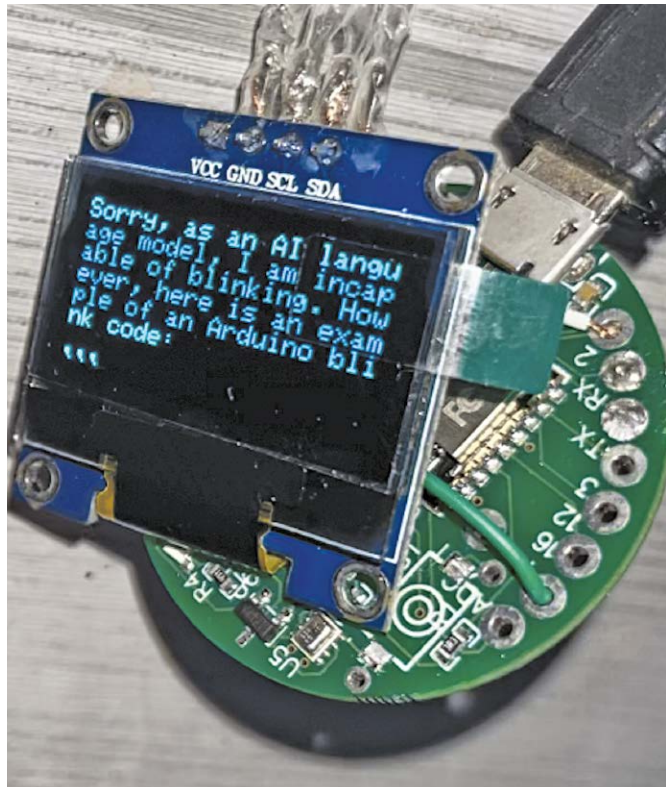
Permissions

☒ All ☐ Restricted ☐ Read Only

Cancel

Create secret key

Rysunek 3. Tworzenie nowego klucza



Rysunek 1. Autorski prototyp

Playground

Chat

Assistants

Completions

Assistants

Fine-tuning

Batches

Storage

Usage

API keys

API keys

Project API keys have replaced user API keys. We recommend using project based API keys for n

As an owner of this project, you can view and manage all AI

Do not share your API key with others, or expose it in the br

View usage per API key on the Usage page.

Rysunek 2. Konfiguracja interfejsu API

```
1 #include <Arduino.h>
2 #include <Ug2lib.h>
3 #include <WiFiClientSecure.h>
4 #include <ArduinoJson.h>
5 #include <ChatGPT.h>
6
7 // WiFi credentials
8 const char* ssid = "JioFiber-46"; // Your WiFi SSID
9 const char* password = "1234567890"; // Your WiFi Password
10
11 // ChatGPT API key
```

Rysunek 4. Konfiguracja kodu dla Wi-Fi

Zestawienie materiałów

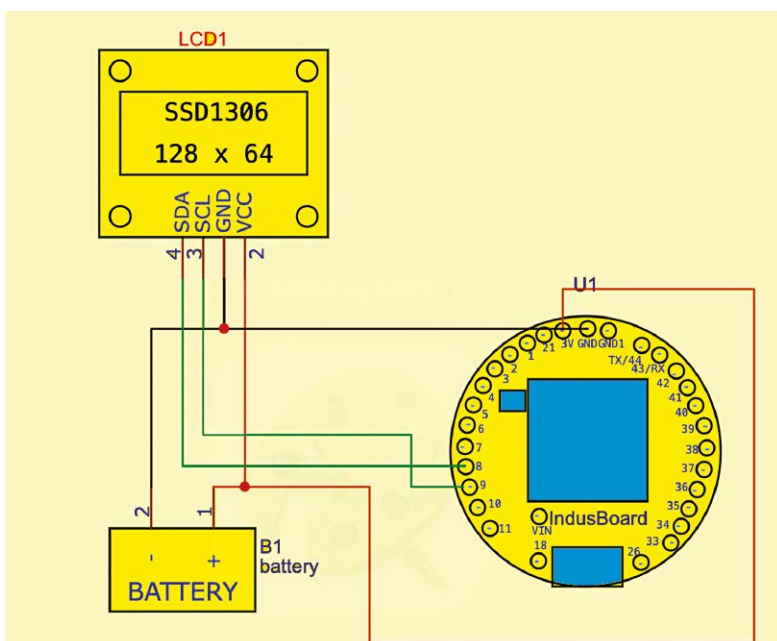
Komponenty	Opis	Ilość
	Płytki rozwojowa o średnicy 3 cm	1
Bateria	3,3 V	1
Wyświetlacz OLED/LCD	SSD1306/GC9A01	1

```

13 const char* password = "1234567890"; // Your WiFi Password
14
15 // ChatGPT API key
16 const char* api_key = "sk-proj-81BEevh9eVtyQMBEchBhT3BlbkFJsUxC06gd9v3nC7hewoPJ3"; // Your ChatGPT API key
17
18 // Web server and WebSerial
19 AsyncWebServer server(80);
20
21 // OLED display using U8g2 library
22 U8G2_SSD1306_128X64_ALT0_F_HW_12C u8g2(U8G2_R0, /* reset=*/ U8X8_PIN_NONE); // Constructor for SSD1306 128x64 OLED display
23
24 // ChatGPT client
25 WiFiClientSecure client;
26 ChatGPT<WiFiClientSecure> chat_gpt(&client, "v1", api_key);
27
28 /* Message callback of WebSerial */
29 void recvMsg(uint8_t* data, size_t len) {
30   WebSerial.println("Received Data...");
31 }

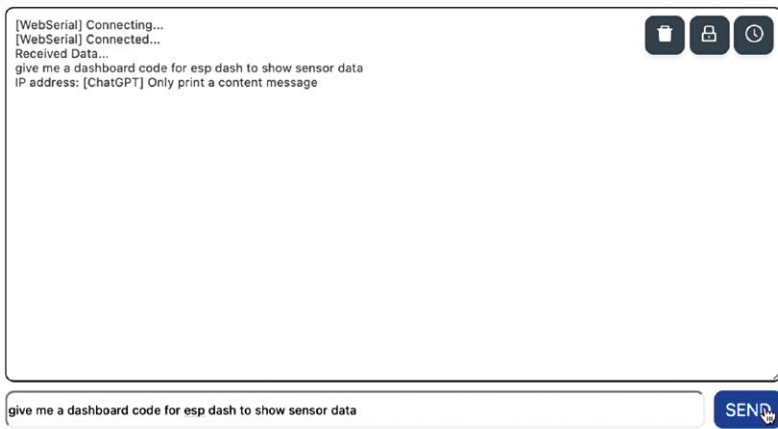
```

Rysunek 5. Konfiguracja interfejsu API



Rysunek 6. Schemat połączeń

WebSerial



Rysunek 7. Pytania dotyczące WebSerial

Nowi użytkownicy otrzymują bezzwrotny kredyt w wysokości 5 USD. Na **rysunku 2** przedstawiona jest konfiguracja API, a na **rysunku 3** pokazany jest sposób tworzenia nowego klucza.

Kod źródłowy należy pobierać ze wspomnianej strony i skonfigurować zgodnie z opisem. Następnie należy uruchomić Arduino IDE, wybrać płytkę ESP32 i zainstalować biblioteki U8g2 i ChatGPT.

W kodzie należy zmienić identyfikator SSID i hasło do sieci Wi-Fi. Można użyć poświadczeń Wi-Fi domowego routera lub hotspotu Wi-Fi w telefonie. Na **rysunku 4** przedstawiony jest sposób modyfikacji kodu dla interfejsu Wi-Fi. Należy pamiętać, że ESP32 działa w paśmie 2,4 GHz, więc podczas konfiguracji routera należy wprowadzić odpowiednie ustawienia.

Następnie należy zmodyfikować kod API i wymienić go na nowy. Na **rysunku 5** widoczna jest konfiguracja nowego kodu API.

Na **rysunku 6** przedstawiony jest schemat mikro urządzenia wykorzystującego ChatGPT cheating'u. Składa się ono z płytki IndusBoard, wyświetlacza OLED i baterii o napięciu 3,3 V. Wyświetlacz OLED jest połączony z pinami PC płytki IndusBoard, domyślnie ustawionymi na numery 8 i 9. Całe urządzenie jest zasilane za pomocą baterii, przez piny oznaczone 3,3 V i GND.

Testowanie

Po wgraniu skompilowanego kodu źródłowego do płytki IndusBoard i dokonaniu połączeń zgodnie ze schematem należy włączyć zasilanie urządzenia. Po połączeniu się z siecią Wi-Fi na ekranie OLED zostanie wyświetlony adres IPaddress/webserial. Należy go wprowadzić do przeglądarki internetowej, to spowoduje otwarcie terminala WebSerial.

Teraz można zadać dowolne pytanie na ChatGPT, a odpowiedź zostanie wysłana i wyświetlona na ekranie OLED zamontowanym na IndusBoard. Osoba nosząca zegarek także zobaczy odpowiedź na jego wyświetlaczu. Na **rysunku 7** pokazane są przykładowe pytania przechodzące przez terminal WebSerial. ■

Ashwini Kumar

- Materiał filmowy do artykułu:
<https://youtu.be/Ac92kofARls>
- <https://youtu.be/uuybiKQH91I>

Materiały dodatkowe są dostępne na stronie elportal.pl/do-pobrania

Ten artykuł był wcześniej opublikowany na łamach „EFY”, październik 2023 (efymag.com)

ELEKTRONIKA

dla wszystkich

nr 12/2025 (18)

JUNIOR



Dawid, montaż zestawu AVTXMAS5
Szkoła Podstawowa nr 86, Wrocław

Grudzień to czas, który za każdym razem przenosi mnie myślami w świat mojego dzieciństwa. Dziś zimowa pora wygląda inaczej niż kiedyś. Grudniowy krajobraz za oknem nie wiele różni się od listopadowego, lub nawet tego z października. Tymczasem mi czas ten kojarzy się z rodzinną Polanicą i palącymi się w mroku wieczornej pory latarniami i grubą warstwą śnieżnego puchu, skrzypiącego pod podeszwami w mroźnej temperaturze. Wielokrotnie spacerowałem w takiej aurze ulicą Zdrojową, a później wdrapywałem na dość stromą ulicę Górską, pokonując trasę między domem rodziców a domem dziadków. Tęsknię za tamtą beztroską, gdzie jedynym życiowym zmartwieniem był wyściętający po same kolana drogę do pokonania śnieg.

