

Przetwornica typu flyback krok po kroku (3)



Projektowanie przetwornicy typu flyback za pomocą programu Webench

Topologia flyback jest najczęściej wybierana dla przetwornic odizolowanych galwanicznie od źródła zasilania, ponieważ umożliwia uzyskanie wielu napięć wyjściowych z użyciem tylko jednego tranzystora kluczującego i stosunkowo niewielkiej liczby komponentów zewnętrznych. Jednak mimo nieskomplikowanej budowy przetwornice o topologii flyback mają pewne cechy szczególne, które mogą ograniczać zakres ich zastosowań, jeśli nie zostaną zrozumiane i dogłębnie przeanalizowane przez konstruktora. Posługując się matematyką oraz programem Webench w kolejnych artykułach odkryjemy tajemnice konstrukcyjne tych przetwornic oraz podamy porady umożliwiające wykonywanie optymalnych konstrukcji.

W ostatniej części artykułu, omówimy metody kompensowania wzmacniacza błędu odpowiedzialnego za stabilizację napięcia wyjściowego, sposoby wykonywania odizolowanej galwanicznie pętli sprzężenia zwrotnego oraz kończąc rozważania teoretyczne, pokażemy praktyczne rozwiązania przetwornicy *flyback*.

Metody kompensowania wzmacniacza błędu

Powodem, dla którego wykonuje się obwód kompensujący wzmacniacz błędu jest przeciwdziałanie nierównomierności charakterystyk amplitudowej i fazowej, która mogłaby doprowadzić do niestabilnej pracy regulatora przetwornicy, której częścią jest ten wzmacniacz. Celem nadrzędnym jest takie ukształtowanie charakterystyki funkcji transferu funkcji kontrolnej, aby spełnić kryteria pracy stabilnej. Zatem funkcja transferu obwodów kompensujących jest dodawana do funkcji wyjściowej funkcji transferu obwodów kontrolnych w taki sposób, aby zostały spełnione wymagania statyczne i dynamiczne oraz była zagwarantowana stabilność pracy.

Teoretycznie, idealna pętla sprzężenia zwrotnego wzmacniacza powinna mieć następujące właściwości:

Krótki czas odpowiedzi osiągany dzięki szerokiemu pasmu przenoszenia (wysoka częstotliwość wzmocnienia jednostkowego).

Nachylenie charakterystyki wzmocnienia pętli 20 dB/dekadę, od sygnału o niskiej częstotliwości do sygnału o częstotliwości równiej połowie częstotliwości kluczowania.

Bardzo dobra dokładność regulacji DC – niewielkie zmiany napięcia DC powstałe na skutek zmiany obciążenia opanowywane muszą być niwelowane dzięki bardzo dużemu wzmocnieniu DC.

Dobra odporność na zaburzenia z małym wzmocnieniem dla sygnałów o wysokiej częstotliwości, zbliżonej do częstotliwości kluczowania.

Plaska charakterystyka fazowa dla sygnałów o częstotliwości wzmocnienia jednostkowego.

Margines fazowy zapewniający dobrą stabilność przy minimalnym przeregulowaniu.

Wzmocnienie DC jest wzmocnieniem dla sygnałów o niskiej częstotliwości. Częstotliwość wzmocnienia jednostkowego jest częstotliwością sygnału, przy którym charakterystyka wzmocnienia przekracza linię 0 dB (wzmocnienie jest równe 1). Margines fazy jest wartością, o którą przesunięcie fazowe układu z otwartą pętlą sprzężenia zwrotnego jest mniejsze od 180°, gdy wzmocnienie wynosi 0 dB (jest równe 1).

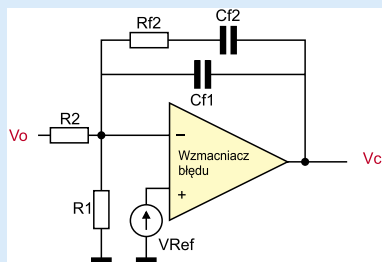
Aby zachować szerokie pasmo przenoszenia przy jednoczesnym ograniczeniu wzmocnienia w obszarze sygnałów o wysokiej częstotliwości, dla przetwornicy *flyback* częstotliwość wzmocnienia jednostkowego powinna być ograniczona do częstotliwości RHPZ po mniejszej o 25% RHPZ.