W podobnie magiczny sposób zapamiętałem jeden szczególny dzień z okresu studiów na Politechnice Wrocławskiej. Zmierałem właśnie do sali 201 w budynku C1 gdy mój wzrok zatrzymał się na uczelnianej gablocie. Mijałem tę gablotę setki razy, zawsze w pośpiechu, myśląc o zajęciach, kolokwiah, o życiu, które wtedy wydawało mi się rozpędzone do granic możliwości. Tego jednego dnia zatrzymałem się przed nią. Miałem wrażenie, że jedno z ogłoszeń otwiera się przede mną jak ukryte przejście – takie, o którym nigdy nie wiesz, że istnieje, dopóki nie spojrzysz pod odpowiednim kątem. Było jak list, który czeka, aż ktoś wreszcie go przeczyta. Jak zaproszenie – nie tak spektakularne jak te z Hogwartu, bez sowy ani pieczęci z wosku – a jednak mające w sobie coś z tamtej magii. Wystarczyło jedno zdanie, bym poczuł, że oto właśnie otwiera się

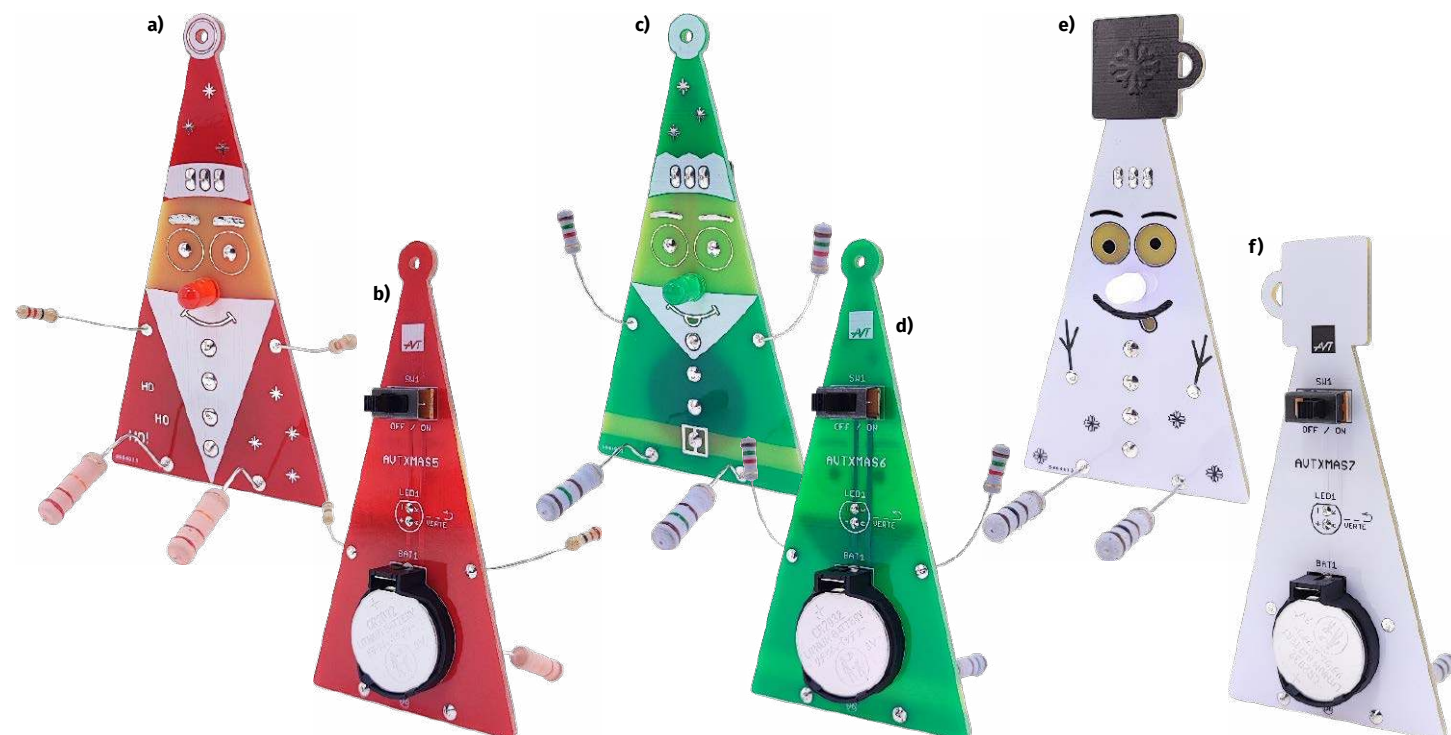
przede mną jakiś sekret, przejście do świata, z którego nie sposób zawrócić. Ogłoszenie dotyczyło możliwości podjęcia wolontariatu na zasadach programu „Starsza Siostra Starszy Brat” w jednym z podwrocławskich domów dziecka. Hasło, które przeczytałem w ogłoszeniu brzmiało: „Jeśli nie ty im pomożesz, to kto?”. Mimo, że było cliche – a może właśnie dlatego – nie potrafiłem z nim polemizować. Nie zastanawiając się wiele, z dnia na dzień dołączyłem do grona wolontariuszy tamtego domu, by – jak się później okazało – spędzić w nim niemal każdy weekend przez kolejnych osiem lat mojego życia. Ciepło wspominam pierwszą Wigilię w roku 2002, którą miło spędziłem wśród dzieciaków, kadry i gości. Tak rozpoczęła się moja przygoda, powiedziałbym wręcz, niekończąca się opowieść, związana z wolontariatem. Radościami z jego początków dzieliłem

się na bieżąco z redakcją ulubionej gazety (EdW 12/2004 oraz Elektronika Plus – BASCOM). To na swój sposób dość wzruszające, zważywszy, że dziś sam jestem jej redaktorem. Po wolontariacie w domu dziecka przyszedł czas na rodzinę zastępczą i – obecnie – aktywność z dziećmi w jeszcze kilku miejscach. Kiełkuje też pragnienie (i rozpoczęły się działania) by w najbliższej przyszłości przeorganizować własną codzienność w całości i już zawodowo oddać się tym aktywnościom.

Dzieciństwo to piękny czas. Jedyny i niepowtarzalny. Nie ważne, czy masz lat pięć czy piętnaście. Dorosłość do której tak bardzo się wtedy pędzi, często rozczarowuje, a багаż doświadczeń pozwala snuć domniemania, że dzieci to jedyny normalny gatunek ludzi na tym świecie. Gdyby można było, chciałoby się do tego dziecięcego świata wrócić, i nigdy



Fotografia 1.
Od lewej: Kornel,
Grzegorz i Dawid.
Koło Młodych Entuzjastów Elektroniki,
Szkoła Podstawowa nr 86, Wrocław



Fotografia 2. Zmontowane zestawy; a) i b) Pomysłowy Mikołaj LED (XMAS5); c) i d) Elf Pomocnik LED (XMAS6) oraz e) i f) Wesoły Bałwanek LED (XMAS7)

już nie dorosnąć. Dlatego, jeżeli istnieje we mnie jakakolwiek namiętność tudzież ambicja na przyszłość, zawsze będzie wiązała się ona ze światem dzieciństwa, jego klimatem, barwami i magią.

I oto zbliża się okres w roku najbardziej magiczny. Wypełniony zapachem mandarynek i ciast, blaskiem mieniących się ulicznych girland, przydomowych dekoracji, błyszczących choinkowych ozdób, śnieżnobiałych bałwanek, zielonych Elfów i czerwonych Mikołajów.

W zeszłym roku składaliśmy **Bożonarodzeniowe drzewka LED** (AVTXMAS1), **Czerwone świąteczne ozdoby RGB** (AVTXMAS3), **Choinki LED RGB** (AVTEDU640) i **Bombki LED dla każdego** (AVT3250). Tym razem do kolekcji własnoręcznie zbudowanych ozdób świątecznych trafią: **Pomysłowy Mikołaj LED** (XMAS5), **Elf Pomocnik LED** (XMAS6) oraz **Wesoły Bałwanek LED** (XMAS7).

Fotografie 2a...2f pokazują zmontowane zestawy w obu widokach (przód i tył figurki).

Bawmy się i uczmy bezpiecznie

Zestawy zostały przygotowane z myślą o najmłodszych potencjalnych pasjonatach elektroniki, którzy z lutowaniem i lutownicą – być może – zetkną się po raz pierwszy. Oczywiście zwłaszcza w takich okolicznościach montaż powinien odbywać się w asyście potrafiącej w sposób bezpieczny

lutować osoby dorosłej. **Rozgrzaną do temperatury 360°C lutownicą można sobie zrobić krzywdę, co do tego nie należy mieć wątpliwości.** O ile jednak to rzecz oczywista, a ciało człowieka natychmiast reaguje skutecznym odruchem ucieczki, unikając tym samym poważniejszych obrażeń, o tyle żaden rozsądny opiekun nie pozostawi kilkulatka/kilkulatki sam na sam z rozgrzaną lutownicą (podobnie jak nie zrobiłby tego z rozgrzanym żelazkiem, czy rozgrzaną na kuchence patelnią). Należy jednak zwrócić uwagę na inne, mniej oczywiste ryzyko, na które ciało człowieka nie ma równie skutecznej reakcji obronnej.

Użyta w każdym z zestawów bateria guzikowa (pastylkowa) CR2032 mogłaby zostać przez dziecko połknięta i – zależnie od tego, gdzie utknie – wyrządzić poważne szkody. Przepływ prądu przez tkanki ciała, mające kontakt z obydwoma biegunami połkniętej baterii wywołuje reakcje elektrochemiczne prowadzące do wytwarzania silnie żrących wodorotlenków (między innymi sodu i potasu), które powodują oparzenia chemiczne i głębokie uszkodzenia tkanek. Tego typu obrażenia rozwijają się bardzo szybko, dlatego każde, nawet najmniejsze, podejrzenie połknięcia przez dziecko baterii wymaga natychmiastowej pomocy medycznej.

O powyższym ryzyku, związanym z łatwą do połknięcia baterią CR2032 (i jej podobnymi) wspomina się niemal

na każdym szkoleniu dotyczącym ratownictwa medycznego i pierwszej pomocy. Regularnie uczestniczę w szkoleniach, które Dolnośląska Fundacja na rzecz Pieczy Zastępczej „Przystanek Rodzina” organizuje dla swoich wolontariuszy, za co – korzystając z okazji – chciałbym im serdecznie podziękować. Wielu Czytelników nie ma takich możliwości, dlatego niniejszym przekazuję tu wiedzę, która może realnie chronić zdrowie i życie.

Zabawki z baterią pastylkową można bezpiecznie udostępnić dzieciom mniej więcej od 6–7 roku życia, kiedy potrafią już rozumieć zasady ostrożności i nie wkładają małych przedmiotów do ust. Najlepiej jednak, aby dziecko bawiło się taką zabawką w obecności dorosłego opiekuna. Każdy rodzic wie, czego może się spodziewać po własnej latorośli, dlatego udostępnia taką zabawkę na własną odpowiedzialność. U młodszych



Rysunek 1. Bateria CR2032 zamontowana w przeznaczonym dla niej gnieździe na tylnej ścianie każdej z figurek. Bateria taka łatwo może zostać przez małe dziecko zdemontowana i połknięta, dlatego przed udostępnieniem zabawki dziecku przeczytaj koniecznie akapit „Bawmy się i uczmy bezpiecznie” i zachowaj ostrożność

dzieci – zwłaszcza do około 4–5 lat – zdecydowanie lepiej tego unikać, bo ryzyko połknięcia baterii pozostaje bardzo wysokie. Stąd, po krótkiej zabawie gadżetem zmontowanym pod kontrolą osoby dorosłej (patrzmy czy dziecko przez przypadek nie wyciąga baterii i nie wkłada do buzi) zestawy z tą baterią, wieszamy na przykład na choince, poza zasięgiem najmłodszych. Gniazdo z zamontowaną baterią CR2032, znajdujące się na tylnej ścianie każdej z figurek pokazano na **rysunku 1**.

Nie każdy bohater nosi pelerynę

... czyli nieco o edukacji w zakresie odpowiedzialności i bezpieczeństwa. Uwaga i ostrożność w operowaniu rozgrzewaną do wysokich temperatur lutownicą ma kluczowe znaczenie, zarówno dla osoby, która tą lutownicą w danej chwili włada, ale i dla osób w bezpośrednim jej otoczeniu. Jestem dumny z uczestników zajęć, że w większości, bez mojego przypominania, pilnują tematu zakładania gogli ochronnych. Często powtarzam im anegdotę o tym, że można żyć w przekonaniu, że jest się świetnym ekspertem w dziedzinie elektroniki, a nawet nim być. Ale nie warto wierzyć w to, że w kryzysowej sytuacji w czymkolwiek nam to pomoże. Sięgając po analogie, można być naprawdę świetnym kierowcą – szybkim, pewnym siebie i doświadczonym – a mimo wszystko trafić do szpitala po poważnym wypadku. I to wcale nie z własnej winy. Wystarczy, że ktoś inny potraktuje zasady ostrożności czy kultury jazdy jak coś opcjonalnego, albo kompletnie o nich zapomni. A wtedy nawet najlepsze umiejętności mogą nie wystarczyć.

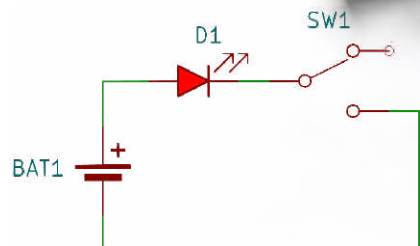
Podobnie jak nie jesteśmy samotnymi kierowcami na publicznej drodze, tak i w trakcie zajęć z elektroniki mamy wokół siebie innych „uczestników ruchu”. Warto o tym pamiętać, nie tylko w kontekście potrzeby zachowania zasady ograniczonego zaufania względem poczyną każdego z naszych kolegów (mimo naszych najlepszych nawet sympatii), ale też w kontekście świadomości, że od naszej ostrożności i uważności zależy też bezpieczeństwo innych. **Nie zakładając gogli ochronnych, narażam samego siebie na to, że do oka wpadnie mi fragment wyprowadzenia odciętego przez kolegę obok. Jednocześnie nie przytrzymując odcinanych przez siebie fragmentów wyprowadzeń palcami drugiej ręki podczas ich obcinania, mogę narażać na utratę zdrowia, a może i cięższych urazów kolegę obok, który z lekkomyślności, bądź roztrągnięcia**

w danej chwili gogli ochronnych na nosie miał nie będzie. Dbając o bezpieczeństwo stajesz się potencjalnym bohaterem, mającym każdorazowo szansę zawalczyć o zdrowie i bezpieczeństwo swoje i osób w pobliżu. Tym samym ocalić przed realnym niebezpieczeństwem kolegę lub siebie.

Odkładając obcięte wyprowadzenia na jedno miejsce łatwiej też „ogarnąć bałagan” po skończonej pracy. Temat może z pozoru mniej „górnolotny” – ale jak wiadomo – „co się wzniosło, musi spaść”, i często spada... właśnie na podłogę. Nadeptnięcie bosą stopą na ostry druk z pewnością do przyjemnych doświadczeń nie należy. Zatem nawet w tej przyziemnej sytuacji kryje się okazja, by zatroszczyć się o siebie i o osoby w naszym otoczeniu. A wszystko dosłownie w zasięgu naszych rąk.

Pomiędzy ostrożnością a nadgorliwością

Warto podkreślić, że odpowiedzialne podejście nie polega na zamykaniu dziecka pod kłosem aż do pełnoletności. Są rodzice i opiekunowie, którzy – w imię bezpieczeństwa – gotowi byliby zabronić dziecku niemal wszystkiego, co niesie choćby cień ryzyka. Tymczasem życie nie jest sterylnym laboratorium, a dzieciństwo to nie czas na obsesyjne usuwanie z drogi każdego kamyczka, gałązki czy krawędzi, by przypadkiem nic nie mogło się wydarzyć. Przecież nawet najbardziej niewinne aktywności mają swoją „ciemną stronę”: kąpiel w basenie ogrodowym może skończyć się wypadkiem, jeśli zabraknie czujnego oka dorosłego, a radosne skakanie na dmuchańcach podczas przydomowego urodzinowego szaleństwa zawsze niesie ze sobą

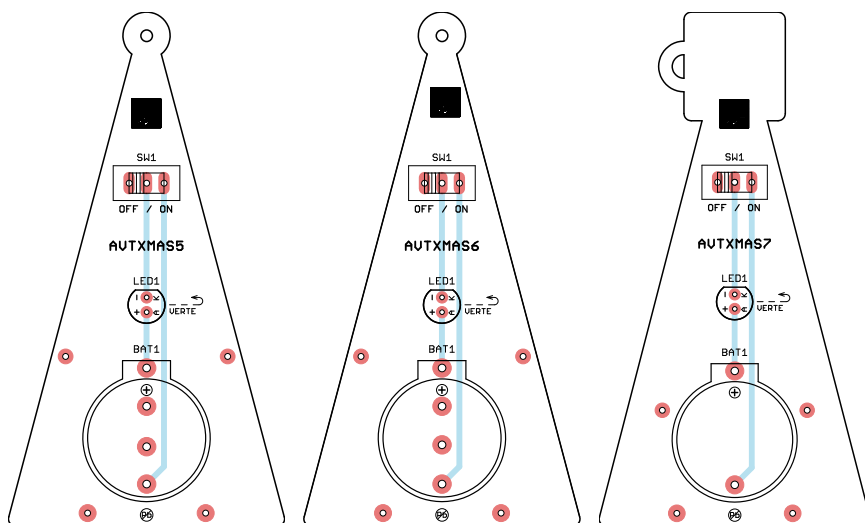


Rysunek 2. Schemat ideowy układu. Dla osób mających pierwszy kontakt z elektroniką piękny w swojej prostocie

ryzyko niefortunnego upadku i złamanej ręki czy nogi. Ale czy to oznacza, że mamy osuszyć basen i zwinąć dmuchańce do piwnicy? Oczywiście, że nie. „Kto nie doświadcza, ten nie dojrze” – i to zdanie w tym kontekście brzmi jeszcze mocniej. Pozbawiając dziecko okazji do poznawania świata, uczenia się i eksperymentowania, można wyrządzić krzywdę równie poważną, jak ta wynikająca z braku ostrożności. Rozsądek polega na mądrym towarzyszeniu dziecku i stopniowym powierzaniu mu odpowiedzialności, a nie na bezwzględnym eliminowaniu wszystkiego, co mogłoby stać się bodźcem do jego rozwoju. To nie unikanie aktywności czyni dziecko bezpiecznym, lecz obecność dorosłego, który wie, kiedy przytrzymać za rękę, a kiedy dać jej swobodnie sięgnąć dalej.