Biegun na początku charakterystyki wzmocnienia kompensatora jest dołączony w celu zagwarantowania nominalnego poziomu błędu w okolicach stabilnego zera. Rozmieszczenie zer i biegunów na charakterystyce wzmacniacza błędu może być łatwo określone z użyciem metody „współczynnika K”. Bazuje ona na prostej koncepcji. Mając daną częstotliwość wzmocnienia jednostkowego f_c , dla której wzmocnienie wynosi 0 dB oraz wymagany margines fazy ϕ_{fc} i początkowy margines fazy nieskompensowanej pętli wzmacniacza ϕ_{mu} przy częstotliwości f_c , konstruktor może obliczyć podbicie fazy ϕ_{pb} , które musi zapewnić obwód kompensujący przy częstotliwości f_c :

$$(34) \quad \phi_{pb} = \phi_{fc} - \phi_{mu} + 90^\circ$$

Dodatkowe przesunięcie fazowe „+90°” wprowadzono po to, aby skompensować przesunięcie fazowe o -90° wynikające z umieszczenia bieguna na początku charakterystyki.

Typowo, przetwornica zasilająca *flyback* pracująca w trybie kontroli prądu szczytowego ma podbicie



Rysunek 12. Obwód kompensujący typu drugiego

fazy zawierające się w zakresie $0 < \phi_{hb} < 70^\circ$, w którym to dla zapewnienia wymaganego przesunięcia fazowego jest niezbędna kompensacja typu drugiego. Obwód kompensacyjny *Type2*, będący częścią schematu z **rysunku 12**, ma jedno zero i jeden biegun ułożone odpowiednio przy $\omega_{z1} = f_c/K$ i $\omega_{p1} = f_c \cdot K$, gdzie $K = \tan(\phi_{hb}/2 + 45^\circ)$.

Funkcja transferu stopnia kompensacji $B_{error}(s)$, przy pominięciu wpływu na jej kształt parametrów rzeczywistego wzmacniacza operacyjnego, ma następującą postać:

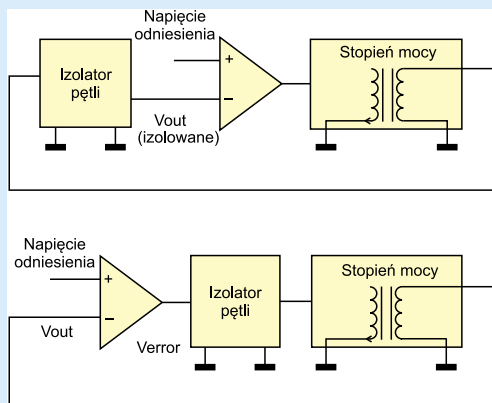
$$(35) \quad B_{error} = \frac{V_c(s)}{V_o(s)} = \frac{1}{R_2 \cdot (Cf_2 + Cf_1)} \cdot \frac{\left(1 + \frac{s}{\omega_{z1}}\right)}{\left(1 + \frac{s}{\omega_{p1}}\right)}$$

Całkowita funkcja transferu zamkniętej pętli sprzężenia zwrotnego jest otrzymywana przez przemnożenie funkcji transferu stopnia kompensacji (wyrażenie 35) przez funkcję transferu funkcji kontrolnej stopnia mocy (wyrażenie 32).

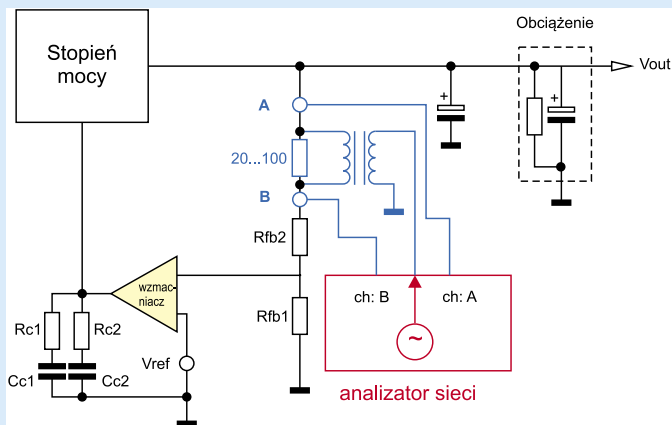
Sposób przekroczenia bariery izolacji

Jedną z głównych zalet przetwornic pracujących w topologii *flyback* jest możliwość wyprowadzenia jednego lub wielu wyjść o masie odizolowanej galwanicznie od masy wejściowej. W dużych systemach zasilania mających wiele szyn, izolacja pomiędzy masami ułatwia uziemienie całego systemu w pojedynczym punkcie i zabezpiecza przed tworzeniem się pętli mas. Często wymagania odnośnie do izolacji są określone przez różne instytucje zajmujące się bezpieczeństwem.