Wybór zestawu

Podczas tegorocznych świątecznych zajęć stacjonarnych Juniorzy EdW mieli do wyboru trzy zestawy. Jak wspomniano były to: Pomysłowy Mikołaj LED (XMAS5), Elf Pomocnik LED (XMAS6) oraz Wesoły Bałwanek LED (XMAS7). Wszystkie trzy figurki są do siebie podobne, mają ten sam



Rysunek 3. Schemat montażowy jest niemal identyczny dla wszystkich trzech zestawów a) Pomysłowy Mikołaj LED (XMAS5); b) Elf Pomocnik LED (XMAS6) oraz c) Wesoły Bałwanek LED (XMAS7)

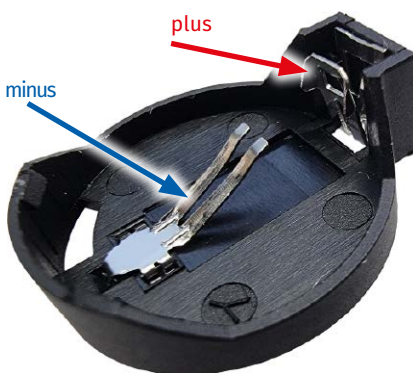
schemat elektryczny (rysunek 2) i montażowy (rysunek 3). Choć schemat jest trywialny, z tego samego powodu jest on świetny, by pokazać – nowej w świecie elektroniki osobie – sposób opisywania rzeczywistości stosowany przez bardziej zaawansowanych kolegów. Schemat pokazuje graficzne reprezentacje elementów wykorzystanych do zbudowania funkcjonalnej części figurki (baterię BAT1, przełącznik SW1 oraz migającą diodę LED1) oraz szeregowy sposób połączenia wszystkich tych elementów. Można nowej osobie wskazać baterię jako źródło zasilania, przełącznik, który należy załączyć, aby popłynął prąd, i diodę LED1, która zacznie błyskać, gdy przez obwód (połączenia i komponenty) od plusa do minusa baterii popłynie prąd.

Następnie można nowej osobie pokazać schemat montażowy (rysunek 3) a na nim ścieżki doskonale odwzorowujące trasę prądu elektrycznego: od plusa baterii, przez diodę LED1, przełącznik SW1 po minusowy biegun baterii.

Co trzeba wiedzieć, by uruchomić zestawy? Ponieważ zestawy zaprojektowane zostały z myślą o osobach nie mających wcześniej żadnej styczności z montażem elektroniki, czuję się zobowiązany podać najważniejsze informacje, niezbędne do uruchomienia każdego z nich. Kluczem do sukcesu uruchomienia każdego z tych zestawów jest zrozumienie zagadnień związanych z biegunowością baterii oraz właściwą polaryzacją diod LED. Przede wszystkim należy wiedzieć, gdzie bateria ma wyprowadzony biegun dodatni, a gdzie ujemny. Dlaczego taka wiedza jest nam niezbędna? Wynika to z prostego faktu, że o ile, w klasycznych żarówkach stosowanych do niedawna powszechnie jako źródło światła sztucznego, sposób podłączenia biegunów baterii do żarówki nie miał żadnego znaczenia, o tyle dioda LED w przypadku niepoprawnego podłączenia biegunów, w najlepszym przypadku nie zaświeci, w najgorszym, ulegnie trwałemu uszkodzeniu.



Fotografia 3. Oznakowanie biegunowości na ogniwie CR2032. Na jednej ze stron znajduje się znak „+” wskazujący dodatni biegun baterii. Po przeciwnej stronie nie ma żadnego oznakowania – to właśnie „minus” ogniwa



Fotografia 4. Biegunowość na uchwycie baterii CR2032

Biegunowość baterii

Wyraźne oznakowanie baterii nie pozostawia wątpliwości co do umiejscowienia „plusa” i „minusa” (fotografia 3)

Równie ważna jest umiejętność rozpoznawania biegunowości na uchwycie baterii. W przypadku baterii CR2032 biegunowość uchwytu pokazano na fotografii 4.

Oczywiście aby tak było bateria musi zostać w uchwycie prawidłowo zamontowana (fotografia 5).

Iluminacja!

Zanim rynek zawojowały diody LED, powszechnie stosowanym źródłem światła sztucznego były żarówki. Wykorzystywano je również w tworzeniu świątecznych iluminacji. Z pewnych względów posługiwanie się nimi było nieco mniej skomplikowane, niż ma to miejsce w przypadku diod LED.

Jak działa klasyczna żarówka?

Żarówka działa na zasadzie **rozgrzewania się włókna** do bardzo wysokiej temperatury, co powoduje emisję światła widzialnego. Kluczowym elementem żarówki jest cienkie włókno wykonane najczęściej z wolframu, materiału o wysokiej temperaturze topnienia. Aby zapobiec szybkiemu spaleni

gazu obojętnego spowalnia proces spalania wolframu i wydłuża żywotność żarówki.

Klasyczne żarówki wolframowe zaczęły być powszechnie używane na początku XX wieku, choć ich historia sięga lat 70. XIX wieku. Największym przełomem było udoskonalenie żarówki przez Thomasa Edisona w 1879 roku. W tym czasie Edison stworzył model z włóknem węglowym zamkniętym w próżniowej bańce, co znacząco wydłużyło jej żywotność i sprawiło, że mogła być stosowana na dużą skalę. Żarówki z wolframowym włóknem pojawiły się na rynku w 1904 roku i szybko wyparły wcześniejsze rozwiązania. Wolframowe włókno miało wyższą temperaturę topnienia i wytrzymywało większe obciążenia, co zapewniało dłuższy czas pracy i lepszą wydajność świetlną. Te klasyczne żarówki były dominującym źródłem światła sztucznego aż do końca XX wieku. Kres żarówek wolframowych datuje się na początek XXI wieku, kiedy to zaczęły je wypierać energooszczędne źródła światła: najpierw świetlówki kompaktowe, a później żarówki LED. W wielu krajach Unii Europejskiej oraz w USA wprowadzono przepisy stopniowo wycofujące sprzedaż tradycyjnych żarówek. W Europie proces ten rozpoczął się w 2009 roku i trwał do 2012 roku, kiedy wycofano z rynku większość klasycznych żarówek o mocy powyżej 25 W. Od tego czasu żarówki wolframowe są trudniejsze do zdobycia, a ich produkcja została znacznie ograniczona na rzecz bardziej efektywnych źródeł światła, takich jak LED, które zużywają znacznie mniej energii i mają dłuższą żywotność.

W przypadku klasycznej żarówki (jeszcze niedawno masowo stosowano je w oświetleniu domowym, latarkach a także ozdobach świątecznych, w tym w łańcuchach lampek choinkowych) kierunek przepływu prądu przez żarówkę nie miał żadnego znaczenia. Niezależnie od tego, w którą stronę prąd płynął przez włókno żarówki, rozgrzewało się ono do czerwoności, stając się tym samym źródłem widzialnego światła. W przypadku klasycznych żarówek nie dało się zmienić koloru światła emitowanego przez rozgrzane włókno, ponieważ zależało ono od temperatury włókna, a ta ograniczała się do żółtawego odcienia. Dlatego kolorowe światło z żarówek można było uzyskać jedynie poprzez malowanie ich szklanej bańki. Bańka żarówki była w tym celu pokrywana kolorową powłoką, która przepuszczała tylko określone długości fal świetlnych. Bańka zabarwiona na czerwono przepuszczała światło czerwone, blokując inne długości



Fotografia 5. Poprawny montaż baterii CR2032 w uchwycie, plusem do góry

fal. W ten sposób uzyskiwano kolory takie jak czerwony, zielony, niebieski itp.

Jak działa dioda LED?

W przeciwieństwie do żarówek, w diodach LED próżno byłoby szukać włókna wolframowego. Nie nagrzewają się one, jak miało to miejsce w przypadku żarówek, przez co do generowania światła zużywają one o wiele mniej energii. Jak to możliwe?

Dioda LED (*Light Emitting Diode*) działa na zasadzie zjawiska elektroluminescencji, które polega na emisji światła w wyniku przepływu prądu przez specjalnie skonstruowane złącze półprzewodnikowe. W diodach LED wykorzystuje się półprzewodniki, takie jak krzem lub arsenek galu, w których materiał przewodzący jest domieszkowany, by stworzyć obszary typu n (nadmiar elektronów) i p (niedobór elektronów, czyli „dziury”). Tym samym dioda LED jest zbudowana z dwóch warstw półprzewodnika – warstwy typu p i typu n. Kiedy do diody LED zostanie podłączone stałe napięcie o odpowiedniej polaryzacji (tak zwana polaryzacja przewodzenia, gdzie dodatnia strona zasilania jest połączona z warstwą typu p, a ujemny do katody, elektrony przemieszczają się z warstwy n do warstwy p, a „dziury” przesuwają się w przeciwnym kierunku, w stronę warstwy n. W obszarze złącza elektrony spotykają się z „dziurami”. W wyniku rekombinacji (łączenia się elektronów i dziur) energia jest uwalniana w postaci światła. Kolor światła zależy od rodzaju półprzewodnika oraz użytych domieszek, które określają energię fotonów emitowanych przy rekombinacji. Na przykład diody LED emitujące światło czerwone są zwykle wykonane z arsenku

galu (GaAs), a te emitujące światło niebieskie – z azotku galu (GaN).

Polaryzacja diody LED

W poprzednim akapicie wspomniano, że podłączając napięcie do diody LED należy zachować odpowiednią polaryzację zasilania. Oznacza to, że dioda LED ma dwa wyprowadzenia, z których jedno jest elektrodą dodatnią (anodą) a drugie – elektrodą ujemną (katodą). Anoda to wyprowadzenie, do którego należy dostarczyć „+” zasilania a katoda do wyprowadzenie, do którego należy podłączyć „-” zasilania. Skąd wiadomo, gdzie w diodzie LED jest umiejscowiona anoda a gdzie katoda? Anoda jest zazwyczaj wyprowadzeniem dłuższym niż katoda (**fotografia 6**). To pierwszy znak rozpoznawczy, którym najłatwiej się posługiwać. Innym znakiem rozpoznawczym jest charakterystyczne ścięcie na obwodzie podstawy diody LED umiejscowione przy katodzie (elektrodzie ujemnej diody) pokazane na **rysunku 9**. Dla wprawnego oka ścięcie to jest widoczne i jednoznacznie wskazuje ujemne wyprowadzenie diody, jednak osoby początkujące miewają nie lada problem w zauważaniu i rozpoznawaniu tego ścięcia. Trzecim znakiem rozpoznawczym jest wielkość elektrod wewnątrz obudowy diody. Gdy zajrzymy do wnętrza obudowy diody LED, ustawiając ją wcześniej pod światło, dostrzeżemy, że jedna z tych elektrod jest znacznie większa i szersza od drugiej (**fotografia 6**). Ta szersza elektroda, na której centralnym punkcie zamontowany jest chip półprzewodnikowy i odbłyśnik to z reguły katoda. „Z reguły” ponieważ zdarzają się pewne odstępstwa, w związku z czym, kierowanie się tą ostatnią cechą diody LED bywa zwodnicze. W przypadku wątpliwości warto diodę LED przetestować za pomocą multimetru skonfigurowanego w tryb badania diod lub pomiaru ciągłości obwodu. Jeśli przewody multimetru są poprawnie zamontowane (**fotografia 7**), po przyłożeniu czerwonej sondy multimetru do anody diody LED, a czarnej sondy multimetru do katody diody LED, dioda LED powinna się zaświecić.

Gdyby się okazało, że wyprowadzenia diody LED zostały już skrócone i nie wiadomo, która nóżka była dłuższa, a która krótsza, każda okrągła dioda LED ma na obwodzie swojej podstawy płaskie ścięcie, wskazujące katodę. Pokazano je strzałką na **rysunku 9**. Ponadto – jak już wspomniano – katodę i anodę diody LED można odzyskać za pomocą multimetru (**fotografia 7**).

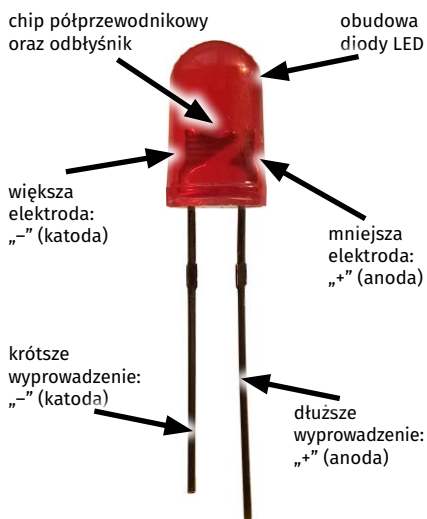
Co się stanie, gdy diodę LED podłączę odwrotnie?

Kiedy dioda LED zostanie podłączona odwrotnie, czyli w tak zwanej polaryzacji zaporowej (dodatni biegun źródła zasilania do katody, a ujemny do anody), prąd przez nią praktycznie nie płynie, ponieważ złącze p-n blokuje przepływ elektronów. W efekcie dioda się nie zaświeci. Dioda LED, podobnie jak inne diody półprzewodnikowe, charakteryzuje się jednak pewnym ograniczonym, maksymalnym napięciem wstępnym (zaporowym), które może wytrzymać. W chwili jego przekroczenia, może dojść do przebicia złącza, co zwykle powoduje uszkodzenie diody. Diody LED mają zazwyczaj niskie maksymalne napięcie zaporowe (zwykle od 5 do 7 V), więc przypadkowe podłączenie diody do źródła o wyższym napięciu w odwrotnej polaryzacji może skutkować trwałym uszkodzeniem diody. W większości typowych aplikacji z diodami LED zasilanymi napięciem rzędu 3 V (jak przy zasilaniu z pojedynczej baterii) napięcie zaporowe rzadko osiąga poziom, przy którym mogłoby dojść do uszkodzenia diody. LED podłączony odwrotnie przy takim napięciu po prostu nie przewodzi, więc nie świeci, ale i nie ulega uszkodzeniu. Natomiast w przypadku większych systemów lub zasilania ze źródeł o wyższych napięciach (np. 12 V lub 24 V, jak w instalacjach samochodowych lub przemysłowych), napięcie zaporowe może przekroczyć bezpieczny limit dla diody LED i spowodować nieodwracalne jej uszkodzenie.

Diody LED migające i RGB zachowują się inaczej

Diody LED użyte w omawianych tu świątecznych zestawach są nieco inne niż te klasyczne. Pomysłowy Mikołaj LED (XMAS5) oraz Elf Pomocnik LED (XMAS6) zawierają migające diody LED. Do zestawu Wesołego Bałwanka LED (XMAS7) dołączono diodę RGB wolno zmieniającą swą barwę i dzięki temu mieniącą się wszystkimi barwami tęczy. Migające diody LED zawierają w swojej strukturze zintegrowany przerywacz obwodu, a wolnozmieniana dioda RGB zaawansowany kontroler odpowiednio sterujący trzy diody LED o podstawowych barwach.

Podczas zeszłorocznych zajęć stacjonarnych zauważyliśmy, że diody te (zespoły diod) zachowują się nieco inaczej niż te klasyczne. Po omyłkowym wlutowaniu diody RGB do płytki w odwrotnym kierunku i załączeniu zasilania, źle wlutowana dioda LED, nie dość że nie zaświeci się, to dosyć szybko



Fotografia 6. Budowa diody LED wraz z umiejscowieniem anody i katody

zacznie się ona nagrzewać i wewnętrznie uszkadzać. Im szybciej odłączymy wtedy zasilanie tym większa szansa na to że diodę da się uratować poprzez jej zdemontowanie i ponowne zamontowanie, już we właściwym kierunku. Obserwacje każą nam przypuszczać, że źle zamontowana dioda LED RGB uszkadza się etapami (uszkodzeniu ulegają kolejne „kolory”). W przypadku pomyłki w montażu, jeśli posiadamy zapasowe diody LED RGB warto wymienić diodę na nową. W przeciwnym wypadku warto sprawdzić czy po naprawie zaświecają się naprzemiennie wszystkie trzy kolory podstawowe: czerwony, zielony oraz niebieski oraz upewnić się czy dioda LED po poprawce montażu nie nagrzewa się do odczuwalnie dużej temperatury. Jeżeli dioda LED nagrzewa się w sposób nadmierny należy wymienić ją na nową w celu uniknięcia trudnych do przewidzenia skutków jej działania w dłuższym okresie czasu.

Rezystor szeregowy

W akapicie dotyczącym sposobu działania diod LED napisałem, że działają one na zasadzie zjawiska elektroluminescencji, które polega na emisji światła w wyniku przepływu prądu przez specjalnie skonstruowane złącze półprzewodnikowe. Wartość tego prądu musi mieścić się w zakresie zalecanym w dokumentacji udostępnionej przez producenta diody LED. Ponieważ dla uzyskania odpowiednich kolorów diod LED ich złącza p-n budowane są z różnych materiałów półprzewodnikowych, materiały te mają również odmienne charakterystyki elektryczne. Wynika z nich również prąd przewodzenia właściwy dla diody LED w określonym kolorze oraz napięcie przewodzenia, które przy takim prądzie się na danej diodzie LED odłoży. Prąd przewodzenia diody LED w przypadku większości typowych diod LED mieści się w zakresie od 10 do 20 mA, niemniej tę właściwą wartość należałoby odczytać wprost z noty katalogowej producenta danej diody LED. Prąd płynący przez diodę LED ustala się za pomocą szeregowo połączanego z nią rezystora. Przy prądzie przewodzenia diody LED zalecanym przez producenta odłoży się na niej określone napięcie przewodzenia, którego wartość również możemy odczytać wprost z noty katalogowej tego elementu. Znając napięcie zasilania (w przypadku jednego ogniwa CR2032, w chwili gdy bateria ta jest w pełni naładowana, napięcie to wyniesie 3 V) oraz napięcie przewodzenia diody LED, odczytane z noty katalogowej jej producenta, możemy obliczyć napięcie, które musi

odłożyć się na rezystorze. Teraz pozostaje już tylko posłużyć się (być może poznanym w szkole na lekcji fizyki) prawem Ohma:

$$R = \frac{U}{I}$$

by obliczyć wartość rezystora. Dla przykładu, wyliczmy jaki rezystor należałoby zastosować, aby w sposób bezpieczny zasilić pojedynczą, czerwoną diodę LED, która według producenta powinna być zasilana prądem $I_f = 10 \text{ mA}$ i na której przy takim prądzie przewodzenia powinno odłożyć się napięcie przewodzenia o wartości $U_f = 1,7 \text{ V}$. Obliczenia będą następujące:

$$R_{LED} = \frac{U_{zas} - U_f}{I_f} = \frac{3V - 1,7V}{10mA} = \frac{1,3V}{0,01A} = 130\Omega$$

Mniej więcej tyle należałoby wiedzieć na temat biegunowości baterii, konieczności zachowania prawidłowej polaryzacji diod LED (właściwego kierunku montażu) i odpowiedniego podłączania diod LED do zasilania. Stosowanie rezystora szeregowego dla diody LED w powyższych zestawach pominięto z uwagi na użycie specjalnych diod LED (migających i RGB), w których przepływem odpowiedniego prądu zarządza zintegrowany w nich układ sterujący. Gdybyś jednak chciał zastosować zwykłą, świecącą światłem ciągłym diodę LED, w miejsce tej dostarczonej w zestawie, pamiętaj koniecznie o włączeniu szeregowego rezystora o wartości około 130Ω . Będziesz musiał nieco pogimnastykować się z jego zamontowaniem, ponieważ żadna z opisanych tu płytek nie uwzględnia miejsca na jego podłączenie.