Podstawowym wyzwaniem dla konstruktora przetwornicy jest rzetelne przekazanie informacji o napięciu wyjściowym lub sygnału błędów – napięcia występującego pomiędzy masą odniesienia a inną masą, jak pokazano na **rysunku 13**.



Rysunek 13. Na górze – separacja galwaniczna sygnału wyjściowego, na dole – izolacja galwaniczna sygnału błędów



Rysunek 14. Izolowany obwód sprzężenia zwrotnego z kompensacją typu drugiego

Sygnał sprzężenia zwrotnego, który wraca poprzez granicę izolacji, może być sygnałem proporcjonalnym do napięcia wyjściowego (rys. 13 na górze) lub sygnałem proporcjonalnym do różnicy pomiędzy napięciem wyjściowym i napięciem odniesienia (rys. 13 na dole). Zazwyczaj przez barierę izolacji przekazuje się sygnał błędów (jest to rozwiązanie preferowane przez większość konstruktorów) ze względu na fakt, że jeśli doprowadzi się przez nią do obwodów regulacji napięcie wyjściowe, to jakkolwiek niedokładność wprowadzana przez obwód izolujący będzie powodowała błędy regulacji napięcia wyjściowego.

Typową implementację pętli sprzężenia zwrotnego wzmacniacza z izolacją galwaniczną pokazano na **rysunku 14**. Zewnętrzny wzmacniacz operacyjny LMV431 lub podobny z obwodem kompensacyjnym typu drugiego jest używany do wytworzenia sygnału błędów, który przekracza barierę izolacji galwanicznej za pomocą transoptora. Rezystor R_{c2} i kondensator C_c są zewnętrzną siecią dołączoną do wyjścia transoptora, aby wspomóc stabilizację. CRT jest ilorazem prądów kolektora fototranzystora $I_c(s)$ oraz prądu diody LED $I_d(s)$ będących elementami składowymi transoptora.

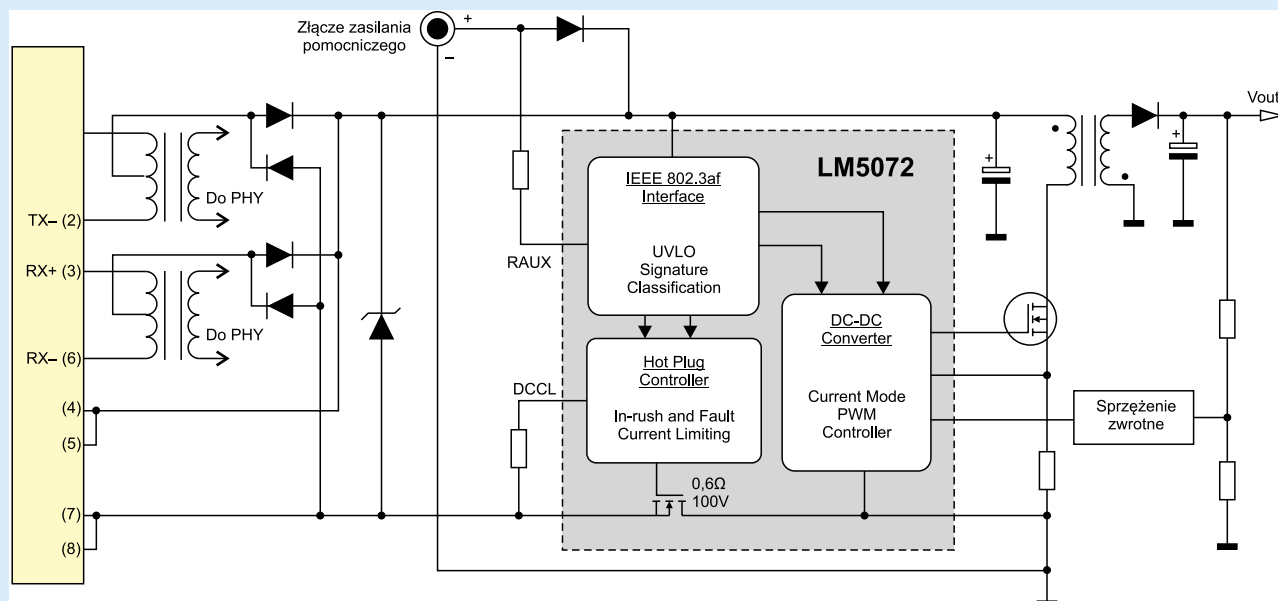
Dodatkowa funkcja transferu sieci optoizolującej (wyjście wzmacniacza błędów do wejścia modulatora) ma postać:

$$(36) \quad B_{opt}(s) = \frac{\hat{v}_e(s)}{\hat{v}_a(s)} = \frac{R_{CTR} \cdot k}{R_d} \cdot \frac{1 + sC_c R_{c2}}{\left(1 + \frac{s}{s_p}\right)(1 + sC_c(R_c + R_{c2}))}$$

Biegun optoizolacji s_p jest zależny od konfiguracji obwodu optycznego (w tym wypadku – wspólny emiter) i może być określony na podstawie znajomości rezystancji małosygnałowej i pojemności transoptora.

Całkowita funkcja transferu sprzężenia zwrotnego jest obwodem kompensacji (równanie 35) omówionym poprzednio przemnożonym przez funkcję transferu obwodu optoizolacji (wyrażenie 36).

Wyjaśniona wyżej metoda, aby modelować całkowity obwód kompensacji, z izolacją i bez niej, nie bierze pod uwagę nieidealnych parametrów wzmacniacza operacyjnego pracującego jako wzmacniacz błędów. Nieidealna odpowiedź częstotliwościowa wzmacniacza operacyjnego mogłaby doprowadzić do zbyt dużej estymacji marginesu fazy i wzmocnienia DC całego systemu. Dlatego też zawsze jest zalecane, aby nie polegać jedynie na modelowaniu matematycznym ale optymalizować wartości komponentów sieci kompensacji za pomocą dokładnych pomiarów laboratoryjnych całkowitego systemu.



Rysunek 15. Typowy obwód aplikacyjny z kontrolerem LM5072

Opracowanie typowego konwertera DC-DC przeznaczonego do zasilania urządzeń PoE

Na rysunku 15 pokazano typową aplikację układu LM5072 zintegrowanego z 100V interfejsem PoE PD i kontrolerem PWM z wejściem zasilania pomocniczego. Kontroler LM5072 zapewnia elastyczność aplikacji przetwornicy akceptując również zasilanie z zewnętrznych źródeł, takich jak zasilacz sieciowy lub bateria akumulatorów. Wysoki stopień integracji kontrolera LM5072 umożliwia spełnienie wymagań normy IEEE 802.3af przy użyciu minimalnej liczby komponentów zewnętrznych.

Sposób pomiaru pętli sprzężenia zwrotnego za pomocą analizatora sieci

Pomimo faktu, że modele matematyczne impulsowych źródeł zasilania są coraz lepsze, nadal mają ograniczoną dokładność. Głównym powodem tych ograniczeń są nieznanne szczegóły związane z funkcjonowaniem systemu. Na przykład: komponenty pasozytnicze, płytka drukowana, oddziaływanie temperatury otoczenia, opóźnienia wynikające z czasów propagacji, nieliniowość komponentów półprzewodnikowych itp. Wyniki pomiarów wykonanych rzeczywistym obwodzie przeważnie odbiegają od matematycznych przewidywań i wysiłki mające na celu otrzymanie dobrych modeli mogą być ekstremalnie trudne oraz czasochłonne. Z tego powodu, zawsze dobrym pomysłem jest zmierzenie funkcji transferu prototypu obwodu. Oczywiście, modele matematyczne są bardzo użyteczne do obliczenia wartości elementów RC przed wykonaniem prototypu i jego pomiarów oraz „dostrojeniem” pętli. Aby mieć jasność – nie jest kwestią czy używać modelu matematycznego, czy wykonać pomiary stabilności rzeczywistego obwodu. Ekspert zajmujący się konstruowaniem przetwornic impulsowych zawsze używa obu metod. Ten tok rozumowania może być zastosowany do jakiegokolwiek przetwornicy impulsowej, niezależnie od jej topologii lub sposobu modulowania sygnału PWM.