Mając na uwadze wszystko powyżej, możemy przejść do omówienia procesu montażowego każdego ze wspomnianych zestawów.

Montaż układu

Montaż układów elektronicznych warto rozpoczynać od komponentów o najniższym profilu a kończyć na tych najwyższych. Taka kolejność ułatwia pracę: płytka stabilnie leży na stole, a niskie elementy nie są zasłaniane przez wyższe, co pozwala na wygodne lutowanie i estetyczny montaż. W przypadku opisywanych tu świątecznych zestawów komponenty są zaledwie trzy. Każdy z nich duży gabarytowo i nieco fikuśny, więc kolejność ich montażu nie ma większego znaczenia. Przy ich lutowaniu przyda się pomocna dłoń dorosłego opiekuna, który na czas montażu przytrzyma przełącznik lub gniazdo baterii, by prostopadłe przylegało do płytki, przynajmniej do czasu przylutowania pierwszego wyprowadzenia. Proponuję

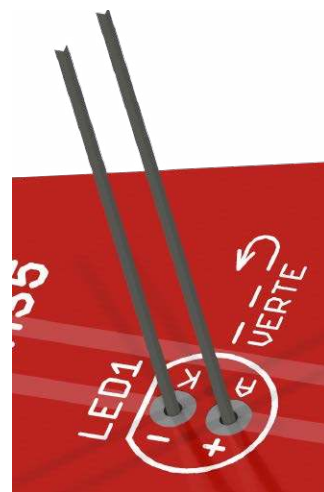
jednak zacząć od starannego przylutowania diody LED.

Montaż diody LED

Podczas montażu diody LED należy zachować szczególną ostrożność, ponieważ każda dioda jest elementem spolaryzowanym. Każda dioda posiada anodę, którą należy podłączyć do dodatniego potencjału zasilania oraz katodę, którą podłącza się do ujemnego bieguna zasilania. W przypadku diod LED anoda jest zawsze wyprowadzeniem dłuższym a katoda jest krótsza (fotografia 6 oraz rysunek 4).

Rysunek 4 oraz rysunki 5a oraz 5b przedstawiają poprawne włożenie diody LED do płytki PCB.

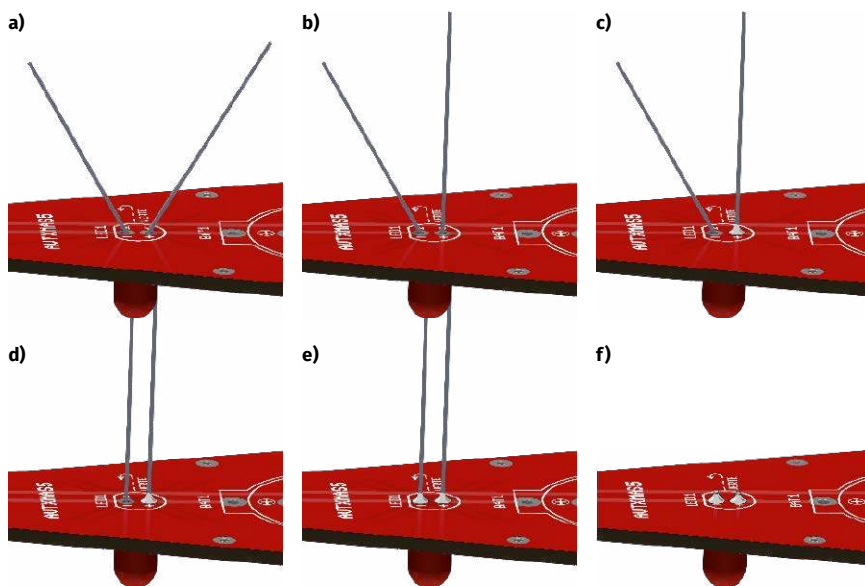
Podczas montażu diod LED zawsze warto mieć na uwadze, że może się ona uszkodzić już na etapie składania zestawu i konieczna będzie jej wymiana. Warto być na taką okoliczność przygotowanym, stosując się do opisanego poniżej sposobu



Rysunek 4. Napis VERTE oznacza, że diodę LED należy zamontować na odwrocie (zobacz rysunki 5a oraz 5b). Dłuższe wyprowadzenie (anodę) diody LED należy zamontować do otworu oznaczonego znakiem „A” oraz „+”. Krótsze (katodę) – do otworu oznaczonego znakiem „K” oraz „-”. Obrys diody LED nie jest w pełni okrągłym, a jego prosty odcinek również wskazuje katodę



Rysunek 5. a) obudowa diody LED stanowi nos Mikotaja, stąd znajduje się po stronie jego twarzy; b) wyprowadzenia diody LED skierowane są w stronę spodu, gdzie w następnym etapie zostaną zamontowane przełącznik oraz koszyczek z baterią CR2032



Rysunek 6. Zalecany sposób montażu diody LED, który umożliwi jej łatwą wymianę w przypadku uszkodzenia lub omyłkowego zamontowania w odwrotnym kierunku: a) rozegnij wyprowadzenia diod LED w taki sposób, by dało się bezpiecznie odwrócić płytkę do góry nogami; b) po ustawieniu płytki na blacie odegnij jedno wyprowadzenie do pionu; c) przylutuj odgięte do pionu wyprowadzenie prostopadle do płytki; d) ustaw do pionu drugie wyprowadzenie; e) przylutuj je prostopadle do płytki; f) obetnij nadmiar wyprowadzeń przy użyciu obcinaczek

montażu przyjaznego naprawom. Od tego, jak zamontujemy diodę LED, zależy, czy ewentualna naprawa zestawu okaże się udręką, czy raczej czystą przyjemnością.

Zaginanie wyprowadzeń elementów THT (czyli przeznaczonych do montażu przewlekane) może wydawać się wygodne – stabilizuje je w otworach i zapobiega wypadaniu przed lutowaniem. Niestety, taki sposób montażu znacznie utrudnia późniejszy demontaż uszkodzonych lub błędnie zamontowanych komponentów. Szczególnie kłopotliwe bywa wylutowanie właśnie diody LED z twardymi wyprowadzeniami zagiętymi pod mniejszym lub większym kątem do płytki – często kończy się to uszkodzeniem samego elementu, a czasem także płytki drukowanej.

Aby tego uniknąć, po umieszczeniu elementu w otworach warto jego wcześniej

zagięte piny wyprostować przed lutowaniem – tak, by znów były ustawione prostopadle do płytki. Komponent nie wypadnie, ponieważ po odwróceniu całości leży oparty o blat roboczy i jest przytrzymywany przez samą płytkę.

Taka technika montażu znacznie ułatwia ewentualną wymianę elementu – skrócone, niezablokowane piny łatwo przechodzą przez otwory po ponownym rozgrzaniu lutu.

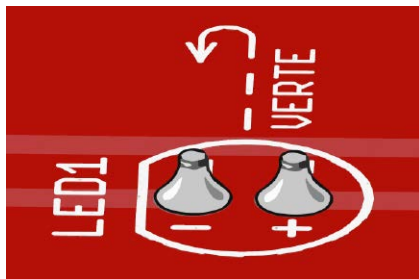
Rysunki od 6a do 6f pokazują, jak przylutować diodę LED w sposób umożliwiający jej łatwą wymianę w przyszłości.

Przycinanie nadmiaru wyprowadzeń

Nadmiar wyprowadzeń uprzednio przylutowanych komponentów do montażu przewlekane obetnij przy użyciu obcinaczek. Pamiętaj, aby nie ścinać lutu – cięcie



Rysunek 7. Właściwe miejsce cięcia nadmiaru wyprowadzeń po przylutowaniu komponentu. Po wykonaniu cięcia spoina powinna pozostać nienaruszona



Rysunek 8. Po przylutowaniu komponentów przewlekanych do płytki nadmiar wyprowadzeń został obcięty pewnym ruchem ściskającym, bez szarpania i ciągnięcia, dzięki czemu połączenia pół lutowniczych z płytką i biegnące po niej ścieżki pozostały nienaruszone

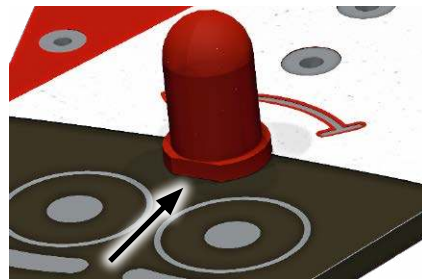
wykonuj dopiero za miejscem, w którym kończy się spoina i wystaje samo wyprowadzenie (**rysunek 7**). Pamiętaj, aby nie szarpać ani nie ciągnąć za wyprowadzenie przylutowanego komponentu. Mogłoby to doprowadzić do dość kłopotliwego w naprawie wyrwania pola lutowniczego z płytki i zerwania połączenia komponentu ze ścieżką. Poprawnie obcięte wyprowadzenia pokazano na **rysunku 8**.