Należy zauważyć, że dla stabilnej pracy pętli jest niezbędne ujemne sprzężenie zwrotne. Odpowiedź

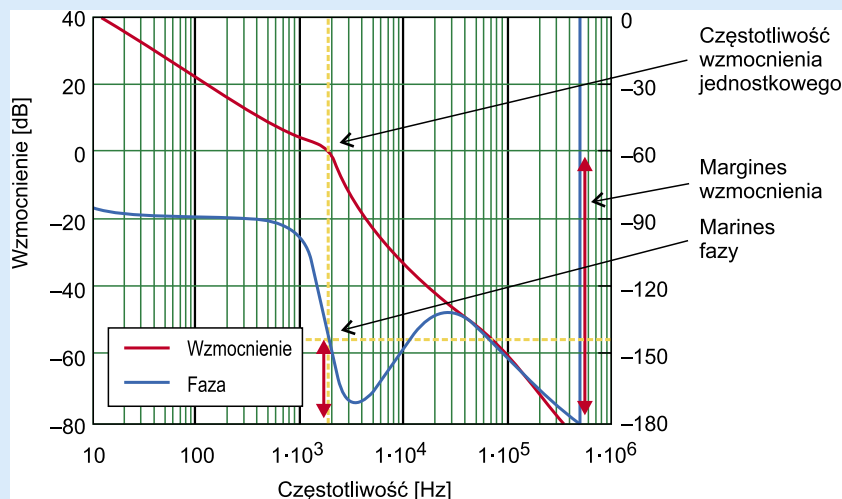
układu regulacji objętego pętlą sprzężenia zwrotnego musi być przeciwna do zmiany na wyjściu. Oznacza to, że jeśli napięcie wyjściowe próbuje wzrosnąć (lub zmaleć), pętla wspólnie z układem regulacji powinny odpowiednio wymusić powrót do wartości nominalnej.

Jeśli fala sinusoidalna (lub szum) zostaną doprowadzone do pętli, sygnał przechodzi przez pętlę i powraca przemnożony przez współczynnik wzmocnienia pętli, z pewnymi opóźnieniami fazowymi zależnymi od doprowadzonego sygnału. Przesunięcie fazowe jest zdefiniowane przez całkowitą liczbę opóźnień fazowych, odnoszonych do punktu początkowego -180° (ujemne sprzężenie zwrotne), które są wprowadzane do sygnału sprzężenia zwrotnego przechodzącego przez pętlę sprzężenia. Wzmocnienie pętli może być wyznaczone jako iloraz amplitudy sygnału, który przechodzi przez pętlę oraz amplitudy sygnału doprowadzonego do pętli: .

$$k_{loop} = 20 \cdot \log \frac{v_a}{v_b}$$

Załóżmy, że wprowadzamy sygnał sinusoidalny do pętli w szerokim zakresie częstotliwości. Sygnały składowe o niskich częstotliwościach wracają z podwyższoną amplitudą, a o wysokich są tłumione. Ich amplitudy i fazy można zarejestrować i wykonać wykres Bodego. Wiele współczesnych, skomputeryzowanych przyrządów pomiarowych wykonuje takie wykresy automatycznie. Przykład rysunku sporządzonego na podstawie wyników pomiarów pokazano na rysunku 16. Wykres mówi wiele o stabilności pętli i zachowaniu się układu w pewnych punktach szczególnego zainteresowania.

Częstotliwość, przy której wzmocnienie wynosi 0 dB (jest równe 1), oznaczana f_c : punkt, w którym wprowadzany sygnał powraca mając tę samą amplitudę (0 dB), jest nazywany częstotliwością wzmocnienia jednostkowego. Margines fazy ϕ_m jest zdefiniowany jako różnica pomiędzy całkowitym przesunięciem sygnału sprzężenia zwrotnego i punktem początkowym -180° , mierzonymi przy częstotliwości wzmocnienia jednostkowego. Margines wzmocnienia G_c jest sumą wzmocnienia ujemnego (tłumienia), przy którym całkowite przesunięcie fazowe wynosi 180° .