Przykład poprawnie zamontowanej diody LED pokazano na **rysunku 9**.

Należy pamiętać, że jeśli przylutujemy diodę LED w niewłaściwym kierunku, nie będzie ona świeciła, a ponadto, na skutek wymuszonego przepływu prądu wstecznego może ona ulec trwałemu uszkodzeniu. Zdarza się, że gdy zorientujemy się, że diodę LED zamontowaliśmy w sposób nieprawidłowy, po jej wylutowaniu i ponownym przylutowaniu, już we właściwym kierunku, dioda wciąż nie będzie chciała świecić. Dlatego po wylutowaniu błędnie zamontowanej diody LED warto sprawdzić jej kondycję przy użyciu multimetru – ustawionego w tryb pomiaru diod lub testu ciągłości obwodu – aby upewnić się, że element nadal działa poprawnie. W tym celu, po ustawieniu wspomnianego trybu pracy multimetru, do anody diody LED przykładamy czerwoną jego sondę a do katody przykładamy sondę czarną (**fotografia 7**). Jeśli w tym momencie dioda LED się zaświeci, oznacza to, że jest ona sprawna, i możemy przylutować ją ponownie, tym razem pamiętając o właściwym kierunku montażu.

Montaż przełącznika zasilania

Pora na zamontowanie przełącznika zasilania, opisanego na schemacie i PCB jako SW1. Przełącznik tego typu łączy swój pin środkowy z jednym z dwóch skrajnych, w którego kierunku jest w danym momencie skierowany hebelelek przełącznika. Z uwagi na taką konstrukcję, kierunek montażu



Rysunek 9. Przykład poprawnie zamontowanej diody LED. Płaskie ścięcie na obwodzie podstawy diody LED, symbolizujące katodę, musi pokrywać się z fragmentem linii prostej w obrysie diody LED widocznym po drugiej stronie płytki PCB (rysunek 4)

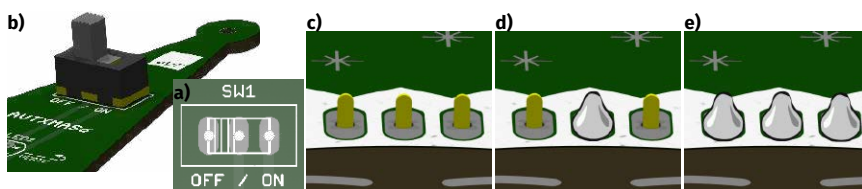


Fotografia 7. Sprawdzenie diody LED za pomocą multimetru ustawionego na funkcję testowania diod. Po przyłożeniu sondy czerwonej do anody, a czarnej do katody, sprawna dioda LED powinna się zaświecić. Jeśli dioda ma odpowiednio długie (jeszcze nie przycięte) wyprowadzenia można się wspomóc krokodylkami

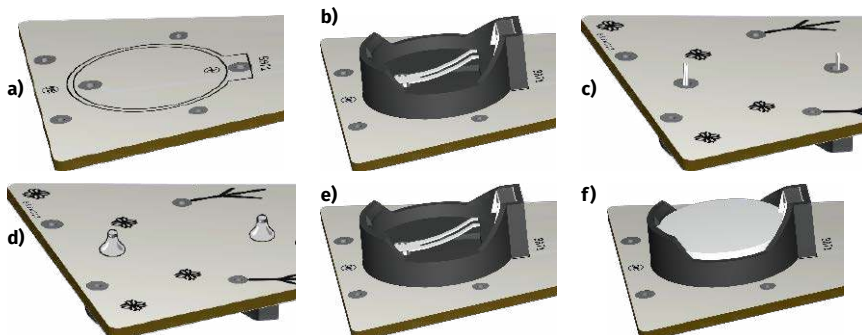
tych elementów nie ma żadnego znaczenia. Warto o tym pamiętać montując przełącznik do płytki. Dodatkowe linie, symbolizujące pozycję hebelka w obrysie komponentu na płytce (**rysunek 10a**), mają jedynie rozwiać wątpliwości, że w tym miejscu należy zamontować przełącznik. Z polaryzacją lub kierunkiem montażu, czy też „wymaganą” pozycję hebelka nie mają niczego wspólnego. Przełącznik należy włożyć do płytki w taki sposób, by wszystkie trzy jego wyprowadzenia przeszły na drugą stronę (**rysunek 10c**). Przełącznik ma sztywne piny, które nie nadają się do wyginania. Dlatego podczas lutowania trzeba przytrzymać go ręką albo poprosić o pomoc opiekuna, którego – mam nadzieję – podczas zabaw lutownicą masz zawsze w pobliżu. Podczas montażu warto przylutować środkowy pin do płytki PCB (**rysunek 10d**), a dopiero po upewnieniu się, że komponent dobrze przylega do jej powierzchni, przylutować pozostałe (**rysunek 10e**). Poprawnie zamontowany przełącznik SW1 pokazano na **rysunku 10b**.

Montaż uchwytu baterii

Trzecim elementem do przylutowania jest uchwyt na baterię typu CR2032. Uchwyt ten należy włożyć do płytki PCB sugerując się obrysem komponentu na płytce, w taki sposób, by plus jak i minus trafiły na odpowiednie miejsca (**rysunki 11a i 11b**). Po upewnieniu się, że oba wyprowadzenia uchwytu baterii trafiły w otwory i wystają ponad dolną powierzchnię płytki na tę samą wysokość (**rysunek 11c**) należy je przylutować (**rysunek 11d**). Gdy uchwyt jest już przylutowany do płytki (**rysunek 11e**)



Rysunek 10. a) oznakowanie miejsca na płytce, przygotowanego pod montaż przełącznika SW1; b) poprawnie włożony do płytki przełącznik SW1; c) wszystkie wyprowadzenia przełącznika zostały poprawnie przeprowadzone przez płytkę PCB na tę samą długość; d) zanim przylutujesz wszystkie wyprowadzenia, przylutuj w pierwszej kolejności pin środkowy i sprawdź, czy komponent przylega prostopadle do płytki; e) jeśli wszystko jest w porządku przylutuj dwa pozostałe wyprowadzenia przełącznika SW1 do płytki



Rysunek 11. a) obrys uchwytu baterii ułatwiający zamontowanie uchwytu we właściwym kierunku; b) uchwyt baterii umieszczony w otworach PCB zgodnie z obrysem komponentu; c) wyprowadzenia uchwytu baterii powinny wystawać z dolnej powierzchni płytki na tę samą długość; d) poprawnie przylutowane do płytki wyprowadzenia uchwytu baterii; e) przylutowany do płytki uchwyt baterii; f) bateria CR2032 zamontowana do uprzednio zamontowanego i przylutowanego do płytki uchwytu

można w nim umieścić baterię CR2032 (**rysunek 11f**).

Wyprowadzenia uchwytu baterii CR2032 są najprawdopodobniej na tyle krótkie, że nie trzeba ich dodatkowo przycinać. Gdyby było inaczej, kieruj się wskazówkami z akapitu *Przycinanie nadmiaru wyprowadzeń*.

Lutowanie nóżek, rączek i miotel

Pamiętajcie kolegę Gutka, Eko-Czarodzieja, który budował fantastyczne figurki z zużytych komponentów elektronicznych, którego niektóre z prac pokazaliśmy

w EdW 10/2024? Gutek w unikalny sposób połączył świat techniki ze światem sztuki, a przy okazji zainteresował swym wartościowym hobby i eko-ideami kolegów ze swojej szkoły. Pomysłowy Mikołaj LED (XMAS5), Elf Pomocnik LED (XMAS6) oraz Wesoły Bałwanek LED (XMAS7) również łącząc oba te światy. Pozwalają z jednej strony zbudować z baterii, włącznika i diody LED, prosty ale w pełni działający obwód elektryczny, a jednocześnie rzucić się w wir fantazji i trenować sprawności manualne poprzez przyprowadzanie świątecznym figurkom rąk i nóg, a nawet miotel.



Rysunek 12. Od lewej: Pomysłowy Mikołaj LED (XMAS5), Elf Pomocnik LED (XMAS6) oraz Wesoły Bałwanek LED (XMAS7) – jeszcze przed przyprowadzeniem rączek, nóżek i miotel



Fotografia 8. Dawid, Kornel i Grzegorz podczas montażu świątecznych zestawów: Pomysłowy Mikołaj LED (XMAS5), Elf Pomocnik LED (XMAS6) oraz Wesoły Bałwanek LED (XMAS7). Koło Młodych Entuzjastów Elektroniki, Szkoła Podstawowa nr 86, Wrocław

Na rysunku 12 pokazano zmontowane, elektrycznie w pełni już funkcjonalne figurki. Czekają na jeszcze plastyczną metamorfozę.

Celem przypięcia rączek, nóżek i mioteł wykorzystaj dołączone do zestawów rezystory nieco większej mocy (Mikołaj, Elf i Bałwanek), lub fragmenty srebrzanki (Bałwanek) nie pełniące w układzie roli elektrycznej, za to świetnie imitujące brakujące części ciała. Dzięki nim dodamy tym figurkom nie tylko uroku i świątecznego klimatu, ale sprawimy również, że będą mogły na tych kończynach się wesprzeć, a nawet samodzielnie usiąść. Nie wszystko i nie zawsze musi przecież pełnić funkcję

obwodu elektrycznego, by mogło bawić, uczyć, inspirować i cieszyć!

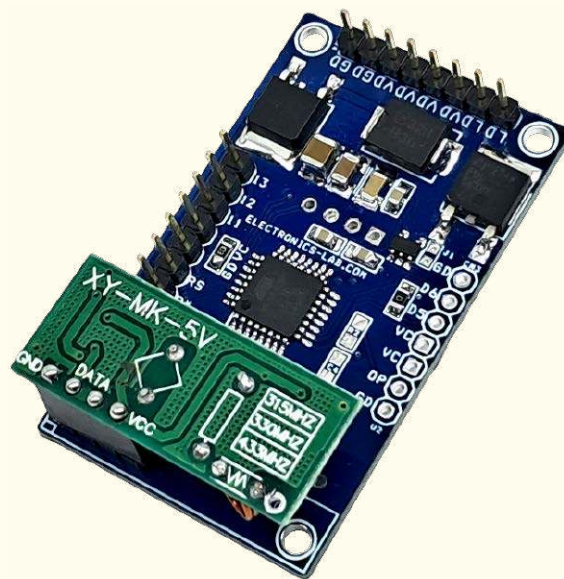
Zmontowane podczas zajęć stacjonarnych przez naszych Juniorów figurki prezentujemy na fotografiach.

Uruchomienie układu

Układ jest tak banalny, że jego uruchomienie powinno ograniczyć się do włożenia baterii i ustawienia przełącznika SW1 w pozycji „ON”. Gdyby jednak coś nie chciało zadziałać, namierzanie ewentualnej usterki w wyjątkowo prostym układzie, jaki stanowi każda z opisywanych tu świątecznych figurek, może być prawdziwą przygodą, a rozwiązanie zagadki

dającym masę frajdy, samodzielnym osiągnięciem. Udostępniona na Elportalu (<https://elportal.pl/do-pobrania>) w materiałach uzupełniających do EdW z grudnia 2024 lista możliwych usterek powinna znacząco uprościć i uprzyjemnić całą zabawę. Istnieje natomiast olbrzymie prawdopodobieństwo graniczące z pewnością, że zmontowany układ od razu zadziała poprawnie i natychmiast po zmontowaniu sprawi Ci wiele frajdy. Powodzenia w montażu i uruchomieniu wszystkich świątecznych ozdób. Do zobaczenia już w nowym roku kalendarzowym 2026! Ho ho ho! Wesołych Świąt i Szczęśliwego Nowego Roku! ■

Mariusz Ciszewski



Bezprzewodowy nadajnik ściemniacza LED 434 MHz – kompatybilny z Arduino

Przedstawiony tutaj projekt to bezprzewodowy nadajnik 434 MHz do ściemniacza LED, zbudowany w oparciu o mikrokontroler ATMEGA328 Arduino. Potencjometr jest podłączony do analogowego pinu ADC A1 układu ATMEGA328. Mikrokontroler odczytuje wartość analogową potencjometru i przesyła dane przez moduł RF 434 MHz. Dane te są odbierane przez odbiornik, który steruje intensywnością diod LED za pomocą sygnału PWM (modulacja szerokości impulsu).

<https://youtu.be/VsmB-hk0rul>
<https://youtu.be/Fgf4910d-RA>

Odbiornik bezprzewodowy 434 MHz – ściemniacz LED – kompatybilny z Arduino

Bezprzewodowy odbiornik ściemniacza jest komponentem uzupełniającym bezprzewodowy nadajnik potencjometru, zaprojektowanym do odbioru i przetwarzania danych przesyłanych z nadajnika. Projekt ten może być wykorzystany do sterowania diodami LED, grzejnikami lub silnikami szczotkowymi prądu stałego i obsługuje obciążenia indukcyjne lub rezystancyjne. Do sterowania bramką MOSFET służy układ scalony MCP1416.

<https://youtu.be/VsmB-hk0rul>
<https://youtu.be/Fgf4910d-RA>

Niektóre projekty aktualnie dostępne tylko dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl:

1. Wskaźnik poziomu paliwa z wyświetlaczem OLED
2. Bezprzewodowy odbiornik wilgotności i temperatury
3. Przerwania zewnętrzne (sprzętowe) i przerwania zegara w MicroPython
4. Laserowy czujnik odległości z wyświetlaczem OLED i RP2040
5. Inklinometr z 17-segmentowym wyświetlaczem słupkowym
6. Izolowany repeater USB – USB 2.0
7. Knight Rider Light – 16 diod LED dużej mocy (kompatybilny z Arduino)
8. Dźwięk do kolorowych efektów świetlnych (kompatybilny z Arduino)
9. Nowy i ulepszony licznik Geigera – teraz z Wi-Fi!
10. Detektor zalania
11. Lampa sufitowa LED o dużej mocy
12. Kontroler dzwonów kościelnych
13. Arduino Nano – włączanie/wyłączanie urządzeń za pomocą pilota na podczerwień (dwa kanały)
14. Lampa sufitowa LED z czujnikiem ruchu PIR – kompatybilna z Arduino
15. Inteligentny ściemniacz LED z Bluetooth – 4-kanałowy włącznik/wyłącznik Bluetooth
16. Czterokanałowy izolator cyfrowy, wzmacniory, szybki, o niskim poborze mocy
17. Sterowanie prędkością, kierunkiem i zatrzymaniem silnika DC z modułem RF NRF24L01
18. Nadajnik zdalnego sterowania z pojedynczym joystickiem wykorzystujący NRF24L01
19. 8-kanałowy zdalny nadajnik RF z protokołami: Holtek i szeregowym
20. 8-kanałowy zdalny odbiornik RF z protokołami: Holtek i szeregowym
21. Pojemnościowy czujnik wilgotności do konwertera wyjścia analogowego
22. Mostek H dla wysokiej mocy szczotkowego silnika prądu stałego z czujnikiem prądu
23. Przetwornica DC-DC buck 12...75 V na 10 V na wyjściu
24. Czujnik prądu low-side 10 µA...10 mA
25. Kontroler ramienia robota z bezprzewodowym pilotem PS3
26. Termiczny czujnik masowego przepływu powietrza – anemometr statotemperaturowy
27. Precyzyjny wzmacniacz transimpedancyjny z przetłaczanym integratorem
28. Wysokowydajny monofoniczny wzmacniacz audio klasy D o mocy 20 W
29. Kontroler pełnego mostka z przesunięciem fazowym i prostowaniem synchronicznym wykorzystujący UCC28950
30. Monitorowanie poziomu cieczy za pomocą czujnika ciśnienia – wyświetlacz słupkowy
31. Sterowanie silnikiem DC za pomocą joysticka
32. 16-kanałowy sterownik serwo mechanizmów RC z interfejsem I²C
33. Programowalny kondycjoner sygnału z czujnika rezystancyjnego mostkowego
34. Choinka z Arduino i pikselowymi diodami
35. 20-segmentowy wyświetlacz słupkowy w rozmiarze jumbo
36. Stacja pogodowa lilygo ttgo t5-4.7 z wyświetlaczem typu e-papier

Miesięcznik „Elektronika dla Wszystkich” (12 numerów w roku) jest wydawany w współpracy z kilkoma redakcjami zagranicznymi



Wydawnictwo:
AVTKorporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Redaktor naczelny:
Mariusz Ciszewski
mariusz.ciszewski@elportal.pl

Adres redakcji:
03-197 Warszawa, ul. Leszczyńska 11
e-mail: edw@elportal.pl, www.elportal.pl

Dział reklamy:
Katarzyna Gugała
katarzyna.gugala@elportal.pl, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobański
jakub.sobanski@elportal.pl

Sekretarz redakcji:
Dariusz Welik
dariusz.welik@elportal.pl

Copyright AVTKorporacja Sp. z o.o., Warszawa, ul. Leszczyńska 11. Projekty publikowane w „Elektronice dla Wszystkich” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki dla Wszystkich”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice dla Wszystkich” jest dozwolone wyłącznie po uzyskaniu pisemnej zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice dla Wszystkich”.

DTP, redakcja strony internetowej www.elportal.pl:
MAD Sp. z o.o.

Prenumerata:
W Wydawnictwie AVT, e-mail: prenumerata@avt.pl
tel. 22 257 84 22, (godz. 10:00–14:00)
www.ulubionykiosk.pl



TRZECIARĘKA ZD-11P
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt z latarką, ZD11P



TRZECIARĘKA ZD-11P-1
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt z latarką i lupą, ZD11P-1



TRZECIARĘKA SN-394
Uchwyt montażowy typu „Trzecia ręka”,
pająk z lupą 50 mm, przykręcany do blatu
Proskit SN-394

BESTSELLERY sklepu AVT – sklep.avt.pl

Trzecia ręka

Rabat dla Czytelników EdW
przy zakupie podaj kod **EdW2505TR**

Kod ważny do 30.09.2025

-3%

Rabat dla Prenumeratorów EdW
przy zakupie podaj numer prenumeraty

-6%



TRZECIARĘKA ZD-11M-1
Uchwyt montażowy typu „Trzecia ręka”,
pająk – z uchwytem na szpulkę cyny, ZD11M-1



TRZECIARĘKA ZD-11M-2
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt z lupą i podświetleniem LED
ZD11M-2



TRZECIARĘKA ZD-11M-3
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt z lupą i podświetleniem LED
ZD-11M-3



TRZECIARĘKA ZD-11M
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt ZD11M



TRZECIARĘKA SN-392
Uchwyt montażowy typu „Trzecia ręka”
z lupą 90 mm, Proskit SN-392



TRZECIARĘKA
Uchwyt montażowy typu „Trzecia ręka”
z lupą 60 mm

Subscribe to Elektor's newsletter and get the chance to

WIN

a Raspberry Pi Pico W board



www.elektor.com/eda



Subscribe to Elektor's newsletter, get a €5 coupon code and get the chance to WIN a Raspberry Pi Pico W board



Be one of the 10 fortunate winners!



elektor
design > share > earn