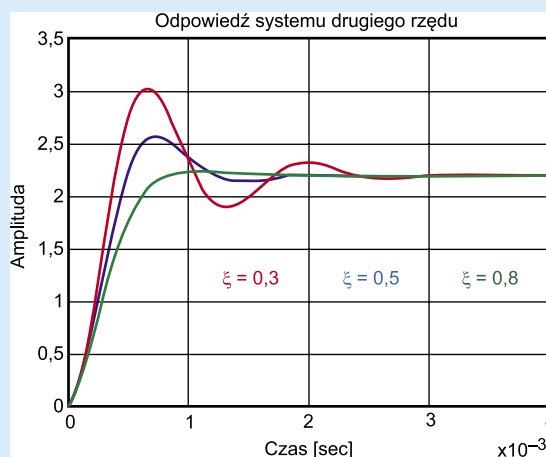
Rysunek 16. Przykładowy wykres Bodego umożliwiający określenie czy obwód będzie stabilny

Odpowiedni margines fazy jest wymagany, aby zapobiec oscylacjom. Odpowiedź krokowa, **rysunek 17**, może być zaobserwowana w systemach drugiego rzędu, w których współczynnik spadku wynosi $\xi \approx \phi_m/100$. Optymalny margines fazy wynosi 52° (niebieska linia). Mniejszy margines fazy doprowadza do zbyt tłumionej odpowiedzi systemu (czerwona linia) a wyższy margines błędu doprowadza do nad tłumionej odpowiedzi systemu (zielona linia).

Jak powiedziano wcześniej, są dwa główne parametry które dają odpowiedź na pytanie o stabilność systemu: margines fazy i margines wzmocnienia. W teorii nawet 20° margines fazy w najgorszych warunkach może być wystarczający do stabilnej pracy, aczkolwiek większy margines zapewni stabilność pętli w różnych warunkach. Kiedy wyspecyfikujemy i zmierzmy margines fazy również musimy rozważyć, jak bardzo zmieni się on w najgorszych warunkach pracy przetwornicy wynikających np. z obciążenia wyjścia oraz zmiany temperatury. Ważne jest również sprawdzenie minimalnej wartości marginesu przesunięcia fazowego poniżej częstotliwości wzmocnienia jednostkowego. Jeśli przesunięcie fazowe zbliży się do 0, system może oscylować, gdy wzmocnienie wzmacniacza błędu zwiększa się, na przykład przy wzroście obciążenia lub spadku napięcia wyjściowego.

Do określenia czy system jest stabilny, bardzo użyteczny jest analizator sieci. Opisane wyżej parametry można sprawdzić za pomocą sygnału wyjściowego analizatora, dzięki wprowadzeniu go do pętli kontrolnej, przy zadanym przemiataniu sygnałem o częstotliwości od kilku dziesiątek Hz do wyższej od częstotliwości kluczowania oraz pomiar sygnałów A i B, jako pokazano na rysunku 17.

Sinusoidalny sygnał wyjściowy z analizatora jest wprowadzany za pomocą transformatora. Aby zbyt nie zakłócić zamkniętej pętli systemu, sygnał wstrzykiwany do pętli ma amplitudę od kilku dziesiątych do setnych części mV. Rezystor jest dołączony równolegle do wyjścia transformatora. Dwa kanały wejściowe są dołączone do punktów połączeniowych na transformatorze



Rysunek 17. Pomiary zamkniętej pętli sprzężenia zwrotnego

w celu pomiarów wejścia pętli (kanał B) oraz wyjścia pętli (kanał A). Dzięki tej metodzie pomiaru, analizator może wykreślić wzmocnienie (iloraz kanału A i B) oraz fazę sygnałów w skali logarytmicznej, częstotliwościowo.

Transformator izolujący jest potrzebny do zapewnienia wprowadzenia „pływającego” napięcia (bez poziomu odniesienia na masie systemu) sinusoidalnego do pętli sprzężenia zwrotnego poprzez rezystor o niewielkiej rezystancji. Jest on włączony równolegle do wyjścia transformatora i szeregowo do dzielnika pętli, więc powinien mieć rezystancję o wiele mniejszą, niż rezystor sprzężenia zwrotnego, aby nie zmieniać napięcia wyjściowego DC. Transformator powinien mieć niewielką pojemność pomiędzy uzwojeniami pierwotnym i wtórnym oraz płaską charakterystykę częstotliwościową. Transformatory przeznaczone do tego celu są dostępne w handlu i zazwyczaj sprzedawane za kilkaset Euro. Niemniej, można samodzielnie wykonać podobny nawijając dwa uzwojenia na rdzeniu toroidalnym.

Michele Sclocchi
Application Engineer
Texas Instruments